



Colloque International Book of Abstracts Edité par K. Boukhetala et J.F Dupuy

Kamal Boukhetala, Jean-François Dupuy

► To cite this version:

Kamal Boukhetala, Jean-François Dupuy. Colloque International Book of Abstracts Edité par K. Boukhetala et J.F Dupuy. USTHB. Modélisation Stochastique et Statistique, Nov 2014, Alger, Algeria. 1, Bibliothèque National d'Alger, 2014, Modélisation Stochastique et Statistique, 978-9931-9222-0-9. hal-01086342

HAL Id: hal-01086342

<https://hal.science/hal-01086342>

Submitted on 24 Nov 2014

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.

Le colloque international Modélisation Stochastique et Statistique, dans sa troisième édition, est une rencontre scientifique de haut niveau regroupant chercheurs universitaires et experts praticiens du calcul stochastique et statistique et ses applications dans les domaines socio économique, industriel et environnemental. Par les différents thèmes abordés, cette rencontre constituera un cadre idéal pour débattre et discuter des développements récents. Un atelier sur le calcul numérique stochastique intensif et ses applications sera animé par des spécialistes en la matière. D'autre part, cette rencontre offrira une occasion de rapprochement entre académiciens et professionnels en matière d'échange d'idées et d'expériences dans le domaine de l'utilisation de l'outil stochastique et statistique en analyse, modélisation, simulation et prospection pour l'aide à la prise de décision.

ISBN 978-9931-9222-0-9



9 789931 922209

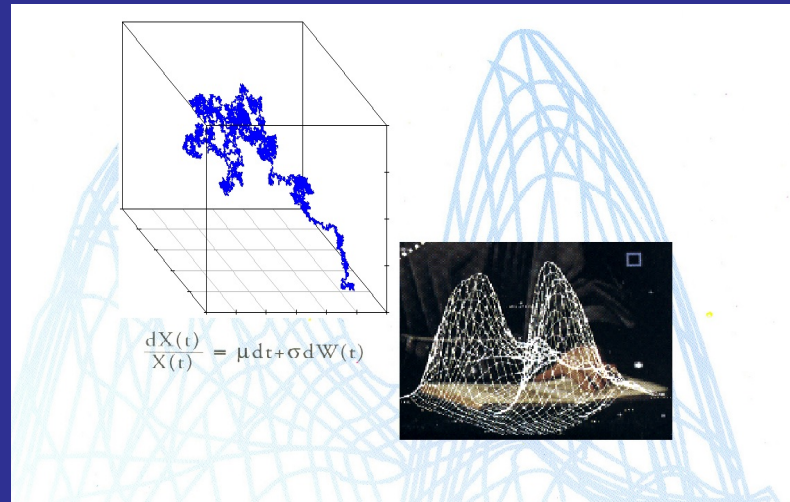
MSS'2014 — Book of Abstracts



MSS'2014

Colloque International Modélisation Stochastique et Statistique

Book of Abstracts



Edité par

Kamal Boukhetala
Jean-François Dupuy

MSS'2014

Colloque International

Modélisation Stochastique et Statistique

23-25 Novembre 2014, 3^{ème} Édition.

Faculté de Mathématiques, USTHB,
BP 32 El-Alia, 16111, Bab-Ezzouar,
Alger, Algérie.

Edité par

Kamal Boukhetala

Jean-François Dupuy

Professeur Kamal Boukhetala
Faculté de Mathématiques, USTHB
BP 32 El-Alia, 16111, Bab-Ezzouar
kboukhetala@usthb.dz
Algérie

Professeur Jean-François Dupuy
INSA de Rennes
20 Avenue des Buttes de Cöesmes
CS 70839, 35708 Rennes cedex 7
Jean-Francois.Dupuy@insa-rennes.fr
France

Preface

Le Département de Probabilités et Statistiques de la Faculté de Mathématiques (USTHB) en collaboration avec la Société Mathématique d'Algérie organise le Colloque International "Modélisation Stochastique et Statistique (MSS'2014)" (3ème édition), avec la participation du laboratoire de Mathématiques et Applications, UMR 7348 CNRS, Futuroscope, Université de Poitiers, (France), l'INSA de Rennes (France), l'Université de Toulouse (France), le laboratoire CREAM de l'Institut de Mathématiques Appliquées de l'université d'Angers (France) et la Faculté des Sciences et de Génie Université Laval, (Canada).

Objectifs du Colloque

Le colloque international Modélisation Stochastique et Statistique, dans sa troisième édition, est une rencontre scientifique de haut niveau regroupant chercheurs universitaires et experts praticiens du calcul stochastique et statistique et ses applications dans les domaines socio économique, industriel et environnemental. Par les différents thèmes abordés, cette rencontre constituera un cadre idéal pour débattre et discuter des développements récents. Un atelier sur le calcul numérique stochastique intensif et ses applications sera animé par des spécialistes en la matière. D'autre part, cette rencontre offrira une occasion de rapprochement entre académiciens et professionnels en matière d'échange d'idées et d'expériences dans le domaine de l'utilisation de l'outil stochastique et statistique en analyse, modélisation, simulation et prospection pour l'aide à la prise de décision.

Thèmes du Colloque

- Analyse des Données, Applications.
- Calcul Numérique Stochastique Intensif et Optimisation.
- Modèles Econométriques, Applications Financières et Actuarielles.
- Processus Aléatoire.
- Statistique Computationnelle, Simulation.

Remerciements

Les organisateurs de la conférence et les éditeurs de ce volume souhaitent d'abord exprimer leur sincère gratitude aux auteurs des documents dans le présent volume, pour leur précieuse contribution et pour leur participation enthousiaste qui a fait de cette conférence un forum utile pour l'échange d'idées. Et remercie aussi les sponsors qui ont bien voulu contribuer à l'aboutissement de ce colloque. Nous tenons à exprimer notre gratitude aux membres du comité scientifique, et à tous les membres du comité local d'organisation.

Alger,
Novembre, 2014

K. Boukhetala
J.F. Dupuy

Comité d'Organisation

Ameraoui A	(USTHB, DZ)	Kadem S	(USTHB, DZ)
Benmansour M	(USTHB, DZ)	Laouar A	(USTHB, DZ)
Bensalloua M	(USTHB, DZ)	Messaci Y	(USTHB, DZ)
El Mahdaoui A	(USTHB, DZ)	Tatachak A	(USTHB, DZ)
Fersi K	(ISG, Sousse, TN)	Yahi M	(USTHB, DZ)
Guessoum Z	(USTHB, DZ)	Yahi F	(USTHB, DZ)
Guidoum A.C	(USTHB, DZ)		

Comité Scientifique

Adjabi S	(U. Béjaia, DZ)	Gamboa F	(U. Toulouse, FR)
Adjaoud F	(U. Ottawa, CA)	Guerbyenne H	(USTHB, DZ)
Aider M	(USTHB, DZ)	Guessoum Z	(USTHB, DZ)
Aïssani A	(USTHB, DZ)	Hamadouche M	(U. Tizi-Ouzou, DZ)
Aïssani D	(U. Béjaia, DZ)	Hamdi F	(USTHB, DZ)
Baggag A	(U. Laval, CA)	Khaldi K	(U. Boumerdes, DZ)
Benchettah A	(USTHB, DZ)	Marion J.M	(U. Angers, FR)
Benhenni K	(U. Grenoble, FR)	Mezerdi B	(U. Biskra, DZ)
Beninel F	(U. Poitiers, FR)	Mohdeb Z	(U. Constantine, DZ)
Boukhetala K	(USTHB, DZ)	Mourid T	(U. Tlemcen, DZ)
Bousbaine A	(Nestlé, SW)	Necir A	(U. Biskra, DZ)
Boutabia H	(U. Annaba, DZ)	Ould-Said E	(U. Littoral, FR)
Chauvet P	(U. Angers, FR)	Oulidi A	(U. Angers, FR)
Coeurjolly J.F	(U. Grenoble, FR)	Rachedi M	(U. Grenoble, FR)
Dahmani A	(U. Béjaia, DZ)	Radjef M.S	(U. Béjaia, DZ)
D'Aubigny G	(U. Grenoble, FR)	Rahmani F.L	(U. Constantine, DZ)
Djedour M	(USTHB, DZ)	Rebbouh A	(USTHB, DZ)
Drouilhet R	(U. Grenoble, FR)	Remita R	(U. Annaba, DZ)
Dupuy J.F	(INSA Rennes, FR)	Sabre R	(U. Dijon, FR)
El Methni M	(U. Grenoble, FR)	Tatachak A	(USTHB, DZ)
Fellag H	(U. Tizi-Ouzou, DZ)	Yousfate A	(U. Sid Bel Abbes, DZ)
Ferraty F	(U. Toulouse, FR)		

Conférenciers Invités

Coquet F	(ENSAI Rennes, FR)	Gamboa F	(U. Toulouse, FR)
Dupuy J.F	(INSA Rennes, FR)	Hanafi M	(Oniris Nantes, FR)

Contents

Preface	i
A. Conférences Plénières	3
Sélection informative d'un échantillon: propriétés asymptotiques <i>F. Coquet</i>	5
Estimation de l'indice des valeurs extrêmes conditionnel en présence de cen- sure <i>A. Diop, J.F. Dupuy, P. Ndao</i>	7
La méthode <i>Pick and Freeze</i> de Sobol à la sauce COSTA BRAVA <i>F. Gamboa</i>	10
Algebra for multiblock data and models <i>M. Hanafi</i>	12
B. Communications Orales	13
Part I: Analyse des Données, Applications	15
The k nearest neighbors estimation of the conditional hazard function for functional data <i>M.K. Attouch, F.Z. Belabed</i>	16
Classification de l'EEG par SVMs pour la reconnaissance des tâches mentales <i>S.A. Belhadj, N. Benmoussat, M. Della Krachai</i>	19
Tests de rupture épidémique dans la variance <i>F. Graiche, D. Merabet, D. Hamadouche</i>	24
Part II: Processus Aléatoire, Modèles Économétriques, Applications Financières et Actuarielles	29
Étude asymptotique de solutions d'équations différentielles stochastiques <i>I. Abi-ayad, T. Mourid</i>	30
Sur la stabilité forte des systèmes d'attente classiques avec distribution générale inconnue <i>A. Bareche, D. Aïssani</i>	34
Solutions Stepanov presque périodique d'une classe d'équations différentielles stochastiques <i>F. Bedouhene, N. Challali, O. Mellah, P. Raynaud De Fitte</i>	38

Borne de stabilité pour la distribution stationnaire du temps d'attente: approche par interaction entre modèles de risque et systèmes d'attente <i>Z. Benouaret, A. Touazi, D. Aïssani</i>	42
Sur l'estimation des retards multiples dans une diffusion non linéaire <i>A. Benyahia, M. Dali-Korso</i>	46
Analyse stochastique du système d'attente $GI/M/s$ avec arrivées négatives <i>L. Berdjoudj, D. Aïssani</i>	50
Analyse stationnaire vaste du modèle d'attente à un seul serveur avec rappels classiques et feedback <i>M. Boualem, M. Cherfaoui, N. Djellab, D. Aïssani</i>	54
Estimation à noyau d'une matrice de transition inconnue du modèle d'attente $M/GI/1/N$ <i>M. Cherfaoui, M. Boualem, D. Aïssani, S. Adjabi</i>	58
Estimation d'un modèle GARCH asymétrique (AGARCH) en présence de données de haute fréquence <i>H. Guerbyenne, L. Messahli</i>	62
On probabilistic properties of a power periodic threshold GARCH model <i>H. Guerbyenne, A. Kessira</i>	66
Normalité asymptotique locale (condition LAN) pour un processus ARH(1) <i>N. Kara Terki, T. Mourid</i>	69
GSPNs modeling and strong stability analysis for the $M/M/1$ queueing system <i>O. Lekadir, L. Ikhlef, D. Aïssani</i>	73
La sensibilité des performances du système d'attente $M/G/1/N$ avec vacances <i>B. Takhedmit, K. Abbas</i>	77
Estimation non paramétrique des densités Heavy-tailed par la méthode du noyau <i>Y. Ziane, S. Adjabi, N. Zougab</i>	81
Part III: Statistique Computationnelle, Simulation	85
Modélisation spatiale de la formation des agglomérations dans la zone algéroise <i>S. Ait Amokhtar, N. EL Saadi</i>	86
Optimal approach on the heavy tailed distribution's mean estimation with right random censoring <i>A. Ameraoui, K. Boukhetala, J.F. Dupuy</i>	90
Bayesian inference about markov switching models using augmented data algorithm <i>I. Bouderbala, G. Saidi, H. Kherchi</i>	96
Application de la modélisation géostatistique dans l'industrie minière: cartographie de la variation spatiale des teneurs en P_2O_5 dans le gisement de phosphate de Kef Essennoun, Algérie <i>M. Dassamiour, H. Mezghache</i>	100
Comparative aspects of multistage designs for phase II clinical trials <i>Z. Djeridi, H. Merabet</i>	105
An R package for modeling stochastic and diffusion processes <i>A.C. Guidoum, K. Boukhetala</i>	108
Validation de modèles de régression linéaire dans le cas non régulier <i>Z. Mohdeb</i>	115
Méthode du noyau dans l'analyse de stabilité forte d'un modèle de risque <i>A. Touazi, Z. Benouaret, S. Adjabi, D. Aïssani</i>	119

Approche bayésienne dans l'estimation de la fonction de régression discrète par noyau binomial <i>N. Zougab, S. Adjabi, K. Célestin C</i>	123
---	-----

C. Posters 127

Comparaison de méthodes d'estimation pour un modèle exponentiel généralisé <i>K. Aidi, N. Seddik-Ameur</i>	129
Bayesian estimation of change point in generalized exponential distribution <i>A. Aknine, K. Nouali</i>	133
Inégalités stochastiques pour le système $M_1; M_2/G_1; G_2/1$ avec rappels à communication bidirectionnelle <i>L.M. Alem, M. Boualem, D. Aïssani</i>	139
Sélection de variables pour réseau bayésien <i>R. Benmakrelouf, W. Karouche, J. Rynkiewicz, M. Djedour</i>	143
Parameter estimation for time dependent drift to stochastic differential equations <i>M. Bensalloua, K. Boukhetala</i>	147
Estimation de la densité des retards dans un processus de type diffusion <i>W. Benyahia, T. Mourid</i>	151
Le temps d'occupation pour les k^{iemes} extrêmes $k \geq 1$ <i>S. Berrouane</i>	155
Asymptotics for distorsion risk measure under dependence <i>Y. Berkoun</i>	160
Some asymptotic normality result of k -Nearest Neighbour estimator of the conditional mode function for independent functional data <i>W. Bouabça, M.K. Attouch</i>	164
Modélisation stochastique de risque de dégradation par processus de diffusion <i>K. Boukhetala, N. Khellouf</i>	168
Analyse bayésienne de la file d'attente $M/M/1$ <i>H. Braham, L. Berdjoudj, M. Boualem</i>	173
Tests d'hypothèses dans un processus de diffusion non linéaire à retard <i>M. Dali-Korso</i>	176
Asymptotic normality of the kernel estimator of the conditional density under association for censored data <i>S. Dhiabi, O. Sadki</i>	181
Simulation de l'estimateur à noyau de la fonction de régression dans un modèle de troncature à gauche <i>F. Hamrani, Z. Guessoum</i>	185
Impact d'une modélisation à volatilité stochastique sur la prime d'un Call européen <i>S. Kadem, K. Boukhetala</i>	189
Analyse des données catégorielles <i>W. Karouche, R. Benmakrelouf, J. Rynkiewicz, M. Djedour</i>	196
Sur l'estimation des processus $SARFIMA-S\alpha S$ <i>R. Lessak</i>	200
Estimation du maximum du taux de hasard dans un modèle associé-tronqué <i>D. Madani, A. Tatachak</i>	204
Estimateur à noyau de la fonction de régression dans un modèle de censure à droite: applications <i>N. Menni, A. Tatachak</i>	208
Estimation récursive à noyau de la densité des variables faiblement dépendantes <i>K.A. Mezhoud, Z. Mohdeb, S. Louhichi</i>	212

Sur l'évolution des performances du protocole DHCP/IP <i>M. Rachedi, N. Djellab</i>	215
Procédures d'inférence Bayésienne séquentielle appliquées aux essais clin- iques <i>A. Zerari, H. Merabet</i>	219
Sur l'étude de l'effet d'allongement des vecteurs dans le cas linéaire <i>D. Zerdazi, A. Chibat</i>	227
Index des Auteurs	231

A. Conférences Plénières

Sélection informative d'un échantillon: propriétés asymptotiques

François Coquet*

*Crest-Ensai
fcoquet@ensai.fr

On s'intéresse dans cette exposé à la distribution d'un caractère Y dans un échantillon issu d'une population finie en situation de sélection informative, c'est-à-dire que la probabilité d'inclure l'individu i dans l'échantillon dépend de la valeur Y_i .

Formellement, on se place dans l'optique d'un modèle de superpopulation : on s'intéresse donc à une population taille finie N , et on suppose que les variables d'intérêt Y_i , $1 \leq i \leq n$ sont indépendantes et identiquement distribuées, de densité commune f , et de fonction de répartition F . Le mécanisme de sélection d'un échantillon dans cette population sera informatif dans le sens où la probabilité de tirer un échantillon donné dépend explicitement du vecteur des réponses de cet échantillon. Dans ce cas, les Y_i observées dans l'échantillon ne sont plus iid de densité f conditionnellement au fait que les individus i correspondants ont été sélectionnés. D'une manière générale, le mécanisme de sélection informative est susceptible d'induire de la dépendance entre les observations sélectionnées.

Un premier objectif de cet exposé est l'étude de cette dépendance sur la fonction de répartition empirique de l'échantillon. Nous donnons un cadre asymptotique et des conditions sur le mécanisme de sélection sous lesquelles la fonction de répartition empirique converge uniformément, dans L^2 et presque-sûrement, vers une version pondérée de la fonction de répartition F associée au modèle de superpopulation. Cela nous donne donc un analogue du théorème de Glivenko-Cantelli. Les hypothèses nécessaires sont raisonnables, et une série d'exemples motivés par des problèmes concrets d'échantillonnage et d'études observationnelles montre qu'elles sont vérifiables pour certains plans de sondages classiques. Nous comblons ainsi dans une certaine mesure le "trou" méthodologique entre une pratique d'échantillonnage sélectif courante dans les enquêtes, et une utilisation des observations, fréquente dans la littérature méthodologique, supposant implicitement qu'elles sont iid selon la loi de l'échantillon.

Dans le même cadre général, nous nous intéressons ensuite à l'inférence sur un paramètre en présence de covariables. Là encore, les méthodes classiques doivent être modifiées pour rendre compte de la sélection informative dont provient l'échantillon dont nous disposons. Nous appliquons cette idée à la technique du maximum de vraisemblance empirique, qui traite les observations comme si elles étaient indépendantes et identiquement distribuées selon une version pondérée de f . Dans ce cadre, nous prouvons la consistance et la normalité asymptotique de l'estimateur du maximum de vraisemblance empirique. Nos résultats restent valides en présence d'un paramètre de nuisance modélisant le caractère informatif du mécanisme de

Sélection informative d'un échantillon: propriétés asymptotiques

sélection. Là encore, les hypothèses formulées sont vérifiables pour des plans de sondages variés et classiques en méthodologie d'enquêtes. Enfin, on illustre la qualité asymptotique de nos estimateurs sur des résultats provenant de simulations.

(travail commun avec Daniel Bonnéry, University of Maryland-College Park, et F. Jay Breidt, Colorado State University)

Estimation de l'indice des valeurs extrêmes conditionnel en présence de censure

A. Diop*, J.-F. Dupuy**, P. Ndao*

* Université Gaston Berger, Saint-Louis, Sénégal
ndao.pathe@yahoo.fr, aliou.diop@ugb.edu.sn

** Institut de Recherche Mathématique de Rennes - INSA de Rennes
20 Avenue des Buttes de Coesmes, CS 70839, 35708 Rennes cedex 7, France
jean-francois.dupuy@insa-rennes.fr,
<http://dupuy.perso.math.cnrs.fr/>

Résumé. Dans ce travail, nous nous intéressons à l'estimation de l'indice des valeurs extrêmes d'une loi conditionnelle à queue lourde, lorsque les données observées sont censurées aléatoirement à droite. Une version de l'estimateur de Hill adaptée à ces données est proposée. Cette construction repose sur la combinaison d'une méthode à noyau et de la méthode dite de pondération inverse par la probabilité de censure. Nous établissons la normalité asymptotique de l'estimateur proposé et en évaluons les propriétés numériques au moyen de simulations. Nous proposons également un estimateur de type Weissman des quantiles extrêmes conditionnels et l'étudions au moyen de simulations. Enfin, nous illustrons les estimateurs proposés sur un jeu de données réelles.

1 Introduction

Le théorème de Fisher-Tippett-Gnedenko constitue l'un des résultats fondamentaux de la théorie des valeurs extrêmes. Notons (Y_n) une suite de variables aléatoires indépendantes et identiquement distribuées, de fonction de répartition F . S'il existe deux suites réelles $(a_n > 0)$, (b_n) et un réel γ tels que

$$P\left(\frac{\max_{1 \leq i \leq n} Y_i - b_n}{a_n} \leq y\right) \longrightarrow H_\gamma(y)$$

lorsque $n \rightarrow \infty$, alors

$$H_\gamma(y) = \begin{cases} \exp\left[-(1 + \gamma y)_+^{-1/\gamma}\right] & \text{si } \gamma \neq 0, \\ \exp(-e^{-y}) & \text{si } \gamma = 0, \end{cases}$$

où $y_+ = \max(0, y)$ (voir Beirlant et al. (2004) par exemple). La fonction de répartition H_γ est la fonction de répartition de la loi des valeurs extrêmes. Cette loi dépend du seul paramètre γ (appelé "indice des valeurs extrêmes de Y "), qui contrôle le comportement asymptotique de la fonction de survie $\bar{F} := 1 - F$ de Y . On définit trois cas selon le signe de γ :

Indice des valeurs extrêmes conditionnel en présence de censure

- si $\gamma > 0$, on dit que F appartient au domaine d'attraction de Fréchet, qui contient les lois dont la fonction de survie décroît comme une fonction puissance. On parle aussi de lois à queue lourde.
- si $\gamma < 0$, on dit que F appartient au domaine d'attraction de Weibull. Ce domaine contient les lois dont le point terminal à droite est fini.
- si $\gamma = 0$, on dit que F appartient au domaine d'attraction de Gumbel, qui regroupe les lois dont la fonction de survie est à décroissance exponentielle.

La connaissance de γ apparaît donc primordiale pour résoudre les problèmes liés aux valeurs extrêmes d'une variable aléatoire, tels que l'estimation des quantiles extrêmes.

Dans les applications, il est fréquent que la loi de la variable Y soit liée à une covariable X . A titre d'exemple, citons l'étude :

- du niveau de production maximal d'une entreprise en fonction du facteur travail disponible (Daouia et al., 2010),
- des températures extrêmes en fonction de paramètres topologiques (Ferrez et al., 2011),
- de la magnitude de séismes exceptionnels en fonction de leur position sur la Terre (Pisarenko et Sornette, 2003),
- de la durée de survie de patients atteints du VIH, en fonction de leur âge au moment de l'infection (Einmahl et al., 2008).

On s'intéresse alors à l'estimation de la fonction $x \mapsto \gamma(x)$, que l'on appelle "indice des valeurs extrêmes conditionnel de Y sachant $X = x$ ". De nombreux travaux ont été récemment consacrés à ce problème (voir Ndao et al. (2014b) et Ndao et al. (2014a) pour une revue de la littérature).

Dans les applications, il est également fréquent que les observations de Y soient censurées à droite. On observe alors seulement une borne inférieure pour Y . Cette situation est fréquente en statistique des durées de vie. Supposons par exemple que Y représente la durée entre le début d'un traitement médical et la guérison, chez des patient suivis dans un essai clinique. Si un patient quitte l'étude en cours (on dit alors qu'il est "perdu de vue") ou atteint la fin de l'étude sans avoir guéri, l'information disponible pour ce patient est : $Y > C$, où C est la durée d'observation du patient. Plusieurs auteurs se sont récemment intéressés à l'estimation de l'indice des valeurs extrêmes γ à partir de durées censurées (*e.g.*, Beirlant et al. (2007); Einmahl et al. (2008); Brahim et al. (2013); Worms et Worms (2014)).

Dans cet exposé, nous nous intéressons au problème de l'estimation de l'indice des valeurs extrêmes conditionnel de Y , lorsque les observations de Y sont censurées aléatoirement à droite. Nous considérons le cas où la loi de Y est à queue lourde. Nous proposons une version modifiée de l'estimateur de Hill adaptée à la présence de covariables et de censure. Nous établissons sa normalité asymptotique et en étudions les propriétés à distance finie au moyen de simulations. Nous proposons également un estimateur des quantiles extrêmes conditionnels en présence de censure. Nous évaluons ses propriétés à l'aide de simulations. Enfin, nous illustrons les estimateurs proposés sur un jeu de données réelles.

Références

Beirlant, J., Goegebeur, J. Y., Teugels, et J. Segers (2004). *Statistics of Extremes: Theory and Applications*. John Wiley & Sons, Ltd.

- Beirlant, J., A. Guillou, G. Dierckx, et A. Fils-Villetard (2007). Estimation of the extreme value index and extreme quantiles under random censoring. *Extremes* 10, 151–174.
- Brahimi, B., D. Meraghni, et A. Necir (2013). On the asymptotic normality of hill’s estimator of the tail index under random censoring. Technical report, arXiv: arXiv-1302.1666.
- Daouia, A., J.-P. Florens, et L. Simar (2010). Frontier estimation and extreme value theory. *Bernoulli* 16, 1039–1063.
- Einmahl, J., A. Fils-Villetard, et A. Guillou (2008). Statistics of extremes under random censoring. *Bernoulli* 14, 207–227.
- Ferrez, J., A. Davison, et M. Rebetez (2011). Extreme temperature analysis under forest cover compared to an open field. *Agricultural and Forest Meteorology* 151, 992–1001.
- Ndao, P., A. Diop, et J.-F. Dupuy (2014a). Nonparametric estimation of the conditional extreme-value index with random covariates and censoring. Technical report, HAL: hal-01056117.
- Ndao, P., A. Diop, et J.-F. Dupuy (2014b). Nonparametric estimation of the conditional tail index and extreme quantiles under random censoring. *Computational Statistics & Data Analysis* 79, 63–79.
- Pisarenko, V. et D. Sornette (2003). Characterization of the frequency of extreme earthquake events by the generalized pareto distribution. *Pure and Applied Geophysics* 160, 2343–2364.
- Worms, J. et R. Worms (2014). New estimators of the extreme value index under random right censoring, for heavy-tailed distributions. *Extremes* 17, 337–358.

Summary

Estimation of the extreme-value index of a heavy-tailed distribution is addressed when some random covariate information is available and the data are randomly right-censored. An inverse-probability-of-censoring-weighted kernel version of Hill’s estimator of the extreme value index is proposed and its asymptotic normality is established. Based on this, a Weissman-type estimator of conditional extreme quantiles is also constructed. A simulation study is conducted to assess the finite-sample behaviour of the proposed estimators. The methodology is illustrated on a real data set.

La méthode *Pick and Freeze* de Sobol à la sauce COSTA BRAVA

Fabrice Gamboa*

*Institut de Mathématiques de Toulouse
Université Paul Sabatier
118 Route de Narbonne
F-31062 Toulouse cedex
fabrice.gamboa@math.univ-toulouse.fr,
<http://www.math.univ-toulouse.fr/~gamboa>

Résumé. Nous présentons ici des résultats récents obtenus sur l'estimation des indices de Sobol dans le projet ANR COSTA BRAVA¹. Ces coefficients apparaissent naturellement dans la décomposition de Hoeffding.

Code numérique, décomposition de Hoeffding et indice de Sobol

On modélise fréquemment un code numérique scalaire, à entrées $X \in \mathbb{R}^d$ incertaines, par un modèle de régression du type

$$Y = F(X).$$

Si l'on suppose Y centré, de carré intégrale et que le vecteur aléatoire des entrées est constitué de composantes indépendantes alors la fonction F possède une décomposition de Hoeffding unique de la forme

$$\begin{aligned} F(X) &= F_1(X_1) + \cdots + F_d(X_d) + F_{1,2}(X_1, X_2) + \cdots + F_{d-1,d}(X_{d-1}, X_d) + \cdots + F_{1,\dots,d}(X) \\ &= \sum_{A \subset [1:d], A \neq \emptyset} F_A(X_A). \end{aligned}$$

Où, pour $A \subset [1:d]$, $A \neq \emptyset$, on note $X_A := (X_j)_{j \in A}$. Les fonctions (F_A) sont deux à deux orthogonales (on a en réalité l'orthogonalité de ces fonctions avec des espaces fonctionnels engendrés par les composantes de X). La décomposition de Hoeffding conduit à la dissociation de la variance de Y à partir des variances des variables aléatoires $F_A(X_A)$:

$$\text{Var}Y = \sum_{A \subset [1:d], A \neq \emptyset} \text{Var}F_A(X_A).$$

Si l'on définit l'indice de Sobol $S_A := \frac{\text{Var}F_A(X_A)}{\text{Var}Y}$, la relation précédente se réécrit

1. http://www.math.univ-toulouse.fr/COSTA_BRAVA/index.php

$$1 = \sum_{A \subset [1:d], A \neq \emptyset} S_A.$$

Dans cet exposé, nous discuterons de l'estimation statistique des indices S_A et de l'extension de ces indices à un cadre d'entrées dépendantes ou vectorielles (sortie Y à valeurs dans $\mathbb{R}^k$ ou dans un espace de Hilbert).

Summary

We will discuss recent results on Sobol' indices estimation obtained in the project sur ANR COSTA BRAVA². These coefficients naturally appear naturally in Hoeffding decomposition.

2. http://www.math.univ-toulouse.fr/COSTA_BRAVA/index.php

Algebra for multiblock data and models

Mohamed Hanafi*

*Unité de Sensométrie et de Chimiométrie,
ONIRIS, site de la Géraudière,
Nantes, France
mohamed.hanafi@oniris-nantes.fr

A challenging problem in multivariate statistics is the study of relationships between several blocks of data, also known as "multi-block data analysis".

After twenty years of involvement around this methodology, and nearly a century of literature in this field, I have arrived to the conclusion that : the framework provided by matrix algebra and matrix factorization principles are not enough to describe the whole richness and the complexity of block matrices, their manipulation, or the factorization strategies behind multiblock methods.

The present talk proposes a new paradigm for multi-block data analysis considered as a new form of thinking about multiblock data structures with new ways of manipulation and statistical treatment. This paradigm is based on both algebraic and geometric framework.

I extend the classical matrix algebra to the case of block matrices and I show some practical interests of this new framework in order to prototyping algorithms for multi-block data analysis. Likewise, I introduce some "multi-block matrix factorization" strategies as a natural extension of conventional matrix factorization. Geometric mechanisms of some multi-block methods are clarified and some differences between these methods will be elucidated

B. Communications Orales

Part I: Analyse des Données, Applications

The k nearest neighbors estimation of the conditional hazard function for functional data

Mohammed Kadi Attouch*, Fatima Zohra Belabed**

*Univ. Djillali Liabès Sidi Bel Abbès, Algeria
attou_kadi@yahoo.fr

**Univ. Djillali Liabès Sidi Bel Abbès, Algeria
zahira_bell@yahoo.fr

Résumé. Dans cet article, on va étudier l'estimateur nonparamétrique fonctionnel de la fonction de hasard conditionnelle en utilisant la méthode des k plus proches voisins (k -NN) avec une variable de réponse réelle. On donne la convergence presque complète de cet estimateur et la normalité asymptotique.

1 Introduction

The conditional hazard function remains an indispensable tool in survival analysis and many other fields (medicine, reliability or seismology).

2 Main results

Let $(X_i, Y_i)_{i=1, n}$ be an independent sequence identically distributed (i.i.d.) as (X, Y) which is a random pair valued in $\mathcal{E} \times \mathbf{R}$. Here $(\mathcal{E}, d)$ is a semi-metric space. $\mathcal{E}$ is not necessarily of a finite dimension, and we do not suppose the existence of a density for the functional random variable X .

Our goal, is to estimate the conditional hazard function defined by :

$$h^X(Y) = \frac{f^X(Y)}{1 - F^X(Y)}, \quad (1)$$

where :

$f^X(Y)$: is the conditional density function of Y given X .

$F^X(Y)$: is the conditional distribution function of Y given X .

For a fixed $x \in \mathcal{E}$, the k -NN kernel estimator of $h^x(Y = y)$ is given by :

$$\hat{h}_{k-NN}^x(Y = y) = \hat{h}^x(y) = \frac{\hat{f}^x(y)}{1 - \hat{F}^x(y)} \quad (2)$$

The k nearest neighbors estimation of the conditional hazard function for functional data

with :

$$\hat{F}^x(y) = \frac{\sum_{i=1}^n K(H_n^{-1}d(x, X_i))R(g_n^{-1}(y - Y_i))}{\sum_{i=1}^n K(H_n^{-1}d(x, X_i))}; \forall y \in \mathbf{R}. \quad (3)$$

And

$$\hat{f}^x(y) = \frac{\sum_{i=1}^n K(H_n^{-1}d(x, X_i))g_n^{-1}R'(g_n^{-1}(y - Y_i))}{\sum_{i=1}^n K(H_n^{-1}d(x, X_i))}. \quad (4)$$

where K is an asymmetrical kernel, H_n is a positive random variable, defined as follows :

$$H_n(x) = \min\{h \in \mathbf{R}^+ / \sum_{i=1}^n \mathbf{I}_{B(x,h)}(X_i) = k\} \quad (5)$$

with :

$$B(x, h) = \{x' \in \mathcal{E}; d(x, x') < h\}.$$

R is a distribution function and $(g_n)_{n \in \mathbf{N}}$ is a sequence of strictly positive real numbers (depending on n).

In parallel, in order to emphasize differences between the k -NN method and the traditional kernel approach, we define the estimator of the conditional hazard function FERRATY et al. (?) by :

$$\hat{h}_{kernel}^x(y) = \frac{\hat{f}_{kernel}^x(y)}{1 - \hat{F}_{kernel}^x(y)} \quad (6)$$

with :

$$\hat{f}_{kernel}^x(y) = \frac{\sum_{i=1}^n K(h_n^{-1}d(x, X_i))g_n^{-1}R'(g_n^{-1}(y - Y_i))}{\sum_{i=1}^n K(h_n^{-1}d(x, X_i))} \quad (7)$$

and :

$$\hat{F}_{kernel}^x(y) = \frac{\sum_{i=1}^n K(h_n^{-1}d(x, X_i))R(g_n^{-1}(y - Y_i))}{\sum_{i=1}^n K(h_n^{-1}d(x, X_i))}. \quad (8)$$

Where K is a kernel, R is a distribution function and $(h_n)_{n \in \mathbf{N}}, (g_n)_{n \in \mathbf{N}}$ are sequences of strictly positive numbers.

Before studying the k -NN estimator, we remind asymptotic properties of $\hat{h}_{kernel}^x$ defined by equation (6). FERRATY et al.(?), showed the almost complete convergence of this estimator.

Theorem 1 – In the "continuity type" model and under some assumptions already posed, we have :

$$\hat{h}_{kernel}^x(y) \longrightarrow h^x(y). \quad a.co.$$

– Under the "Lipschitz type" model, we have :

$$\hat{h}_{kernel}^x(y) - h^x(y) = O(h_n^\alpha) + O(g_n^\beta) + O\left(\sqrt{\frac{\log n}{n\varphi_x(h)}}\right).$$

Theorem 2 In the "continuity type" model and under the hypotheses already posed, suppose that $k = k_n$ is a sequence of positive real numbers such that $\frac{k_n}{n} \rightarrow 0$ and $\frac{\log n}{k_n} \rightarrow 0$, then we have :

$$\lim_{n \rightarrow \infty} \hat{h}^x(y) = h^x(y). \quad a.co.$$

Theorem 3 The hypotheses of 1 imply

$$\hat{h}^x(y) - h^x(y) = O\left(\varphi_x^{-1}\left(\frac{k}{n}\right)^\alpha\right) + O(g_n^\beta) + O\left(\sqrt{\frac{\log n}{k_n g_n}}\right). \quad a.co.$$

Theorem 4 Under the hypotheses already posed, then for any $x \in \mathcal{A}$, we have :

$$\left(\frac{k_n g_n}{\sigma_h^2(x, y)}\right)^{1/2} [\hat{h}^x(y) - h^x(y)] \xrightarrow{\mathcal{D}} \mathcal{N}(0, 1) \quad \text{as } n \rightarrow \infty \quad (9)$$

where

$$\sigma_h^2(x, y) = \frac{\alpha_2 h^x(y)}{\alpha_1^2 (1 - F^x(y))} \quad (\text{with : } \alpha_j = K^j(1) - \int_0^1 (K^j)'(s) \alpha_x(s) ds \text{ for } j = 1, 2) \quad (10)$$

$$\mathcal{A} = \{x \in \mathcal{E}, f^x(y)[1 - F^x(y)] \neq 0\}$$

$\xrightarrow{\mathcal{D}}$ means the convergence in distribution.

Références

Laksaci, A. et B. Mechab (2012). Conditional hazard estimate for functional random fields.

Summary

In this paper, we study the nonparametric estimator of the conditional hazard function using the k nearest neighbors (k -NN) estimation method for a scalar response variable given a random variable taking values in a semi-metric space. We give the almost complete convergence (its corresponding rate) of this estimator and we establish the asymptotic normality. Then the effectiveness of this method is exhibited by a comparison with the kernel method estimation given in FERRATY et al.(?) and LAKSACI and MECHAB (Laksaci et Mechab (2012)) in both cases simulated data and real data.

Classification de l'EEG par SVMs pour la reconnaissance des tâches mentales

Sid Ahmed BELHADJ*, Nawal BENMOUSSAT**
, Mohamed DELLA KRACHAI***

Université des Sciences et de la Technologie d'Oran (USTO-MB),
B.P 1505 El M'Naouer, Algérie

* sidahmed.belhadj.univ@gmail.com

** nawalbb@yahoo.com

*** mohamed.dellakrachai@univ-usto.dz

Résumé. Ce papier présente une approche pour la classification des données électroencéphalogrammes (EEG) dans le contexte des Interfaces Cerveau-Ordinateur (ICOs). L'algorithme utilisé est basée sur l'apprentissage automatique en exploitant les performances d'un mélange de classifieurs SVMs. Chaque classifieur est entraîné sur une partie des données d'apprentissage dérivant d'une même session d'acquisition d'ICO, adapté à une méthode permettant la sélection automatique des canaux EEG. Ainsi, cette technique vise à pallier le problème d'"apprentissage-intersession" dû à la variabilité des données EEG durant les sessions d'acquisition du même utilisateur. L'évaluation des performances de cet algorithme est appliquée sur des données EEG provenant de l'institut Wadsworth Center.

1 Introduction

Certains patients atteints par un accident vasculaire cérébral grave restent dans un état de paralysie musculaire complète. Un handicap moteur très sévère qui n'atteint pas ses facultés mentales (le cerveau), mais lui interdit toute communication avec son entourage Farwell et Donchin (1988). Ceci a poussé les chercheurs à trouver une approche (système) pour contrôler le milieu environnant seulement par la pensée. Actuellement, les interfaces Cerveau-Ordinateur (ICOs) sont l'approche qui semble la plus prometteuse pour pallier cet handicap). Par ailleurs, l'ICO peut être décrite comme un système de communication et de contrôle direct s'appuyant uniquement sur l'activité cérébrale sans aucune intervention musculaire entre une personne et un système électrique ou mécanique (un ordinateur, un fauteuil roulant...) (Wolpaw et al., 2002).

Dans cet article, nous intéressons aux ICOs destinées aux applications de la communication palliative permettant d'épeler un mot par la pensée. Le principe de ce type d'interface est basé sur l'apparition d'un potentiel évoqué dans l'EEG provoqué lors d'un stimuli visuel auquel le sujet supposé répondre. Cependant, le problème consiste à la détection de ces potentiels dans un signal EEG.

Ainsi, dans cette étude, nous présentons une méthodologie pour la reconnaissance de ces signaux, sachant que les signaux EEG sont connus par leur faible rapport signal sur bruit (RSB) et leur variabilité durant le temps. Pour cela, nous avons recours à des méthodes d'apprentissage automatique, en exploitant les performances d'un mélange de classifieurs SVMs pour traiter ces problèmes.

2 Méthodologie

Avant de développer le système permettant la connaissance des potentiels évoqués noyés dans les signaux EEG, nous allons décrire succinctement l'Interface Cerveau-Ordinateur utilisée, ainsi la base de données des signaux EEG.

2.1 L'Interface Cerveau-Ordinateur

L'Interface Cerveau-Ordinateur utilisée dans cette étude est de type BCI synchrone "P300Speller". Cette interface est la communication la plus répandue depuis 1988, originalement développée par Farwell et Donchin Farwell et Donchin (1988) qui permet d'épeler des mots par la saisie d'un texte sur un écran. Ce traitement remplace le clavier traditionnel par un autre virtuel (voir figure 1). L'interface est basée sur le protocole suivant : il est demandé à l'utilisateur de fixer son regard sur une matrice $[6 * 6]$ soit 36 caractères. Plusieurs stimulus visuels sont alors provoqués par l'illumination d'une manière aléatoire des 12 lignes et colonnes. L'utilisateur doit alors compter le nombre de fois d'intensification de la ligne ou la colonne du caractère qu'il souhaite épeler. Cette tâche cognitive de comptage entraîne la génération d'un potentiel évoqué de type P300 détecté par l'ICO, ce qui permet à cette dernière de sélectionner la ligne et la colonne contenant le symbole désiré.

L'objectif de cette étude est de discriminer parmi les 12 ondes cérébrales la présence ou non de P300 dans les signaux EEG.

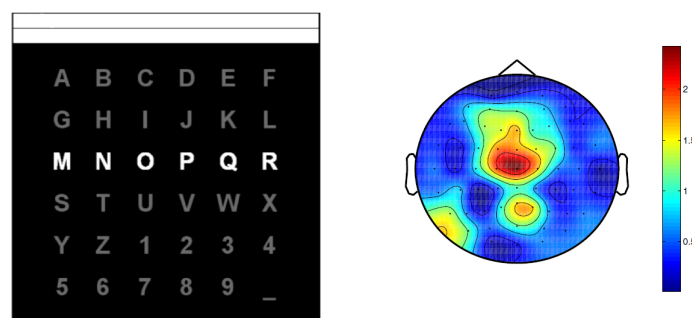


FIG. 1 – Écran P300 Speller qui permet de saisir un texte par la pensée (gauche). L'emplacement et la distribution des canaux EEG à la réponse d'un stimuli (droite).

2.2 Prétraitement de données EEG

Les ondes cérébrales EEG utilisées dans cet article sont issues d'expérience d'épellation de mot réalisée au sein de l'institut Wadsworth center (Schalk et al., 2004). Ces ondes ont été acquises grâce à un scalp de 64 électrodes échantillonnées à 240 Hz correspondant à l'épellation par le même sujet, pendant 3 sessions d'acquisition, respectivement 5, 6 et 8 mots, avec bien entendu de répéter l'épellation 15 fois du même caractère avant de passer au suivant.

Dans le fait que le potentiel évoqué qu'on le cherche apparaît après 300 ms d'un stimulus, bien que la fenêtre temporelle qui nous intéresse correspond généralement autour 600 ms. Ainsi, pour améliorer le RSB un filtre passe-bande d'ordre 8 de type Tchebychev (Labbé et al., 2010) est appliqué avec une bande passante [0.1-20] Hz, la figure 2 illustre la distribution de données d'un exemple après ce prétraitement.

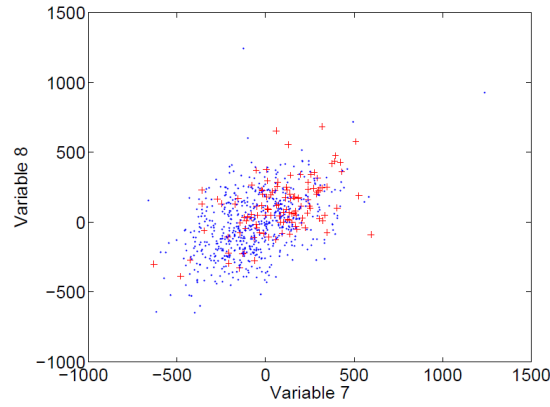


FIG. 2 – Distribution de données avec P300 en ("+" rouge) et sans P300 en ("." bleu) en fonction des variable 7 et 8.

2.3 Discrimination de données EEG

Après l'étape précédente de prétraitement, les données peuvent maintenant s'introduire dans l'entrée du classifieur, notre choix s'est porté sur un classifieur discriminant de type SVM Vapnik (1995). Cette technique est celle introduite par (Kaper et al., 2004). Pour traiter le problème d'"apprentissage-intersession" dû à la variabilité des données cérébrales EEG durant le temps (voir figure 3(a))

La méthode que nous avons utilisée consiste à attribuer une fonction discriminante pour chaque session d'acquisition du même utilisateur. Ainsi, les différents classifieurs provenant des différentes sessions d'acquisition seront fusionnés pour former finalement un classifieur qui va s'adapter à discriminer les données cérébraux du l'utilisateur concerné pour décider la classe finale des donnée.

En plus, pour discriminer la colonne ou la ligne correspondante au caractère désiré nous avons créé une fonction score $S_{Col(i)}$ et $S_{lin(i)}$ qui correspond respectivement au score de la

i -ème colonne et ligne de la matrice d'épellation.

$$S_{Col(i)} = S_{Col(i)} + \sum_k f_k(x)$$

avec $f_k(x)$ représente la fonction de décision liée à un mot de la base d'entraînement avec $k \in [1, 11]$ qui correspond ainsi à 11 mots de la base d'entraînement.

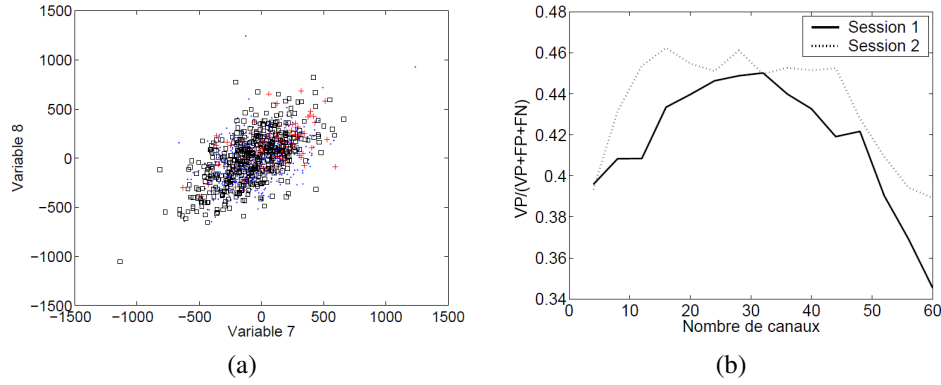


FIG. 3 – (a) Distribution de données avec P300 en ("+" rouge) et sans P300 en ("." bleu) en fonction des variable 7 et 8. D'autres données sans P300 d'une autre session d'acquisitions d'ICO sont également illustrées en carré noir: (b) Variations du critère d'évaluation C_{sel} pour deux sessions d'acquisition différentes d'ICOS.

2.4 Résultat

Pour l'évaluation de notre classifieur, nous avons calculé le critère de sélection C_{sel} (voir figure 3(b)). Ainsi, nous constatons la variations de ce critère C_{sel} pour deux sessions d'acquisition différentes d'ICOS en fonction du nombre de canaux pendant la phase d'élimination de canal. Dans le premier cas, le critère est maximisé en utilisant 32 canaux et 16 dans le deuxième cas.

$$C_{sel} = \frac{VP}{VP + FP + FN}$$

Avec VP , FP et FN sont respectivement le taux de vrai positifs, faux positifs et faux négatifs.

Références

Farwell, L. A. et E. Donchin (1988). Talking off the top of your head : toward a mental prosthesis utilizing event-related brain potentials. *Electroencephalography and clinical Neurophysiology* 70(6), 510–523.

- Kaper, M., P. Meinicke, U. Grossekhoefer, T. Lingner, et H. Ritter (2004). Bci competition 2003-data set iib : support vector machines for the p300 speller paradigm. *Biomedical Engineering, IEEE Transactions on* 51(6), 1073–1076.
- Labbé, B., X. Tian, et A. Rakotomamonjy (2010). Mlsp competition, 2010 : Description of third place method. In *Machine Learning for Signal Processing (MLSP), 2010 IEEE International Workshop on*, pp. 116–117. IEEE.
- Schalk, G., D. J. McFarland, T. Hinterberger, N. Birbaumer, et J. R. Wolpaw (2004). Bci2000 : a general-purpose brain-computer interface (bci) system. *Biomedical Engineering, IEEE Transactions on* 51(6), 1034–1043.
- Vapnik, V. N. (1995). *The Nature of Statistical Learning Theory*. New York, NY, USA : Springer-Verlag New York, Inc.
- Wolpaw, J. R., N. Birbaumer, D. J. McFarland, G. Pfurtscheller, et T. M. Vaughan (2002). Brain–computer interfaces for communication and control. *Clinical neurophysiology* 113(6), 767–791.

Summary

In this paper, we present an approach to analyze EEG-based Brain-Computer Interface (BCI) using the machine learning technique mixed Support Vector Machines (SVMs). Each SVM classifier is trained on a part of the training set from the same session BCI data acquisition, combined with a channel selection algorithm. Hence, the mixture of SVM able to copes the problem of "intersession-variability" which is caused by non-stationarity of brain patterns over time. Performances of this algorithm have been evaluated on EEG data provided by Wadsworth Center.

Tests de rupture épidémique dans la variance

Farid Graiche, Dalila Merabet et Djamel Hamadouche

Laboratoire de Mathématiques Pures et appliquées
Faculté des Sciences
Université Mouloud Mammeri de Tizi-Ouzou
Algérie

email : faridgraiche@yahoo.fr, d_merabet@yahoo.fr and djhamad@mail.ummo.dz

Résumé. Pour la détection de ruptures épidémiques dans la moyenne d'un échantillon de taille n , Račkauskas et Suquet (2004) ont introduit des nouvelles statistiques basées sur des sommes partielles. Sous l'hypothèse nulle, on obtient leurs lois limites pour des variables aléatoires indépendantes non stationnaires et dépendantes stationnaires (α -mélangeantes) et on donne une application pour un test de rupture de variances.

1 Introduction

Soit $(X_1, X_2, \dots, X_n)$ un échantillon de variables aléatoires de moyennes $m_1, m_2, \dots, m_n$ respectivement. On veut tester l'hypothèse nulle

$(H_0) : m_1 = m_2 = \dots = m_n,$

contre l'hypothèse alternative

$(H_A) : \exists 1 < k^* < m^* < n$ tel que

$m_1 = \dots = m_{k^*} = m_{m^*+1} = \dots = m_n, m_{k^*+1} = \dots = m_{m^*} \text{ et } m_{k^*} \neq m_{k^*+1}.$

On note ξ_n le processus de sommes partielles basé sur l'échantillon $(X_1, X_2, \dots, X_n)$. Donsker et Prokhorov ont prouvé que si les X_i sont i.i.d. de variance finie, $\sigma^{-1}n^{-\frac{1}{2}}\xi_n$ converge dans $C[0, 1]$ vers le mouvement brownien W . Ceci implique que $g(\sigma^{-1}n^{-\frac{1}{2}}\xi_n)$ converge en loi vers $g(W)$ pour une fonctionnelle continue g . Ceci procure de multiples applications statistiques. Un exemple classique est la détection de rupture épidémique de l'espérance d'un échantillon (hypothèses H_0 et H_A). Levine et Kline (1985) ont proposé la statistique

$$Q = \max_{1 \leq i < j \leq n} |S(j) - S(i) - S(n)(\frac{j}{n} - \frac{i}{n})|, \quad (1)$$

Tests de rupture épidémique dans la variance

où $S(k) = \sum_{i \leq k} X_i$. Si on note B le pont brownien, alors sous l'hypothèse nulle d'un échantillon i.i.d. de variance 1 et d'espérance fixé, on a $n^{-\frac{1}{2}}Q$ converge en loi vers $\sup_{0 < t < 1} |B(t)|$. Sous l'hypothèse alternative, on peut utiliser la statistique proposé par Račkauskas et Suquet (2004)

$$UI(n, \alpha) = \max_{1 \leq i < j \leq n} \frac{|S(j) - S(i) - S(n)(\frac{j-i}{n})|}{(\frac{j-i}{n})^\alpha}. \quad (2)$$

La fonctionnelle correspondante est

$$h(x) = \sup_{0 < |t-s| < 1} \frac{|x(t) - x(s)|}{|t-s|^\alpha}.$$

h n'est pas continue sur $C[0, 1]$, mais elle est continue sur l'espace de Hölder H_α^0 des fonctions $x : [0, 1] \rightarrow \mathbb{R}$ tel que $|x(t+h) - x(t)| = o(h^\alpha)$, uniformément en t . Lamperti a prouvé que si $E|X_1|^p < \infty$, où $p = (\frac{1}{2} - \alpha)^{-1}$, alors $n^{-\frac{1}{2}}\sigma^{-1}\xi_n$ converge dans H_α^0 pour tout $0 < \alpha < \frac{1}{2}$. Ceci nous permet d'avoir la limite de $h(n^{-\frac{1}{2}}\sigma^{-1}\xi_n) = n^{-\frac{1}{2}}\sigma^{-1}UI(n, \alpha)$. Une autre statistique proposée par Račkauskas et Suquet (2004) est

$$DI(n, \alpha) = \max_{1 \leq 2^j \alpha \leq n} \frac{1}{2^{-j\alpha}} \max_{r \in D_j} \left| S(nr) - \frac{1}{2}S(nr^+) - \frac{1}{2}S(nr^-) \right|, \quad (3)$$

où D_j est l'ensemble des dyadiques sur $[0, 1]$ de niveau j .

L'intérêt des statistiques hölderiennes est dans la détection de courtes épidémies. La statistique Q détecte seulement des épidémies de longueur l^* de l'ordre d'au moins $n^{\frac{1}{2}}$, alors que $UI(n, \alpha)$ détecte des épidémies de longueur l^* d'ordre n^δ , $0 < \delta < \frac{1}{2}$.

Notre objectif dans ce travail est de proposer des statistiques de type DI basées sur des variables aléatoires indépendantes ou dépendantes, d'étudier leurs comportements asymptotiques et donner une application pour des tests de rupture de variances.

2 Lois limites de $UI(n, \alpha)$ et $DI(n, \alpha)$

On considère les variables aléatoires

$$UI(\alpha) = \sup_{0 < t-s < 1} \frac{|B(t) - B(s)|}{[(t-s)(1-t+s)]^\alpha}, \quad (4)$$

et

$$DI(\alpha) = \sup_{j \geq 1} \frac{1}{2^{-j\alpha}} \max_{r \in D_j} \left| W(r) - \frac{1}{2}W(r^+) - \frac{1}{2}W(r^-) \right|. \quad (5)$$

La loi de $DI(\alpha)$ est connue contrairement à celle de $UI(\alpha)$ (cf Račkauskas et Suquet (2004)).

3 Convergence et consistance de $DI(n, \alpha)$

Cas indépendant non stationnaire

Soit $(X_i)_{i \geq 1}$ une suite de variables aléatoires de variances $(\sigma_i^2)_{i \geq 1}$ et soit $s_n^2 = \sum_{i=1}^n \sigma_i^2$.

Sous l'hypothèse

(H'_0) : Les variables aléatoires X_i sont indépendantes de moyenne 0. On a le résultat suivant.

Théorème 1 Supposons que

$$\lim_{n \rightarrow +\infty} \frac{\max_{1 \leq k \leq k_n} \sigma_k^2}{s_n^2} = 0$$

et pour $q > \frac{1}{1/2 - \alpha}$

$$\lim_{n \rightarrow +\infty} s_n^{q(2\alpha-1)} \sum_{k=1}^n \sigma_k^{-2\alpha q} \mathbb{E} |X_k|^q = 0.$$

Alors

$$s_n^{-1} DI(n, \alpha) \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} DI(\alpha), \forall 0 < \alpha < \frac{1}{2} - \frac{1}{q}.$$

On a montré que le test basé sur la statistique $DI(n, \alpha)$ est consistant.

Cas dépendant stationnaire (α -mélangeante)

Sous l'hypothèse

(H'_0) : La suite $(X_n)_n$ est strictement stationnaire, de moyenne 0 et α -mélangeante avec un coefficient de mélange α_n . On a le résultat suivant

Théorème 2 : Supposons qu'il existe $\gamma > 2$ et $\epsilon > 0$ tel que

$E |X_1|^{\gamma+\epsilon} < +\infty, P(|X_1| > t) = o(t^{-p}),$ où $p = (\frac{1}{2} - \alpha)^{-1}$ et

$\sum_{n=1}^{\infty} (n+1)^{\frac{\gamma}{2}-1} (\alpha_n)^{\frac{\epsilon}{\gamma+\epsilon}} < +\infty.$ Alors pour tout $\alpha < \frac{1}{2} - \frac{1}{\gamma}$, on a

$$\sigma^{-1} n^{-\frac{1}{2}} DI(n, \alpha) \xrightarrow[n \rightarrow \infty]{\mathcal{L}} DI(\alpha),$$

où $\sigma^2 = E(X_1^2) + 2 \sum_{j=2}^{\infty} cov(X_1, X_j) < +\infty.$

Le test basé sur $DI(n, \alpha)$ est consistant.

4 Exemple d'application pour un test de rupture de variance

Soit $(X_1, X_2, \dots, X_n)$ un échantillon de variables aléatoires de variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$ respectivement. On veut tester l'hypothèse nulle

$(H_0) : \sigma_1^2 = \sigma_2^2 = \dots = \sigma_n^2 = \tilde{\sigma}^2,$

Tests de rupture épidémique dans la variance

contre l'hypothèse alternative

$(H_A) : \exists 1 < k^* < m^* < n$ tels que
 $\sigma_1^2 = \dots = \sigma_{k^*}^2 = \sigma_{m^*+1}^2 = \dots = \sigma_n^2, \sigma_{k^*+1}^2 = \dots = \sigma_{m^*}^2$ et $\sigma_{k^*}^2 \neq \sigma_{k^*+1}^2$.

On définit alors

$$V_n(s) = \sum_{1 \leq k \leq ns} (X_k^2 - \tilde{\sigma}^2), s \in [0, 1].$$

On considère la statistique de test :

$$\nu(n, \alpha) = \max_{1 \leq j \leq \log n} \frac{1}{2^{-j\alpha}} \max_{r \in D_j} |\lambda_r(V_n)|.$$

Dans les deux cas précédents, on a montré sous quelques conditions de moments la convergence de $\nu(n, \alpha)$ vers $DI(\alpha)$ et que le test basé sur $\nu(n, \alpha)$ est consistant.

Enfin, on a présenté un exemple numérique de rupture de variance sur des données simulées.

Références

- Graiche F., Merabet D. and Hamadouche D. (2014), Testing epidemic change in the variance. *A paraître dans Communications in Statistics-Theory and Methods*.
- Hamadouche D. (2000), Invariance principles in Hölder spaces. *Portugaliae Mathematica*, 57, 127–151.
- Levine, B. and Kline, J. (1985), CUSUM tests of homogeneity, *Statistics in medicine* 4 , 469-488.
- Račkauskas A., Suquet Ch. (2004), Hölder norm test statistics for epidemic change. *J. Planning and Inference*. 126,495-520.
- Račkauskas A., Suquet Ch. (2006), Testing epidemic changes of infinite dimensional parameters. *Statistical Inference for Stochastic processes*. 9,111-134.

Summary

To detect epidemic change point in the mean, Račkauskas and Suquet (2004) introduced a new statistics based on partial sums. We obtain their limit distributions under the null hypothesis for independent not identically distributed or α -mixing random variables and we present application for change point in variates.

Part II: Processus Aléatoire, Modèles Économétriques, Applications Financières et Actuariales

Étude Asymptotique de Solutions d'Équations Différentielles Stochastiques

ABI-AYAD Ilham*, MOURID Tahar**

* Université Abou Bekr Belkaïd, Tlemcen
ilham.abiayad@yahoo.fr, ** Université Abou Bekr Belkaïd, Tlemcen
t_mourid@univ-tlemcen.dz

Résumé. Nous étudions le comportement asymptotique de solution d'équation différentielle stochastique avec drift positif en se basant sur un article de G.Keller, G. Kersting et U. Rösler (1984). Sous certaines conditions, la solution de l'EDS possède un comportement déterministe donné par une équation différentielle ordinaire. Nous illustrons cette étude par des simulations numériques par comparaison des trajectoires selon certains critères.

1 Introduction

Soit $(X_t, t \geq 0)$ défini par l'équation différentielle stochastique suivante :

$$dX_t = g(X_t)dt + \sigma_1(X_t)dW_t, \quad X_0 = 1, \quad (1)$$

où les fonctions g et σ_1 sont strictement positives et (W_t) est un mouvement brownien standard. Notons par $\mu(t)$ la solution de l'équation différentielle ordinaire donnée par

$$d\mu_t = g(\mu_t)dt, \quad \mu_0 = 1, \quad (2)$$

Dans l'article de G.Keller, G. Kersting et U. Rösler (1984) les auteurs ont donné des conditions sur le drift g et le coefficient de diffusion σ_1 pour que le rapport $X_t/\mu_t \rightarrow 1$. Soit $(Y_t, t \geq 0)$ le processus transformé défini par

$$Y_t = G(X_t), \quad (3)$$

où

$$G(t) = \int_1^t \frac{ds}{g(s)}, \quad (4)$$

Dans cette étude le coefficient de diffusion σ_1 est donné par $\sigma_1(X_t) = g(X_t)\sigma(X_t)$ avec $\sigma(X_t) > 0$. On introduit les fonctions suivantes pour $t > 0$:

Pour une fonction f on note

$$\begin{aligned} \tilde{f}(t) &= f(\mu(t)) = f(\mu_t) \\ h(t) &= \frac{1}{2}g'(t)\sigma^2(t), \end{aligned} \quad (5)$$

Étude Asymptotique de Solutions d'Équations Différentielles Stochastiques

$$\psi(t) = \int_1^t \frac{\sigma^2(s)}{g(s)} ds, \quad \tilde{\psi}(t) = \int_0^t \tilde{\sigma}^2(s) ds, \quad (6)$$

Soit $\alpha_t, t \geq 0$, la solution de l'équation différentielle ordinaire

$$\alpha'(t) = 1 - \tilde{h}(\alpha_t) \quad (7)$$

de valeur initiale α_0 . Posons enfin

$$Z_t = \int_0^t \sigma(X_s) dW_s = \int_0^t \tilde{\sigma}(Y_s) dW_s \quad (8)$$

Nous imposons les conditions suivantes : (A1) $g : R^+ \rightarrow R^+$ est strictement positive, deux fois continuellement dérivable et

$$G(\infty) = \int_1^\infty \frac{ds}{g(s)} = \infty$$

(A2) $h(t) \rightarrow 0$, quand $t \rightarrow \infty$

(A3) $\sigma : R^+ \rightarrow R^+$ est strictement positive, deux fois continuellement dérivable et

$$\int_0^\infty t^{-2} \tilde{\sigma}^2(t) dt < \infty$$

(A4) Les fonctions $g, g', \tilde{\sigma}^2$ et $\tilde{h}$ sont soit concaves ou convexes pour t assez grand.

En utilisant la formule d'Itô, Y_t vérifie l'équation suivante

$$dY_t = (1 - \tilde{h}(Y_t)) dt + \tilde{\sigma}(Y_t) dW_t, \quad Y_0 = 0 \quad (9)$$

ou encore sous la forme intégrale

$$Y_t = t - \int_0^t \tilde{h}(Y_s) ds + Z_t. \quad (10)$$

Nous avons le théorème suivant sur le comportement asymptotique du processus (Y_t) :

Théorème 1.1 *Supposons (A1)-(A4).*

Alors sur $\{X_t \rightarrow \infty\}$ on a $Y_t/t \rightarrow_{t \rightarrow \infty} 1$ p.s.

De plus

- i) *Si $\psi(\infty) < \infty$, alors sur $\{X_t \rightarrow \infty\}$ on a $Y_t - \alpha_t$ converge p.s. et la loi limite conditionnelle possède une densité unimodale qui est strictement positive partout.*
- ii) *Si $\psi(\infty) = \infty$, alors p.s. sur $\{X_t \rightarrow \infty\}$*

$$Y_t = \alpha_t + Z_t + o(\tilde{\psi}_t^{1/2})$$

De plus, il existe un mouvement brownien $B(t)$ tel que p.s. sur $\{X_t \rightarrow \infty\}$

$$Z_t = B(\tilde{\psi}_t) + o(\tilde{\psi}_t) \quad (11)$$

et pour tout réel b , on a

$$P(Z_t \leq b \tilde{\psi}_t^{1/2} | X_t \rightarrow_{t \rightarrow \infty} \infty) \rightarrow P(B(1) \leq b) \quad (12)$$

Sur le comportement asymptotique du processus (X_t) nous avons les deux résultats suivant :

Théorème 1.2 *Supposons (A1)-(A4).*

Nous avons les équivalences suivantes :

- i) $g(t)/t = o(\psi^{-1/2}(t))$
- ii) sur $\{X_t \rightarrow \infty\}$ on a en probabilité $X_t/\mu_t \rightarrow_{t \rightarrow \infty} 1$
- iii) il existe une fonction positive β_t telle que sur $\{X_t \rightarrow \infty\}$ on a en probabilité $X_t/\beta_t \rightarrow_{t \rightarrow \infty} 1$

Théorème 1.3 *Supposons (A1)-(A4). Nous avons les équivalences suivantes :*

- i) Il existe $0 \leq c < \infty$ tel que $\psi^{1/2}(t)g(t)/t \rightarrow_{t \rightarrow \infty} c$
- ii) il existe des fonctions γ_t, δ_t telles que sur $\{X_t \rightarrow \infty\}$, $\gamma_t X_t + \delta_t$ converge en loi vers une v.a. non dégénérée.

De plus, sous ces conditions nous avons

- a) Si $\psi(\infty) < \infty$, alors sur $\{X_t \rightarrow \infty\}$ on a $(X_t - \mu_t)/\tilde{g}(t)$ converge p.s. vers une limite non dégénérée
- b) Si $\psi(\infty) = \infty$ et $c = 0$, alors sur $\{X_t \rightarrow \infty\}$ on a en probabilité

$$(X_t - \mu_t)/\tilde{g}(t) = B(\tilde{\psi}_t + o_p(\tilde{\psi}_t^{1/2}))$$

- c) Si $\psi(\infty) = \infty$ et $c > 0$, alors sur $\{X_t \rightarrow \infty\}$ on a en probabilité

$$\log(X_t/\mu_t) = c\tilde{\psi}_t^{-1/2}B(\tilde{\psi}_t) - c^2 + o_p(1)$$

2 Simulations numériques

Pour illustrer les résultats précédents sur le comportement asymptotique du processus (X_t) et des fonctions μ_t nous choisissons le drift $g(t)$ et le coefficient de diffusion sous les formes suivantes :

$$g(t) = t^\alpha \quad \sigma^2(t) = t^{\alpha'} \text{ ou } \sigma^2(t) = e^{-\beta t}$$

pour des valeurs de α, α' et $\beta > 0$ de sorte que les conditions (A1)-A(4) soient vérifiées.

Plusieurs trajectoires de (X_t) et les solutions μ_t seront présentées.

Références

- Gihman, I., Skorohod, A.V. *Stochastic differential equations*. Berlin, Heidelberg, New York : Springer 1972.
- Ibragimov, I. *On the composition of unimodal distributions*. Theor. Probability Appl. 1, 255-260(1956)
- Keller, G., Kersting G., Rösler U. *On the asymptotic behaviour of solutions of stochastic differential equations* Z. Wahrscheinlichkeitstheorie verw. Springer-Verlag Gebiete 68,163-189.(1984)
- Rösler U. *The tail σ -field of time-homogeneous one-dimensional diffusion-processes*. Ann. Probability 7,847-857 (1979)
- Rösler U. *Unimodality of passage times for one-dimensional strong Markov processes*. Ann. Probability 8, 853-859 (1980)

Summary

In this work, we study the asymptotic behavior of solutions of stochastic differential equations. We see under which conditions the solution of the SDE behaves like the ordinary differential equation one

Sur la stabilité forte des systèmes d'attente classiques avec distribution générale inconnue

Aicha Bareche*, Djamil Aïssani*

*Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
aicha_bareche@yahoo.fr, lamos_bejaia@hotmail.com
<http://www.lamos.org>

Résumé. L'objet principal de ce travail est d'effectuer une synthèse sur l'applicabilité de la méthode de stabilité forte à l'étude de quelques systèmes de files d'attente classiques lorsque la distribution des arrivées ou celle de la durée de service est générale et inconnue et doit être préalablement estimée par une méthode non paramétrique.

1 Introduction

L'un des problèmes qui apparaît lors de la conception et de l'exploitation des systèmes complexes est précisément l'analyse de stabilité de leur fonctionnement. Par stabilité d'un système, nous comprenons une dépendance continue des caractéristiques de fonctionnement de ce dernier par rapport à ses paramètres. Un système est dit stable si une petite déviation dans ses paramètres entraîne une petite déviation dans ses caractéristiques. Ainsi l'écart entre ces caractéristiques peut être obtenu en fonction de l'écart entre les paramètres. L'étude de la stabilité consiste à délimiter le domaine dans lequel le modèle idéal peut être utilisé comme une bonne approximation du modèle réel.

À la différence des autres approches de stabilité, la méthode de stabilité forte (Aïssani et Kartashov, 1983; Kartashov, 1996) suppose que la perturbation du noyau de transition est petite par rapport à une certaine norme d'opérateurs. Cette condition, beaucoup plus stricte que les conditions habituelles, permet d'obtenir essentiellement de meilleures estimations pour les distributions stationnaires perturbées.

Le but de ce travail est de résumer sous forme d'une synthèse certains problèmes liés à la mise en oeuvre de la méthode de stabilité forte en pratique. Notre contribution concerne, plus précisément, la discussion et la combinaison de quelques aspects statistiques et numériques dans le but de prouver, d'améliorer et d'étendre le champ d'applicabilité de la méthode de stabilité forte à l'étude de certains systèmes d'attente classiques lorsque l'une des lois les régissant est générale et inconnue et doit être préalablement estimée par une méthode non paramétrique (Parzen, 1962; Chen, 2000; Schuster, 1985).

En effet, dans la pratique, on est généralement confronté à des situations où la fonction densité d'une des lois régissant un système de files d'attente donné est générale et inconnue. En outre, les distributions utilisées dans les systèmes d'attente sont pour la plupart bornées à gauche, ce qui crée un problème d'effet de bord lors de l'estimation des fonctions densité.

Par ailleurs, une des conditions liées à l'application de la méthode de stabilité forte aux systèmes d'attente est que la perturbation faite doit être petite, dans le sens que la densité de la loi générale régissant un système d'attente doit être proche de celle de la loi exponentielle qui est définie sur un support borné $[0, +\infty[$. Or, quand on applique la méthode du noyau dans ce cas on est confronté au problème des effets de bord, donc on doit recourir à certaines techniques permettant leur correction (Bareche et Aïssani, 2007, 2008, 2011, 2013, 2014).

2 Méthodes et matériel

Nous décrivons dans cette section un ensemble de méthodologies utilisées pour l'accomplissement de notre travail.

2.1 Stabilité forte

Nous rappelons la définition de base donnant le critère de stabilité forte.

Définition 2.1 (Kartashov, 1996) *Une chaîne de Markov X de noyau de transition P et de distribution stationnaire π est dite v -fortement stable par rapport à la norme $\|\cdot\|_v$ (où $\|\alpha\|_v = \sum_{j \geq 0} v(j) |\alpha_j|$ pour toute mesure α), si $\|P\|_v < \infty$ et chaque noyau stochastique Q dans un certain voisinage $\{Q : \|Q - P\|_v < \epsilon \text{ pour } \epsilon > 0\}$ admet une unique distribution stationnaire $\mu = \mu(Q)$ telle que $\|\pi - \mu\|_v \rightarrow 0$ lorsque $\|Q - P\|_v \rightarrow 0$.*

2.2 Estimation non paramétrique de la densité de probabilité

Nous présentons dans cette section, un panorama très réduit des techniques d'estimation par noyaux. Soit $X_1, \dots, X_n$ un échantillon issu d'une variable aléatoire X de densité de probabilité f et de distribution F . L'estimateur à noyau de Parzen-Rosenblatt (Parzen, 1962) de la densité f au point $x \in \mathbb{R}$, $f_n(x)$, est donné par :

$$f_n(x) = \frac{1}{nh_n} \sum_{j=1}^n K\left(\frac{x - X_j}{h_n}\right), \quad (1)$$

où K est une densité symétrique appelée noyau, h_n est dit paramètre de lissage.

Plusieurs résultats sont connus dans la littérature quand la fonction densité est définie sur $\mathbb{R}$. Lorsque cette dernière est définie sur un support borné, sans correction, l'estimateur à noyau souffre d'effets de bord puisque il présente un biais au bord. Ce problème est dû à l'usage d'un noyau fixe qui assigne un poids en dehors du support quand le lissage est considéré près du bord.

Pour remédier à ce problème, Schuster (1985) suggère de créer l'image miroir des données de l'autre côté du bord puis d'appliquer l'estimateur (1) pour l'ensemble de données initiales et de nouvelles données créées. $f(x)$ est alors estimée, pour $x \geq 0$, par :

$$\tilde{f}_n(x) = \frac{1}{nh_n} \sum_{j=1}^n \left[K\left(\frac{x - X_j}{h_n}\right) + K\left(\frac{x + X_j}{h_n}\right) \right]. \quad (2)$$

Une autre simple idée pour éviter les effets de bord, est l'usage d'un noyau flexible, qui n'assigne jamais un poids en dehors du support de la fonction densité et qui corrige automatiquement et implicitement les effets de bord. On cite les estimateurs à noyaux asymétriques (Chen, 2000) donnés par la forme :

$$\hat{f}_b(x) = \frac{1}{n} \sum_{i=1}^n K(x, b)(X_i), \quad (3)$$

où b est le paramètre de lissage et le noyau asymétrique K peut être pris comme la densité d'une loi Gamma K_G avec les paramètres $(\frac{x}{b} + 1, b)$ donnée par :

$$K_G\left(\frac{x}{b} + 1, b\right)(t) = \frac{t^{x/b} e^{-t/b}}{b^{x/b+1} \Gamma(x/b + 1)}. \quad (4)$$

2.3 Mesures de performance

Pour mesurer les performances d'un système d'attente par la méthode de stabilité forte, on utilise une approche générale basée sur la simulation à événements discrets et le test statistique de Student (Bareche et Aïssani, 2011, 2014). En effet, l'un des plus importants usages de la simulation est la comparaison des performances de deux ou plusieurs systèmes. Pour faire cette comparaison sur des échantillons obtenus à partir de la simulation des systèmes observés, on doit recourir aux méthodes statistiques (test de Student).

3 Application aux systèmes d'attente classiques

Nous avons employé les techniques décrites dans la section précédente pour l'étude de quelques systèmes d'attente classiques, à savoir :

3.1 Systèmes markoviens : Exemple du système $M/M/1$

Dans un premier temps, nous avons considérée l'étude de la stabilité forte du système élémentaire $M/M/1$. On s'est intéressé à deux types de perturbation :

- Perturbation du flot des arrivées : dans (Bareche et Aïssani, 2007, 2008), on a évalué la proximité des systèmes $G/M/1$ et $M/M/1$.
- Perturbation de la durée de service : dans (Bareche et Aïssani, 2008), on a évalué la proximité des systèmes $M/G/1$ et $M/M/1$.

3.2 Systèmes non markoviens : Exemple du système $G/G/1$

Nous nous sommes penché par la suite sur le problème de la recherche d'approximations numériques pour le système complexe $G/G/1$. Là aussi, nous avons considéré les deux même types de perturbation :

- Dans (Bareche et Aïssani, 2013, 2014), on s'est intéressé à l'approximation du système $G/G/1$ par le système $M/G/1$.
- Dans (Bareche et Aïssani, 2011), on s'est intéressé à l'approximation du système $G/G/1$ par le système $G/M/1$.

Références

- Aïssani, D. et N. V. Kartashov (1983). Ergodicity and stability of Markov chains with respect to operator topology in the space of transition kernels. *Doklady Akademii Nauk Ukrainskoi SSR 11 (seriya A)*, 3–5.
- Bareche, A. et D. Aïssani (2007). Interest of kernel density in the use of strong stability method to precise the proximity of $G/M/1$ and $M/M/1$ systems. *Proceedings of the Second International Conference Valuetools'2007 (Performance Evaluation, Methodologies and Tools)*, 23-25 Octobre 2007, Nantes, France, ACM International Conference Proceedings Series. ICST (Institute for Computer Sciences Social Informatics and Telecommunications Engineering), ICST, Brussels, ISBN: ICST 978-963-9799-00-4, 1–5.
- Bareche, A. et D. Aïssani (2008). Kernel density in the study of the strong stability of the $M/M/1$ queueing system. *Operations Research Letters* 36, 535–538.
- Bareche, A. et D. Aïssani (2011). Statistical techniques for a numerical evaluation of the proximity of $G/G/1$ and $G/M/1$ queueing systems. *Computers and Mathematics with Applications* 61, 1296–1304.
- Bareche, A. et D. Aïssani (2013). Combinaison des méthodes de stabilité forte et d'estimation non paramétrique pour l'approximation de la file d'attente $G/G/1$. *Journal Européen des Systèmes Automatisés* 47, 155–164.
- Bareche, A. et D. Aïssani (2014). Statistical Methodology for Approximating $G/G/1$ Queues by the Strong Stability Technique. In *Proceedings of the 3rd International Conference on Operations Research and Enterprise Systems*, DOI: 10.5220/0004834002410248, Copyright ©SCITEPRESS, 241–248.
- Chen, S. X. (2000). Probability density function estimation using gamma kernels. *Ann. Inst. Statist. Math.* 52, 471–480.
- Kartashov, N. V. (1996). *Strong Stable Markov chains*. Utrecht: TbiMC Scientific Publishers, VSPV.
- Parzen, E. (1962). On estimation of a probability density function and mode. *Ann. Math. Stat.* 33, 1065–1076.
- Schuster, E. F. (1985). Incorporating support constraints into nonparametric estimation of densities. *Commun. Statist. Theory Meth.* 14, 1123–1136.

Summary

The main purpose of this work is to make a synthesis on the applicability of the strong stability method to the study of some classical queueing systems when the distribution of arrivals or that of service time is general and unknown and must be first estimated by a nonparametric method.

Solutions Stepanov presque périodique d'une classe d'équations différentielles stochastiques

Fazia BEDOUHENE*, Nouredine CHALLALI*
Omar MELLAH *, Paul RAYNAUD De FITTE**

* Laboratoire LMPA, Université Mouloud Mammeri, Tizi-Ouzou, BP No. 17 RP 15000

E-Mail: fbedouhene@yahoo.fr; challalin@yahoo.fr; omellah@yahoo.fr

**Laboratoire Raphaël Salem, UMR CNRS 6085, Université of Rouen, France

E-Mail: prf@univ-rouen.fr

Résumé. Dans cet article, on discutera la presque périodicité au sens de Stepanov d'une classe d'équation d'évolution stochastique non-autonome. On montre que, contrairement à ce qu'ont affirmé quelques auteurs, les solutions de certaines classes d'équations différentielles stochastiques ne sont jamais Stepanov presque périodiques en moyenne quadratique, et qu'elles peuvent être Bohr presque périodique en loi infinidimensionnelle.

1 Introduction

Durant les 20 dernières années, le problème d'existence de solutions presque périodiques des équations différentielles stochastiques a été intensément étudié notamment par Arnold et Tudor (1998), Da Prato et Tudor (1995). Le type de presque périodicité considérée par ces auteurs est celle de la loi infini-dimensionnelle des trajectoires du processus solution. Récemment, Bezandry et Diagana (2007, 2009); Bezandry (2008) ont affirmé que certaines EDS à coefficients presque périodiques admettent des solutions presque périodiques en moyenne quadratique. L'étude comparative des différents concepts de presque périodicité d'un processus stochastique établie par Bedouhene et al. (2012) montre que la presque périodicité de Bohr en moyenne quadratique est une propriété assez forte que la presque périodicité en loi (fini-dimensionnelle). Cette étude est complétée par Mellah et Raynaud de Fitte (2013) en donnant des contre exemples qui illustrent l'existence de solutions stationnaires (donc presque périodique en loi) d'EDS vérifiant les hypothèses de Bezandry et Diagana (2007), Bezandry et Diagana (2009) sans qu'elles soient presque périodiques en moyenne quadratique.

La presque périodicité au sens de Stepanov semble une notion très voisine de celle de Bohr. Toute fonction Bohr presque périodique est Stepanov presque périodique, inversement, toute fonction Stepanov presque périodique et uniformément continue est Bohr presque périodique. La question d'existence de solution Stepanov presque périodique dans le cas déterministe ou stochastique se pose donc naturellement.

Dans le cas déterministe, il existe une littérature abondante dédiée à la question, notamment, les travaux de Levitan et Žikov (1978), Corduneanu (1997), Hu (2005), Hu et Mingarelli (2008), Andres et Pennequin (2012). Toutefois, un récent résultat démontré par Andres et Pennequin (2012) sur les solutions des équations différentielles ordinaires à coefficients Stepanov

presque périodiques indique que ces solutions sont Bohr presque périodiques et qu'il n'existe pas de solutions purement Stepanov presque périodiques.

Concernant le cas stochastique, à notre connaissance, le seul travail dédié à la question d'existence de solutions (mild) Stepanov presque périodiques en moyenne quadratique des équations d'évolution stochastique

$$dX(t) = A(t)X(t)dt + f(t)dt + g(t)dW(t) \quad (1)$$

$$dX(t) = A(t)X(t)dt + f(t, X(t))dt + g(t, X(t))dW(t) \quad (2)$$

est établi par Bezandry et Diagana (2008) moyennant les conditions d'Acquistapace-Terreni sur le générateur $A(t)$ et la presque périodicité au sens de Stepanov en moyenne quadratique des coefficients f et g .

Le but du présent travail est, d'une part, de donner un contre-exemple aux affirmations de Bezandry et Diagana (2008) et, d'autre part, de montrer que les d'équations (1)-(2) admettent une solution mild (unique) Bohr presque périodique en loi infinidimensionnelle.

2 Définitions et préliminaires

Soit $(\mathbb{E}, d)$ un espace métrique (séparable). Rappelons qu'une fonction continue $f : \mathbb{R} \rightarrow (\mathbb{E}, d)$ est dite Bohr presque périodique (resp. Stepanov presque périodique) si pour tout $\varepsilon > 0$, il existe $\ell = \ell(\varepsilon) > 0$ tel que tout intervalle de longueur ℓ rencontre $T(f, \varepsilon) = \{\tau \in \mathbb{R}, \sup_{t \in \mathbb{R}} d(f(t), f(t + \tau)) < \varepsilon\}$ (resp. $\{\tau \in \mathbb{R}, \sup_{x \in \mathbb{R}} \int_x^{x+1} d(f(t), f(t + \tau)) dt < \varepsilon\}$). Soit $(\Omega, \mathcal{F}, \mathbf{P})$ un espace de probabilités et $X : \mathbb{R} \times \Omega \rightarrow (\mathbb{E}, d)$ un processus stochastique. On note par $\text{law}(X(t))$ la loi de la variable aléatoire $X(t)$. On dit que X est *presque périodique en loi unidimensionnelle* si l'application $t \mapsto \text{law}(X(t))$ de $\mathbb{R}$ dans $(\mathcal{P}(\mathbb{E}), d_{BL})$ est presque périodique, où

$$d_{BL}(P, Q) = \sup \left\{ \int_E f d(P - Q) ; \|f\|_{BL} \leq 1 \right\} \text{ où } \|f\|_{BL} = \sup_{x \in \mathbb{E}} |f(x)| + \sup_{\substack{x, y \in E \\ x \neq y}} \frac{|f(x) - f(y)|}{d(x, y)}.$$

Si X est à trajectoires continues, on dit que X est *presque périodique en distribution ou en loi infinidimensionnelle* si l'application $t \mapsto \text{law}(X(t + \cdot))$ de $\mathbb{R}$ dans $\mathcal{P}(C(\mathbb{R}; \mathbb{E}))$ est presque périodique, où $C(\mathbb{R}; \mathbb{H}_2)$ est muni de la topologie de la convergence uniforme sur les parties compactes et $\mathcal{P}(C(\mathbb{R}; \mathbb{E}))$ est muni de la distance d_{BL} . Soit $L^2(\mathbf{P}, \mathbb{H}_2)$ l'espace des variables aléatoires à valeurs dans $\mathbb{H}_2$ de carré intégrable. Un processus X est dit *Bohr presque périodique en moyenne quadratique* (resp. *Stepanov presque périodique en moyenne quadratique*) si l'application $X : \mathbb{R} \rightarrow L^2(\mathbf{P}, \mathbb{H}_2)$ est Bohr presque périodique (resp. Stepanov presque périodique).

Dans ce qui suit, on adoptera les notations de Bezandry et Diagana (2008) : soient $(\mathbb{K}, \|\cdot\|_K)$ et $(\mathbb{H}, \|\cdot\|)$ sont deux espaces de Hilbert séparables. L'espace $L_2(\mathbb{K}, \mathbb{H})$ désigne l'ensemble des opérateurs de Hilbert-Schmidt définis sur $\mathbb{K}$ à valeurs dans $\mathbb{H}$ muni de la norme classique de Hilbert-Schmidt, notée $\|\cdot\|_2$. Pour un opérateur $Q \in L_2(\mathbb{K}, \mathbb{H})$ semi-défini positif à trace finie, on suppose que $\{W(t) : t \in \mathbb{R}\}$ est un Q -processus de Wiener défini sur la base stochastique $(\Omega, \mathcal{F}, (\mathcal{F}_t)_{t \in \mathbb{R}}, \mathbf{P})$ à valeurs dans $\mathbb{K}$. Soit $\mathbb{K}_0 = Q^{\frac{1}{2}}(\mathbb{K})$ et $L_2^0 = L_2(\mathbb{K}_0; \mathbb{H})$ muni de la norme

$$\|\Phi\|_{L_2^0}^2 = \left\| \Phi Q^{\frac{1}{2}} \right\|_2^2 = \text{Tr}(\Phi Q \Phi^*).$$

On considère la classe d'EDS (2). On suppose que :

- (H0) Pour chaque $t \in \mathbb{R}$, $A(t) : D(A(t)) \subset L^2(\mathbb{P}; \mathbb{H}) \rightarrow L^2(\mathbb{P}; \mathbb{H})$ est une famille d'opérateurs linéaires fermés à domaine dense dans $L^2(\mathbb{P}; \mathbb{H})$ et que $A(t)$ vérifie les conditions dites d'Acquistapace-Terreni : il existe des constantes $\lambda_0 \geq 0$, $\theta \in (\frac{\pi}{2}, \pi)$, $L, K \geq 0$ et $\alpha, \beta \in (0, 1]$ avec $\alpha + \beta = 1$ telles que

$$\Sigma_\theta \cup \{0\} \subset \rho(A(t) - \lambda_0), \|R(\lambda, A(t) - \lambda_0)\| \leq \frac{K}{1 + |\lambda|} \quad (3)$$

et

$$\|(A(t) - \lambda_0)R(\lambda, A(t) - \lambda_0)[R(\lambda_0, A(t)) - R(\lambda_0, A(s))]\| \leq L|t - s|^\alpha |\lambda|^\beta$$

pour tout $s, t \in \mathbb{R}$, $\lambda \in \Sigma_\theta := \{\lambda \in \mathbb{C} - \{0\} : |\arg \lambda| \leq \theta\}$. Notons que les conditions d'Acquistapace-Terreni garantissent l'existence d'une famille d'évolution $\{U(t, s) : t \geq s, t, s \in \mathbb{R}\}$ associée à $A(t)$.

- (H1) Les opérateurs $A(t)$ et $U(r, s)$ commutent et la famille d'évolution $U(t, s)$ générée par $A(t)$ est asymptotiquement stable, à savoir, il existe des constantes $M, \omega > 0$, tel que

$$\|U(t, s)\| \leq M \exp -\delta(t - s).$$

- (H2) $u \mapsto U(t + u, s + u)$ est Stepanov presque périodique pour $t, s \in \mathbb{R}$ avec $|t - s| \geq h > 0$.

- (H3) et $F : \mathbb{R} \times L^2(\mathbb{P}; \mathbb{H}) \mapsto L^2(\mathbb{P}; \mathbb{H})$, $G : \mathbb{R} \times L^2(\mathbb{P}; \mathbb{H}) \mapsto L^2(\mathbb{P}; L_2^0)$ deux fonctions continues Stepanov presque périodiques.

Le premier résultat que nous avons démontré concerne l'existence et l'unicité de la solution mild Bohr presque périodique en loi infinidimensionnelle de l'équations (1), nous avons :

Theorem 2.1. *Sous les conditions (H0) – (H3), l'EDS (1) possède une unique solution solution mild bornée en moyenne quadratique, donnée par*

$$X(t) = \int_{-\infty}^t U(t, s)f(s)ds + \int_{-\infty}^t U(t, s)g(s)dW(s), \text{ a.e } t \in \mathbb{R}. \quad (4)$$

Cette solution est de plus Bohr presque périodique en loi infini-dimensionnelle.

Références

- Andres, J. et D. Pennequin (2012). On the nonexistence of purely Stepanov almost-periodic solutions of ordinary differential equations. *Proc. Amer. Math. Soc.* 140(8), 2825–2834.
- Arnold, L. et C. Tudor (1998). Stationary and almost periodic solutions of almost periodic affine stochastic differential equations. *Stochastics Stochastics Rep.* 64, 177–193.
- Bedouhene, F., O. Mellah, et P. Raynaud de Fitte (2012). Bochner-almost periodicity for stochastic processes. *Stoch. Anal. Appl.* 30(2), 322–342.

- Bezandry, P. H. (2008). Existence of almost periodic solutions to some functional integro-differential stochastic evolution equations. *Statist. Probab. Lett.* 78(17), 2844–2849.
- Bezandry, P. H. et T. Diagana (2007). Existence of almost periodic solutions to some stochastic differential equations. *Appl. Anal.* 86(7), 819–827.
- Bezandry, P. H. et T. Diagana (2008). Existence of S^2 -almost periodic solutions to a class of nonautonomous stochastic evolution equations. *Electron. J. Qual. Theory Differ. Equ.*, No. 35, 19.
- Bezandry, P. H. et T. Diagana (2009). Existence of quadratic-mean almost periodic solutions to some stochastic hyperbolic differential equations. *Electron. J. Differential Equations*, No. 111, 14.
- Corduneanu, C. (1997). Almost periodic solutions to differential equations in abstract spaces. *Rev. Roumaine Math. Pures Appl.* 42(9-10), 753–758. Collection of papers in honour of Academician Radu Miron on his 70th birthday.
- Da Prato, G. et C. Tudor (1995). Periodic and almost periodic solutions for semilinear stochastic equations. *Stochastic Anal. Appl.* 13(1), 13–33.
- Hu, Z. (2005). Boundedness and Stepanov's almost periodicity of solutions. *Electron. J. Differential Equations*, No. 35, 7.
- Hu, Z. et A. Mingarelli (2008). Bochner's theorem and Stepanov almost periodic functions. *Annali di Matematica Pura ed Applicata* 187(4), 719–736.
- Levitan, B. M. et V. V. Žikov (1978). *Pochti-periodicheskie funktsii i differentsialnye uravneniya*. Moskov. Gos. Univ., Moscow.
- Mellah, O. et P. Raynaud de Fitte (2013). Counterexamples to mean square almost periodicity of the solutions of some SDEs with almost periodic coefficients. *Electron. J. Differential Equations*, No. 91, 7.

Summary

In this paper, we discuss the existence and uniqueness of Stepanov (quadratic-mean) almost periodic solutions to a class of nonautonomous stochastic evolution equations on a separable real Hilbert space. We show that contrarily to what is claimed by some authors, the solutions of some EDS with Stepanov almost periodic coefficients in mean square are never mean square Stepanov almost periodic, but they can be Bohr almost periodic in distribution.

Borne de stabilité pour la distribution stationnaire du temps d'attente: Approche par interaction entre modèles de risque et systèmes d'attente

Zina BENOURET, Atik TOUAZI,
Djamil AISSANI

Unité de Recherche LaMOS Université de Béjaïa, Targa-Ouzamour, 06000 (Algérie)
benouaret_z@yahoo.fr,
<http://www.lamos.org>

Résumé. Le but de l'interaction connue entre la théorie de risque et celle des systèmes d'attente est la translation des résultats qui permet d'éviter la complexité d'une étude dans certains modèles. D'autre part, la méthode de stabilité forte possède un large champ d'application en théorie de files d'attente et récemment en modèles de risque.

Dans ce travail, notre objectif est la translation des inégalités de stabilité forte obtenues dans les modèles de risque afin d'estimer la déviation de la distribution stationnaire du temps d'attente qui ne peut être déterminée explicitement et exactement que dans des cas très simple.

1 Introduction

C'est l'académicien V. Kalashnikov, qui a initié l'étude de stabilité forte (voir Aissani et Kartashov (1983)) dans les modèles de risque (voir Kalashnikov (2000)). Par la suite, plusieurs autres applications ont été réalisé afin d'estimer la déviation des probabilités de ruine dans des modèles qui cernent mieux la réalité (voir Benouaret et Aissani (2010)). D'autres part, d'énormes progrès ont été réalisés sur la performance de la méthode de stabilité forte dans l'extension des fondements théoriques et dans l'applicabilité (du point de vue théorique) à plusieurs classes spécifiques de systèmes d'attente en analysant la stabilité de la distribution stationnaire du nombre de clients. Cependant, la distribution stationnaire du temps d'attente, qui ne peut aussi être estimée explicitement et exactement que dans des cas très simple, représente une caractéristique très importante à analyser. (voir Hêche et al. (2003))

Dans ce travail, notre objectif est de déterminer des inégalités de stabilité forte pour la distribution stationnaire du temps d'attente en utilisant interaction qui existe entre la théorie de risque et celle des files d'attente (voir Janssen (1982)). Autrement dit, nous présentons le concept de translation de borne de stabilité obtenue d'un modèle de risque vers un système de file d'attente qui lui est équivalent.

2 Interaction entre théorie de risque et files d'attente

Dans cette section, nous présentons quelques relations précises sur l'interaction qui existe entre la théorie de risque et celle des files d'attente. (voir Janssen (1982))

Système d'attente $GI/G/1$

Dans une file d'attente $GI/G/1$, l'étude du temps d'attente est possible à l'aide d'une relation classique : l'équation de Lindley. Afin d'établir cette relation, nous considérons une séquence d'arrivées de client indexée par n :

- C_n le n ème client arrivé dans le système,
- t_n l'intervalle de temps entre l'arrivée de C_{n-1} et C_n ($t_n = a_n - a_{n-1}$),
- s_n le temps de service de C_n ,
- w_n le temps d'attente de C_n .

Notons $W_n(y)$ la fonction de répartition du temps d'attente w_n de C_n . La file étant supposée stable, le processus $\{w_n, n \geq 0\}$ est ergodique et admet une distribution invariante $W(y)$ vérifiant l'équation fonctionnelle

$$W(y) = \int_0^\infty U(y-w) dW(w), \quad y \geq 0$$

où U est la fonction de répartition de $u_n = s_n - t_{n+1}$, indépendante de n .

La distribution stationnaire du temps d'attente satisfait donc l'équation intégrale de Lindley. Cette équation de type Wiener-Hopf ne peut être résolue explicitement et exactement que dans des cas très simple.

Modèles de risque

D'une manière générale, un modèle de risque peut être entièrement décrit à l'aide des éléments suivant :

1. $\{A_n, n \geq 1\}$ est le processus des inter-arrivées successives des réclamations.
2. $\{B_n, n \geq 1\}$ est le processus des montants successifs des réclamations.
3. Nous supposons que les revenus de la compagnie d'assurance ont un taux constant $c > 0$.

Associions au modèle de risque considéré les processus stochastiques suivants :

$$\{X_n, n \geq 0\}, \text{ avec } X_n = B_n - cA_n, \quad n \geq 0 \quad (1)$$

$$\{S_n, n \geq 0\}, \text{ avec } S_n = \sum_{k=1}^n X_k, \quad n \geq 0, \quad S_0 = 0 \quad (2)$$

$$\{T_n, n \geq 0\}, \quad T_n = \sum_{i=1}^n A_i, \quad n > 0, \quad T_0 = 0 \quad (3)$$

$$\{M_n, n \geq 0\}, \text{ avec } M_n = \sup\{S_0, S_1, \dots, S_n\} \quad (4)$$

Nous considérons aussi le processus $\{N(t), t \geq 0\}$ où $N(t)$ est le nombre total de réclamations dans l'intervalle $[0, t]$ et le processus $\{X(t), t > 0\}$, où $X(t)$ représente le montant cumulé des réclamations sur l'intervalle $[0, t]$.

Autrement dit,

$$X(t) = \begin{cases} 0 & \text{si } t < A_1 \\ \sum_{k=1}^{N(t)} B_k, & t \geq A_1 \end{cases} \quad (5)$$

$\{Z(t), t > 0\}$, où $Z(t)$ représente la fortune de l'assureur à l'instant t , avec

$$Z(t) = u + ct - X(t) \quad (6)$$

En utilisant le processus des réserves $\{Z(t), t > 0\}$, on peut définir deux types de probabilités de non-ruine ; en temps fini et en temps infini, définies comme suit :

$$\Psi(u, t) = P(Z(t') \geq 0, t' \in [0, t] / Z(0) = u), \quad t > 0, \quad u \geq 0 \quad (7)$$

$$\Psi(u) = P(Z(t') \geq 0, \forall t' \geq 0 / Z(0) = u), \quad u \geq 0 \quad (8)$$

Parallélisme théorie de risque- théorie de files d'attente

Les trajectoires du processus de temps d'attente et celle du processus des réserves ont la même structure d'un point de vue géométrique. Le problème fondamental est de montrer que cette technique conduit à une équivalence entre les deux modèles d'un point de vue probabiliste. Nous avons dans les équations suivantes des relations précises présentées par Janssen (1982) sur cette interaction :

$$P(W_n \leq x) = P(M_n \leq x) = \Psi_n(x) \quad (9)$$

$$P(W_{N(t)} \leq x) = P(M_{N(t)} \leq x) = \Psi(x, t) \quad (10)$$

3 Translation de bornes de stabilité forte

Dans cette partie, nous présentons le concept de translation des inégalités de stabilité forte obtenues dans un modèle de risque vers le système de file d'attente qui lui est équivalent.

Définition : La chaîne de Markov homogène X est dite fortement stable par rapport à la norme $\|\cdot\|$ si

1. $\|P\| < \infty$.
2. Chaque noyau de transition Q dans un certain voisinage $\{Q : \|Q - P\| < \epsilon\}$, admet une mesure invariante unique $\nu = \nu(Q)$.
3. Il existe une constante $C = C(P)$, telle que $\|\nu - \pi\| \leq C\|P - Q\|$.

Le principe de l'application de la méthode de stabilité forte dans les modèles de risque est basé sur la représentation de la probabilité de ruine en fonction de la distribution stationnaire d'un processus dit "processus inverse" noté $\{V_n\}_{n \geq 0}$ tel que

$$\Psi(u) = \lim_{n \rightarrow \infty} P(V_n > u) \quad (11)$$

Par conséquent, une estimation de la déviation $\|\Psi - \Psi'\|_v$ de la probabilité de ruine a été obtenue avec un calcul explicite des constantes. Par ailleurs, en utilisant les relations d'interaction, nous déterminons une borne de stabilité forte qui estime la déviation de la distribution stationnaire du temps d'attente. La figure suivante illustre le concept de la translation.

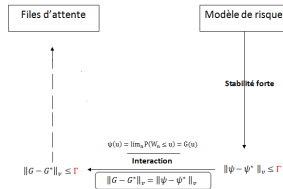


FIG. 1 – Concept de la translation : modèles de risque vers systèmes d'attente

En résumé, afin d'éviter la complexité de l'application directe du critère de stabilité forte dans l'analyse de la distribution stationnaire du temps d'attente, nous avons exploité l'interaction qui existe entre la théorie de risque et celle des files d'attente. Nous avons donc présenté des inégalités de stabilité forte pour cette caractéristique qui ne possède pas de formules explicites que dans des cas très simple.

Références

- Aissani, D. et N.-V. Kartashov (1983). Ergodicity and stability of Markov chains with respect to operator topology in the space of transition kernels. *Dokl. Akad. Nauk Ukr. S.S.R., ser. A 11*, 3–5.
- Benouaret, Z. et D. Aissani (2010). Strong stability in a two-dimensional classical risk model with independent claims. *Scandinavian Actuarial Journal 2010*, 83–93.
- Hêche, J.-F., T.-M. Liebling, et D. Werra (2003). *Recherche Opérationnelle pour ingénieurs II*. Lausanne : Presses polytechniques et universitaire romandes, CH-1015.
- Janssen, J. (1982). On the interaction between risk and queueing theories. *Paper presented at the first Tagung fiber Risikotheorie at the Mathematics Research Center, Oberwolfach*.
- Kalashnikov, V. (2000). The stability concept for stochastic risk models. *Working Paper Nr 166. Lab. of Actuarial Mathematics. University of Copenhagen*.

Summary

The aim of the interaction well-known between risk and queueing theories is the translation of results which allow avoiding the complexity study in some models. On the other hand, there exists numerous application of the strong stability method in queueing theory and recently in risk models.

In this work, Our objective is to translate strong stability bounds obtained in risk models in order to estimate the deviation of the waiting time stationary distribution which can be determined explicitly and exactly only in very simple cases.

Sur l'estimation des retards multiples dans une diffusion non linéaire

Awatif BENYAHIA*, Malika Dali-Korso**

*Laboratoire de Statistiques et Modélisations Aléatoires(LASMA)
Département de Mathématique, Université de Tlemcen, 13000, Algérie
awatif.ben@outlook.fr

Résumé. Dans ce travail, nous traitons le problème d'estimation multidimensionnelle de retard pour des processus de type diffusion non linéaires. Nous montrons sous certaines conditions de régularité que l'EMV des paramètres de retard est consistant, asymptotiquement normal et asymptotiquement efficace lorsque le coefficient de diffusion tend vers zéro.

1 Introduction

En 1988, Kutoyants traite le problème d'estimation paramétrique du retard dans le modèle suivant où la dérive présente la difficulté de non dérivabilité par rapport au paramètre θ

$$dX_t = X_{t-\theta}dt + \varepsilon dW_t, \quad X_s = x_0, \quad s \leq 0.$$

Il résoud en 1994 le problème à retard (si $\theta \in (0, 1)$) de la forme

$$dX_t = X_{\theta t}dt + \varepsilon dW_t, \quad X_s = x_0, \quad s \leq 0$$

et dans un but de généralisation, il cite le modèle suivant comme solvable du point de vue estimation de θ

$$dX_t = S(X_{t-\theta})dt + \varepsilon dW_t, \quad X_s = x_0, \quad s \leq 0.$$

En 1992, Kutoyants, Mourid et Bosq résolvent un problème d'estimation de deux paramètres : θ et γ dans $\mathbb{R}^k$ par l'EMV dans le processus avec k retards

$$dX_t = \sum_{l=1}^k \gamma_l X_{t-\theta_l}dt + \varepsilon dW_t.$$

Ensuite en 1994 Kutoyants et Mourid reprennent la même e.d.s. et proposent un estimateur de la distance minimale pour ces mêmes paramètres et étudient les propriétés de cet estimateur. Enfin en 2007, Mourid propose un estimateur pour le paramètre k représentant le nombre de retards dans ce même processus.

Estimation Multidimensionnelle

Nous considérons l'équation différentielle stochastique de type diffusion non linéaire suivante

$$\begin{aligned} dX_t &= (S_1(X_{t-\theta_1}) + S_2(X_{t-\theta_2}))dt + \varepsilon dW_t, \quad 0 \leq t \leq T \\ X_s &= x_0, \quad s \leq 0. \end{aligned}$$

où $S_1(\cdot)$ et $S_2(\cdot)$ sont deux fonctions régulières données, le paramètre à estimer est $\vartheta = (\theta_1, \theta_2)^t \in (0, T)^2$ et x_0 est fixé. Nous nous proposons d'estimer le paramètre ϑ à partir des observations $X^\varepsilon = \{X_t, 0 \leq t \leq T\}$ sur l'intervalle de temps fixé $[0, T]$, et d'étudier les propriétés asymptotiques de l'estimateur du maximum de vraisemblance du paramètre dans le cadre des petites diffusions ($\varepsilon \rightarrow 0$). Nous montrons que cet estimateur est consistant, asymptotiquement normal et asymptotiquement efficace.

2 Notations et hypothèses

Soit l'e.d.s. de type diffusion non linéaire suivante

$$\begin{aligned} dX_t &= (S_1(X_{t-\theta_1}) + S_2(X_{t-\theta_2}))dt + \varepsilon dW_t, \quad 0 \leq t \leq T \\ X_s &= x_0, \quad s \leq 0. \end{aligned}$$

- $S_1(\cdot)$ et $S_2(\cdot)$ sont deux fonctions régulières données,
- x_0 est fixé,
- l'espace paramétrique Θ est défini par

$$\Theta = (\alpha_1, \beta_1) \times (\alpha_2, \beta_2), \quad 0 < \alpha_1 < \beta_1 < T, \quad 0 < \alpha_2 < \beta_2 < T$$

avec $\alpha_1, \beta_1, \alpha_2, \beta_2$ connus et $\vartheta = (\theta_1, \theta_2)^t \in \Theta$ représente deux retards.

Nous associons à cette e.d.s. l'équation déterministe

$$\begin{aligned} \frac{dx_t}{dt} &= S_1(x_{t-\theta_1}) + S_2(x_{t-\theta_2}), \quad 0 \leq t \leq T \\ x_s &= x_0, \quad s \leq 0. \end{aligned}$$

Nous ferons l'hypothèse :

H₁ : S_1 et S_2 sont des fonctions deux fois continûment dérivables de dérivées bornées sur $\mathbb{R}$.

Le système dynamique admet alors une solution unique et les mesures induites par des solutions correspondant à des valeurs différentes du paramètre ϑ , sont équivalentes sur l'espace $(\mathcal{C}_T, \mathcal{B}_T)$ (R.S.Lipster, A.N.Shiryayev).

Nous imposons aussi une hypothèse d'identifiabilité pour ce problème :

H₂ : les fonctions S_1 et S_2 sont telles que pour tout $\delta > 0$, et tout compact $\mathcal{K}$ de Θ ,

$$\inf_{\vartheta \in \mathcal{K}} \inf_{|\vartheta - \tilde{\vartheta}| > \delta} \int_a^b (S_1(x_{t-\theta_1}) + S_2(x_{t-\theta_2}) - S_1(x_{t-\tilde{\theta}_1}) - S_2(x_{t-\tilde{\theta}_2}))^2 dt > 0$$

où $\vartheta = (\theta_1, \theta_2)^t$, $\tilde{\vartheta} = (\tilde{\theta}_1, \tilde{\theta}_2)^t$ et x est solution de l'équation déterministe pour la valeur ϑ .

• Le rapport de vraisemblance pour deux valeurs du paramètre $\vartheta_0 = (\theta_{0,1}, \theta_{0,2})^t$, $\vartheta = (\theta_1, \theta_2)^t \in \Theta$ est

$$\begin{aligned} L(\vartheta, \vartheta_0; X) &= \frac{dP_{\vartheta}^{(\varepsilon)}}{dP_{\vartheta_0}^{(\varepsilon)}}(X) \\ &= \exp \left\{ \frac{1}{\varepsilon^2} \int_0^T \left(S_1(X_{t-\theta_1}) + S_2(X_{t-\theta_2}) \right. \right. \\ &\quad \left. \left. - S_1(X_{t-\theta_{0,1}}) - S_2(X_{t-\theta_{0,2}}) \right) dX_t \right. \\ &\quad \left. - \frac{1}{2\varepsilon^2} \int_0^T \left((S_1(X_{t-\theta_1}) + S_2(X_{t-\theta_2}))^2 \right. \right. \\ &\quad \left. \left. - (S_1(X_{t-\theta_{0,1}}) + S_2(X_{t-\theta_{0,2}}))^2 \right) dt \right\}. \end{aligned}$$

• L'estimateur du maximum de vraisemblance EMV $\hat{\vartheta}_\varepsilon$ est défini par

$$L(\hat{\vartheta}_\varepsilon, \vartheta_0; X) = \sup_{\vartheta \in \bar{\Theta}} L(\vartheta, \vartheta_0; X)$$

où $\bar{\Theta}$ est l'adhérence de Θ et ϑ_0 est une valeur fixée dans Θ .

• La matrice de l'Information de Fisher associée à ce problème est

$$I(\vartheta) = \begin{pmatrix} I_{11} & I_{12} \\ I_{12} & I_{22} \end{pmatrix}$$

où

$$\begin{aligned} I_{11} &= \int_0^T (S_1(x_{t-2\theta_1}) + S_2(x_{t-\theta_1-\theta_2}))^2 S_1'^2(x_{t-\theta_1}) dt, \\ I_{12} &= \int_0^T (S_1(x_{t-2\theta_1}) + S_2(x_{t-\theta_1-\theta_2})) S_1'(x_{t-\theta_1}) S_2'(x_{t-\theta_2}) \\ &\quad \times (S_1(x_{t-\theta_1-\theta_2}) + S_2(x_{t-2\theta_2})) dt, \\ I_{22} &= \int_0^T (S_1(x_{t-\theta_1-\theta_2}) + S_2(x_{t-2\theta_2}))^2 S_2'^2(x_{t-\theta_2}) dt, \end{aligned}$$

Nous ferons enfin l'hypothèse suivante :

H₃ : $I(\vartheta)$ est une matrice définie positive uniformément par rapport à $\vartheta \in \Theta$: pour tout $\lambda \in \mathbb{R}^2$, $\lambda \neq 0$

$$\inf_{\vartheta \in \Theta} (I(\vartheta)\lambda, \lambda) > 0$$

Nous montrons que sous **H₁**, **H₂** et **H₃** ce problème d'estimation devient régulier au sens classique, Et nous obtenons les résultats suivants :

Condition LAN

Si les conditions $\mathbf{H}_1$ et $\mathbf{H}_3$ sont satisfaites, alors la famille des lois $\{P_\vartheta^{(\varepsilon)}, \vartheta \in \Theta\}$ est uniformément LAN, de matrice de normalisation dans la représentation du rapport de vraisemblance $\varphi_\varepsilon(\vartheta) = \varepsilon I^{-1/2}(\vartheta)$ et de vecteur aléatoire

$$\Delta_\varepsilon(\vartheta, X) = \varepsilon^{-1} I^{-1/2}(\vartheta) \int_0^T \begin{pmatrix} (S_1(x_{t-2\theta_1}) + S_2(x_{t-\theta_1-\theta_2})) S'_1(x_{t-\theta_1}) \\ (S_1(x_{t-\theta_1-\theta_2}) + S_2(x_{t-2\theta_2})) S'_2(x_{t-\theta_2}) \end{pmatrix} \\ \times [dX_t - (S_1(X_{t-\theta_1}) + S_2(X_{t-\theta_2})) dt]$$

qui admet la loi $\mathcal{N}(0, Id)$ sous $P_\vartheta^{(\varepsilon)}$.

De plus pour tout estimateur $\tilde{\vartheta}_\varepsilon$ et tout $l \in \mathbf{W}_{e,2}$ nous avons l'inégalité de Hajek-Le Cam

$$\lim_{\delta \rightarrow 0} \liminf_{\varepsilon \rightarrow 0} \sup_{|\vartheta - \vartheta_0| < \delta} \mathbb{E}_\vartheta l \left(\varphi_\varepsilon^{-1}(\vartheta_0)(\tilde{\vartheta}_\varepsilon - \vartheta) \right) \geq \frac{1}{2\pi} \int_{\mathbb{R}^2} l(x) \exp \left(-\frac{|x|^2}{2} \right) dx$$

Pour montrer le Théorème nous suivons I.A.Ibragimov, R.Z.Khasminskii et Y.Kutoyants et nous avons besoin de trois Lemmes techniques.

Propriétés de l'estimateur du maximum de vraisemblance

Si les hypothèses $\mathbf{H}_1$, $\mathbf{H}_2$ et $\mathbf{H}_3$ sont réalisées, alors uniformément sur tout compact $\mathcal{K}$ de Θ , l'estimateur du maximum de vraisemblance $\hat{\vartheta}_\varepsilon$ vérifie

- $P_\vartheta - \lim_{\varepsilon \rightarrow 0} \hat{\vartheta}_\varepsilon = \vartheta$.
- $\varphi_\varepsilon(\vartheta)^{-1}(\hat{\vartheta}_\varepsilon - \vartheta) \implies \mathcal{N}(0, Id)$ où Id est la matrice unité.
- La convergence des moments $\mathbb{E}_\vartheta |\varphi_\varepsilon(\vartheta)^{-1}(\hat{\vartheta}_\varepsilon - \vartheta)|^p$ vers les moments $\mathbb{E} |\Delta|^p$ pour tout $p > 0$, où Δ est un vecteur aléatoire de loi $\mathcal{N}(0, Id)$.
- L'EMV est localement asymptotiquement minimax au sens de Hajek pour les fonctions de pertes $l(\cdot) \in \mathbf{W}_{e,2}$ i.e

$$\lim_{\delta \rightarrow 0} \liminf_{\varepsilon \rightarrow 0} \sup_{|\vartheta - \vartheta_0| < \delta} \mathbb{E}_\vartheta l \left(\varphi_\varepsilon^{-1}(\vartheta_0)(\hat{\vartheta}_\varepsilon - \vartheta) \right) = \frac{1}{2\pi} \int_{\mathbb{R}^2} l(x) \exp \left(-\frac{|x|^2}{2} \right) dx$$

La preuve de ce Théorème est basée sur deux Lemmes qui montrent que la famille des lois des processus du rapport de vraisemblance $(Z_\varepsilon(u), u \in \varepsilon^{-1}(\Theta - \vartheta))$ est tendue dans l'espace $\mathcal{C}_0(\mathbb{R}^2)$.

Summary

in this work, we study the problem of multidimensional estimation of delay with non linear diffusion type process. we show under certain conditions of regularity that the MLE is consistent, asymptotically normal and asymptotically efficient

Analyse stochastique du système d'attente GI/M/s avec arrivées négatives

Louiza Berdjoudj*,
Djamil Aissani*

*Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
l_berdjoudj@yahoo.fr, lamos_bejaia@hotmail.com
<http://www.lamos.org>

Résumé. Dans ce travail, nous considérons l'analyse stochastique du système d'attente GI/M/s avec arrivées négatives (RCE) en utilisant la chaîne de Markov induite. Nous avons démontré l'existence du régime stationnaire de cette chaîne incluse. Nous avons également obtenu quelques mesures de performance.

1 Introduction

Recemment, il est apparu dans la littérature des files d'attente, des publications des travaux portant sur les systèmes d'attente caractérisés par la présence simultanée de deux types d'arrivées. D'un côté, les arrivées positives ou régulières qui ont pour l'objectif l'occupation du service. Elles sont traitées normalement par le serveur. De l'autre côté, les arrivées négatives dont la présence dans un système de file d'attente affecte ce dernier de différentes manières. L'intérêt porté à cette nouvelle famille de réseaux de files d'attente avec arrivées positives et négatives, introduite par (Gelenbe, 1989), était motivée initialement, par la modélisation des réseaux de neurones où les arrivées négatives et positives représentent des signaux excitateurs, qui font croître le potentiel du neurone et sa tendance à produire une impulsion, et inhibiteurs, qui diminuent le potentiel du neurone et sa tendance à produire une impulsion, respectivement. Puis leurs domaines d'applications se sont étendus pour toucher d'autres systèmes plus complexes comme les réseaux informatiques avec infection par un virus (Artalejo et Gomez-Corral, 1999), élimination des transactions dans les bases de données, les systèmes inventaires, les systèmes de télécommunication, les systèmes de production, ... etc. Une revue sur ce thème peut-être trouvée dans (Artalejo, 1994). Dans notre travail, nous allons nous intéresser au cas où une arrivée négative élimine un seul client positif. Le nombre de publications sur les systèmes de type GI/M/1 avec arrivées négatives est réduit comparé au nombre de publications disponibles dans la littérature des systèmes de files d'attente de type M/G/1 avec arrivées négatives, avec rappels et arrivées négatives. Le système d'attente M/M/1 à arrivées positives et négatives a été étudié par Harisson et Pitel (Harrison et Pitel, 1993). Cette file d'attente M/M/1 a été ensuite étendue à la file d'attente M/G/1 par Harisson et Pitel (Harrison et Pitel, 1996) et à la file d'attente GI/M/1 par Yang et Chae (Yang et Chae, 2001). En outre les clients négatifs ont été pris en compte dans une file d'attente avec rappels par Artalejo et Gomez-Corral (Artalejo et Gomez-Corral, 1996). (Kyung et al., 2010) ont considéré la version temps discret du

système GI/M/1 avec arrivées négatives.

Le but de notre travail est de faire une analyse stationnaire du système d'attente GI/M/s avec arrivées négatives.

2 Le modèle d'attente

On considère un système de files d'attente à plusieurs serveurs indépendants montés en parallèle avec deux types de clients où les intervalles entre deux arrivées consécutives ξ sont des variables aléatoires positives, indépendantes et toutes de même loi de densité de probabilité $g(t)$ et de transformée de Laplace Stieljes $G^*(\theta)$. Les arrivées négatives RCE (remove the customer in the end), sous la discipline FCFS, et dans le cas où tous les serveurs sont occupés éliminent un client (le dernier) de la file. On suppose que le processus générant les arrivées négatives est poissonien de taux $\delta \geq 0$ et que une arrivée négative n'a aucun effet sur un système dont la file vide (il ya au moins un serveur libre). Les durées de service sont des variables aléatoires indépendantes de loi exponentielle de moyenne $1/\mu$. Finalement, on suppose que les inter-arrivées sont indépendantes des durées de service et la discipline de service est FIFO.

3 Chaîne de Markov induite

Soit $X(t)$ le nombre de clients dans le système à l'instant t , $t_1 < t_2 < \dots < t_n < \dots$ les instants d'arrivées des clients $1, 2, \dots, n$, soit $X_n = X(t_n^-)$ le nombre de clients dans le système just avant l'instant d'arrivée du $n^{\text{ème}}$ client (positif).

On a pour tout $n \geq 1$

$$X_{n+1} = X_n + 1 - D_{n+1} \quad (1)$$

où D_{n+1} est le nombre de départs entre les instants d'arrivées consécutifs t_n et t_{n+1} . Or D_{n+1} est une variable aléatoire qui ne dépend que de la durée de l'intervalle (t_n, t_{n+1}) , et $(t_{n+1} - t_n)$ est aussi une variable aléatoire indépendante de X_n , de l'état du processus avant t_n ainsi que de n . Donc $(X_n)_{n \geq 1}$ est bien une chaîne de Markov homogène et possède les probabilités de transitions suivantes : $P_{ij} = P(X_{n+1} = j / X_n = i)$ tel que :

$$P_{ij} = \begin{cases} 0, & j > i + 1; \\ \int_0^{+\infty} C_{i+1}^{i+1-j} (1 - e^{-\mu\xi})^{i+1-j} e^{-j\mu\xi} g(\xi) d\xi, & j \leq i + 1 \leq s; \\ \int_0^{+\infty} e^{-(s\mu+\delta)t} \frac{[(s\mu+\delta)t]^{i+1-j}}{(i+j-1)!} g(t) dt, & s \leq j \leq i + 1; \\ \int_0^{+\infty} \int_0^{+\infty} d_{s-j}(v, t) (s\mu + \delta) \frac{[(s\mu+\delta)v]^{i-s}}{(i-s)!} e^{-(s\mu+\delta)v} g(t) dv dt, & j < s < i + 1. \end{cases}$$

Où $d_{s-j}(v, t) = C_s^{s-j} (1 - e^{-\mu(t-v)})^{s-j} e^{-j\mu(t-v)}$ avec v est la durée nécessaire pour servir les $(i + 1 - s)$ premiers clients.

La chaîne de Markov induite $(X_n)_{n \geq 1}$ est ergodique si et seulement si

$$E(\xi) = \int_0^{+\infty} t g(t) dt > \frac{1}{s\mu + \delta}.$$

Admet donc une unique distribution stationnaire α donnée par

$$\alpha_j = \begin{cases} C\beta_j, & \text{pour } j = \overline{0, s-2}; \\ C\sigma^{j-s+1}, & \text{pour } j \geq s-1. \end{cases}$$

et $C = \frac{1}{\sum_{j=0}^{s-2} \beta_j + \frac{1}{1-\sigma}}$ et $0 < \sigma < 1$ et l'unique solution de l'équation
 $x = G^*[(s\mu + \delta)(1 - x)]$

4 Mesures de performance du système

Les arrivées étant non poissonniennes, la probabilité α_n est différente de la probabilité que le système contienne n clients à un instant arbitraire ($\alpha_n \neq \pi_n$) donc on ne peut calculer le nombre moyen de clients dans le système par la formule $L = \sum_{n=1}^{\infty} n\pi_n$, les π_n étant a priori inconnus. Soit t_f la durée d'attente d'un client dans la file. Sa fonction de répartition est $W(x) = P(t_f \leq x)$. On a

$$\begin{aligned} P(0 < t_f \leq x) &= \int_0^x (s\mu + \delta) C \sigma e^{-(s\mu + \delta)t} \left(\sum_{n=s}^{\infty} \frac{\sigma^{n-s} [(s\mu + \delta)t]^{n-s}}{(n-s)!} \right) dt \\ &= \int_0^x C(s\mu + \delta) \sigma e^{-(s\mu + \delta)t(1-\sigma)} dt \\ &= \frac{C(s\mu + \delta)\sigma}{(s\mu + \delta)(1-\sigma)} [1 - e^{-(s\mu + \delta)x(1-\sigma)}] \end{aligned}$$

On trouve après tout calcul fait $W(x) = 1 - \frac{\sigma e^{-(s\mu + \delta)x(1-\sigma)}}{1 + (1-\sigma) \sum_{j=0}^{s-2} \beta_j}$. La densité de la loi de t_f est donc

$$w(x) = W'(x) = \frac{\sigma(s\mu + \delta)(1-\sigma)e^{-(s\mu + \delta)x(1-\sigma)}}{1 + (1-\sigma) \sum_{j=0}^{s-2} \beta_j}.$$

Les mesures de performance du système sont alors données par :

La durée moyenne de séjour dans la file est :

$$\bar{W}_q = \frac{C\sigma}{(s\mu + \delta)(1-\sigma)^2}. \quad (2)$$

Puis le **temps moyen de séjour d'un client dans le système** est

$$\bar{W} = \bar{W}_q + 1/\mu.$$

Le nombre moyen de clients dans le système est

$$\bar{L} = \bar{W}/E(\xi).$$

Le nombre moyen de clients rencontrés par un client qui arrive dans le système est :

$$\bar{L}_a = (s\mu + \delta)\bar{W}_q \left[1 + \frac{K(1-\sigma)^2}{C\sigma} \right]. \quad (3)$$

$$\text{où } K = (s-1) - \sum_{k=0}^{s-2} (s-1-k)\alpha_k$$

5 Cas limites particuliers

- Si $\delta \rightarrow 0$, le système considéré n'est rien d'autre que le système classique GI/M/s sans arrivées négatives.
- Lorsque $s \rightarrow \infty$; $\alpha \rightarrow C(\beta_0, \beta_1, \dots, \beta_s, \dots)$, et on ne peut pas généraliser. Le problème est de déterminer les β_j et ceux-ci dépendent des P_{ij} définis uniquement dans les cas :
 - (1) $j > i + 1$ où $P_{ij} = 0$
 - (2) $j \leq i + 1 \leq s$ (voir b)) alors que la formalisation établie dans la section 2 avec σ portait sur le cas
 - (3) $s \leq j \leq i + 1$ où $P_{ij} = d_{i+1-j}$

Références

- Artalejo, J. R. (1994). G-networks: A versatile approach for work removal in queueing networks works. *European Journal of Operational Research* 15, 401–417.
- Artalejo, J. R. et A. Gomez-Corral (1996). On a single server queue with negative arrivals and request repeated. *Journal of applied Probability* 36, 907–918.
- Artalejo, J. R. et A. Gomez-Corral (1999). Computation of the limiting distribution in queueing systems with repeated attempts and disasters. *RAIRO Operations Research* 33, 371–382.
- Gelenbe, E. (1989). Random neural network with negative and positive signals and product form solution. *Neural Computation* 1, 502–510.
- Harrison, P. G. et E. Pitel (1993). Sojourn times in single-server queues with negative customers. *Advanced in Applied Probability* 28, 540–546.
- Harrison, P. G. et E. Pitel (1996). The M/G/1 queue with negative customers. *Journal of Applied Probability* 30, 943–963.
- Kyung, C. C., H. M. Park, et W. S. Yang (2010). A GI/Geo/1 queue with negative and positive customers. *Applied Mathematical Modelling* 34, 1662–1671.
- Yang, W. S. et K. C. Chae (2001). A note on the GI/M/1 queue with poisson negative arrivals. *Journal of Applied Probability* 38, 1081–1085.

Summary

In this work, we consider the stochastic analysis of a multiserver GI/M/s queue with negative arrivals under the RCE discipline using the imbedded Markov chain. We carried out the existence of stationary regime of the embedded Markov chain. We also derived some performance measures.

Analyse stationnaire vaste du modèle d'attente à un seul serveur avec rappels classiques et feedback

Mohamed Boualem*, Mouloud Cherfaoui*,**
Natalia Djellab*** Djamil Aissani*

*Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
robertt15dz@yahoo.fr, lamos_bejaia@hotmail.com
<http://www.lamos.org>

**Département de Mathématiques, Université de Biskra, 07000 Algérie
mouloudcherfaoui2013@gmail.com

***Département de Mathématiques, Université de Annaba, 23000 Algérie
djellab@yahoo.fr

Résumé. L'objectif de ce papier est de faire une analyse stationnaire vaste, du système d'attente à un seul serveur avec rappels classiques et feedback, en utilisant la technique de la chaîne de Markov induite. Pour cela, nous avons dérivé certaines mesures de performance, importantes et intéressantes, pour le système considéré, telles que: le nombre moyen de clients dans le système et en orbite, le temps moyen d'attente et le nombre de rappels ainsi que le nombre moyen de clients qui demandent un service supplémentaire. Des exemples numériques sont donnés pour illustrer les résultats théoriques obtenus.

1 Introduction

La majorité des études sur les systèmes d'attente avec rappels considère le modèle sans feedback. Néanmoins, plusieurs situations réelles peuvent être modélisées comme des systèmes de files d'attente avec rappels et feedback. La notion de feedback a été initialement introduite par (Takas, 1963) pour l'étude de certains systèmes d'attente classiques, et depuis plusieurs papiers sont apparus sur ce sujet en considérant d'autres types de systèmes avec différentes variantes (Krishna Kumar et al., 2010). Le phénomène de feedback dans les systèmes d'attente avec rappels peut apparaître dans plusieurs situations pratiques. Par exemple, dans les systèmes MATS (*Multiple Access Telecommunication Systems*) où des messages s'avérant comme erreurs à la destination (messages perdus ou corrompus) sont renvoyés. En particulier, les systèmes d'attente avec rappels et feedback sont utilisés pour modéliser le protocole ARQ (*Automatic Repeat Request*) dans un réseau de communication à haute fréquence. Le modèle $M/G/1$ avec rappels et feedback a été étudié dans (Djellab, 2005), l'auteur a obtenu la distribution du nombre de clients dans le système à l'état stationnaire, en utilisant la chaîne de Markov auxiliaire. Pour le même modèle, (Boualem et al., 2012) ont obtenu la propriété de la décomposition stochastique du nombre de clients dans le système au régime stationnaire ainsi que différents résultats de monotonie, en utilisant l'analyse qualitative.

Notre objectif est focalisé sur l'analyse stationnaire vaste du système d'attente $M/G/1$ avec rappels classiques et Bernoulli feedback. Pour cela, en exploitant le résultat obtenu dans (Boualem et al., 2012), concernant l'expression de la fonction génératrice de la distribution stationnaire du nombre de clients dans le système aux instants de départs, nous avons déterminer certaines performances, à savoir : le nombre moyen de clients dans le système, le nombre moyen de clients dans en orbite, le temps moyen d'attente, le nombre de rappels et nombre moyen de clients qui demandent un service supplémentaire. Des exemples numériques sont donnés pour confirmer les résultats théoriques obtenus.

2 Le modèle d'attente

Les clients primaires arrivent suivant un flux poissonnien de taux λ . Un client qui arrive et trouve le serveur occupé, quitte la zone du service pour rejoindre l'orbite (clients bloqués). Après un certain temps aléatoire, il renouvelle sa tentative d'entrer en service jusqu'à ce qu'il le trouve disponible. Une fois servi, le client doit décider, soit de rejoindre l'orbite pour un autre service avec une probabilité c où de quitter le système définitivement avec une probabilité $\bar{c}$. Les intervalles de temps inter-rappels suivent une distribution exponentielle de taux $k\mu$ ($k \in \mathbb{N}$) (rappels classiques) (Boualem et al., 2011). La durée de service τ est de loi générale, de fonction de distribution $B(x)$, de TLS $\tilde{B}(s)$ et des deux premiers moments finis β_1 et β_2 , respectivement. Toutes les variables aléatoires introduites sont mutuellement indépendantes.

L'état du système est décrit par le processus $X(t) = (C(t), N_o(t), \xi(t))_{t \geq 0}$, où $C(t) = 1$ (si le serveur est oisif) ou 0 (si le serveur est occupé), $N_o(t)$: le nombre de clients en orbite à l'instant t et $\xi(t)$ représente le temps de service écoulé du client en service à l'instant t , si $C(t) = 1$.

Soit $\{t_n, n \in \mathbb{N}\}$ une suite d'instant de la complétion d'un service. La suite de variables aléatoires $Y_n = \{q_n = N_o(t_n^+), n \in \mathbb{N}\}$ forme une chaîne de Markov induite, dont l'équation fondamentale est $q_{n+1} = q_n - \delta_{q_n} + v_{n+1} + \eta$, où,

- v_{n+1} : représente le nombre de clients qui arrivent pendant un temps de service qui se termine à l'instant t_{n+1} . Sa distribution est donnée par : $k_i = P(v_{n+1} = i) = \int_0^\infty \frac{(\lambda x)^i}{i!} e^{-\lambda x} dB(x)$, avec la fonction génératrice $K(z) = \sum_{i \geq 0} k_i z^i = \tilde{B}(\lambda - \lambda z)$.

- $\delta_{q_n} = \begin{cases} 1, & \text{si le } (n+1)^{\text{ème}} \text{ client provient de l'orbite,} \\ 0, & \text{sinon.} \end{cases}$

De plus, $P(\delta_{q_n} = 1 | q_n = k) = \frac{k\mu}{\lambda + k\mu}$, et $P(\delta_{q_n} = 0 | q_n = k) = \frac{\lambda}{\lambda + k\mu}$.

- La distribution de la variable aléatoire η est : $P[\eta = 1] = c$ et $P[\eta = 0] = \bar{c} = 1 - c$.

Sous la condition d'ergodicité $\rho = \lambda\beta_1 + c < 1$, l'expression de la fonction génératrice de la distribution stationnaire du nombre de clients dans le système aux instants de départs est donnée par (Boualem et al., 2012) :

$$\begin{aligned} \pi(z) &= \left[\frac{(1 - \rho)\tilde{B}(\lambda - \lambda z)(1 - z)}{(\bar{c} + cz)\tilde{B}(\lambda - \lambda z) - z} \right] \\ &\times \left[(\bar{c} + cz) \exp \left(\frac{\lambda}{\mu} \int_1^z \frac{(1 - \tilde{B}(\lambda - \lambda u)(\bar{c} + cu))}{(\tilde{B}(\lambda - \lambda u)(\bar{c} + cu) - u)} du \right) \right]. \end{aligned} \quad (1)$$

3 Mesures de performance

En exploitant le résultat précédent (voir l'équation (1)), on obtient :

Nombre moyen de clients dans le système :

$$\bar{n} = \pi'(1) = \rho + c(1 - \lambda\beta_1) + \frac{\bar{c}\lambda\rho}{\mu(1 - \rho)} + \frac{c\lambda\beta_1}{(1 - \rho)} + \frac{\lambda^2\beta_2}{2(1 - \rho)}. \quad (2)$$

Nombre moyen de clients en orbite :

$$\bar{n}_0 = \bar{n} - \rho = c(1 - \lambda\beta_1) + \frac{\bar{c}\lambda\rho}{\mu(1 - \rho)} + \frac{c\lambda\beta_1}{(1 - \rho)} + \frac{\lambda^2\beta_2}{2(1 - \rho)}. \quad (3)$$

Temps moyen d'attente :

$$\bar{w} = \frac{\bar{n}_0}{\lambda} = \frac{c(1 - \lambda\beta_1)}{\lambda} + \frac{\bar{c}\rho}{\mu(1 - \rho)} + \frac{c\beta_1}{(1 - \rho)} + \frac{\lambda\beta_2}{2(1 - \rho)}. \quad (4)$$

Nombre moyen de rappels par client :

$$\bar{d} = \mu\bar{w} = \frac{c\mu(1 - \lambda\beta_1)}{\lambda} + \frac{\bar{c}\rho}{(1 - \rho)} + \frac{c\mu\beta_1}{(1 - \rho)} + \frac{\lambda\mu\beta_2}{2(1 - \rho)}. \quad (5)$$

Nombre moyen de clients qui demandent un service supplémentaire :

$$\bar{F} = c\pi'(1) = c\rho + c^2(1 - \lambda\beta_1) + \frac{c\bar{c}\lambda\rho}{\mu(1 - \rho)} + \frac{c^2\lambda\beta_1}{(1 - \rho)} + \frac{\lambda^2c\beta_2}{2(1 - \rho)}. \quad (6)$$

4 Aspect pratique

Afin de confirmer les résultats analytiques (voir la Section 3), nous avons élaboré un simulateur, sous environnement Matlab, qui reproduira le comportement du modèle considéré et de faire les comparaisons entre les caractéristiques estimées à celles obtenues analytiquement. Pour cela, on considère les trois cas suivants, dont les résultats sont résumés dans le tableau 1 :

1° **cas** : on considère la file $M/M/1$ avec rappels et feedback avec les paramètres suivants :

$\lambda = 2, \beta_1 = 1/4, \mu = 2$ et $c = 0, 3$.

2° **cas** : on considère la durée de temps de service suit une distribution de Cox d'ordre 2, ayant les paramètres $\nu_1 = 4, \nu_2 = 7, 5$ et $\gamma = 0, 05$. De plus, on fixe $\lambda = 2, \mu = 2$ et $c = 0, 3$.

3° **cas** : on considère le cas particulier $c = 0$ (pas de demande d'un service supplémentaire); avec les paramètres : $\lambda = 2, \beta_1 = 1/4$ et $\mu = 5$.

On constate que toutes les statistiques de Student calculées (voir le tableau 1), pour un seuil $\alpha = 0, 05$, sont inférieures à $t_{(0,025; 99)} = 1, 96$. Du ce fait, on ne rejette pas H_0 au seuil $\alpha = 0, 05$, ce qui signifie que les résultats théoriques ne sont pas contradictoires avec ceux de simulation au seuil $\alpha = 0, 05$.

En guise de conclusion, les résultats obtenus dans ce travail sont considérés comme une extension des résultats obtenus par Takas (Takas, 1963).

	Caractéristiques	$\bar{n}$	$\bar{n}_0$	$\bar{w}$	$\bar{d}$	$\bar{F}$
1 ^{ier} cas	Résultats théoriques	3,725	2,925	2,1 625	4,325	1,117
	Résultats de simulation	3,731	2,931	1,9 716	3,9 432	1,119
	La statistique $t_{(n,\alpha)}$	0,363	0,408	0,628	0,973	0,082
2 ^{ième} cas	Résultats théoriques	4,077	3,266	0,726	1,452	1,223
	Résultats de simulation	4,060	3,246	0,835	1,669	1,218
	La statistique $t_{(n,\alpha)}$	0,804	0,927	0,931	1,872	0,236
3 ^{ième} cas	Résultats théoriques	1,250	0,750	0,325	1,625	—
	Résultats de simulation	1,261	0,761	0,3 663	1,206	—
	La statistique $t_{(n,\alpha)}$	1,930	1,890	1,035	1,734	—

TAB. 1 – Comparaison des résultats de simulation avec les résultats théoriques.

Références

- Boualem, M., N. Djellab, et D. Aïssani (2011). Approche régénérative de la file d'attente M/G/1 avec rappels classiques et vacances exhaustives du serveur. *Journal Européen des Systèmes Automatisés* 45, 253–267.
- Boualem, M., N. Djellab, et D. Aïssani (2012). Stochastic approximations and monotonicity of a single server feedback retrial queue. *Mathematical Problems in Engineering* 2012, 1–13.
- Djellab, N. (2005). On the M/G/1 retrial queue with feedback. *Proceedings of the International Conference on Mathematical Methods of Optimization of Telecommunication Networks, Minsk, Byelorussia*, 32–35.
- Krishna Kumar, B., G. Vijayalakshmi, A. Krishnamoorthy, et S. S. Basha (2010). A single server feedback retrial queue with collisions. *Computer and Operations Research* 37, 1247–1255.
- Takas, I. (1963). A single server queue with feedback. *Bell System technical Journal* 42, 505–519.

Summary

In this paper, we consider the embedded Markov chain to study the performance measures of a single server retrial queues and feedback, using steady state analysis. Some interesting and important performance measures of the system are obtained. Finally, numerical illustrations are provided.

Estimation à noyau d'une matrice de transition inconnue du modèle d'attente $M/GI/1/N$

Mouloud Cherfaoui*,**, Mohamed Boualem**
Djamil Aïssani**, Smail Adjabi**

*Département de Mathématiques, Université de Biskra, 07000 Algérie
mouloudcherfaoui2013@gmail.com

**Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
robertt15dz@yahoo.fr, lamos_bejaia@hotmail.com, adjabi@hotmail.com
<http://www.lamos.org>

Résumé. Le papier se situe dans le cadre de l'estimation à noyau d'une matrice de transition inconnue d'une chaîne de Markov associée au modèle d'attente $M/GI/1/N$. L'objectif est de comparer, par une étude de simulation, quatre procédures de choix de la fenêtre de lissage (intervenant dans l'estimation de la matrice de transition). L'estimateur du paramètre de lissage choisi, par la minimisation d'une certaine norme matricielle, fournit des meilleurs résultats en terme de la vitesse de convergence et d'erreurs.

1 Introduction

Tout au long de l'histoire, nous disposons de plusieurs méthodes de l'estimation paramétrique d'une matrice de transition associée à une chaîne de Markov (la méthode du maximum de vraisemblance, la méthode de régression matricielle, etc.). Ces méthodes ont l'avantage d'être simple à utiliser, mais il est bien difficile d'estimer avec précision des matrices de transition modélisant des phénomènes complexes. Pour palier à ces situations difficiles, on fait souvent appel aux méthodes d'estimation non paramétrique. (Roussas, 1969) est le premier à aborder la question de l'estimation fonctionnelle en utilisant des observations markoviennes. L'auteur a établi la convergence de l'estimateur à noyau de la densité de la loi de transition au sens de L^2 . Les résultats obtenus par ce dernier ont été complétés dans (Basu et Sahoo, 1998). Plus tard, (Laksaci et Yousfate, 2002) s'intéressent à la construction et aux propriétés d'un estimateur fonctionnel de la densité de l'opérateur de transition. Récemment, plusieurs auteurs ont tendance à prouver l'efficacité d'utilisation de l'estimateur à noyau dans la théorie des file d'attente. (Bareche et Aïssani, 2008) ont évalué l'approximité des systèmes $GI/M/1$ et $M/M/1$ lorsque la densité des temps inter-arrivées est estimée par la méthode du noyau.

Dans ce travail on s'intéresse à l'estimation de la matrice de transition inconnue $\mathbb{P} = (P_{ij})$, associée au modèle d'attente $M/GI/1/N$, par la méthode de noyau. Essentiellement, notre étude se focalise sur la qualité de cet estimateur par rapport aux procédures de sélection du paramètre de lissage. Pour ce faire, nous allons effectuer une étude comparative des résultats

obtenus par les méthodes classiques de sélection du paramètre de lissage et d'autres méthodes qui se basent sur les normes matricielles. Contrairement aux méthodes classiques, ces dernières prennent en considération la loi des temps des inter-arrivées, qui joue le rôle d'une pondération de la distribution générale des temps de service dans les éléments de la matrice de transition considéré.

2 L'estimateur à noyau

Soient $T_1, T_2, \dots, T_n$ un n -échantillon, issu de la variable aléatoire positive T , d'une densité de probabilité inconnue g et d'une fonction de répartition G . Parmi les estimateurs proposés pour cette famille de distributions, on cite l'estimateur de (Chen, 2000) qui suggère de remplacer l'estimateur classique de (voir (Parzen, 1962) et (Rosenblatt, 1956)) par :

$$g_h(t) = \frac{1}{n} \sum_{i=1}^n \frac{T_i^{(t/h)} e^{-(T_i/h)}}{h^{(t/h)+1} \Gamma((t/h) + 1)}, \quad (1)$$

où, h est le paramètre de lissage.

La décision du choix de ce paramètre de lissage repose sur la spécification d'un critère d'erreur qui puisse être optimisé. Cette optimalité est étroitement liée au choix d'un critère qui fait intervenir à la fois la densité inconnue g et l'estimateur g_h . Concernant l'estimateur donné par la formule (1), le paramètre de lissage optimal au sens du *MISE* a été élaboré par (Chen, 2000) dont sa forme explicite est donnée par :

$$h_1^* = \left[\frac{1}{2\sqrt{\pi}} \int_0^\infty t^{-1/2} g(t) dt \right]^{2/5} \left[\int_0^\infty \left\{ g'(t) + \frac{1}{2} t g''(t) \right\}^2 dt \right]^{-2/5} 4^{-2/5} n^{-2/5}. \quad (2)$$

3 Modèle d'attente $M/GI/1$ ($FIFO, N$)

Considérons le modèle d'attente $M/GI/1$ à capacité finie (N) où les clients sont servis selon la discipline *FIFO*. Les temps d'inter-arrivées des clients sont indépendants et identiquement distribués selon une loi exponentielle de paramètre λ . Les temps de service est une suite de variables aléatoires distribuées selon une loi générale commune G de moyenne $1/\mu$. On suppose que toutes les variables introduites sont mutuellement indépendantes.

Soit $X_n = X(t_n)$ la chaîne de Markov définie aux instants de départ du $n^{\text{ème}}$ client et qui représente le nombre de clients dans le système juste après le départ du $n^{\text{ème}}$ client. Il est clair que X_n est une chaîne de Markov induite d'espace d'état $\{0, 1, \dots, N-1\}$ et dont l'équation fondamentale est donnée par

$$X_n = \max(X_{n-1} - 1, 0) + A_n, \quad (3)$$

où A_n représente le nombre d'arrivées entre les instants de deux départs consécutifs. L'opérateur de transition $\mathbb{P} = (P_{ij})$ de la chaîne de Markov X_n s'écrit comme suit :

$$P_{ij} = \begin{cases} \int_0^\infty \frac{e^{-\lambda t} (\lambda t)^j}{j!} g(t) dt, & \text{si } i = 0 \text{ et } 0 \leq j+1 \leq N-1, \\ \int_0^\infty \frac{e^{-\lambda t} (\lambda t)^{j-i+1}}{(j-i+1)!} g(t) dt, & \text{si } 1 \leq i \leq j+1 \leq N-1, \\ 1 - \sum_{j=0}^{N-2} P_{ij}, & \text{si } j = N-1, \\ 0, & \text{sinon.} \end{cases} \quad (4)$$

4 Estimateur $\hat{\mathbb{P}}$ et choix du h

Dans cette Section, on s'intéresse au choix du paramètre de lissage dans l'estimation de la matrice de transition de la chaîne de Markov induite, donnée par l'équation (4), qui minimise un certain critère d'erreur. Pour ce faire, on suppose qu'on dispose d'un n -échantillon $T_1, T_2, \dots, T_n$ qui représentent les temps de service ayant comme densité de probabilité inconnue g . En outre, d'une part, l'estimation de la matrice de transition, notée $\mathbb{P} = (P_{ij})$, consiste à estimer la densité inconnue g et substituer son estimateur, g_h , dans la matrice de transition. Dans notre cas, on opte pour l'estimateur à noyau Gamma donné dans l'expression (1) et le paramètre de lissage, h_1^* , qui minimise le critère *MISE* donné par la formule (2). D'autre part, nous proposons l'utilisation des normes matricielles afin de prendre en considération les pondérations qui ont un impact sur la qualité de l'estimateur $\hat{\mathbb{P}} = (\hat{P}_{ij})$. En effet, l'utilisation de ces dernières nous permettent d'inclure les pondérations, qui sont d'une loi de Poisson du paramètre $\lambda * t$, de la quantité $g(t)$ dans l'expression des P_{ij} lors de l'estimation. Dans ce cas, le paramètre de lissage optimal est calculé selon l'une des trois normes suivantes :

$$h_2^* = \arg \min_h \|\hat{\mathbb{P}} - \mathbb{P}\|_1 = \arg \min_h \left[\max_j \left(\sum_{i=0}^N |\hat{P}_{ij} - P_{ij}| \right) \right]. \quad (5)$$

$$h_3^* = \arg \min_h \|\hat{\mathbb{P}} - \mathbb{P}\|_2 = \arg \min_h \left[\sum_{i=0}^N \sum_{j=0}^N (\hat{P}_{ij} - P_{ij})^2 \right]^{1/2}. \quad (6)$$

$$h_4^* = \arg \min_h \|\hat{\mathbb{P}} - \mathbb{P}\|_\infty = \arg \min_h \left[\max_i \left(\sum_{j=0}^N |\hat{P}_{ij} - P_{ij}| \right) \right]. \quad (7)$$

Où, $\hat{P}_{ij}$ est l'estimateur de P_{ij} lorsque on remplace $g(t)$ par son estimateur $g_h(t)$.

5 Application numérique et discussions

Les étapes de l'algorithme de simulation, conçu pour répondre à notre objectif, sont :

Étape 1 : Fixer les paramètres du système (loi des temps de service, λ, μ, N),

Étape 2 : Générer un échantillon de taille n de la loi des temps de service,

Étape 3 : Estimer h par la formule (2), calculer $\hat{\mathbb{P}}$ et l'erreur associée,

Étape 4 : Estimer h par les formules (5)–(7), calculer $\hat{\mathbb{P}}$ et les erreurs associées.

Pour analyser l'impact de variation de taux d'arrivées λ (poids de g dans les P_{ij}) et la taille de l'échantillon n sur la qualité de l'estimateur $\hat{\mathbb{P}}$, lorsque h est calculé par l'une des formules (2)–(7), dans le système $M/GI/1/N$. Nous considérons les situations suivantes : $g(t) = \mu \exp(-\mu t)$ avec $\mu = 2/3$, $N = 10$, $\lambda = [0.2, 0.4, 0.6, 0.8, 1.0]$ et $n = [100, 200, 300]$. À l'aide des étapes de l'algorithme précédent et en utilisant le premier noyau gamma (donné par (1)), les résultats obtenus nous permet de constater :

- Tous les estimateurs considérés convergent en fonction de la taille de l'échantillon (n). De plus, la vitesse de convergence des estimateurs de $\mathbb{P}$, au sens des normes considérées, est plus considérable pour le h_4^* , mais elle est plus faible pour h_3^* .
- L'estimateur de $\mathbb{P}$ dépend du taux des arrivées, lorsque le paramètre de lissage est estimé par la norme $\|\cdot\|_\infty$, contrairement au cas où il est estimé par la norme $\|\cdot\|_2$.

Il est préférable dans l'estimation d'une matrice de transition, par la méthode du noyau, d'utiliser les normes matricielles, pour le choix du paramètre de lissage, que les méthodes classiques de sélection. Les résultats obtenus peuvent être porteurs de nombreuses applications.

Références

- Bareche, A. et D. Aïssani (2008). Kernel density in the study of the strong stability of the M/M/1 queueing system. *Operations Research Letters* 36, 535–538.
- Basu, A. K. et D. Sahoo (1998). On berry–esseen theorem for nonparametric density estimation in markov sequences. *Bull. Inform. Cybernet.* 30, 25–39.
- Chen, S. (2000). Probability density functions estimation using gamma kernels. *Annals of the Institute of Statistical Mathematics* 52, 471–480.
- Laksaci, A. et A. Yousfate (2002). Estimation fonctionnelle de la densité de l'opérateur de transition d'un processus de markov à temps discret. *C. R. Acad. Sci. Paris, Ser. I* 334, 1035–1038.
- Parzen, E. (1962). On estimation of a probability density function and mode. *Ann. Math. Stat.* 33, 1065–1076.
- Rosenblatt, M. (1956). Remarks in some nonparametric estimates of a density function. *Ann. Math. Stat.* 27, 832–837.
- Roussas, G. G. (1969). Nonparametric estimation in markov processes. *Ann. Inst. Statist. Math.* 21, 73–87.

Summary

The paper is situated in the framework of the bandwidth choice in Kernel estimation of transition matrix of an embedded Markov chain associate to an $M/GI/1/N$ queue. The aim is to compare, by a simulation study, four selection procedures of the smoothing parameter (involved in the estimation of the transition matrix). The smoothing parameter chosen by minimizing a certain matrix norm provides the best results in terms of errors and convergence speed.

Estimation d'un modèle GARCH asymétrique (AGARCH) en présence de données de haute fréquence

Hafida Guerbyenne *, Leïla Messahli**

*Université des Sciences et de la Technologie, U. S. T. H. B.

**Ecole Nationale Supérieure Agronomique, E.N.S.A.

Résumé. Nous nous intéressons à l'estimation des paramètres d'un modèle *AGARCH*, en présence de données de haute fréquence. La fonction de vraisemblance est basée sur des proxies de volatilité H_n qui exploitent l'information contenue dans le processus infra-journalier. Dans notre étude, nous nous limitons à un modèle *AGARCH*(1, 1). Une étude de simulation intensive a été entreprise. Elle montre un gain certain en efficacité dans l'estimation des paramètres. Une application sur une série de données réelles est proposée.

1 Introduction

L'asymétrie de la volatilité par rapport aux signes des chocs passés est connue depuis les articles de Black (1976) et Christie (1982). Ces articles expliquent par l'effet de levier l'impact plus fort sur la volatilité des baisses des cours que des hausses de même ampleur. L'asymétrie a été mise en évidence dans de nombreuses études empiriques. En 1986, Taylor a constaté le phénomène de la corrélation sérielle des rendements en valeur absolue. Par suite, ce même auteur et Schwert (1989), ont introduit le modèle *GARCH* asymétrique dans sa version la plus élémentaire *i.e.*, ils décrivent l'écart-type standard conditionnel des rendements journaliers comme une fonction linéaire des rendements en valeur absolue. Ces travaux ont été généralisés par Ding et al. (1993) à une classe des modèles dite *GARCH* en puissance asymétrique (*AGARCH*).

Par essence, la volatilité est inhérente à la plupart des modèles financiers. Comme elle ne s'observe pas naturellement, alors, il est nécessaire de construire un certain nombre de statistiques H_n dites *proxies* sur lesquelles on se base lors de spécification, d'estimation et d'évaluation des modèles de volatilité.

Mesurer la volatilité à une fréquence très élevée, de façon presque continue en utilisant des données infra-journalières est possible. L'idée de la volatilité réalisée est initialement proposée par French, Schwert et Stambaugh (1987). Andersen et Bollerslev (1997) montrent que la volatilité journalière réalisée peut être construite en additionnant les rendements infra-journaliers au carré, en prenant la racine carrée de cette dernière. Visser (2008) a incorporé des données de haute fréquence dans le modèle (journalier) en temps discret *GARCH* (1, 1).

Nous nous intéressons à l'estimation des paramètres d'un modèle *AGARCH*, en présence de données de haute fréquence. Les propriétés probabilistes du modèle *AGARCH* classique ont été étudiées par Straumann et Mikosch (2006). nous développons une théorie sur l'estimation de ce modèle fondée sur la méthode du quasi maximum de vraisemblance (QMV). L'idée

est de construire la fonction de vraisemblance en se basant sur des proxies H_n , qui exploitent l'information contenue dans le processus infra-journalier (modèle d'échelle). Dans notre étude, nous nous limitons à un modèle $AGARCH(1, 1)$. Une étude de simulation et une application sur une série de données réelles ont été faites. Elles confortent les résultats théoriques.

2 Présentation du modèle

Le modèle $AGARCH(p, q)$ proposé indépendamment par Ding et al. (1993) et Zakoian (1994) est défini par

$$X_n = \sigma_n \eta_n$$

$$\sigma_n^2 = \alpha_0 + \sum_{i=1}^p \alpha_i (|X_{n-i}| - \gamma X_{n-i})^2 + \sum_{j=1}^q \beta_j \sigma_{n-j}^2, \quad (2.1)$$

où $\alpha_0 > 0$, $\alpha_i, \beta_j \geq 0$ et $|\gamma| \leq 1$. Cette dernière condition garantit l'identifiabilité du modèle.

3 Représentation markovienne et stationnarité

La représentation markovienne $Y_{n+1} = A_n Y_n + B_n$ du modèle a permis à Straumann et Mikosch (2006) d'établir, à partir du plus grand exposant de Lyapounov, les propriétés de stationnarité du processus $AGARCH$.

$$Y_n = (\sigma_n^2, \sigma_{n-1}^2, \dots, \sigma_{n-q+1}^2, (|X_{n-1}| - \gamma X_{n-1})^2, \dots, (|X_{n-p+1}| - \gamma X_{n-p+1})^2)'$$

Un calcul simple permet d'obtenir A_n et B_n .

3.1 Existence d'une solution strictement stationnaire

L'équation (3.1) admet une solution unique, strictement stationnaire et ergodique de la forme :

$$Y_n = \sum_{k=0}^{\infty} \prod_{l=1}^k A_{n-l} B_{n-k}, \quad (3.1)$$

où la série (3.4) converge presque sûrement, pourvu que le plus grand exposant de Lyapounov ρ

$$\rho = \lim_{t \rightarrow \infty} \frac{1}{t+1} \mathbb{E}(\log \|A_n \dots A_{n+1}\|) < 0, \quad (3.2)$$

soit strictement négatif. Avec la convention $\prod_{l=1}^0 \cdot \equiv 1$, $\|\cdot\|$ désigne la norme matricielle.

3.2 Existence d'une solution stationnaire au second ordre

Une condition suffisante de stationnarité au second ordre est que $\sum_{i=1}^p \alpha_i \mathbb{E}(|\eta_{n-1}| - \gamma \eta_{n-1})^2 + \sum_{j=1}^q \beta_j < 1$, cette condition implique que $\mathbb{E}(X_0) < \infty$. Par contre, $\sum_{j=1}^q \beta_j < 1$, est une condition nécessaire de stationnarité du modèle $AGARCH$.

4 Estimation du modèle en présence de données de haute fréquence

Soit $\{X_n, n \in \mathbb{N}\}$, un processus *AGARCH* (p, q)

Le paramètre $\theta_0 = (\alpha_0^0, \alpha_1^0, \dots, \alpha_p^0, \beta_1^0, \dots, \beta_q^0, \gamma)'$ est tel que

1. Le processus est strictement stationnaire.
2. $\alpha_0^0 > 0, \alpha_1^0 > 0$ pour au moins un indice $i \geq 1$ et $(\alpha_p^0, \beta_q^0) \neq (0, 0)$.
3. Les polynômes $a^0(z) = \sum_{i=1}^p \alpha_i^0 z^i$ et $b^0(z) = 1 - \sum_{j=1}^q \beta_j^0 z^j$ n'ont pas de racine commune.
4. De plus la distribution de η_0 n'est pas concentrée en deux points.
5. Soit $\Theta \subset]0, \infty[\times [0, \infty[\times B \times [-1, 1]$ un espace compact et $\theta_0 \in \Theta$, où $B = \{(\beta_1, \dots, \beta_q)' \in [0, 1]^q \mid \sum_{j=1}^q \beta_j < 1\}$.

Alors, l'estimateur du QMV est fortement consistant et asymptotiquement gaussien si, de plus, $E(\eta_0^4) < \infty$. (Straumann et Mikosch (2006)).

Nous proposons le modèle suivant

$$R_n(u) = \sigma_n \Psi_n(u), 0 \leq u \leq 1, \quad (4.1a)$$

$$\sigma_n^2 = \alpha_0 + \alpha_1(|X_{n-1}| - \gamma X_{n-1})^2 + \beta_1 \sigma_{n-1}^2, \quad (4.1b)$$

pour le processus des prix infra-journaliers où, $R_n(\cdot)$ est le processus des rendements. La suite de processus $(\Psi_n(\cdot))_n$ est supposée *i.i.d.*, et à trajectoires *càdlàg*, tel que $\mathbb{E}(\Psi_n(1)^2) = 1$. $R_n(0)$ et $R_n(1)$ donnent, respectivement, le rendement *overnight* et le rendement *close-to-close* X_n , respectivement.

La journée ouvrable est représentée par l'intervalle de temps unité $[0, 1]$. Pour $\gamma = 0$, nous retrouvons le modèle proposé par Visser (2008).

Une variable aléatoire $H_n = H(R_n)$ est une *proxy*, si la fonctionnelle H est positive et possède la propriété d'homogénéité positive en R_n . Nous supposons que la variable aléatoire $H(\Psi)$ n'est pas identiquement nulle ; donc $\mu_2^H = \sqrt{\mathbb{E}(H^2(\Psi))} > 0$. Notons η_H , les innovations normalisées *i.e.*, $\eta_H \equiv H(\Psi) / \mu_2^H < \infty$. Ainsi, $\mathbb{E}(\eta_H^2) = 1$.

$$H_n = H(R_n) = \sigma_n H(\Psi_n). \quad (4.2)$$

ou encore $H_n = \sigma_{n,H} \eta_{H,n}$ avec $\sigma_{n,H} = \sigma_n \mu_2^H$. Alors, (4.1) devient

$$H_n = \sigma_{H,n} \eta_{H,n}, \quad (4.3a)$$

$$\sigma_n^2 = \alpha_0 + \alpha_1(|X_{n-1}| - \gamma X_{n-1})^2 + \beta_1 \sigma_{n-1}^2, \quad (4.3b)$$

et $\mathbb{E}(H_n^2 \mid \text{tciFourier}_{n-1}) = \sigma_{H,n}^2$, où *tciFourier*_{*n*-1} est l'information dont on dispose jusqu'à l'instant $n - 1$.

Étant donné des observations $(y_1, y_2, \dots, y_N)$ issues du modèle (4.3) ($y_n = \text{sgn}_n H_n$ où sgn_n est une variable aléatoire uniforme sur $\{-1, +1\}$) et sous les conditions 1 – 5 nous montrons la consistance et la normalité asymptotique des estimateurs du QMV gaussien si de plus $E(\eta_{H,0}^4) < \infty$,

5 Simulation et application

Une étude de simulation a été entreprise respectivement pour différentes tailles d'échantillon. Les estimations des paramètres basées sur la *proxy* de volatilité (variance réalisée) sont meilleures au sens du critère de l'erreur quadratique moyenne (*RMSE*). Une application sur une série de données réelles est proposée.

Références

- Andersen, T. et T. Bollerslev (1997). Intraday periodicity and volatility persistence in financial markets. *Journal of Empirical Finance* 4, Issues 2-3, 69–293.
- Ding, Z., C. W. J. Granger, et R. Engle (1993). A long memory property of stock market returns and a new model. *Journal of Empirical Finance* 1, 83–106.
- Straumann, D. et T. Mikosch (2006). Quasi-maximum-likelihood estimation in conditionally heteroscedastic time series : a stochastic recurrence equations approach. *The Annals of Statistics* 34, number 5, 2449–2495.
- Visser, M. (2008). Garch parameter estimation using high-frequency data. *MPRA paper*, 9076.

Summary

We estimate the parameters of an *AGARCH* model using intraday high frequency data in addition to the daily returns. The likelihood function is based on volatility proxies H_n which uses the intraday information. In this work, we are interested in an *AGARCH*(1, 1) model. A simulation study shows that the parameter estimation is improved. An application on a real time-series data is proposed.

On Probabilistic properties of a Power Periodic Threshold GARCH Model

Hafida Guerbyenne *, Abderrahim Kessira *

* Université des Sciences et de la Technologie, U. S. T. H. B.
BP 32 El Alia, Bab Ezzouar, 16111, Algiers, Algeria

Résumé. Un modèle *GARCH* en puissance δ , périodique et à seuil est défini. L'existence d'une solution périodiquement strictement stationnaire, périodiquement stationnaire au second ordre, périodiquement géométriquement ergodique and β -mélangeante est établie. Nous donnons une condition nécessaire et suffisante pour l'existence des moments d'ordre supérieur. Ces propriétés sont importantes pour l'inférence statistique. L'analyse est fondée sur la théorie des chaînes de Markov à espace d'état général.

1 Introduction

The Power Periodic Threshold *GARCH*(1, 1) process (*PP* – *TGARCH*(1, 1)) is a class of *GARCH*-type models which features periodicity, regime effects and asymmetry in the conditional variance.

To the best of our knowledge, it seems that the probabilistic properties of the *PP* – *TGARCH*(1, 1) has not yet been studied. Our aim is to find conditions under which the given process is strictly stationary, second order stationary, with finite higher order moments, geometrically ergodic and β -mixing with exponential decay.

2 The *PP* – *TGARCH*(1, 1) Model

The Power Periodic Threshold *GARCH*(1, 1) process is a class of *GARCH*-type models which features periodicity and regime effects in the conditional variance.

Define the *PP* – *TGARCH*(1, 1) by :

$$\varepsilon_{s+S\tau} = \sqrt{h_{s+S\tau}} \eta_{s+S\tau} \quad (1)$$

$$h_{s+S\tau}^\delta = \omega_s + (\alpha_s + \gamma_s I_{\{\varepsilon_{s-1+S\tau} \leq 0\}}) \varepsilon_{s-1+S\tau}^{2\delta} + \beta_s h_{s-1+S\tau}^\delta, \delta > 0 \quad (2)$$

where $(\eta_t)_{t \in \mathbb{Z}}$ is a sequence of independent and identically distributed (*i.i.d*) random variables with mean 0 and finite 2δ -th moment. $\varepsilon_{s+S\tau}$ refers to the observation during season s , $1 \leq s \leq S$, in the period τ , $\tau \in \mathbb{Z}$; we assume also that ε_t is independent of $\eta_{t'}$ when $t < t'$. The coefficients are such that $\omega_s > 0$, $\alpha_s > 0$ and $\alpha_s + \gamma_s \geq 0$, $\beta_s \geq 0$, $1 \leq s \leq S$ and are periodic functions of time.

3 The Markovian Representation and Stationarity

In this section, we give the markovian representation of the *PP* – *TGARCH*(1, 1) that will be needed later, a necessary and sufficient condition for the existence of the strictly stationary solution and a necessary and sufficient condition for the existence of the second order stationary solution.

3.1 Markovian representation

Rewrite (2) as follows

$$\begin{aligned}\varepsilon_{s+S\tau} &= \sqrt{h_{s+S\tau}}\eta_{s+S\tau} \\ h_{s+S\tau}^\delta &= \omega_s + ((\alpha_s + \gamma_s I_{\{\eta_{s-1+S\tau} \leq 0\}})\eta_{s-1+S\tau}^{2\delta} + \beta_s)h_{s-1+S\tau}^\delta.\end{aligned}\quad (3)$$

Set $\phi_s(\eta_{s-1+S\tau}) = (\alpha_s + \gamma_s I_{\{\eta_{s-1+S\tau} \leq 0\}})\eta_{s-1+S\tau}^{2\delta} + \beta_s$, so model (2) is equivalent to

$$h_{s+S\tau}^\delta = \omega_s + \phi_s(\eta_{s-1+S\tau})h_{s-1+S\tau}^\delta$$

which is a temporarily non-homogeneous Markov chain. Seasons in a periodic time varying coefficients process can be embedded to lead to a multivariate stationary process; this approach goes back to Gladyshev (1961).

Define

$$\mathbf{h}_\tau = (h_{1+S\tau}^\delta, h_{2+S\tau}^\delta, \dots, h_{S+S\tau}^\delta)', \delta_\tau = (\eta_{S+S(\tau-1)}, \eta_{1+S\tau}, \dots, \eta_{S-1+S\tau})'.$$

So,

$$\mathbf{h}_\tau = A(\delta_\tau)\mathbf{h}_{\tau-1} + B(\delta_\tau) \quad (4)$$

which is a generalized autoregressive process. $A(\delta_\tau)$ and $B(\delta_\tau)$ are obtained from simple calculation.

3.2 Existence of the strictly periodically stationary solution

The process $(A(\delta_\tau), B(\delta_\tau))_{\tau \in \mathbb{Z}}$ is an *i.i.d* sequence of $S \times S$ random matrices $A(\delta_\tau)$ and S -dimensional vectors $B(\delta_\tau)$. Following Bougerol and Picard (1992), we state now our first result.

Theorem 3.1 *Equation (4) has a unique strictly stationary and ergodic solution if and only if γ is negative. The unique strictly stationary and ergodic solution of (4) in its causal expansion is given by and so, the s -th component of $\mathbf{h}_\tau$ is the unique strictly periodically stationary and periodically ergodic solution of (2).*

In order to avoid the calculation of $\gamma(A)$ which is difficult to obtain explicitly, we give a sufficient condition which is more practical to check and deal with.

$$E(\phi_s(\eta_{s+S\tau})) = \alpha_s E(\eta_{s+S\tau}^{2\delta+}) + (\alpha_s + \gamma_s) E(\eta_{s+S\tau}^{2\delta-}) + \beta_s, \text{ for all } s = 1, \dots, S.$$

We give the following theorem.

Theorem 3.2 *Equation (4) has a unique strictly stationary and ergodic solution if*

$$\prod_{s=1}^S (\alpha_s E(\eta_{s+S\tau}^{2\delta+}) + (\alpha_s + \gamma_s) E(\eta_{s+S\tau}^{2\delta-}) + \beta_s) < 1,$$

and for all $\tau \in \mathbb{Z}$, the series converges a.s. and $(\mathbf{h}_\tau)_{\tau \in \mathbb{Z}}$ is the unique strictly stationary and ergodic solution of (4).

3.3 Existence of the second order periodically stationary solution

Theorem 3.3 *The process given by (4) has a stationary solution in $\mathbb{L}^1$ if and only if*

$$\prod_{s=1}^S (\alpha_s E(\eta_{s+S\tau}^{2\delta+}) + (\alpha_s + \gamma_s) E(\eta_{s+S\tau}^{2\delta-}) + \beta_s) < 1, \quad (5)$$

further, this solution is unique, causal and ergodic.

4 Higher order moments existence

The existence of moments for a stationary time series model is fundamental for statistical applications. Therefore, it is significant to find necessary and sufficient conditions for the existence of higher order moments for a stationary process $PP - TGARCH(1, 1)$.

Theorem 4.1 *Let $(h_\tau)_{\tau \in \mathbb{Z}}$ be a stationary solution for (4), we suppose that $E(\eta_{s+S\tau}^{2\delta p}) < +\infty, p \geq 2$, then h_τ is in $\mathbb{L}^p$ if and only if*

$$\prod_{j=1}^S E(\phi_{S-j+1}^p(\eta_{S-j+S\tau})) < 1.$$

5 Geometric ergodicity and β -mixing

(E). The sequence $(\delta_\tau)_{\tau \in \mathbb{Z}}$ is *i.i.d*. Its probability distribution function is absolutely continuous with respect to the Lebesgue measure on $\mathbb{R}^S$. The support of δ_τ is defined by its strictly positive density and contains an open set and zero.

Theorem 5.1 *Under condition (E), suppose that there exists $0 < r \leq 1$ such that $E(\|A(\delta_\tau)\|^r) < 1$ and $E(\|B(\delta_\tau)\|^r) < \infty$. if*

$$\prod_{s=1}^S (\alpha_s E(\eta_{s+S\tau}^{2\delta+}) + (\alpha_s + \gamma_s) E(\eta_{s+S\tau}^{2\delta-}) + \beta_s) < 1.$$

then the process $(h_\tau)_{\tau \in \mathbb{Z}}$ defined by (4) is geometrically ergodic, and if initializing from the invariant measure, then $(h_\tau)_{\tau \in \mathbb{Z}}$ is strictly stationary and β -mixing with exponential decay.

References

- Bibi, A. and Aknouche, A. (2008) On Periodic *GARCH* Processes : Stationarity, Existence of Moments and Geometric Ergodicity. *Mathematical Methods of Statistics*, Vol. 17, No. 4, 305-316.
- Bougerol, P. and Picard, N. (1992) Stationarity of *GARCH* processes and of some non-negative time series, *Journal of Econometrics* 52, 115-127.
- Gladyshev, E. G. (1961) Periodically Correlated Random Sequences. *SovietMath.* 2, 385-388.
- Glosten, L.R., Jagannathan, R. and Runkle, D., 1993. On the relation between the expected value and the volatility of the nominal excess return on stocks. *Journal of Finance* 48, 1779-1801.

Summary

A power periodic threshold *GARCH*(1, 1) process is introduced and studied. The existence of a strictly periodically stationary, second order periodically stationary, geometrically periodically ergodic and β -mixing solution is established. We give a necessary and sufficient condition for the existence of moments. These properties are important for statistical inference. The analysis relies on Markov chain theory for a general state space.

Normalité asymptotique locale (Condition LAN) pour un processus ARH(1)

Nesrine Kara Terki*, Tahar Mourid**

* kara.nesrine@yahoo.fr,

**Laboratoire de statistiques et modélisations aléatoires, Université de Tlemcen
t_mourid@mail.univ-tlemcen.dz

Résumé. Nous considérons les processus stochastiques autorégressifs à valeurs dans un espace de Hilbert séparable ARH(1) avec un opérateur dépendant d'un paramètre. Nous montrons la condition de normalité asymptotique locale (LAN condition) et la normalité asymptotique locale uniforme (ULAN condition) sous certaines conditions. Nous en déduisons une borne minimax de Hajek, la convergence et la normalité asymptotique des estimateurs du maximum de vraisemblance.

1 Introduction

Soit $(\Omega, \mathcal{A}, P)$ un espace de probabilité et $(H, \langle \cdot, \cdot \rangle)$ un espace de Hilbert séparable muni de sa tribu borélienne $\mathcal{B}_H$.

Nous considérons un processus stochastique $X = (X_n, n \geq 0)$ défini sur $(\Omega, \mathcal{A}, P)$ et à valeurs dans H défini par

$$X_n = \rho_\theta(X_{n-1}) + \varepsilon_n \quad (1)$$

où ρ_θ un opérateur linéaire borné de H vers H , θ un paramètre dans Θ (Θ un ouvert de $\mathbb{R}$), et $(\varepsilon_n, n \in \mathbb{Z})$ un H bruit blanc fort d'opérateur de covariance B . Sous la condition $\forall \theta, \exists j_0 \geq 1$ tel que $\forall j \geq j_0 \|\rho_\theta^j\| < 1$ le processus défini dans (1) admet une solution strictement stationnaire (voir Bosq.D. (2000)). Dans ce travail nous allons considérer la condition de Normalité asymptotique locale (LAN condition), une notion très utile en statistiques des processus. Plus particulièrement la normalité asymptotique locale uniforme (ULAN condition) (voir Ibragimov.I.A et Hasminski.R.Z. (1981)) on peut voir aussi ROUSSAS.G.G et PHILIPPOU.A.N (1973) ou Swensen.A.R. (1985)). Si on note $P_{n,\theta}$ la loi de $(X_1, \dots, X_n)$ de (1) et $\theta \in \Theta$ l'espace des paramètres un ouvert de $\mathbb{R}$, le rapport de vraisemblance est donné par

$$\tilde{\Lambda}_n = \log(p_{n,\theta_n}/p_{n,\theta}), \quad \theta_n = \theta + h/\sqrt{n}$$

où $p_{n,\theta}$ est la densité de $P_{n,\theta}$ par rapport à une mesure σ -finie et $h \in \mathbb{R}$.

Rappelons la définition de la condition LAN et ULAN.

Définition 1 La famille des lois $\{P_{n,\theta}, \theta \in \Theta\}$ vérifie la normalité asymptotique locale (condition LAN) au point θ s'il existe une suite de v.a. $\{S_n(\theta)\}$ et un nombre positif $\Gamma(\theta)$ tels que sous $P_{n,\theta}$:

Condition LAN pour un processus ARH(1)

1. pour tout h

$$\tilde{\Lambda}_n = hS_n(\theta) - 1/2h^2\Gamma(\theta) + o_p(1)$$

2.

$$S_n(\theta) \rightarrow^d N(0, \Gamma(\theta))$$

où $S_n(\theta), \Gamma(\theta)$ représentent respectivement la fonction score et l'information.

Définition 2 La famille $\{P_{n,\theta}, \theta \in \Theta\}$ vérifie la condition de la normalité asymptotique locale uniforme (condition ULAN) si pour tout $\theta \in \Theta$ la condition LAN est satisfaite et pour tout M , sous P_θ :

$$\sup_{|h| \leq M} |\tilde{\Lambda}_n - hS_n(\theta) + 1/2h^2\Gamma(\theta)| = o_p(1)$$

On suppose que $(\varepsilon_n, n \in \mathbb{Z})$ est un H bruit blanc fort gaussien d'opérateur de covariance B.

X est un processus de Markov homogène et ergodique (voir Mourider (1995)) de noyau de transition Q_θ , défini par

$$Q_\theta(x, dy) := P(X_1 \in dy / X_0 = x) = P_{\varepsilon_1 + \rho_\theta(x)}(dy)$$

Le log du rapport de vraisemblance Λ_n conditionnel à X_0 est défini par :

$$\Lambda_n = \log(L_n(\theta_n)/L_n(\theta))$$

où $L_n(\theta)$ est la vraisemblance de $(X_1, \dots, X_n)$ sous P_θ conditionnel à X_0 et $\theta_n = \theta + h/\sqrt{n}$, θ un point fixé de Θ

Notons $\{e_i, i \geq 1\}$ une base hilbertienne de l'espace H formée de vecteurs propres de l'opérateur de covariance B de $(\varepsilon_n, n \in \mathbb{Z})$ et $(\sigma_k^2, k \geq 1)$ ses valeurs propres.

L'opérateur de covariance B étant symétrique positif et à trace, on a :

$$Bx = \sum_{k=1}^{\infty} \sigma_k^2 \langle x, e_k \rangle e_k \quad \forall x \in H$$

En gardant les notations précédentes, le log du rapport de vraisemblance conditionnel à X_0 du processus défini par 1 s'écrit :

$$\begin{aligned} \Lambda_n &= \sum_{i=1}^n \sum_{k=1}^{\infty} \frac{-1}{2\sigma_k^2} \langle \rho_{\theta_n} X_{i-1}, e_k \rangle [\langle \rho_{\theta_n} X_{i-1}, e_k \rangle - 2 \langle X_i, e_k \rangle] \\ &+ \sum_{i=1}^n \sum_{k=1}^{\infty} \frac{1}{2\sigma_k^2} \langle \rho_\theta X_{i-1}, e_k \rangle [\langle \rho_\theta X_{i-1}, e_k \rangle - 2 \langle X_i, e_k \rangle] \end{aligned}$$

Nous avons les quatre résultats principaux suivants :

Théorème 1 Sous des conditions, la famille des lois $\{P_{n,\theta}, \theta \in \Theta\}$ du processus X (ARH(1)) satisfait la condition LAN de la définition 1 avec

$$S_n(\theta) = \frac{2}{\sqrt{n}} \sum_{i=1}^n \sum_{k=1}^{\infty} \frac{-1}{2\sigma_k^2} \frac{d}{d\theta} \langle \rho_{\theta} X_{i-1}, e_k \rangle [\langle \rho_{\theta} X_{i-1}, e_k \rangle - \langle X_i, e_k \rangle]$$

et

$$\Gamma(\theta) = \sum_{k=1}^{\infty} \frac{1}{\sigma_k^2} E_{\theta} \left[\frac{d}{d\theta} \langle \rho_{\theta} X_{i-1}, e_k \rangle \right]^2$$

Théorème 2 Sous des conditions , la famille des lois $\{P_{n,\theta}, \theta \in \Theta\}$ du processus X (ARH(1)) satisfait la condition ULAN de la définition 2 sur tout compact $K \subset \Theta$ et $S_n(\theta), \Gamma(\theta)$ sont définis dans le théorème 1.

Si on définit l'estimateur du maximum de vraisemblance $\hat{\theta}_n$ de θ alors ,nous obtenons le théorème suivant :

Théorème 3 Sous des conditions la famille des lois $\{P_{n,\theta}, \theta \in \Theta\}$ du processus X (ARH(1)) nous avons :

1. sur tout K compact de Θ :

$$\sqrt{n\Gamma(\theta)}(\hat{\theta}_n - \theta) \Rightarrow \zeta$$

où $\zeta \hookrightarrow N(0, 1)$

2. Pour une classe de fonctions de risque w

$$\lim_{n \rightarrow +\infty} E_{\theta}[w(\sqrt{n\Gamma(\theta)}(\hat{\theta}_n - \theta))] = Ew(\zeta).$$

Extention Dans le cas où l'espace des paramètres Θ est un ouvert de $\mathbb{R}^p, p \geq 1$ les quantités $S_n(\theta), \Gamma(\theta)$ se définissent par :

$$S_n(\theta) = \frac{2}{\sqrt{n}} \sum_{i=1}^n \sum_{k=1}^{\infty} \frac{-1}{2\sigma_k^2} \nabla \langle \rho_{\theta} X_{i-1}, e_k \rangle [\langle \rho_{\theta} X_{i-1}, e_k \rangle - \langle X_i, e_k \rangle]$$

et

$$\Gamma(\theta) = \sum_{k=1}^{\infty} \frac{1}{\sigma_k^2} E_{\theta} [\nabla \langle \rho_{\theta} X_{i-1}, e_k \rangle \nabla^T \langle \rho_{\theta} X_{i-1}, e_k \rangle]$$

$\nabla \langle \rho_{\theta} X_{i-1}, e_k \rangle$ désigne le vecteur des dérivées partielles.

Nous avons le résultat suivant :

Théorème 4 Sous des conditions , la famille des lois $\{P_{n,\theta}, \theta \in \Theta\}$ du processus X (ARH(1)) satisfait pour tout $\theta \in \Theta$,les propriétés suivantes : sous $P_{n,\theta}$,

1. pour tout h

$$\Lambda_n = h^T S_n(\theta) - 1/2 h^T \Gamma(\theta) h + o_p(1)$$

Condition LAN pour un processus ARH(1)

2.

$$S_n(\theta) \rightarrow^d N(0, \Gamma(\theta))$$

3. Pour tout $M > 0$:

$$\sup_{|h| \leq M} |\Lambda_n - h^T S_n(\theta) + 1/2 h^T \Gamma(\theta) h| = o_p(1)$$

Références

- Bosq.D. (2000). *Linear Processes in Function Spaces: Theory and Applications*. Springer.
- Ibragimov.I.A et Hasminski.R.Z. (1981). *Statistical Estimation .Asymptotic Theory*. New York: Springer.
- Mourid, T. . (1995). *Contribution à la statistique des processus autorégressifs à temps continu*. Thèse de doctorat, université Paris 6.
- ROUSSAS.G.G et PHILIPPOU.A.N (1973). Asymptotic distribution of the likelihood function in the independent not identically distributed case. *Statist,Anal* 3, 454–471.
- Swensen.A.R. (1985). The asymptotic distribution of the likelihood ratio for autoregressive time series with a regression trend. *J.Multivariate Anal* 16, 54–70.

Summary

We consider an Hilbert valued random autoregressive process ARH(1). we give the local asymptotic normality property (LAN condition) and uniform local asymptotic normality property (ULAN condition).we then deduce an asymptotically minimax Hajek- bound for the risks , convergence and asymptotic normality of maximum likelihood estimators.

GSPNs modeling and stability analysis for the $M/M/1$ queue

O. Lekadir*, L. Ikhlef** D. Aïssani***

Unité de recherche LAMOS, Université Abderrahmane Mira de Bejaia, 06000 Bejaia, Algérie.

* ouizalekadir@gmail.com,

**ikhlefilyes@gmail.com

***lamos_bejaia@hotmail.com

Abstract. In this paper, the modeling and the stability problem of the $M/M/1$ system after perturbing the $M/G/1$ service distribution are established. The modeling of the markovian $M/M/1$, resp. non Markov $M/G/1$ system is obtained via the formalism GSPNs (generalized stochastic Petri net), resp. MR-SPNs (Markovian regenerative SPNs) and the stability of the $M/M/1$ -GSPN is established by the strong stability approach. The main contribution of this work consists in combining strong stability theory with PNs models to give a complete solution to the stability problem for queuing systems. The presented methodology, which is the first attempt in the field of analysis of stability, results to be innovative in contrast to the classical stochastic approach usually used.

1 Introduction

The investigation of the Markovian queueing systems is simpler due to the memoryless property of their arrival and service distributions. However, the analytical treatment of the non markovian queueing systems is generally difficult to obtain because of the complexity of these systems due to the absence of the above memoryless property. These non markovian queueing systems are known of increasing interest in the literature and requires the development of suitable modeling tools, approximation methods and approaches. So, recently the new Petri nets formalism MRSPNs (Markovian regenerative Stochastic Petri Nets) have been introduced in the literature by H. Choi et al in Choi et al. (1994) with the aim of combining exponential and non-exponential firing times into a single model. The MRSPNs new formalism, gives the stochastic Petri nets in which the underlying marking process is a Markov regenerative process (MRGP) which is a continuous-time discrete state stochastic process with an embedded sequence of regenerative time points (RTPs), at which the process enjoys the Markov property. Once the Markov property is acquired, the investigation of these non-Markov systems becomes possible. After modeling a queueing system, the most important performance issues to be considered in this system is its stability. A lot of tools, approaches and methods are established to address the stability problem of queueing systems. In this paper, we are interested in the stability problem for queueing systems modeled with timed Petri nets. The early works on this field used the Lyapunov stability approach combining with others theories can be found in Retchkiman (2013) and the references in it. This paper is a first attempt for using the strong

stability approach in Petri nets (PNs). Indeed, the main contribution of this paper consists in combining the strong stability theory Kartashov (1996) with PNs modeling to give a complete and precise solution to the stability problem for the $M/M/1$ queuing system (QS) modeled with GSPNs.

Our objectif is to approximate an $M/G/1$ QS for which the service distribution is close to the exponential one by the simpler queue $M/M/1$. It is therefore very important to justify these approximation and estimate the resulting error using the strong stability approach. Thus, we consider first as nominal model the markovian QS $M/M/1$, and as perturbed model the non markovian QS $M/G/1$. In the $M/G/1$ QS the examined process (queue length, waiting time, etc ...) are not necessarily of the Markov type since the service time distribution may not have the memoryless property, so its investigation requires different methods. In the $M/G/1$ QS no conditions ensures that the queue length $L(t)$ is a Markov process (MP), but one can make it a MP with an other state space, that is what we will do in this work using the MRSPNs formalism which makes $M/G/1$ as MP using the theory of Markov regenerative processes. After the description of the considered models, we prove the v -strong stability of the considered nominal model $M/M/1$.

2 The models

• The nominal model “M/M/1-GSPN”:

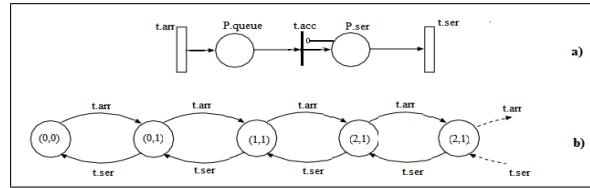


FIG. 1 – The $M/M/1$ queue: a- The GSPNs representation; b- The reachability graph of the $M/M/1$ -GSPN.

From the corresponding reachability graph of $M/M/1$ -GSPN reported in (FIG.1-b), we extract the infinitesimal generator $\bar{Q}$. then we define the transition matrix of this $M/M/1 - GSPN$:

$$\bar{P} = (p_{ij}) = \begin{cases} 1, & i = 0, j = 1; \\ a_j = \frac{\mu}{1+\mu} \left(\frac{\lambda}{1+\mu} \right)^j, & i > 0, j \geq i - 1; \\ 0; & \text{otherwise.} \end{cases}$$

• The real model “M/G/1-MRSPN”:

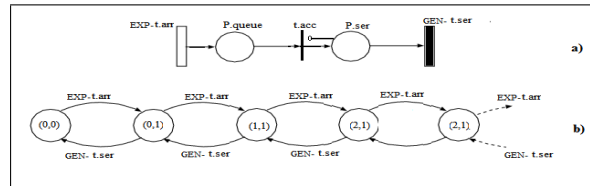


FIG. 2 – The $M/G/1$ queue: a- The MRSPNs representation; b- The reachability graph of the $M/G/1$ -MRSPN.

The regenerative moments of the $M/G/1$ -MRSPN correspond to the firing of $EXP - t.arr$ where the marking of $P.queue$ is zero-tail, otherwise the firing of $GEN - t.arr$. Consider then the subordinate process to the transition $GEN - t.arr$ which gives the local and global kernels needed for the $M/G/1$ -MRSPN analysis.

$$\begin{aligned} \text{Local kernel:} \quad E(t) &= \begin{cases} e^{-\lambda t} & \text{if, } i = j = 0; \\ \frac{(\lambda t)^{j-i}}{(j-i)!} e^{-\lambda t} [1 - F^g(t)] & \text{if, } i \geq 1, j \geq i; \\ 0 & \text{otherwise.} \end{cases} \quad (1) \\ \text{Global kernel:} \quad K(t) &= \begin{cases} 1 - e^{-\lambda t} & \text{if, } i = 0, j = 1; \\ \int_0^t \frac{(\lambda x)^{j-i+1}}{(j-i+1)!} e^{-\lambda x} dE(x), & \text{if, } i \geq 1, j \geq i-1; \\ 0 & \text{otherwise.} \end{cases} \quad (2) \end{aligned}$$

From the global kernel we obtain the one step transition matrix P , where:

$$P = K(\infty) = \begin{cases} 1 & \text{if, } i = 0, j = 1; \\ \int_0^\infty \frac{(\lambda x)^{j-i+1}}{(j-i+1)!} e^{-\lambda x} dE(x), & \text{if, } i \geq 1, j \geq i-1; \\ 0 & \text{otherwise.} \end{cases}$$

3 The strong stability approach

In this section we give the theorems of the theory of strong stability (see Kartashov (1996)) which we will use later.

Theorem 3.1 *The Markov chain X with the transition kernel $\mathbf{P}$ and invariant measure π is strongly v -stable with respect to the norm $\|\cdot\|_v$ if and only if there exists a measure α and a nonnegative measurable function h on $\mathbb{N}$ such that the following conditions hold:*

1. $\pi\alpha > 0$, $\alpha\mathbf{I} = 1$, $\alpha h > 0$.
2. The operator $\mathbf{T} = \mathbf{P} - h \circ \alpha$ is a nonnegative; $\circ$ denotes the convolution.
3. $\exists \rho < 1$ such that $Tv(k) \leq \rho v(k)$ for all $k \in \mathbb{N}$,
4. $\|\mathbf{P}\| < \infty$.

Theorem 3.2 *Let X be a Markov chain satisfying the conditions 1-4 of theorem 3.1. If ν is the invariant measure of an operator $\mathbf{Q}$, then, for $\|\Delta\|_v < C^{-1}(1 - \rho)$, we have the estimate:*

$$\|\nu - \pi\|_v \leq \|\Delta\|_v \|\pi\|_v C(1 - \rho - C\|\Delta\|_v)^{-1}, \quad (3)$$

where: $\Delta = \mathbf{Q} - \mathbf{P}$, $C = 1 + \|\mathbf{I}\|_v \|\pi\|_v$, and $\|\pi\|_v \leq (\alpha v)(1 - \rho)^{-1}(\pi h)$.

4 The strong v -Stability of “ $M/M/1$ -GSPN”

Using the theorem 3.1, the v -stability of the “ $M/M/1$ -GSPN” is stated in the following theorem.

Theorem 4.1 *We suppose that $\frac{\lambda}{\mu} < 1$, then for all β such that $1 < \beta < \frac{\mu}{\lambda}$ the embedded markov chain of $M/M/1$ -GSPN is strongly v -stable for the test fonction $v(l) = \beta^l$.*

Lemma 4.1 We suppose that: $\int x | F(x) - G(x) | dx < \frac{w}{\lambda}$, then $\exists \beta > 1$, $\int e^{(\lambda\beta-\lambda)x} | F(x) - G(x) | dx < w\beta$,

Theorem 4.2 Let $\bar{P}$ (resp P) be the transition operator of the embedded Markov chain of $M/M/1$ -GSPN (resp: $M/G/1$ -MRSPN), suppose that $\frac{\lambda}{\mu} < 1$, then for all $1 < \beta < \frac{\mu}{\lambda}$, and under the assumptions of the lemma 4.1, we have: $\|\bar{P} - P\|_v < w$, for $v(l) = \beta^l$.

Using the theorems 3.2 and 4.2, the deviation of the stationary distribution is stated in the following theorem.

Theorem 4.3 Let $\bar{\pi}$ (resp π) be the transition operator of the embedded Markov chain of $M/M/1$ -GSPN (resp: $M/G/1$ -MRSPN). provided that the conditions of the theorems 4.1 and 4.2 hold, then for all $1 < \beta < \frac{\mu}{\lambda}$ and under the condition $w < \frac{(1-\rho)(\mu-\lambda\beta)}{2\mu-\lambda(1+\beta)}$, we have the following estimate:

$$\|\pi - \bar{\pi}\| \leq \frac{w \left(\frac{\mu-\lambda}{\mu-\lambda\beta} \right) \left(\frac{2\mu-\lambda(1+\beta)}{\mu-\lambda\beta} \right)}{1 - \left(\frac{\mu}{\lambda+\mu} \right) \left(\frac{1}{\beta(1-\frac{\beta\lambda}{\lambda+\mu})} \right) - w \left(\frac{2\mu-\lambda(1+\beta)}{\mu-\lambda\beta} \right)}. \quad (4)$$

References

- Choi, H., V. G. Kulkarni, and K. S. Trivedi (1994). Markov regenerative stochastic petri nets. *Performance Evaluation* 20, 337–356.
- Gharbi, N. and M. Ioualalen (2010). Numerical investigation of finite-source multiserver systems with different vacation policies. *Journal of Computational and Applied Mathematics* 234(3), 625–635.
- Kartashov, N. V. (1996). *Strong Stable Markov Chains*. Kiev: Edition VSP, Utrecht, TBIMC Scientific Publishers.
- Retchkiman, Z. (2013). Timed petri nets modeling and, lyapunov/max-plus algebra stability analysis for a type of queueing systems. *International Journal of Pure and Applied Mathematics* 86(2), 301–323.

Résumé

Dans ce travail, la modélisation et le problème de stabilité du système $M/M/1$, après perturbation de la distribution de service du système $M/G/1$, sont établies. La modélisation des systèmes $M/M/1$, resp. $M/G/1$, est obtenue via le formalisme GSPNs (generalized stochastic Petri nets), resp. MRSPNs (Markovian regenerative SPNs) et la stabilité du $M/M/1$ -GSPN est établie par l'approche de stabilité forte. La principale contribution de ce travail est la combinaison de la théorie de la stabilité forte et de l'outil de modélisation "réseaux de Petri", pour donner une solution précise du problème de stabilité aux systèmes de files d'attente. Cette combinaison est une première tentative innovante dans le domaine d'analyse de stabilité.

La sensibilité des performances du système d'attente M/G/1/N avec vacances

Baya Takhedmit*, Karim Abbas**

LaMOS- Université Béjaia, Targa Ouzamour, 06000, Béjaia, Algérie.

*bayatakh@gmail.com,

**kabbas.dz@gmail.com.

Résumé. Les systèmes de files d'attente avec vacances sont des modèles très adéquats à la modélisation de plusieurs situations réelles. Pendant la dernière décennie, des techniques basées sur les développements en série de Taylor sont devenues un outil important pour l'analyse des systèmes stochastiques pouvant être décrits par des chaînes de Markov, tout en mettant l'accent sur l'analyse de perturbation. Dans cet article, nous considérons une analyse de sensibilité de quelques mesures de performance du modèle d'attente M/G/1/N avec vacances du serveur et discipline exhaustive, et ce en utilisant l'approche des développements en séries de Taylor, où les coefficients du polynôme de Taylor sont donnés en terme de la matrice de déviation relative à la chaîne de Markov décrivant l'état du système étudié. Quelques illustrations numériques sont donnés afin de montrer l'efficacité de l'approche utilisée.

1 Introduction

Depuis les travaux des Takagi (4), plusieurs auteurs ont investi l'étude des modèles de files d'attente avec vacance et service exhaustif. Les issues générales liées aux systèmes de files d'attente avec vacance et service exhaustif sont discutées dans (2). Ces dernières années, de nombreuses approches d'analyse des systèmes de files d'attente ont été développées dans le cadre d'évaluation des performance des systèmes de files d'attente avec vacances du serveur (2; 4). Entres autres, on peut citer celle des développements en séries de Taylor (1). Celle-ci est devenue un outil important pour l'analyse des systèmes stochastiques pouvant être décrits par des chaînes de Markov, tout en mettant l'accent sur l'analyse de perturbation. Plus particulièrement, elle est utilisée d'une manière très efficace dans le cadre d'analyse des performances des systèmes de files d'attente (voir par exemple (1)). En effet, en calculant un nombre fini de dérivées d'ordre supérieur, les développements en séries de Taylor permettent d'évaluer les fonctions de performance d'un certain modèle comme étant une fonction du paramètre d'intérêt. Dans ce cas, une telle mesure de performance est représentée sous forme polynômial, ce qui facilite leur manipulation mathématique (optimisation, analyse de sensibilité, etc).

Dans cet article, nous considérons une analyse de sensibilité de quelques mesures de performance du modèle d'attente M/G/1/N avec vacances du serveur et discipline exhaustive. Pour ce faire, on utilise l'approche des développements en séries de Taylor, où les coefficients

du polynôme de Taylor sont donnés en terme de la matrice de déviation relative à la chaîne de Markov décrivant l'état du système étudié. Celle-ci nous permet de considérer une analyse fonctionnelle de ces performances par rapport au paramètre du contrôle (taux de vacances). Enfin, quelques illustrations numériques sont donnés afin de montrer l'efficacité de l'approche utilisée.

2 Description du modèle

Considérons le modèle de files d'attente M/G/1/N avec vacances du serveur et service exhaustif noté M/G/1/N - (V, E). Supposons que :

- Le flux des arrivées est poissonnien de paramètre λ ;
- La distribution de la durée de service est générale, de fonction de répartition S ;
- La durée V de chaque vacance est i.i.d de taux $\theta > 0$.

Soit $X = \{X_n, n \in N\}$ le nombre de clients présents dans le système M/G/1/N(V, E) à l'instant de fin de service du n^{me} client. Ce processus forme une chaîne de Markov de matrice de probabilités de transition définie par :

$$P = \begin{pmatrix} b_0 & b_1 & b_2 & b_3 & \cdots & b_{N-2} & 1 - \sum_{k=0}^{N-2} b_k \\ a_0 & a_1 & a_2 & a_3 & \cdots & a_{N-2} & 1 - \sum_{k=0}^{N-2} a_k \\ 0 & a_0 & a_1 & a_2 & \cdots & a_{N-3} & 1 - \sum_{k=0}^{N-3} a_k \\ 0 & 0 & a_0 & a_1 & \cdots & a_{N-4} & 1 - \sum_{k=0}^{N-4} a_k \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \cdots & a_0 & 1 - a_0 \end{pmatrix},$$

où,

$$a_j = \int_0^\infty \frac{(\lambda t)^j}{j!} e^{-\lambda t} dS(t), \quad v_j = \int_0^\infty \frac{(\lambda t)^j}{j!} e^{-\lambda t} dV(t) \text{ et } b_j = \sum_{i=1}^{j+1} \frac{v_i}{1-v_0} a_{j-i+1}, \quad j \in N.$$

2.0.1 Analyse de sensibilité de la file d'attente M/G/1/N-(V,E)

Dans ce qui suit, nous intéressons à l'effet de perturbation du taux de vacances sur les performances du système d'attente M/G/1/N-(V,E). Cela nous permet de représenter π_θ par un polynôme de Taylor, où ces coefficients sont définis en fonction de la matrice de déviation D_θ définie par (3) :

$$D_\theta = \sum_{n \geq 0} (P_\theta^n - \Pi_\theta),$$

où Π_θ est le projecteur stationnaire de la chaîne de Markov X . Supposons que les composantes de P_θ sont n fois continûment dérivables par rapport au paramètre de contrôle θ . Ainsi, on définit la matrice dérivée d'ordre k de P par :

$P^{(k)} = (p_{ij}^{(k)})_{0 \leq i, j \leq N-1}$, où $p_{ij}^{(k)} = \frac{d^k}{d\theta^k} p_{ij}$, $0 \leq i, j \leq N-1$. Par la suite, considérons la représentation polynomiale de la distribution stationnaire suivante :

$$\pi_{\theta+\Delta} = \underbrace{\sum_{n=0}^k \frac{\Delta^n}{n!} \pi_\theta^{(n)}}_{H_\theta(k, \Delta)} + r_\theta(k, \Delta), \quad (1)$$

où $\pi_\theta^{(n)}$ est la sensibilité d'ordre k de la distribution π_θ . Celle-ci est donnée dans (3) :

$$\pi_\theta^{(n)} = \pi_\theta K_\theta(n),$$

avec

$$K_\theta(n) = \sum_{\substack{1 \leq m \leq n; \\ 1 \leq l_k \leq n; \\ l_1 + \dots + l_m = n}} \frac{n!}{l_1! \dots l_m!} \prod_{k=1}^m \left(P_\theta^{(l_k)} D_\theta \right).$$

Supposons que la loi de la durée de service des clients est Hyper-exponentielle d'ordre 2 (notée H_2) de paramètres μ_1 et μ_2 , donc sa fonction de densité est donnée par :

$$b(x) = \gamma \mu_1 e^{-\mu_1 x} + (1 - \gamma) \mu_2 e^{-\mu_2 x}, \quad \text{avec} \quad 0 \leq \gamma \leq 1. \quad (2)$$

Dans ce qui suit, nous appliquons l'approche décrite ci-dessus et nous fixons les valeurs paramètres du modèle comme suit :

$$\lambda = 1, \theta = 2, \mu_1 = 2, \mu_2 = 3, \gamma = 0.5, N = 4 \text{ et } \Delta \in [0, 0.1],$$

avec Δ étant l'intervalle des valeurs de perturbation du paramètre θ .

Les valeurs obtenues des composantes de la distribution stationnaire via le polynôme (1) d'ordre 3 sont données dans le tableau suivant :

Δ	0.02	0.04	0.06	0.08	0.1
π_0	0.5453	0.5511	0.5514	0.5453	0.5511
π_1	0.2849	0.2772	0.2756	0.2849	0.2772
π_2	0.1177	0.1186	0.1214	0.1192	0.1202
π_3	0.0500	0.0507	0.0507	0.0506	0.0515

TAB. 1 – Distributions stationnaires du système : $M/H_2/1/4 - (V, E)$

Pour toute fonction de coût $f : N \rightarrow R$, on peut évaluer la sensibilité de la mesure de performance, $\eta_\theta = \pi_\theta f$ de la chaîne de Markov par rapport à une perturbation de θ . Pour ce faire, nous utilisons l'estimation suivante :

$$\eta_{\theta+\Delta} = \pi_{\theta+\Delta} f = \sum_{i \geq 0} f(i) \times \pi_\theta(i).$$

Nous illustrons graphiquement les erreurs relatives absolues dues au calcul des composantes de la distribution stationnaire par la série de Taylor d'ordre 3 (Figure 1). Celles-ci sont données par :

$$\left| \frac{H_\theta(3, \Delta)(i) - \pi_{\theta+\Delta}(i)}{\pi_{\theta+\Delta}(i)} \right|$$

pour tout $i = 0, \dots, 3$.

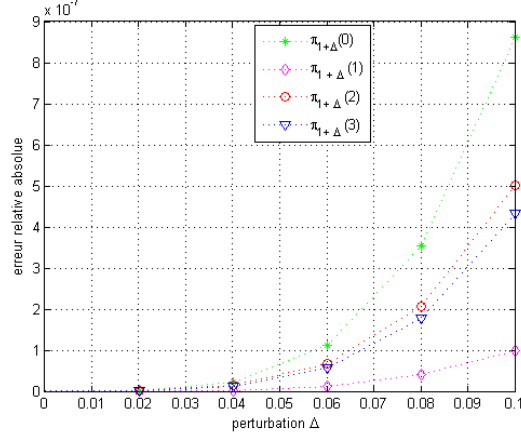


FIG. 1 – Courbes représentatives des erreurs relatives absolues sur le calcul de la distribution stationnaire du modèle M/H₂/1/N-(V,E) par la série de Taylor d'ordre 3.

K. Abbas, B. Heidergott and D. Aïssani, (2013), A functional approximation for the M/G/1/N queue ; Discrete Event Dyn Syst 23 : 93-104.

B.T. Doshi, (1986), Queueing systems with vacations A survry ; Queueing Syst, Vol. 1, 129-166.

B. Heidergott and A. Hordijk(2003), Taylor series expansions for the stationary Markov chains ; Adv. Appl. Prob. 35, 1046-1070 .

H. Takagi, (1993), Queueing Analysis a Foundation of Performance Evaluation, Finite Systems, vol. 2. North-Holland ; New York.

Summary

Queueing models with with vacations are very suitable models for modeling many real situations. During the past decade, Taylor series expansion based techniques have become an important tool for the analysis of stochastic systems that can be described by Markov chains, while focusing on the perturbation analysis. In this paper, we consider a sensitivity analysis of some performance measures of the M/G/1/N queue with server vacation and exhaustive discipline, by using the Taylor series expansions approach, where the Taylor polynomial coefficients are given in terms of the deviation matrix corresponding to the Markov chain that describes the state of the system under study. This allows us to consider a functional analysis of such performance measures with respect to the control parameter. Some numerical illustrations are given to show the effectiveness of the used approach.

Estimation non paramétrique des densités Heavy-tailed par la méthode du noyau

Yasmina ZIANE*, Smail ADJABI**, Nabil ZOUGAB***

*Unité de Recherche laMOS, Université de Béjaia, Algérie
yasmina-ziane@hotmail.com

** Unité de Recherche laMOS, Université de Bejaia, Algérie
Adjabi@hotmail.com

*** Unité de Recherche laMOS, Université de Bejaia, Algérie
nabilzougab@yahoo.fr

Résumé. Un des problèmes de la statistique non paramétrique consiste à estimer la densité de probabilité à partir d'un échantillon d'observations particulièrement celles qui se caractérisent par des événements rares et qui sont à support compact par exemple: $[0, \infty)$. L'objet de ce travail est d'estimer les fonctions de densité du type heavy tailed (queues lourdes) par la méthode du noyau. Pour éviter le problème du biais aux bornes, on a choisi des noyaux asymétriques du type RIG(Réciproque Inverse Gaussien) et gamma. Le paramètre de lissage qui intervient dans l'estimateur de la densité est sélectionné par deux méthodes: la méthode validation croisée et l'approche bayésienne globale. Une étude de simulation est réalisée sur des densités cibles pour comparer les performances de l'estimateur par rapport aux différents noyaux et aux différentes méthodes de sélection du paramètre de lissage selon le critère ISE (Integrate Square Error).

1 Introduction

Dans ce travail, on s'intéresse à l'analyse statistique du phénomène "Heavy-Tailed" à support $[0, \infty)$ dont on se focalise sur l'estimation de la densité de probabilité. Comme les méthodes paramétriques ne répondent pas aux caractéristiques de ce type de données, il est préférable d'utiliser les méthodes non paramétriques, comme la méthode du noyau de Parzen-Rosenblatt. Cette méthode, très importante pour la compréhension des propriétés de répartition des données est fonction de deux paramètres, la fonction noyau K et le paramètre de lissage h . L'efficacité de cette méthode dépend du choix de ces deux paramètres. Les noyaux les plus répandus dans la littérature en raison de leurs simplicité, sont les noyaux symétriques : noyau gaussien et noyau d'Epanechnikov qui minimise le MISE (Mean Integrate Square error). Cependant, ces noyaux ne s'adaptent pas aux données à support compact, car ils causent un sérieux problème aux bords : allocation d'un poids en dehors du support "biais" lorsque le lissage est appliqué près du bord. Cependant, pour éliminer ou réduire le biais, plusieurs auteurs ont proposé des noyaux asymétriques qui correspondent mieux aux support compact citons : noyau Beta Chen (1999), noyau Gamma et Gamma modifié Chen (2000), noyau inverse gaussien (IG) et réciproque inverse gaussien (RIG) Scaillet (2004), noyau log-normal (LN) et le

noyau Birnbaum-Saunders (BS) Jin et Kawczak (2003). Le choix du paramètre de lissage h est crucial, car il influe sur la qualité de l'estimation. Pour cela, on a opté pour estimer le paramètre de lissage par l'approche bayésienne qui se base sur une des méthodes d'approximation Monte Carlo par Chaîne de Markov (MCMC). Le principe de cette technique MCMC est de construire une chaîne de Markov de sorte que sa distribution stationnaire soit similaire à la loi à posteriori du paramètre. Cette approche bayésienne sera comparée à la méthode classique de validation croisée.

2 Choix du noyau

Dans le cas de où le support de la densité à estimer est infini, le noyau n'a aucune influence sur la qualité et les performances de l'estimateur. Lorsqu'il s'agit de données à support borné, le choix de noyau est important à cause du biais engendré aux bords de la densité. Nous utilisons les noyaux asymétriques : gamma modifié et Réciproque Inverse Gaussienne (RIG) définis respectivement par :

$$K_{Gam(\rho_h(x),h)}(\mu) = \frac{u^{\rho_h(x)-1} e^{-\frac{u}{h}}}{h^{\rho_h(x)} \Gamma(\rho_h(x))} \quad (1)$$

$$K_{RIG(\frac{1}{x-h}, \frac{1}{h})}(u) = \frac{1}{\sqrt{2\pi h u}} \exp\left(-\frac{x-h}{2h} \left(\frac{u}{x-h} - 2 + \frac{x-h}{u}\right)\right) \quad (2)$$

où

$$\rho_h(x) = \begin{cases} \frac{x}{h}, & si \quad x \geq 2h; \\ \frac{1}{4}(\frac{x}{h})^2 + 1, & si \quad x \in [0, 2h]. \end{cases}$$

3 Choix du paramètre de lissage

3.1 Validation Croisée

Le paramètre de lissage optimal sélectionné par la méthode de validation croisée est :

$$h_{ucv} = \arg \min_h UCV(h),$$

où

$$\begin{aligned} UCV(h) &= \int \hat{f}_h^2(x) dx - \frac{2}{n} \sum_{i=1}^n \hat{f}_{h,i}(X_i) \\ &= \int \left(\frac{1}{n} \sum_{i=1}^n K_{x,h}(X_i) \right)^2 dx - \frac{2}{n(n-1)} \sum_{i=1}^n \sum_{i \neq j} K_{X_i,h}(X_j) \end{aligned}$$

où $\hat{f}_h$ est l'estimateur de f défini dans (3) et $\hat{f}_{h,i}$ est l'estimateur de f en éliminant l'observation X_i

3.2 approche bayésienne

Cette section présente l'approche bayésienne pour estimer le paramètre de lissage pour une fonction de noyau K donnée. Soit $X_1, X_2, \dots, X_n$ une séquence de variable aléatoire iid de fonction de densité inconnue f . L'estimateur de la densité par la méthode du noyau pour tout $x \in [0, \infty)$ est donné par :

$$\hat{f}_h(x) = \frac{1}{n} \sum_{i=1}^n K_{x,h}(X_i), \quad h > 0. \quad (3)$$

Les estimateurs à noyau gamma modifié et RIG sont donnés respectivement par :

$$\hat{f}_{GAM}(x) = \frac{1}{n} \sum_{i=1}^n \frac{X_i^{\rho_h(x)-1} e^{-\frac{X_i}{h}}}{h^{\rho_h(x)} \Gamma(\rho_h(x))}$$

$$\hat{f}_{RIG}(x) = \frac{1}{(n)\sqrt{2\pi h}} \sum_{i=1}^n \frac{1}{\sqrt{X_i}} \exp\left(-\frac{X_i}{2h} + \frac{x-h}{h} - \frac{(x-h)^2}{2hX_i}\right),$$

Pour une loi a priori inverse gamma (IG) donnée par $\pi(h) = \frac{1}{\beta^\alpha \Gamma(\alpha) h^{\alpha+1}} \exp(-\frac{1}{\beta h})$, h optimal sera estimé par la moyenne de la loi à posteriori de h ,

$$h_{GAM}^* = \frac{1}{\beta^\alpha \Gamma(\alpha)(n-1)} \int \frac{1}{h^\alpha} \prod_{i=1}^n \sum_{j=1, j \neq i}^n \frac{x_j^{\rho_h(x_i)-1} e^{-\frac{x_j}{h}}}{h^{\rho_h(x_i)} \Gamma(\rho_h(x_i))} \exp(-\frac{1}{\beta h}) dh \quad (4)$$

et

$$h_{RIG}^* = \frac{1}{\beta^\alpha \Gamma(\alpha)(n-1)\sqrt{2\pi}} \int \frac{1}{h^\alpha \sqrt{h}} \prod_{i=1}^n \sum_{j=1, j \neq i}^n \frac{1}{\sqrt{x_j}} \exp\left(-\frac{x_j}{2h} + \frac{x_i-h}{h} - \frac{(x_i-h)^2}{2hx_j} - \frac{1}{h\beta}\right) dh \quad (5)$$

La complexité de cette intégrale rend la tâche difficile, voir même impossible avec les méthodes standards. La méthode Monte Carlo par Chaîne de Markov (MCMC) offre une solution indirecte à ce problème. Elle est basée sur l'idée de construction d'une chaîne de Markov ergodique, qui converge vers la loi stationnaire. Cette construction est possible grâce à l'algorithme de Metropolis Hastings introduit par Metropolis et al. (1953) et généralisé par Hastings (1970).

4 Application

Une application numérique a été réalisée pour évaluer les performances des méthodes citées dans lesquelles nous comparons les performances des estimateurs à noyau gamma modifié et RIG, obtenus en utilisant différentes méthodes de sélection de paramètre de lissage (méthodes classique (CV) et l'approche bayésienne). Cette comparaison sera basée sur les données simulées pour des densités cibles qui se caractérisent par une queue lourde à support $[0, \infty)$: log-normal(1,1), gamma(2,1) et weibull(1,1) et cela sur différentes tailles d'échantillon $n=10, 50, 100$ et 500 avec le nombre de simulation $N=50$. Nous examinons les performances en utilisant le critère de comparaison ISE (Integrate Square Error).

Nous présentons ici les résultats obtenus pour le modèle log-normal qui sont donnés dans le tableau 1 et dans la figure 1.

noyaux	n	$\bar{h}_{mcmc}$	$Var(h_{mcmc})$	$\bar{h}_{cv}$	$Var(h_{cv})$	$\overline{ISE}_{Bayes}$	$\overline{ISE}_{cv}$
<i>RIG</i>	10	0.48003	0.23522	0.14040	0.01828	0.10769	0.20814
	50	0.14399	0.01761	0.03088	0.00097	0.03317	0.14495
	100	0.14089	0.01000	0.01853	0.00029	0.01960	0.08007
	500	0.04460	0.00090	0.00352	$1.08805 \cdot 10^{-5}$	0.00518	0.02479
<i>Gam</i>	10	0.72719	0.17729	0.22969	0.04784	0.08250	0.33369
	50	0.16753	0.01086	0.03126	0.00088	0.03748	0.08726
	100	0.14547	0.03130	0.01752	0.00024	0.02468	0.08180
	500	0.06224	0.00260	0.00416	$1.75877 \cdot 10^{-5}$	0.00787	0.01891

TAB. 1 – Résultats de $\bar{h}$ et $\overline{ISE}$ pour la densité cible log-normal par les noyaux RIG et Gamma modifié.

Estimation non paramétrique des densités Heavy-tailed

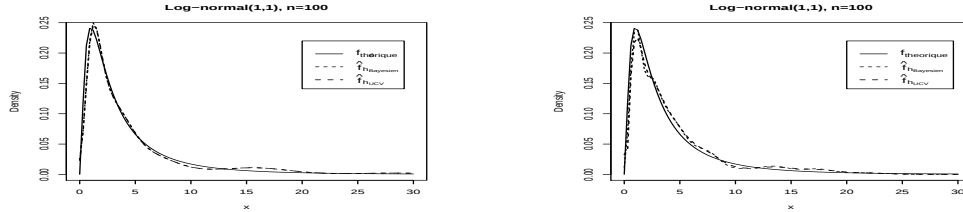


FIG. 1 – Densité cible et les estimateurs de densité log-normal avec le noyau RIG à gauche et noyau gamma modifié à droite.

5 Conclusion

D'après les résultats numériques obtenus pour la densité cible log-normal, l'estimateur obtenu en utilisant l'approche bayésienne pour la sélection de paramètre de lissage avec les deux noyaux (RIG et gamma modifié) est meilleur que celui obtenu par la méthode classique validation croisée et ceci quelque soit la taille de l'échantillon au sens du critère ISE (Integrate Square Error). Par contre, pour les autres densités cibles, l'approche bayésienne est plus performante pour des échantillons de petites ou moyennes tailles.

Références

- Chen, S. X. (1999). Beta kernels estimators for density functions. *Computational Statistics and Data Analysis* 31, 131–145.
- Chen, S. X. (2000). Gamma kernel estimators for density functions. *Annals of the Institute of Statistical Mathematics* 52, 471–480.
- Hastings, W. K. (1970). Monte carlo sampling methods using markov chains and their application. *Biometrika* 57, 97–109.
- Jin, X. et J. Kawczak (2003). Birnbaum-saunders and lognormal kernel estimators for modelling durations in high frequency financial data. *Ann. Econ. Fin* 4, 103–124.
- Metropolis, N., A. W. Rosenbluth, M. N. Rosenbluth, A. H. Teller, et E. Teller (1953). Equations of state calculations by fast computing machines. *Journal of Chemical Physics* 21, 1087–1092.
- Scaillet, O. (2004). Density estimation using inverse and reciprocal inverse gaussian kernels. *Journal of Nonparametric Statistics* 16, 217–226.

Summary

One of the problems of nonparametric statistics is to estimate the probability density from a sample of observations particularly those characterized by rare events and are compact support such: $[0, \infty)$. the subject of this work is to estimate the density functions of the type heavy tailed by the kernel method. To avoid the bias problem at the edges, we choose asymmetric kernels type RIG (Reciprocal Inverse Gaussian) and gamma. The smoothing parameter which intervenes in the density estimator is selected by two methods: Cross Validation method and global bayesian approach. A simulation study is realized on target densities to compare the performances with respect to different kernels and different methods for selecting the smoothing parameter according to the criterion (ISE Integer Square Error).

Part III: Statistique Computationnelle, Simulation

Modélisation spatiale de la formation des agglomérations dans la zone algéroise

Smicha AIT AMOKHTAR*, Nadjia EL SAADI**

*Ecole Nationale Supérieure Agronomique
s.aitamokhtar@ensa.dz,

**Ecole Nationale Supérieure de Statistique et Economie Appliquée
enadjia@gmail.com

Résumé. L'objectif de notre travail est d'analyser les dynamiques sous-jacentes de la formation des agglomérations sur la zone algéroise en se référant aux théories de la Nouvelle Économie Géographique (NEG) et plus précisément au travail de Paul Krugman; lauréat du prix nobel d'économie (2008); "increasing returns and economic geography" qui explique les mécanismes de la concentration des activités économiques à travers deux types de forces: des forces centripètes qui encouragent la concentration des activités économiques et des forces centrifuges qui freinent le processus d'agglomération. Ces mécanismes sont traduits en système d'équations non linéaires, dont la résolution par les méthodes analytiques est considérée comme une tâche très rude d'où le recours aux méthodes numériques. Notre application s'inscrit dans la lignée des travaux empiriques de la NEG dont l'objectif est de mener une analyse prospective de l'armature urbaine de la zone algéroise.

1 Introduction

Maints modèles ont été construits afin de reproduire la réalité économique dont les questions de concentration des activités économiques. La formation des agglomérations a fait l'objet du papier séminal de Krugman (1991). Ce dernier s'est référé à la théorie de l'auto-organisation pour appréhender les mécanismes d'agglomération (KRUGMAN (1998)). Les éléments du système sont constitués d'agents économiques rationnels, ces derniers échangent des biens et des services, chacun poursuivant un objectif de maximisation de son utilité individuelle. L'ensemble de ces agents sont soumis à un processus d'auto-renforcement que l'on peut identifier aux économies d'agglomération.

2 Modèle de Krugman :

Ce modèle est basé sur certaines hypothèses : l'économie est dotée de deux facteurs de production (la main d'œuvre qualifiée utilisée par le secteur industriel et la main d'œuvre non qualifiée exploitée par le secteur agricole). L'espace est considéré unidimensionnel et composé de n régions (KRUGMAN (1991)).

Comportement des consommateurs Les préférences des consommateurs sont considérées les mêmes, elles sont de type Cobb-Douglas : $U = C_M^\mu C_A^{1-\mu}$ où $C_M = (\sum c(i)^{\frac{\sigma-1}{\sigma}})^{\frac{\sigma}{1-\sigma}}$

- C_M : la consommation des biens manufacturés qui sont diversifiés horizontalement ; en outre μ est la part de l'industrie dans l'économie,
- C_A : la consommation des biens agricoles considérés comme bien homogène ,
- $c(i)$: la quantité de la i ème variété consommée,
- σ : le taux de substitution entre les différentes variétés.

Par ailleurs, la dimension spatiale est intégrée dans ce modèle en introduisant les coûts de transport que nous admettons qu'ils sont du type d'iceberg où une partie du bien transporté entre deux localisations se perd, au cours du chemin.

Comportement du producteur L'économie est dotée de L^A travailleurs non qualifiés et de L^M travailleurs qualifiés actifs respectivement dans le secteur agricole et industriel. Nous notons φ_j la part des agriculteurs dans la région j et λ_j est la proportion des ouvriers dans la région j . Le secteur agricole produit un seul produit homogène sous des rendements d'échelles constants selon l'équation suivante : $L_j^A = q_j^A$. Le secteur industriel produit les biens manufacturiers sous des économies d'échelles croissantes. La quantité produite est décrite par la technologie suivante : $L_j^M = \alpha + \beta q(i)_j$ où α et β mesure respectivement les économies d'échelles internes et les coûts fixes quant à $q(i)_j$, elle correspond à la quantité produite de la variété i dans la région j .

L'équilibre spatiale Krugman suppose que les ouvriers sont parfaitement mobiles et ils sont prêts à se déplacer vers les régions où le salaire réel offert est meilleur en se référant au salaire réel moyen. La fonction de mobilité des firmes est considérée comme une équation différentielle de 1er ordre et elle est résumée dans cette équation : $\frac{d\lambda(t)}{dt} = \rho \lambda_j (\omega_j(t) - \bar{\omega}_j(t))$. La localisation des ouvriers est déterminée par ce système d'équations :

$$\begin{aligned} \frac{d\lambda(t)}{dt} &= \rho \lambda_j (\omega_j(t) - \bar{\omega}_j(t)) \\ \omega_j &= W_j(t) T_j^{-\mu} \\ W_j(t) &= \left(\sum Y_k(t) (T_k(t) e^{(-\tau D_{jk})})^{(s-1)} \right) \left(\frac{1}{\sigma} \right) \\ Y_j(t) &= (1 - \mu) \varphi_j + \mu \lambda_j W_j \\ T_j(t) &= \left[\sum \lambda_k (W_k(t) e^{(\tau D_{jk})})^{(1-\sigma)} \right]^{1/(1-\sigma)} \end{aligned}$$

3 Simulation du modèle et application

Stelder (2005) développe ce modèle dans le cas bidimensionnel où il a testé sa robustesse dans la prédiction de la formation des agglomérations sur des structures réelles (Dirk (2005)). Notre application s'inscrit dans la lignée de ce travail (Dirk (2005) et Hakan (2010)) dont l'objectif est de mener une analyse prospective de l'armature urbaine de la zone algéroise. Selon le modèle de Krugman (1991), trois paramètres sont considérés capitaux à l'étude des structures urbaines émergentes à savoir la part des biens industriels dans les dépenses de consommation

(μ), l'élasticité de substitution entre les variétés (σ) et les coûts de transport (τ)(Fujita Masahisa (1999)). Le choix des valeurs est justifié par des études empiriques antérieures (Catherine (2006)). Nous proposons 12 scénarios afin de couvrir des situations différentes, chaque simulation est comparée avec une situation initiale qui est celle de l'année de référence 2008(AIT-AMOKHTAR (2012)).

Par ailleurs, la simulation du modèle nécessite la répartition initiale des travailleurs qualifiés que nous admettons par hypothèse proportionnelle à la répartition des entreprises. En outre, nous avons supposé que le nombre de firmes actives sur le marché algérien est égal au nombre de firmes souscrites au niveau de la société nationale d'électricité et du gaz SONELGAZ voire que cette dernière est la seule compagnie de la distribution d'électricité qui est active en Algérie. Une autre variable que Krugman suppose fixe ; la proportion des agriculteurs dans les différentes communes φ_j qui correspond à l'emploi dans le secteur agricole. Ces données sont issues de l'office national de statistiques (2008). L'apport fondamental de la NEG est l'incorporation de la dimension spatiale dans la conception de ses modèles. La matrice des distances intercommunales est calculée à partir des coordonnées polaires (latitude et longitude) de chaque commune, La résolution du système d'équations requiert de poser des valeurs initiales de W_0 qui peuvent être assimilées aux valeurs du taux d'urbanisation.

Scénarios	$\tau = 0,01$	Scénarios	$\tau = 0,1$
Scénario A.1	$(\mu, \sigma) = (0.3; 2)$	Scénario B.1	$(\mu, \sigma) = (0.3; 2)$
Scénario A.2	$(\mu, \sigma) = (0.3; 5)$	Scénario B.2	$(\mu, \sigma) = (0.3; 5)$
Scénario A.3	$(\mu, \sigma) = (0.5; 5)$	Scénario B.3	$(\mu, \sigma) = (0.5; 5)$
Scénario A.4	$(\mu, \sigma) = (0.5; 2)$	Scénario B.4	$(\mu, \sigma) = (0.5; 2)$
Scénario A.5	$(\mu, \sigma) = (0.1; 2)$	Scénario B.5	$(\mu, \sigma) = (0.1; 2)$
Scénario A.6	$(\mu, \sigma) = (0.1; 5)$	Scénario B.6	$(\mu, \sigma) = (0.1; 5)$

TAB. 1 – Scénarios simulés.

Résultats des simulation Pour quantifier l'effet de la variation des paramètres (μ, σ, τ) sur la dynamique de la concentration du secteur industriel, nous recourons au calcul d'un indice de concentration qui est l'indice de Gini. Cet indice donne une mesure de la concentration par rapport à une région de référence qui est la distribution uniforme.

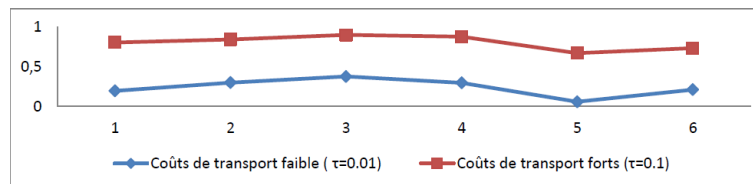


FIG. 1 – L'évolution de l'indice de GINI

La figure (1) résume nos résultats de simulation par le calcul de l'indice de Gini, cet indice est calculé sur la base des valeurs de λ_j associé à l'équilibre spatial. Pour $\tau = 0.01$, nous

obtenons des valeurs de G faibles, ce qui induit une faible concentration spatiale de l'activité industrielle au niveau des communes. L'indice de Gini atteint des valeurs proches de 1 pour des coûts de transport forts, ce qui signifie une concentration de l'activité industrielle dans un nombre limité de régions.

4 Conclusion

Ils sont nombreux les modèles qui cherchent à analyser et à prédire les inégalités économiques spatiales ; le modèle de Krugman est l'un de ces modèles. En simulant le modèle de Krugman sur une structure réelle, nous avons obtenu un résultat majeur : la baisse des coûts de transport permet de rétrécir les inégalités spatiales existantes entre les régions de la zone étudiée.

Références

- AIT-AMOKHTAR, S. (2012). Modélisation spatiale de la formation des agglomérations. application à des villes algériennes. Mémoire de magister, École Nationale Supérieure de Statistique et Économie Appliquée.
- Catherine, B. (2006). *Dépenses publiques, localisation des capitaux et concurrence fiscale : une modélisation et économie géographique*. thèse de doctorat es sciences économiques, Université Paris I.
- Dirk, S. (2005). *Regions and Cities: Five Essays on Interregional and Spatial Agglomeration Modeling*. thèse de doctorat es sciences économiques, université de Groningue.
- Fujita Masahisa, Paul Krugman, A. J. (1999). *The spatial economy*. Princeton University Press.
- Hakan, A. (2010). *L'impact des politiques de transport sur la concentration spatiale des activités*. Phd, université LAVAL, Québec.
- KRUGMAN, P. (1991). Increasing returns and economic geography. *Journal of Political Economy*.
- KRUGMAN, P. (1998). *L'économie auto-organisatrice*. De Boeck Université.

Summary

The goal of our study is to analyze the underlying dynamics of Algiers urban area formation with reference to New Economic Geography (NEG) theories and more precisely to the work, "Increasing returns and economic geography" by Paul Krugman work; a Nobel Prize Laureate in Economics (2008) which explains the mechanisms of economic activities concentration through two types of forces: centripetal forces encouraging the concentration of economic activities and centrifugal forces hindering the agglomeration process. In fact, these mechanisms are translated into a system of nonlinear equations whose resolution analytically is deemed to be a very hard task. As a consequence, the use of numerical methods is highly advocated. We present some numerical simulations using real Algerian data.

Optimal approach on the heavy tailed distribution's mean estimation with right random censoring

Abdelkader Ameraoui *, Kamal Boukhetala*
Jean-François Dupuy**

*Faculty of mathematics PoBox 32, Al alia Bab ezzouar, USTHB - Algiers, Algeria
aameraoui@usthb.dz,
kboukhetala@usthb.dz

**INSA de Rennes (centre de mathématiques) - 20 Avenue des Buttes de Coësmes -
35708 Rennes cedex 7, France
Jean-Francois.Dupuy@insa-rennes.fr

Résumé. Many papers were dedicated to the mean estimation of heavy tailed distributions with tail index $0 < \gamma < 1$, in the absence of censorship. When random censoring is present, an alternative is to introduce a new kind of tail index estimator such as the adaptive Hill's estimator under censoring data. In this paper, we propose new estimator's classes, of heavy tailed distribution position parameter, with right random censoring model. The new estimator is compared with the exact mean of a simulated patterns. We discusses also the optimal fraction level detection procedure in the estimation of the mean.

Keywords : High quantile ; Hall's class ; heavy tailed distributions ; extreme value, random right censoring.

1 Introduction

1.1 Classical extreme value approach

Let $X_1, X_2, \dots, X_n$ be independent and identically distributed (denoted iid) random variables with a common distribution function (denoted df) F . We assume that F is a regularly varying function at infinity with index $\frac{1}{\gamma}$, and that a tail balance condition holds :

$$\lim_{t \rightarrow \infty} \frac{1 - F(tx)}{1 - F(t)} = x^{-\frac{1}{\gamma}} \quad (1)$$

for any $x > 0$ and $0 < \gamma < 1$. The condition in (1) means that F is in the domain of attraction of a Fréchet law.

As Peng (2001), we propose the following expression to estimate the mean $E(X)$, based on the division of the quantile function variation interval into :

$$E(X) = \int_0^1 F^{-1}(u)du = \int_0^{1-\frac{k}{n}} F^{-1}(u)du + \int_{1-\frac{k}{n}}^1 F^{-1}(u)du = \mu_n^{(1)} + \mu_n^{(2)} \quad (2)$$

Mean estimation for heavy tailed distribution's

where $F^{-1}(y) = \inf\{x : F(x) \geq y\}$, for $0 < y < 1$, denotes the inverse function of F and $k = k(n)$ is an intermediate integer sequence such that $k \rightarrow \infty$ and $k/n \rightarrow 0$ as $n \rightarrow \infty$. Let consider now the order statistics $X_{n,1} \leq X_{n,2} \leq \dots \leq X_{n,n}$. For the estimation of the mean, Peng [4] proposed to write $\hat{E}(X) = \hat{\mu}_n^{(1)} + \hat{\mu}_n^{(2)}$, where $\hat{\mu}_n^{(1)}$ is the classical empirical estimator of the trimmed mean defined by :

$$\hat{\mu}_n^{(1)} = \frac{1}{n} \sum_{i=1}^{n-k} X_{n,i}.$$

It remains us to give the expression for $\hat{\mu}_n^{(2)}$ for $0 < \gamma < 1$. This can be done by defining first the Hill EVI estimator :

$$H(k) = \frac{1}{k} \sum_{i=1}^k \log(X_{n,n-i+1}) - \log(X_{n,n-k})$$

Now, since $X_{n,n-k} \sim C.(n/k)^\gamma$ we have $C \sim X_{n,n-k}.(k/n)^\gamma$ The C-estimator is given then by :

$$C_{H(k)}(k) = X_{n,n-k}.(k/n)^{H(k)}$$

and the Weissman quantile estimator at the level p is :

$$Q_{H(k)}^{(p)}(k) = C_{H(k)}(k) \times \left(\frac{1}{p}\right)^{H(k)} \quad (3)$$

It follows that :

$$\begin{aligned} \hat{\mu}_n^{(2)} &= \int_{1-\frac{k}{n}}^1 \left(\frac{k}{n}\right)^{H(k)} X_{n,n-k+1} (1-u)^{-H(k)} du \\ &= \frac{k}{n} X_{n,n-k+1} \frac{1}{1-H(k)} \end{aligned}$$

The expression of a heavy tail distribution mean estimator is thus given by :

$$\hat{E}(X) = \frac{1}{n} \sum_{i=1}^{n-k} X_{n,i} + \frac{k}{n} X_{n,n-k+1} \frac{1}{1-H(k)}$$

1.2 Reduced bias tail index estimation

Let $X_1, X_2, \dots, X_n$ be iid random variables with common df F . If we assume that F is in the Hall's class (see [1]), and if we note $U(t) = F^{-1}(1 - \frac{1}{t})$ then $U(t)$ is a regular variation function at infinity with index $\frac{1}{\gamma}$ if and only if :

$$\lim_{t \rightarrow \infty} \frac{U(tx)}{U(t)} = x^\gamma$$

for any $x > 0$ and $0 < \gamma < 1$.

In order to derive the asymptotic non-degenerate behaviour of semi-parametric estimators

of extreme events parameters, we need more than a first order condition, and we expand the problem to a sub-class of the Hall's class such as :

$$U(t) = Ct^\gamma(1 + \frac{\gamma\beta t^\rho}{\rho} + o(t^\rho)) \quad (4)$$

where $\gamma > 0, \beta \neq 0, \rho > 0$. Then the following condition holds :

$$\lim_{t \rightarrow \infty} \frac{\frac{U(tx)}{U(t)} - x^\gamma}{A(t)} = x^\gamma \frac{x^\rho - 1}{\rho}$$

for all $x > 0$, where A is defined as $A(t) = \gamma\beta t^\rho$. The reduced bias estimator of the heavy tail index, proposed by Caeiro in 2005, is given by :

$$\overline{H}(k) = H(k)(1 - \frac{\beta(k) \left(\frac{n}{k}\right)^{\rho(k)}}{1 - \rho(k)}) \quad (5)$$

The expression of the new reduced bias tail index estimator $\overline{H}(k)$ requires the knowledge (or the estimation) of the second order parameters ρ and β . Caeiro and Gomes (2011) state that, if the second order condition in (5) holds, and for $k \rightarrow \infty, \frac{k}{n} \rightarrow 0$ and $\sqrt{k}A(n/k) \rightarrow \lambda$, finite and non necessarily null, as $n \rightarrow \infty$, then :

$$\sqrt{k} \left(\overline{H}_{\hat{\rho}, \hat{\beta}}(k) - \gamma \right) \stackrel{d}{=} \gamma Z_k + o_p(\sqrt{k}A(n/k)) \quad (6)$$

where Z_k is asymptotically distributed as a standard normal r.v.

1.3 Application to the tail mean estimation

Now, if we replace the Hill estimator of the tail index $H(k)$ by the Hill reduced bias estimator, in the Weissman quantile estimation in (3), we obtain a new Weissman quantile estimator, which is given by :

$$Q_{\overline{H}(k)}^{(p)}(k) = C_{\overline{H}(k)}(k) \times (p)^{-\overline{H}(k)}$$

As a consequence, we also get a new estimator of the tail mean $\mu^{(2)}$:

$$\begin{aligned} \hat{\mu}_n^{(2)} &= \int_{1-\frac{k}{n}}^1 \left(\frac{k}{n}\right)^{\overline{H}(k)} X_{n,n-k+1} (1-u)^{-\overline{H}(k)} du \\ &:= \frac{k}{n} X_{n,n-k+1} \frac{1}{1 - \overline{H}(k)} \end{aligned}$$

In 2006, Caeiro proposed a new reduced bias estimator of C :

$$C_\gamma(k, \theta) = \frac{X_{n,n-[\theta k]} - X_{n,n-k}}{\theta^{-\gamma} - 1} \left(\frac{k}{n}\right)^\gamma$$

Mean estimation for heavy tailed distribution's

where θ is a tuning parameter such that $0 < \theta < 1$, and developed in the same year (see [1]) a new semi-parametric C-estimator :

$$\overline{C}_\gamma(k, \theta) = \frac{X_{n, n-[\theta k]} - X_{n, n-k}}{\theta^{-\gamma} - 1} \left(\frac{k}{n} \right)^\gamma (1 - B_\theta(\gamma, \beta, \rho))$$

where $B_\theta(\gamma, \beta, \rho) = \frac{\theta^{-(\gamma+\rho)} - 1}{\theta^{-\rho} - 1} \times \frac{\gamma\beta(n/k)^\rho}{\rho}$.

A further new reduced bias tail mean estimator $\overline{\mu}^{(2)}$ is thus given by :

$$\begin{aligned} \overline{\mu}^{(2)} &= \int_{1-\frac{k}{n}}^1 \overline{C}_\gamma(k, \theta) (1-u)^{-\overline{H}(k)} du \\ &:= \frac{X_{n, n-[\theta k]} - X_{n, n-k}}{\theta^{-\gamma} - 1} \left(\frac{k}{n} \right)^\gamma (1 - B_\theta(\overline{H}(k), \beta, \rho)) \frac{1}{1 - \overline{H}(k)} \end{aligned}$$

Proposition 1.1. Let F be a df satisfying the second order condition (4). For an integer sequence $k = k(n)$ such that $k \rightarrow \infty$ and $k/n \rightarrow 0$ when $n \rightarrow \infty$, then :

$$\widehat{E}(X) = \frac{1}{n} \sum_{i=1}^{n-k} X_{n,i} + \frac{X_{n, n-[\theta k]} - X_{n, n-k}}{\theta^{-\gamma} - 1} \left(\frac{k}{n} \right)^\gamma (1 - B_\theta(\overline{H}(k), \beta, \rho)) \frac{1}{1 - \overline{H}(k)}$$

2 Estimation under random censoring

Let $X_1, X_2, \dots, X_n$ be iid sample of df F in the domain of attraction of the Fréchet law with extreme value index (EVI) γ_1 , and let $Y_1, Y_2, \dots, Y_n$ be iid sample of df G also in the Fréchet domain of attraction, with EVI γ_2 , and independent of $(X_1, X_2, \dots, X_n)$. Assume that we observe the censored sample :

$$(Z_i, \delta_i), 1 \leq i \leq n, \text{ with } Z_i = \min(X_i, Y_i), \delta_i = I_{\{X_i \leq Y_i\}}$$

Gomes et al. (2011) proposed an adaptive EVI estimator :

$$\hat{\gamma}_{k,Z}^{(C)} = \frac{\overline{H}(k)}{\hat{p}}, \text{ where } \hat{p} = \frac{1}{k} \sum_{i=1}^k \delta_{[n-i+1]},$$

is the proportion of uncensored observations in the k largest observation of Z .

Many papers in the recent literature were dedicated to the asymptotic normality of the censoring adaptive EVI (Gomes et al. (2011), Einmahl et al. (2008), Gomes (2003), ...).

We consider now the estimation of the extreme quantile $Q(p) = F^{-1}(1-p)$. The classical estimator of Q was adopted to censoring as :

$$\hat{Q}_{p,k}^{(C)} = Z_{n, n-k} + \hat{a}_{Z,k}^{(C)} \frac{((1 - \hat{F}_n(Z_{n, n-k}))/p)^{\hat{\gamma}_{k,Z}^{(C)}} - 1}{\hat{\gamma}_{k,Z}^{(C)}}, \quad (7)$$

where $\hat{F}_n$ is the Kaplan-Meier (1958) product-limit estimator,

$$\hat{a}_{Z,k}^{(C)} = \frac{Z_{n, n-k} M_{Z,k}^{(1)} (1 - S_{Z,k})}{\hat{p}},$$

$$S_{Z,k} = 1 - \frac{1}{2} \left(1 - \frac{(M_{Z,k}^{(1)})^2}{M_{Z,k}^{(2)}} \right)^{-1},$$

and $M_{Z,k}^{(j)} = \frac{1}{k} \sum_{i=1}^k (\log Z_{n,n-i+1} - \log Z_{n,n-k})^j, j = 1, 2, \dots$. By introducing (7) in the expression of $\mu_n^{(2)}$ in (2), we get the following estimator :

$$\hat{\mu}_n^{(2,C)}(k) = \left[\left(Z_{n,n-k} - \frac{\hat{a}_{Z,k}^{(C)}}{\hat{\gamma}_{k,Z}^{(C)}} \right) (k/n) \right] + \frac{\hat{a}_{Z,k}^{(C)}}{\hat{\gamma}_{k,Z}^{(C)}} (1 - \hat{F}_n(Z_{n,n-k}))^{\hat{\gamma}_{k,Z}^{(C)}} \left[\frac{(k/n)^{1-\hat{\gamma}_{k,Z}^{(C)}}}{1 - \hat{\gamma}_{k,Z}^{(C)}} \right].$$

Hence, for the first term $\mu_n^{(1)}$ in (2), we propose the estimator of the trimmed mean under censoring as follows :

$$\hat{\mu}_n^{(1,C)} = \frac{1}{n_c} \sum_{i=1}^{n-k} \delta_{[n,i]} Z_{n,i},$$

where n_c is the number of uncensored observations in the ordered sample $Z_{n,1}, Z_{n,2}, \dots, Z_{n,n}$. We can now give our expression of the estimator of EX under right censoring as follow

Definition 2.1. We assume that the df F and G satisfy the second order condition (4). Let $k = k(n)$ be an integer sequence such that $k \rightarrow \infty$ and $k/n \rightarrow 0$ when $n \rightarrow \infty$. Then :

$$\begin{aligned} \hat{E}_C(X) &= \hat{\mu}_n^{(1,C)} + \hat{\mu}_n^{(2,C)}(k) \\ &:= \frac{1}{n_c} \sum_{i=1}^{n-k} \delta_{[i,n]} Z_{n,i} + \left[\left(Z_{n,n-k} - \frac{\hat{a}_{Z,k}^{(C)}}{\hat{\gamma}_{k,Z}^{(C)}} \right) (k/n) \right] \\ &\quad + \frac{\hat{a}_{Z,k}^{(C)}}{\hat{\gamma}_{k,Z}^{(C)}} (1 - \hat{F}_n(Z_{n,n-k}))^{\hat{\gamma}_{k,Z}^{(C)}} \left[\frac{(k/n)^{1-\hat{\gamma}_{k,Z}^{(C)}}}{1 - \hat{\gamma}_{k,Z}^{(C)}} \right]. \end{aligned}$$

Conclusion

Finally, we can deduce that the estimation procedure presented in this paper adapts to the needs of a heavy-tailed distribution's model, in the presence of right random censoring and the estimation seems to have an asymptotic behaviour towards a Gaussian law. These results are confirmed by numericals and graphical approaches and deserves a theoretical reflection on the convergence in probability of the proposed estimator to the mean of the distribution.

Références

- [1] Caeiro F., Gomes M.I., A new class of estimators of a scale second order parameter. *Extremes*, 2006. **(9)**, p. 193–211.
- [2] Caeiro F., Gomes M.I., Asymptotic comparaison at optimal levels of reduced-bias extreme value index estimators. *Statistica Neerlandica*, 2011. **65(4)**, p. 462–488.

Mean estimation for heavy tailed distribution's

- [3] Einmahl John, H.J., Fils-Villetard, A., Guillou, A., Statistics of extremes under random censoring. *Bernoulli*, 2008. 14(1) : p. 207–227.
- [4] Peng, L. (2001). Estimating the mean of a heavy tailed distribution. *Statistics & Probability Letters* **52(2001)**, 255–264.

Bayesian Inference About Markov Switching Models Using Augmented Data Algorithm

Ilhem BOUDERBALA*, Ghania SAIDI**
Hanya KHERCHI**

*USTHB, Faculté des Mathématiques, BP 32 El alia Bab Ezzouar, 16111, Algiers, Algeria,
bouderbala.ilhem@yahoo.fr,

** ENSSEA, 11, Chemin Doudou Mokhtar Ben-Aknoun, Algiers, Algeria,
ghsaidi@yahoo.fr,
hanya.kherchi@gmail.com.

Résumé. Les modèles autorégressifs à changement de régimes Markoviens sont considérés parmi les modèles non linéaires qui permet de capter au mieux les effets non linéaires qui existent dans la plupart des séries réelles. La méthode d'estimation retenue dans notre étude est basée sur une méthodologie Bayésienne, faisant appel aux algorithmes Bayésiens de Monte Carlo par Chaines de Markov, plus précisément en utilisant l'algorithme d'augmentation des données.

1 Introduction

The presence of the stylized facts, which are the major causes of non-linearity in most real world phenomena such as : *asymmetry, heavy tails distributions, multi-modality, the sudden change of dynamic, clustering volatility...* Make their treatment by linear models becomes inappropriate and the use of the nonlinear models to take care of these facts becomes indispensable. We focus here on the treatment of the Markov Switching Autoregressive models (MS-AR) (Hamilton, 1989, 1990).

Our aim through this study is the construction of a Bayesian methodology to estimate the parameters of the MS-AR model using a Bayesian iterative algorithms known as MCMC (Markov Chains Monte Carlo) algorithms (Geman et Geman, 1984), (Tanner et Wong, 1987), (Gelfand et Smith, 1990). More specifically using Augmented Data algorithm which is considered as a general case of the Gibbs sampler in the presence of missing data problem.

2 Markov Switching Autoregressive Models

We define $MS(K)\text{-}AR(p)$ model through a bivariate process $(Y_t, S_t)_{t \in T}$ such as :

- $\{S_t\}_1^T$ is an unobservable process defined as a homogeneous Markov chain in a finite and discrete stats space $E = \{1, 2, \dots, K\}$, K is the number of states in the model, equipped with transition matrix $\Pi = \{\pi_{ij}\} (i, j) \in E$ such as $\pi_{ij} = P(S_t = j | S_{t-1} = i)$;

- $\{Y_t\}_1^T$ is a piecewise linear process defined by

$$Y_t = \phi_0(S_t) + \sum_{i=1}^p \phi_i(S_t)Y_{t-i} + \varepsilon_t(S_t)$$

- $\varepsilon_t(k)$ characterize a white noise under each stat :

$$\varepsilon_t(k) \rightarrow N(0, \sigma^2(k)) \quad \forall k = 1, \dots, K$$

- θ represents the parameters to be estimated :

$$\theta = \{Q_1, \dots, Q_K, \Pi\} = \{Q, \Pi\}, Q_k = \{\phi(k), \sigma^2(k)\}, \phi'(k) = [\phi_0(k), \dots, \phi_p(k)]$$

3 Bayesian estimation

In our methodology (Bouderbala, 2012), (McCulloch et Tsay, 1994) we use *Augmented Data* algorithm initialized by Tanner and Wong (Tanner et Wong, 1987), Gelfand and Smith (Gelfand et Smith, 1990).

AD algorithm : The number of steps used in this algorithm are two. The first step is called *Imputation step* (noted I-step) and the second one is called *Posterior step* (noted P-step). Initialize $\theta^{(0)} \in \Theta$ and switch between these two steps for $l=1, 2, \dots, N$:

I-step : generate $S^{(l)} \rightarrow \pi(S|y, \theta^{(l-1)})$,

P-step : generate $\theta^{(l)} \rightarrow \pi(\theta|y, S^{(l)})$.

Since *AD* is a special case of MCMC algorithms, therefore the empirical mean of the vector θ calculated on a long trajectory converges almost surely to its theoretical counterpart given by

$$\frac{1}{N} \sum_{l=1}^N \theta^{(l)} \xrightarrow{a.s} E_{\pi}(\theta).$$

To make inference on the parameters of the *MS-AR* model using algorithm *AD*, we proceed with the following steps :

- Initialize the algorithm by assigning initial values $\theta^{(0)}$ to θ ,
 - Associate a values to the hyperparameters $\{\psi_{k,0}, A_{k,0}^{-1}, \alpha_k, \beta_k, \alpha_{k1}, \dots, \alpha_{kK}\} \forall k = 1, \dots, K$;
 - For $l=1, \dots, M, \dots, M+N$:
 - Generate the sequence $S^{(l)} = \{S_1^{(l)}, \dots, S_T^{(l)}\}$ using the *Multi-Move algorithm*;
 - Generate the vector of the parameters $\theta^{(l)}$ using the sequence S generated at the same iteration
- $S^{(l)}$. For $k=1, \dots, K$ we have

$$\phi_k^{(l)} | \sigma^{2(l-1)}(k), y_{(k)}^{(l)} \sim N_{p+1}(\psi_{k,*}^{(l)}, A_{k,*}^{-1(l)}), \quad \sigma^{2(l)}(k) | \phi_k^{(l-1)}, y_{(k)}^{(l)} \sim IG(\alpha_{k,*}^{(l)}, \beta_{k,*}^{(l)}),$$

$$\pi_k^{(l)} | S^{(l)} \sim Dir(\alpha_{k1} + n_{k1}^{(l)}, \dots, \alpha_{kK} + n_{kK}^{(l)}),$$

With

$$\begin{aligned}
 A_{k,*} &= \frac{1}{\sigma^2(k)} x'_{(k)} x_{(k)} + A_{k,0}, & \psi_{k,*} &= A_{k,*}^{-1} \left[\frac{1}{\sigma^2(k)} x'_{(k)} x_{(k)} \hat{\phi}_{(k)} + A_{k,0} \psi_{k,0} \right], \\
 \alpha_{k,*} &= \alpha_k + \frac{n_k - p}{2}, & \beta_{k,*} &= \frac{1}{2} [y_{(k)} - x_{(k)} \hat{\phi}_{(k)}]' [y_{(k)} - x_{(k)} \hat{\phi}_{(k)}] + \beta_k, \\
 n_{jk} &= \sum_{t=1}^T (S_{t-1} = j, S_t = k), & & j, k \in \{1, 2, \dots, K\}, \\
 y_{(k)} &= \begin{bmatrix} y_{(k)}^{p+1} \\ y_{(k)}^{p+2} \\ \vdots \\ y_{(k)}^j \\ \vdots \\ y_{(k)}^{n_k} \end{bmatrix} & x_{(k)} &= \begin{bmatrix} 1 & y_{(k)}^p & y_{(k)}^{p-1} & \cdots & y_{(k)}^1 \\ 1 & y_{(k)}^{p+1} & y_{(k)}^p & \cdots & y_{(k)}^2 \\ \vdots & \vdots & \vdots & \cdots & \vdots \\ 1 & y_{(k)}^{n_k-1} & y_{(k)}^{n_k-2} & \cdots & y_{(k)}^{n_k-p} \end{bmatrix}, \phi_k = \begin{bmatrix} \phi_0(k) \\ \phi_1(k) \\ \vdots \\ \phi_p(k) \end{bmatrix},
 \end{aligned}$$

n_k represent the number of observations belonging to the regime k ; $\hat{\phi}_{(k)}$ represents the Ordinary Least Squares (OLS) estimator of ϕ_k , $\hat{\phi}_{(k)} = (x'_{(k)} x_{(k)})^{-1} x'_{(k)} y_{(k)}$.

These posterior distributions dispose respectively the following independents prior :

$$\phi_k \sim N_{p+1}(\psi_{k,0}, A_{k,0}^{-1}); \sigma^2(k) \sim IG(\alpha_k, \beta_k); \pi_k \sim Dir(\alpha_{k1}, \dots, \alpha_{kK}).$$

Once the algorithm is complete, we exclude the first M simulations (*burn-in period*) and we keep the last values N . The final estimations are obtained for each parameter as follows for $k = 1, \dots, K$ and $t = 1, \dots, T$:

$$\begin{aligned}
 \hat{\phi}_k &= \frac{1}{N} \sum_{l=M+1}^{N+M} \phi_k^{(l)}; \sigma^2(k) = \frac{1}{N} \sum_{l=M+1}^{N+M} \sigma^{2(l)}(k); \hat{\pi}_k = \frac{1}{N} \sum_{l=M+1}^{N+M} \pi_k^{(l)}; \\
 P(S_t = k | Y_T, \theta) &= \frac{1}{N} \sum_{l=M+1}^{N+M} P(S_t^{(l)} = k | Y_T, \theta^{(l)}).
 \end{aligned}$$

We choose the values of the sequence $\{S_t\}_1^T$ that maximizes the smoothed probabilities $P(S_t = k | Y, \hat{\theta})$. In order to calculate the smoothed probabilities, we use *Multi-Move* algorithm (Carter et Kohn, 1996), (Jong et Shephard, 1995), (Schnatter, 2006) which enables the simulation of the sequence S in one iteration from the posterior conditional distribution $P(S | Y, \theta)$.

4 Conclusion

Our goal in this work is to estimate the parameters of the Markov Switching Autoregressive models by the Bayesian approach. To do this, we presented a methodology for estimating these parameters by Augmented Data algorithm which is a part of the MCMC algorithms. The purpose is to present an inference, which insures a good representative estimation of the theoretical values of the parameters of the MS-AR model.

Références

- Bouderbala, I. (2012). *Estimation Bayésienne des modèles Autorégressifs à changement de régimes Markoviens (MS-AR). Application : modélisation du Cycle des Affaires Algérien*. Mémoire de magister, ENSSEA. Algeria.
- Carter, C.-K. et R. Kohn (1996). Markov chains monte carlo in conditionally gaussian state space models. *Biometrika* 83, 598–601.
- Gelfand, A.-E. et A.-M. Smith (1990). Sampling based approaches to calculating marginal densities. *Journal of the American Statistical Association* 85, 398–409.
- Geman, S. et D. Geman (1984). Stochastic relaxation, gibbs distribution and the bayesian restoration of image. *IEEE*, 721–741.
- Hamilton, J.-D. (1989). A new approach to the economic analysis of nonstationary time series and the business cycle. *Econometrica* 57, 357–384.
- Hamilton, J.-D. (1990). Analysis of time series subject to change in regime. *Econometrica* 45, 39–70.
- Jong, P. et N. Shephard (1995). The simulation smoother for time series models. *Biometrika* 82, 339–350.
- McCulloch, R.-E. et S. Tsay (1994). Statistical analysis of economic time series via markov switching models.
- Schnatter, S.-F. (2006). *Finite Mixture and Markov Switching Models*. USA: Springer.
- Tanner, N.-A. et H.-W. Wong (1987). The calculation of posterior distribution by data augmentation. *Journal of the American Statistical Association* 82, 528–540.

Summary

The Markov switching autoregressive models are considered one of the nonlinear models which permit to catch some nonlinear effects that are present in a large category of phenomenon. In order to insure the reliability of the parameters we have chosen the Bayesian approach, by using the Markov Chains Monte Carlo algorithms, more specifically using the Augmented Data algorithm.

Application de la modélisation géostatistique dans l'industrie minière : cartographie de la variation spatiale des teneurs en P_2O_5 dans le gisement de phosphate de Kef Essennoun, Algérie

Mohamed Dassamiour*, Hamid Mezghache**

*Département des Sciences de la Terre, Institut d'Architecture et des Sciences de la Terre,
Université Sétif 1, Sétif 19000, Algérie

Dassamiour.m@univ-setif.dz

<http://www.univ-setif.dz/>

**Département des Sciences de la Terre, Faculté des Sciences de la Terre,
Université Badji Mokhtar, Annaba 23000, Algérie

mezghache.hamid@univ-annaba.org

<http://www.univ-annaba.dz/>

Résumé. Le gisement de phosphate de Kef Essennoun se présente sous forme d'une couche d'environ de 35 m d'épaisseur. Afin de localiser les zones les plus riches en minerai de phosphate, un modèle en blocs a été construit à l'aide de méthodes géostatistiques. Le gisement a été subdivisé selon la hauteur de la couche en quatre niveaux et chaque niveau est partagé en 1456 blocs d'exploitation. L'estimation des teneurs en P_2O_5 dans chaque bloc a été effectuée à l'aide de méthodes géostatistiques. Le variogramme expérimental moyen en trois dimensions a été construit à partir des données des échantillons des sondages et ajusté par un modèle de régionalisation exponentiel. Le système de krigeage ordinaire a été utilisé pour l'estimation des teneurs dans les blocs d'exploitation. Les résultats obtenus ont été représentés par des cartes de classes.

1 Introduction

Le glossaire de géostatistique (Olea, 1991) définit la géostatistique comme l'application de méthodes statistiques dans le domaine des sciences de la terre, notamment en géologie. La géostatistique fournit un ensemble d'outils pour le géologue qui les utilisera dans l'analyse des données et permet de construire le modèle numérique correspondant. Elle est largement reconnue aujourd'hui dans l'industrie minière comme étant un outil indispensable pour l'estimation des ressources et la quantification des incertitudes. La géostatistique est l'étude des variables régionalisées, Elle suppose que les caractéristiques spatiales voisines (variables régionalisées) sont spatialement corrélés les uns aux autres. Elle fournisse des informations sur la structure spatiale d'une variable, d'aider à déterminer la stratégie d'échantillonnage spatial optimale et de fournir une base quantitative pour interpoler des points inconnus aux endroits non échan-

tillonnés. Un outil couramment utilisée est la semivariogramme appelée aussi le variogramme (Matheron, 1963).

L'estimation de la variation spatiale des teneurs d'un minerai dans un gisement permet d'orienter l'exploitation selon la demande du marché et permet aussi de réduire les coûts de l'exploitation.

2 Géologie et minéralisation du gisement

Le gisement de Kef Essennoun fait partie du bassin phosphaté de Djebel Onk. Ce dernier est situé au Sud-Est de l'Algérie à 100 Km au Sud de la ville de Tébessa et à 20 Km de la frontière algéro-tunisienne, à environ 09 Km au Sud-Ouest de Bir El Ater. Les phosphates algériens sont liés aux dépôts marins du Tertiaire (Visse, 1952). La série sédimentaire du gisement est exprimée par une succession stratigraphique allant du Crétacé supérieur (Maestrichtien) à l'Eocène moyen (Lutétien), superposée par une série sablo-argileuse continentale datée du Miocène et du Quaternaire (Cielensky et al., 1987).

La formation phosphatée du gisement de Kef Essennoun est une couche d'âge Thanétien supérieur. Elle se présente comme une table monoclinale avec un pendage moyen de 15° vers le sud et une épaisseur d'environ 35 m. Elle s'étend sur plus de 2,7 km de longueur et 1,8 km de largeur.

3 Théorie et données utilisées

3.1 Théorie de géostatistique

La théorie de la géostatistique a été formulée par (Matheron, 1962) qui introduisit un outil pour analyser la continuité spatiale d'une variable régionale, appelé "variogramme" et une méthode d'estimation basée sur le variogramme appelé "krigeage". Le variogramme théorique est défini comme étant l'espérance quadratique de la variable aléatoire $[z(x) - z(x + h)]$ (Chauvet, 1999) :

$$2\gamma(h) = E\{[z(x) - z(x + h)]^2\} \quad (1)$$

L'équation habituelle pour calculer le variogramme est :

$$\gamma(h) = \frac{1}{2N(h)} \times \sum_{i=1}^n [z(x_i) - z(x_i + h)]^2 \quad (2)$$

Où $z(x_i)$ est la valeur de la variable Z dans le point x_i , h est le pas ou la distance entre les points de mesure et $N(h)$ est le nombre de couples entre les points d'échantillonnage ou de mesure distant de h .

Le variogramme expérimental est calculé pour plusieurs pas. Par conséquent celui-ci est ajusté par un modèle théorique, tel que le modèle sphérique et exponentiel. Ces modèles tiennent compte de l'information de la structure de la variation spatiale et qui servent comme

des paramètres utilisés pour l'estimation par krigeage.

La géostatistique désigne sous le nom de krigeage une méthode de construction d'un estimateur. Ainsi il convient de distinguer le krigeage en tant que méthode et les nombreux estimateurs qui résultent de son application, parmi lesquels le krigeage ordinaire. Le krigeage est considéré comme une méthode optimale pour l'estimation de la variation spatiale ou temporelle d'une variable.

L'estimation locale par krigeage ordinaire consiste à trouver le meilleur estimateur linéaire de la valeur moyenne d'une variable régionalisée ; sur un domaine de petite dimension ; vis-à-vis des zones de stationnarité du gisement. Dans notre cas c'est la teneur moyenne d'un bloc d'exploitation ($\bar{z}_i$). C'est une estimation pondérée de la moyenne :

$$\bar{z}_i = \sum_{i=1}^n \lambda_i z(x_i) \quad (3)$$

Où $\bar{z}_i$ est la valeur qui doit être estimée au point ou à l'espace limité i (dans notre cas, c'est le volume du bloc d'exploitation), $z(x_i)$ est la valeur connue au site échantillonné x_i et n est le nombre des échantillons utilisé pour l'estimation. Le formalisme de Lagrange fournit le système d'équations linéaires dit "système de krigeage ordinaire" à $n + 1$ équations linéaires et $n + 1$ inconnus qui contiennent les n pondérateurs et le paramètre de Lagrange μ (Matheron, 1963) :

$$\begin{cases} \sum_{i=1}^n \lambda_i \bar{\gamma}(x_i, x_j) + \mu = \bar{\gamma}(x_j, x_1) \\ \sum_{i=1}^n \lambda_i = 1 \end{cases} \quad (4)$$

pour $i=1$ à n et $j=1$ à n

Les pondérateurs dépendent des paramètres du modèle du variogramme et la configuration des échantillons et qui ont résolu sous les conditions de non biais et la minimisation de la variance d'estimation.

3.2 Données utilisées

Le gisement de Kef Essennoun a été exploré en détail par l'Office de Recherche Géologique et Minière par 29 sondages carottés réalisés à la maille de 250 x 300 m². La couche phosphatée du gisement est une couche inclinée avec un pendage moyen d'environ 15°. Les sondages réalisés par l'EREM sont implantés avec un angle vertical. Ces sondages montrent que l'épaisseur de cette couche n'est pas constante, elle varie entre 30 m (sondages S-20, S-41, S-42) et 53 m (sondage S-7), avec une épaisseur moyenne d'environ 39 m. Des échantillons de 1 m de longueur en moyenne ont été prélevés sur les tronçons de carottes, sur lesquels la teneur en P₂O₅ a été déterminée. Au total 966 échantillons ont été analysés.

Le gisement a été subdivisé en fonction de la géométrie et la composition lithologique et chimique de la couche phosphatée en quatre niveaux d'exploitation. Chaque niveau est divisé en 1456 blocs d'exploitation de dimensions (100 m de longueur x 20 m de largeur x 9 m de hauteur) pour le premier niveau (situé en haut de la couche phosphatée) et de dimensions (100 m de longueur x 20 m de largeur x 10 m de hauteur) pour les trois autres niveaux.

4 Résultats obtenus

4.1 Calcul et ajustement du variogramme expérimental de la variable "teneur en P_2O_5 "

Le variogramme expérimental moyen à trois dimensions des teneurs en P_2O_5 a été construit à partir des données des échantillons de 29 forages carottés. Il a été ajusté par un modèle de régionalisation exponentiel de portée (a) égale à 800 m, d'un effet de pépité (C_0) égale à 1,2%² et de palier (C) égale à 6%² :

$$\gamma(h) = 1,2 + 6 \times (1 - \exp(-h/800)) \quad (5)$$

Le krigeage ordinaire a été utilisé pour l'estimation des teneurs moyennes des blocs d'exploitation en P_2O_5 dans les quatre niveaux d'exploitation.

4.2 Paramètres et plan type de krigeage

Le plan type de krigeage ordinaire et ses paramètres utilisés lors de l'estimation de la teneur d'un bloc d'exploitation de support V ($100 \times 20 \times 10 \text{ m}^3$), à partir de l'ensemble des données des échantillons des sondages de support v_i sont :

- le nombre d'information (nombre des échantillons) retenus pour estimer la teneur d'un bloc V est de 20 à 56 échantillons
- le voisinage restreint de forme ellipsoïdale en 3D permet de mettre le krigeage à l'abri de tout risque de biais et pour faciliter le regroupement de l'information à l'intérieur du voisinage (zone d'influence). Vu que la portée du variogramme en 3D est de 800 m, la zone d'influence ne dépasse pas 800 m

4.3 Cartographie des résultats de krigeage

Pour visualiser les résultats de krigeage ainsi que les zones les plus riches dans le gisement, la méthode de cartographier par classes a été utilisée pour représenter les teneurs en P_2O_5 , ce qui a permis d'obtenir une carte de classes des teneurs en P_2O_5 dans les 1456 blocs de chaque niveau d'exploitation.

Références

- Chauvet, P. (1999). *Aide mémoire de géostatistique linéaire*. Paris : Les Presses de l'Ecole des mines de Paris.
- Cielensky, S., N. Benchermine, et T. Watkowski (1987). *Travaux de prospection et d'évaluation des phosphates dans la région de Bir El Ater*. Rapport interne, Entreprise de Recherche et d'Exploration Minière.
- Olea, R. (1991). *Geostatistical glossary and multilingual dictionary*. Oxford : Oxford University Press.
- Matheron, G. (1962). *Traité de géostatistique appliquée*. Paris : Technip.
- Matheron, G. (1963). *Traité de géostatistique appliquée : le krigeage*. Paris : Technip.

M. Dassamiour et H. Mezghache

Visse, L. (1952). *Genèse des gîtes phosphatés du sud-est Algéro-Tunisien*. Alger : XIXe Congrès Géologique International.

Summary

the purpose of this paper is to demonstrate the usefulness of the geostatistic application in the field of mining exploitation, which helps to take the decision and guide the exploitation works , as well as the reduction of the cost .

Comparative aspects of multistage designs for phase II clinical trials.

Zohra Djeridi*, Hayet Merabet**

*Mathematics Department, Jijel University, Jijel, Algeria
zdjeridi2002@yahoo.fr,

**Laboratoire de mathématiques appliqués et modélisation, Constantine1, Algeria
merabethammadi@outlook.fr

Abstract. The main objective of phase II clinical trials is to estimate treatment efficacy on a relatively small number of patients in order to determine whether it has sufficient biological activity against the disease under study to warrant more extensive development. Such trials are, often, conducted in a multi-institution setting where designs of more than two stages are difficult to manage. They play a primordial role in the drug development process, since the results determine whether or not to proceed to phase III trials. Since phase II clinical trial protocols and reports should include a description of the design selected with a justification for the particular choice, we will describe and compare the index of satisfaction with Simon's design and others respecting the statistical considerations of the protocol of a phase II trials.

1 Introduction

Two or three stage designs are commonly used in phase II cancer clinical trials. The aim of such phase is to determine whether a new treatment is promising for further testing in confirmatory clinical trials. Most exploratory clinical trials are designed as single-arm trials with or without interim monitoring for early stopping. A commonly used primary end point in phase II cancer clinical trials is the clinical response to a treatment, which is a binary end point such as tumor shrinkage defined as the patients achieving complete or partial response within a pre-defined treatment course. Overviews of statistical methods can be found in the work of Lee and Feng (2005). In this paper, we focus on the one arm phase II trials with binary endpoints. Multi-stage designs are often implemented in such settings to increase the study efficiency by allowing early termination if the treatment is deemed inefficacious. In those designs if no responses are observed in the first stage, the new treatment is abandoned. Zohar and *al.* (2008) emphasis that Bayesian approaches or designs have been proposed for single arm clinical trials as they take into account previous information about the quantity of interest as well as accumulated data during a trial. In this context, various approaches or designs have been proposed for single -arm clinical trials, like the PP design proposed by Lee and Liu (2008) and The single threshold design due to Tan and Machin (2002) which was extended into a design incorporate informative prior distribution by Mayo and Gajewsky (2004). We don't forget the recently work of Sambucini (2008) who suggested a Bayesian adaptive two-stage design in which the

sample size for the second stage depends on the results of the first stage. Motivated by all this studies, in this paper we propose a Bayesian adaptive design for single-arm exploratory clinical trials with binary outcomes based on the index of satisfaction created by Merabet (2004) and based on the p -value and on his Bayesian prediction. The proposed design consists of the calculation of the index of satisfaction and sample size selection at any required stage following interim analysis during the course of trial. Discussion: The aim of our work was to propose a simple, coherent and unified Bayesian adaptive design. The design with the index of satisfaction approach provides an excellent alternative for conducting multistage phase II trials. It is sufficient and flexible. It is based on the predictive probability and the sample size can be determined by choosing the smallest N_{max} among all designs that satisfy the design criteria. In the other hand, the interim monitoring using the predictive probability is like the method of Lee and Liu (2008) but we will use the analysis and design prior mentioned in Sambucini (2008).

As they emphasized, the predictive probability approach for interim monitoring is consistent efficient and flexible method that more closely resembles the clinical decision making process. Unlike the PP approach by Lee and Liu (2008) and PSSD approach by Teramukay and al. (2012), the index of satisfaction does not require intensive computation and comprehensive simulations may be required at the design phase to evaluate the operating characteristics (including type I and type II error probabilities) from the frequentist point of view. The effect of varying parameters or the prior distributions can be evaluated accordingly.

2 Statistic method

We recall (see Merabet (2004)), that the experimental context consists of two successive experimentations, of results $\omega' \in \Omega'$ and $\omega'' \in \Omega''$, which are in general carried out independently. Their distributions built in the framework of a well established model, depend on a parameter $\theta \in \Theta$, only ω'' is used to found the official conclusion of the study and to determine the user satisfaction denoted $\phi(\omega'')$ (and on the choice of L about which we will come back in 3). But, on the basis of the result ω' of first step clinical trial, it is useful to anticipate what the satisfaction will be well after the second step. In our study, this prediction is carried out in a bayesian context, i.e., based on the choice of a prior probability on Θ .

2.1 Index of satisfaction

This concept is important when the statistician, who carries out a test, "wishes" to observe a significant result, that is to reject the null hypothesis H_0 . Its satisfaction will be thus larger in the event of rejection, and even in general as much larger as the observation that leads to this rejection is more significant. Its what even the users put in an obvious place, at the end of a test, not only the conclusion "all or no thing" (significant or not significant) but also the smaller value of threshold for which the result y observed will be considered significant that is in theory of the test, the p -value. Being α fixed, let a test of level α be defined by a critical region $\Omega_1^{''(\alpha)}$. It is more interesting to take into account to what level will be the results always appear significant. We will use a new index of satisfaction, that was study by Merabet (2004) , and defined for the bayesian tests, based on the same prior P_Θ .

We propose to calculate the prediction of satisfaction in the case of binary outcomes for $L(p)=(1-p)$ when the prior distribution of the unknown parameter θ is uniform non-informative analysis prior or the design prior proposed by Sambucini (2008) and in the case of a test of threshold, α where the null hypothesis is of type $\theta \leq \theta_0$.

3 Discussion

The aim of our work was to propose a simple, coherent and unified Bayesian adaptive design. The design with the index of satisfaction approach provides an excellent alternative for conducting multistage phase II trials. It is sufficient and flexible. It is based on the predictive probability and the sample size can be determined by choosing the smallest N_{max} among all designs that satisfy the design criteria. In the other hand, the interim monitoring using the predictive probability is like the method of Lee and Liu (2008) but we will use the analysis and design prior mentioned in Sambucini (2008). As they emphasized, the predictive probability approach for interim monitoring is consistent efficient and flexible method that more closely resembles the clinical decision making process. Unlike the PP approach by Lee and Liu (2008) and PSSD approach by Teramukai and al. (2012), The index of satisfaction does not require intensive computation and comprehensive simulations may be required at the design phase to evaluate the operating characteristics (including type I and type II error probabilities) from the frequentist point of view. The effect of varying parameters or the prior distributions can be evaluated accordingly.

Bayesian methods provide many appealing properties for the clinical trials designs and analysis. This include the ability to incorporate the information obtained before the study design to use the information obtained during the trial for monitoring the study, flexibility in trial conduct a consistent way for making inference.

Résumé

L'objectif principal de phase II des essais cliniques est d'estimer l'efficacité du traitement sur un nombre relativement restreint de patients afin de déterminer si elle a une activité biologique suffisante contre la maladie à l'étude pour justifier un développement plus extensif. Ces essais sont souvent menés dans un cadre multi-organismes où les designs de plus de deux étapes sont difficiles à gérer. Ils jouent un rôle primordial dans le processus de développement de médicaments, car les résultats déterminent si on procède ou pas à la phase III des essais. Comme les protocoles et les rapports de phase II d'essais cliniques doivent inclure une description du design choisi avec une justification du choix particulier, nous allons décrire et comparer le design d'indice de satisfaction avec le design de Simon et d'autres en respectant les considérations statistiques du protocole de phase II d'un essai.

An R package for modeling stochastic and diffusion processes

Arsalane Chouaib Guidoum*, Kamal Boukhetala*

*Faculty of mathematics-USTHB, PoBox 32, Bab Ezzouar, Algiers, Algeria
acguidoum@usthb.dz
kboukhetala@usthb.dz

Résumé. The package `Sim.DiffProc` is an object created in the R language for simulating of stochastic differential equations (SDEs), and statistical analysis of diffusion processes solution of SDEs. This package contains many objects (code/function), for example the numerical methods to find the solutions of SDEs (one, two and three dimensional), which can be simulate a flows trajectory, with good accuracy. Statistical analysis can be done to estimated the transformations of diffusion processes and its laws in a fixed time T , or in a fixed time interval $[t_0, T]$, and all others random variables witches can be generated by this process, on the basis of the statistical information in the flow trajectories, as for studying various properties of diffusion.

1 Introduction

Many theoretical problems on the SDEs have become the object of practical research, enabled many searchers in different domains to use these equations to modeling and to analyse practical problems. We seek to motivate further interest in this specific field by introducing the `Sim.DiffProc` package (Guidoum et Boukhetala, 2014) to simulate the solution of a user defined Itô or Stratonovich uni- and multi-dimensional SDEs, estimate parameters from data and visualize statistics and other features, for example the determination of the first passage time in SDEs using the Monte Carlo method; the package (currently version 2.9) is implemented using S3 classes and freely available on the Comprehensive R Archive Network (CRAN) at <http://CRAN.R-project.org/package=Sim.DiffProc>. There already exist a number of packages that can perform for stochastic calculus in R; see `sde` (Iacus, 2014) and `yuima` project package for SDEs (Brouste et al., 2014) a freely available on CRAN, this packages provides functions for simulation and inference for stochastic differential equations. It is the accompanying package to the book of Iacus (2008).

2 Using the Sim.DiffProc package

We can write an d -dimensional SDE in Itô form as :

$$d\mathbf{X}_t = \mathbf{F}(t, \mathbf{X}_t)dt + \mathbf{G}(t, \mathbf{X}_t)d\mathbf{W}_t \quad (1)$$

or in Stratonovich form as :

$$d\mathbf{X}_t = \mathbf{F}(t, \mathbf{X}_t)dt + \mathbf{G}(t, \mathbf{X}_t) \circ d\mathbf{W}_t \quad (2)$$

where $\mathbf{F}(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is called the *drift* of the SED's, $\mathbf{G}(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^{d \times m}$ is called the *diffusion* of the SDE's, and $\mathbf{W}_t$ is an m -dimensional process having independent¹ scalar Wiener process components. It is possible to convert from one interpretation to the other in order to take advantage of one of the approaches as appropriate. For multidimensional SDE's the relationship between the two representations is given by :

$$\underline{\mathbf{F}}_i(t, X_t) = \mathbf{F}_i(t, X_t) - \frac{1}{2} \sum_{j=1}^d \sum_{k=1}^m \mathbf{G}_{jk}(t, X_t) \frac{\partial \mathbf{G}_{ik}}{\partial X_j}(t, X_t), \quad i = 1, \dots, d.$$

More in detail, the user can specify :

- The Itô or the Stratonovich SDE's to be simulated.
- The SDE's structural parameter value. i.e., the drift and diffusion coefficient of SDE's.
- The number of the SDE's solution trajectories to be simulated.
- The numerical integration method : Euler-Maruyama, Predictor-corrector, Milstein, Second Milstein, Itô Taylor order 1.5, Heun order 2 ; Runge-Kutta 1,2 and 3-stage. There a rich literature on simulation of solutions of the SDE's.
- The time interval $[t_0, T]$ to be considered.
- The integration stepsize (discretization).

To obtain :

- Numerical solution of SDE's.
- Plot(s) of the solution trajectories.
- Plot(s) of the trajectories empirical mean, together with their $\alpha\%$ confidence bands.
- Monte-Carlo statistics of the solution process at the end time T , i.e. mean, median, quantiles, moments, skewness, kurtosis, $\alpha\%$ confidence bands,...

2.1 The snssde1d() function

Assume that we want to describe the following SDE in Itô² form :

$$dX_t = \frac{1}{2}\mu^2 X_t dt + \mu X_t dW_t, \quad X_0 = x_0 \quad (3)$$

in Stratonovich form :

$$dX_t = \frac{1}{2}\mu^2 X_t dt + \mu X_t \circ dW_t, \quad X_0 = x_0 \quad (4)$$

1. In this version of the package, multidimensional SDE's need to have diagonal noise.
 2. The equivalently of (3) the following Stratonovich SDE : $dX_t = \mu X_t \circ dW_t$.

In the above $F(t, x) = \frac{1}{2}\mu^2 x$ and $G(t, x) = \mu x$, according to the notation of the (1) in the case $d = 1$ and W_t is a standard Wiener process ($m = 1$). This can be described in `Sim.DiffProc` by specifying the drift and diffusion coefficients as plain R expressions passed as strings which depends on the state variable `x` and time variable `t`, by specifying only one trajectory (`M=1`) in $[t_0, T] = [0, 1]$, with integration stepsize $\Delta t = 0.001$ (by default : `Dt=(T-t0)/N`), $\mu = 0.5$ and $X_0 = 10$. specifying the type of SED by `type="ito"` or `type="str"` (by default `type="ito"`), and the numerical method used (by default `method="euler"`).

```
> f <- expression( (0.5*0.5^2*x) )
> g <- expression( 0.5*x )
> mod1 <- snssde1d(drift=f,diffusion=g,x0=10,M=50,N=1000,type="ito")
> mod2 <- snssde1d(drift=f,diffusion=g,x0=10,M=50,N=1000,type="str")
> mod1
```

Ito Sde 1D:

```
| dX(t) = (0.5 * 0.5^2 * X(t)) * dt + 0.5 * X(t) * dW(t)
```

Method:

```
| Euler scheme of order 0.5
```

Summary:

```
| Size of process          | N = 1000.
| Number of simulation     | M = 50.
| Initial value           | x0 = 10.
| Time of process         | t in [0,1].
| Discretization          | Dt = 0.001.
```

```
> mod2
```

Stratonovich Sde 1D:

```
| dX(t) = (0.5 * 0.5^2 * X(t)) * dt + 0.5 * X(t) o dW(t)
```

Method:

```
| Euler scheme of order 0.5
```

Summary:

```
| Size of process          | N = 1000.
| Number of simulation     | M = 50.
| Initial value           | x0 = 10.
| Time of process         | t in [0,1].
| Discretization          | Dt = 0.001.
```

If we simulate 50 trajectories and let the settings above unchanged (except for the number of simulations, of course); Using Monte-Carlo simulations, the following statistical measures (`S3 method` for class '`snssde1d`') can be approximated for the X_t process at the end time T , i.e. X_T :

1. the expected (mean) value $\mathbb{E}(X_T)$; using the command `mean`.
2. the variance $\text{var}(X_T)$.
3. the median $\text{Med}(X_T)$; using the command `median`.

4. the quartile of X_T ; using the command `quantile`.
5. the skewness and the kurtosis of X_T ; using the command `skewness` and `kurtosis`.
6. the moments of X_T ; using the command `moment`.
7. the $\alpha\%$ confidence bands of X_T ; using the command `bconfint`.

Can be use the `summary` function to produce result summaries of the results of class `'snssde1d'`. The flow of trajectories can be seen in Figure 1.

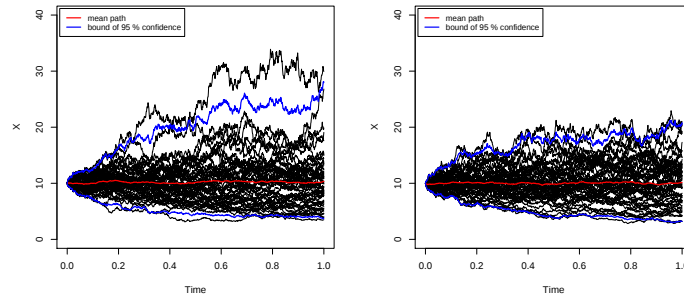


FIG. 1 – 50 paths of Itô SDE 'mod1' (Left), and Stratonovich SDE 'mod2' (Right).

2.2 The `snssde2d()` function

Assume that we want to describe the following SDE (2d) in Itô form :

$$\begin{cases} dX_t = 4(-1 - X_t)Y_t dt + 0.2dW_{1,t} \\ dY_t = 4(1 - Y_t)X_t dt + 0.2dW_{2,t} \end{cases} \quad (5)$$

for (5), we simulate a flow of 50 trajectories, with integration stepsize $t = 0.001$, and using stochastic Runge-Kutta methods 3-stage,

```
> fx <- expression(4*(-1-x)*y)
> gx <- expression(0.2)
> fy <- expression(4*(1-y)*x)
> gy <- expression(0.2)
> mod2d <- snssde2d(driftx=fx,diffx=gx,drifty=fy,diffy=gy,x0=1,y0=-1,
+                  M=50,Dt=0.001,method="rk3")
> mod2d
```

Itô Sde 2D:

$$\begin{aligned} | \quad dX(t) &= 4 * (-1 - X(t)) * Y(t) * dt + 0.2 * dW_1(t) \\ | \quad dY(t) &= 4 * (1 - Y(t)) * X(t) * dt + 0.2 * dW_2(t) \end{aligned}$$

Method:

| Runge-Kutta method of order 3

Summary:

Size of process	N = 1000.
Number of simulation	M = 50.
Initial values	(x0,y0) = (1,-1).
Time of process	t in [0,1].
Discretization	Dt = 0.001.

for plotted (with time) using the command `plot`, and in the plane (O, X, Y) using the command `plot2d`. The result is shown in Figure 2,

```
> plot(mod2d,pos=2)
> plot2d(mod2d)
```

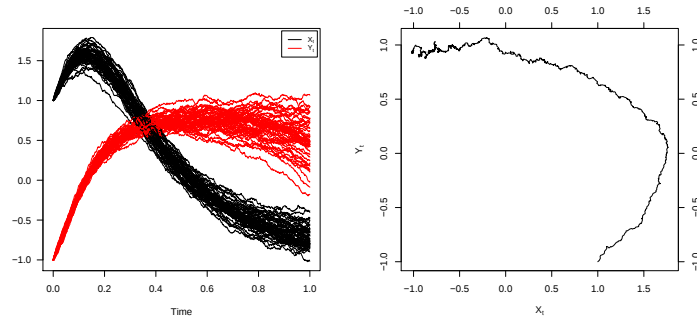


FIG. 2 – 50 paths of (5) (Left), representation of (5) in a plane (O, X, Y) (Right).

2.3 The `snssde3d()` function

We propose the following dispersion models family (Boukhetala, 1996) :

$$\begin{cases} dX_t = \left(\frac{-KX_t}{\sqrt{X_t^2 + Y_t^2 + Z_t^2}} \right) dt + \sigma dW_{1,t} \\ dY_t = \left(\frac{-KY_t}{\sqrt{X_t^2 + Y_t^2 + Z_t^2}} \right) dt + \sigma dW_{2,t} \\ dZ_t = \left(\frac{-KZ_t}{\sqrt{X_t^2 + Y_t^2 + Z_t^2}} \right) dt + \sigma dW_{3,t} \end{cases}, t \in [0, T] \quad (6)$$

Which can easily be implemented (6) in R as follows :

```
> K = 4; sigma = 0.2
> fx <- expression( (-K*x/sqrt(x^2+y^2+z^2)) )
> gx <- expression(sigma)
```

Sim.DiffProc: Simulation of Diffusion Processes

```
> fy <- expression( (-K*y/sqrt(x^2+y^2+z^2)) )
> gy <- expression(sigma)
> fz <- expression( (-K*z/sqrt(x^2+y^2+z^2)) )
> gz <- expression(sigma)
> mod3d <- snssde3d(driftx=fx,diffx=gx,drifty=fy,diffy=gy,driftz=fz,
+                  diffz=gz,N=10000,x0=1,y0=1,z0=1)
> mod3d
```

Ito Sde 3D:

```
| dX(t)=(-K*X(t)/sqrt(X(t)^2+Y(t)^2+Z(t)^2))*dt+sigma*dW1(t)
| dY(t)=(-K*Y(t)/sqrt(X(t)^2+Y(t)^2+Z(t)^2))*dt+sigma*dW2(t)
| dZ(t)=(-K*Z(t)/sqrt(X(t)^2+Y(t)^2+Z(t)^2))*dt+sigma*dW3(t)
```

Method:

```
| Euler scheme of order 0.5
```

Summary:

```
| Size of process      | N = 10000.
| Number of simulation | M = 1.
| Initial values      | (x0,y0,z0) = (1,1,1).
| Time of process     | t in [0,1].
| Discretization      | Dt = 1e-04.
```

for plotted (with time) using the command `plot`, and in the space (O, X, Y, Z) using `plot3D` with two display types ("`rgl`", "`persp`"), the first with `rgl` package and the second display with `scatterplot3d` package. The result is shown in Figure 3,

```
> plot3D(mod3d,display="persp",col="blue") ## in space
> plot(mod3d,union=TRUE,pos=2)             ## with time
```

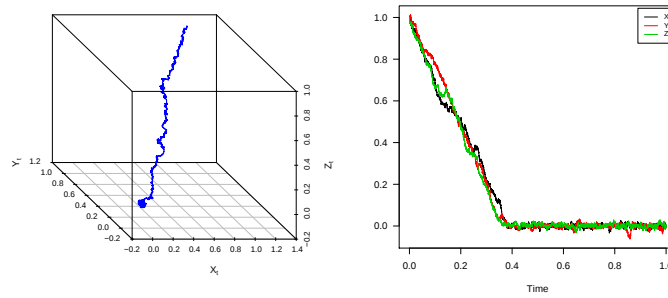


FIG. 3 – 3-dimensional attractive model $\mathcal{M}(K = 4, \sigma = 0.2)$.

3 Summary

This work is about ready to be used `Sim.DiffProc` package for simulation of stochastic differential equations and some related estimation methods based on discrete sampled observations from such models. We hope that the package presented here and the updated survey on the subject might be of help for practitioners, postgraduate and PhD students, and researchers in the field who might want to implement new methods and ideas using R as a statistical environment. The simulation studies implemented in R language seem very preferment and efficient, because it is a statistical environment, which permits to realize, to visualize and validate the simulations.

Références

- Boukhetala, K. (1996). *Modelling and Simulation of a Dispersion Pollutant with Attractive Centre*, Volume 3, pp. 245–252. Boston, USA : Computer Methods and Water Resources, Computational Mechanics Publications.
- Brouste, A., M. Fukasawa, H. Hino, S. M. Iacus, K. Kamatani, Y. Koike, H. Masuda, R. Nomura, T. Ogihara, Y. Shimuzu, M. Uchida, et N. Yoshida (2014). The yuima project : A computational framework for simulation and inference of stochastic differential equations. *Journal of Statistical Software* 57(4), 1–51.
- Guidoum, A. C. et K. Boukhetala (2014). *Sim.DiffProc : Simulation of Diffusion Processes*. R package version 2.9.
- Iacus, S. M. (2008). *Simulation and Inference for Stochastic Differential Equations : with R examples*. New York : Springer-Verlag.
- Iacus, S. M. (2014). *sde : Simulation and Inference for Stochastic Differential Equations*. R package version 2.0.13.

Validation de modèles de régression linéaire dans le cas non régulier

Zaher Mohdeb

Département de Mathématiques, Université Constantine 1, Algérie
zaher.mohdeb@umc.edu.dz

Résumé. Une procédure de test d'hypothèse linéaire sur la fonction de régression f dans un modèle de régression non paramétrique est proposée. Plus précisément, on teste l'hypothèse que f est un élément de E_p , où E_p est un espace vectoriel de dimension finie. En supposant que les fonctions considérées sont Riemann-intégrables et on obtient le comportement asymptotique du test proposé, on a donc ainsi le niveau et la puissance asymptotique du test.

1 Introduction

L'objet du présent travail est de construire des tests d'hypothèses linéaires dans le modèle de régression non paramétrique sans condition de régularité sur la fonction de régression aussi bien sous l'hypothèse nulle que sous l'alternative. On considère donc le modèle de régression suivant

$$Y_{i,n} = f(t_{i,n}) + \varepsilon_{i,n}, \quad i = 1, \dots, n, \quad (1)$$

où f est une fonction réelle inconnue, définie sur l'intervalle $[0, 1]$ et $t_{i,n}, i = 1, \dots, n$, est un échantillonnage fixé de l'intervalle $[0, 1]$. Les erreurs $\varepsilon_{i,n}$ forment un tableau triangulaire de variables aléatoires d'espérance nulle et de variance finie σ^2 .

Soient $x_1(t), \dots, x_p(t)$ des fonctions définies sur $[0, 1]$ et linéairement indépendantes et soit E_p l'espace vectoriel engendré par $x_1, \dots, x_p$. On veut tester l'hypothèse :

$$H_0 : f \in E_p \quad \text{contre} \quad H_1 : f \notin E_p. \quad (2)$$

La plupart des travaux sur les tests d'hypothèses dans le modèle (1) supposent des conditions de régularité sur $f, x_1, \dots, x_p$; généralement ces fonctions sont supposées höldériennes. On peut citer Eubank et Spiegelman (1990), Eubank et Hart (1992) et Härdle et Mammen (1993). Dette et Munk (1998) ont abordé le test (2) avec l'hypothèse f höldérienne d'ordre $\gamma > 1/2$. Mohdeb et Mokkadem (2004) proposent également une autre statistique de test basée sur une autre estimation du carré de la distance de f à E_p .

Dans ce travail, on suppose que $f, x_1, \dots, x_p$ sont Riemann-intégrables ; sous cette seule condition sur les fonctions, on établit un théorème de convergence en loi qui permet de construire des tests avec des hypothèses non régulières.

2 Résultat principal

On considère le modèle de régression (1) et E_p est l'espace vectoriel engendré par des fonctions fixées $x_1(t), \dots, x_p(t)$ définies sur $[0, 1]$ et linéairement indépendantes. Nos hypothèses sont les suivantes :

- (A1) $\max_{i=2, \dots, n} \left| (t_{i,n} - t_{i-1,n}) - \frac{1}{n} \right| = o\left(\frac{1}{n}\right)$;
- (A2) $\forall n, \varepsilon_{1,n}, \dots, \varepsilon_{n,n}$ sont indépendantes et $\exists C \in \mathbb{R}^+$ tel que $E(\varepsilon_{i,n}^4) < C, \forall i, n$;
- (A3) Les fonctions $f, x_1, \dots, x_p$ sont Riemann-intégrables.

Les fonctions que nous considérons, sont aussi dans $L^2(dt)$ muni de son produit scalaire usuel

$$\langle u, v \rangle = \int_0^1 u(t) v(t) dt \quad \text{et} \quad \|u\|_2^2 = \int_0^1 u^2(t) dt, \quad u, v \in L^2(dt).$$

On pose

$$\mathcal{D}^2(f) = \min_{v \in E_p} \|f - v\|_2^2$$

le carré de la distance de f à E_p .

Pour construire notre statistique de test, on utilise une estimation empirique de $\mathcal{D}^2(f)$.

Pour cela, on pose $Y_n = (Y_{1,n}, \dots, Y_{n,n})'$, $\varepsilon_n = (\varepsilon_{1,n}, \dots, \varepsilon_{n,n})'$, $f_n = (f(t_{1,n}), \dots, f(t_{n,n}))'$, $x_{k,n} = (x_k(t_{1,n}), \dots, x_k(t_{n,n}))'$, $k = 1, \dots, p$, et $X = (x_{1,n}, \dots, x_{p,n})$ c'est-à-dire la matrice dont le $(i, j)^{ieme}$ élément est $x_j(t_i)$, $i = 1, \dots, n$ et $j = 1, \dots, p$.

On note aussi $E_{p,n}$, le sous-espace de $\mathbb{R}^n$ engendré par $\{x_{1,n}, \dots, x_{p,n}\}$.

Dans $\mathbb{R}^n$, on utilise le produit scalaire usuel et on pose

$\Pi_n = X(X'X)^{-1}X'$, la matrice de projection sur $E_{p,n}$ et

$\Pi_n^\perp = I_n - X(X'X)^{-1}X'$, la matrice de projection sur l'espace orthogonal de $E_{p,n}$, où I_n est la matrice identité $n \times n$.

Posons aussi

$$\mathcal{D}_n^2 = \frac{1}{n} Y_n' \Pi_n^\perp Y_n \quad \text{et} \quad \tilde{\mathcal{D}}_n^2 = \frac{1}{n} f_n' \Pi_n^\perp f_n.$$

Notons que si $\mathcal{D}_n^2$ est suffisamment petit, on en déduit que f est proche de sa projection sur E_p , ainsi notre statistique de test sera basée sur cette expression empirique.

En remplaçant Y_n par $f_n + \varepsilon_n$, on vérifie que

$$E(\mathcal{D}_n^2) = \tilde{\mathcal{D}}_n^2 + \frac{n-p}{n} \sigma^2.$$

On est amené ainsi à considérer $\mathcal{D}_n^2 - \frac{n-p}{n} \sigma^2$, mais σ^2 est inconnu. On l'estime à l'aide de l'estimateur introduit par Gasser, Sroka, et Jennen-Steinmetz (1986)

$$\hat{\sigma}^2 = \frac{1}{6(n-2)} \sum_{i=2}^{n-1} (Y_{i+1} + Y_{i-1} - 2Y_i)^2.$$

On obtient la statistique de test

$$\widehat{\mathcal{D}}_n^2 = \mathcal{D}_n^2 - \frac{n-p}{n} \widehat{\sigma}^2;$$

on rejette l'hypothèse $H_0 : "f \in E_p"$ si

$$\widehat{\mathcal{D}}_n^2 > c_\alpha,$$

où c_α est un nombre réel positif.

Notre résultat principal est le suivant.

Theoreme 1 *Si les conditions (A1), (A2) et (A3) sont satisfaites, alors*

$$\sqrt{n} \left\{ \widehat{\mathcal{D}}_n^2 - \widetilde{\mathcal{D}}_n^2 + B_n(f) \right\} \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N} \left(0, \frac{17}{9} \sigma^4 + 4\sigma^2 \mathcal{M}^2 \right),$$

$$\text{où } B_n(f) = \frac{1}{6n} \sum_{i=2}^{n-1} \left(f(t_{i+1,n}) + f(t_{i-1,n}) - 2f(t_{i,n}) \right)^2 \text{ et } \mathcal{D}^2 = \mathcal{D}^2(f).$$

Remarque. Sous H_0 , on a $\widetilde{\mathcal{D}}_n^2 = 0$ mais $\sqrt{n} B_n(f)$ n'est pas nécessairement nul, ni même négligeable en général.

3 Application

3.1 Test dans un modèle de régression localement höldérienne

On suppose que, dans le modèle de régression étudié, f est une fonction localement höldérienne d'ordre inconnu, (ou seulement Riemann-intégrable) et on considère un espace vectoriel E_p tel que

- (A4) les fonctions $x_1, \dots, x_p$ sont localement höldériennes d'ordre $\gamma > 1/2$.

On a donc la proposition.

Proposition 1 *Si les conditions (A1), (A2) et (A4) sont satisfaites on a, sous H_0 ,*

$$\sqrt{n} \widehat{\mathcal{D}}_n^2 \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N} (0, \sigma^4) .$$

Cette proposition donne le niveau asymptotique du test; le théorème 1 donne la puissance pour des alternatives qui peuvent être höldériennes d'ordre $\delta < 1/2$, (ou seulement Riemann-intégrable).

Bibliographie

- [1] Dette, H. and Munk, A. (1998). Validation of linear regression models. *Ann. Stat.*, **26**, 778-800.
- [2] Eubank, R. L. and Hart, J. D. (1992). Testing goodness-of-fit in regression using nonparametric via order selection criteria. *Ann. Stat.*, **20**, 1412-1425.
- [3] Eubank, R. L. and Spiegelman, C. H. (1990). Testing the goodness-of-fit of a linear model via nonparametric regression techniques. *J. Amer. Stat. Assoc.* **85**, 410, 387-392.
- [4] Härdle, W. and Mammen, E. (1993). Comparing nonparametric versus parametric regression fits. *Ann. Stat.*, **21**, 1926-1947.
- [5] Gasser, T., Sroka L. and Jennen-Steinmetz, C. (1986). Residual variance and residual pattern in nonlinear regression. *Biometrika*, **73**, 625-633.
- [6] Mohdeb, Z. and Mekkadem, A. (2004). Average squared residuals approach for testing linear hypotheses in nonparametric regression. *J. Nonparametr. Stat.* **16**, no. 1-2, 3-12.

Summary

A procedure for testing linear hypothesis on the regression function f in a nonparametric regression model. More precisely, we test that f is an element of E_p , where E_p is a finite dimensional vector space. We assume that the functions are Riemann-integrable, and we obtain the asymptotic weak behaviour of the proposed test, then we have the level and the asymptotic power of the test.

Méthode du noyau dans l'analyse de stabilité forte d'un modèle de risque

Atik Touazi, Zina Benouaret
Smail Adjabi, Djamil Aissani

Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes),
Université de Béjaia, Algérie
touazi_atik@hotmail.fr
<http://www.lamos.org>

Résumé. Dans ce travail, nous estimons la densité de probabilité f du montant de réclamation, en utilisant le noyau Gamma modifié, pour l'étude de stabilité forte du modèle de risque classique P/P (loi des inter-sinistres et des montants de réclamation exponentielles). L'approche simulation sera utilisée afin d'évaluer numériquement l'erreur d'approximation entre les probabilités de ruine du modèle de risque idéal P/P et du modèle de risque perturbé P/G (loi exponentielle des inter-sinistres et loi générale du montant de réclamation).

1 Introduction

Dans les modèles de risque, la probabilité de ruine est la mesure de risque la plus étudiée dans la littérature. En général, cette mesure est très difficile à évaluer d'une manière explicite, c'est pourquoi on a recours à différentes méthodes d'approximation. La méthode de stabilité forte, qui a été élaborée par Aissani et Kartashov (1983), connaît un large champs d'application en théorie de ruine après le travail de l'académicien Kalashnikov (2000). Ce dernier a présenté de nouvelles bornes de stabilité des probabilités de ruine. Dans ce sens, plusieurs travaux ont été réalisés sur différents modèles : le modèle de risque avec investissement (Rusaityte, 2001) ; les modèles de risque semi-markovien sans investissement (Enikeeva et al., 2001) et le modèle de risque classique à deux dimensions (Benouaret et Aissani, 2010).

Cet article étudie, en utilisant l'estimation non paramétrique par la méthode du noyau, la stabilité forte des probabilités de ruine dans un modèle de risque. En supposant que le nombre de réclamations suit une loi de Poisson, nous clarifions, par l'application de la méthode de stabilité forte, les conditions d'approximation d'un modèle de risque de distribution inconnue des montants de réclamations par le modèle de risque où le montant de réclamations suit une loi exponentielle, avec une estimation de l'erreur de cette approximation.

2 Stabilité forte du modèle de risque classique

Le modèle de risque classique à une dimension est décrit par le processus suivant :

$$X(t) = u + ct - \sum_{i=1}^{N(t)} Z_i, \quad t \geq 0, \quad (1)$$

où, u est le surplus initial, c est le taux de prime constant, $\{N(t), t \geq 0\}$ est un processus de Poisson de paramètre λ , $\{Z_i\}_{i \in \mathbb{N}^*}$ est une suite de variable aléatoire indépendante et identiquement distribuées où Z_i est Le montant du $i^{\text{ème}}$ sinistre, de fonction de distribution F et de moyenne μ finie.

Définition 2.1 Nous appelons probabilité de ruine en temps fini t , la fonction donnée par

$$\Psi(u, t) = P(\exists s \in [0, t] / X(s) < 0), \quad \forall u \geq 0.$$

– En temps infini, elle est définie comme suit

$$\Psi(u, \infty) = P(\exists s \geq 0 / X(s) < 0) = \Psi(u), \quad \forall u \geq 0.$$

Le modèle de risque classique (P/P) est fortement stable par rapport à la fonction poids $v(x) = e^{\epsilon x}$, $x \geq 0$.

Notons par $a = (\lambda, c, F)$, [respectivement $a' = (\lambda', c', F')$] le vecteur de paramètres du modèle de risque idéal (respectivement perturbé).

L'estimation de la déviation entre les noyaux de transition, obtenue par Kalashnikov (2000) est donnée par la formule suivante :

$$\|P - P'\|_v \leq 2 E e^{\epsilon Z} \ln \frac{\lambda c'}{\lambda' c} + \|F - F'\|_v. \quad (2)$$

Sous la condition suivante qui représente le domaine de perturbation des paramètres :

$$u(a, a') = 2 E e^{\epsilon Z} \ln \frac{\lambda c'}{\lambda' c} + \|F - F'\|_v \leq (1 - \rho)^2, \quad (3)$$

Nous avons l'inégalité de stabilité suivante :

$$\|\Psi - \Psi'\|_v \leq \Gamma = \frac{\mu(a, a')}{(1 - \rho)((1 - \rho)^2 - \mu(a, a'))} \quad (4)$$

où $\rho(\epsilon) = E \exp(\epsilon(Z - c \theta))$, θ est une variable aléatoire qui représente les inter-arrivées des réclamations suivant une loi exponentielle de paramètre λ .

3 Méthode du noyau pour l'estimation de la densité du montant de réclamation

Soit $X_1, \dots, X_n$ un n-échantillon issu d'une variable aléatoire X de la fonction de densité inconnue f sur l'ensemble $\mathcal{R} \subseteq R$ borné au moins d'un côté. Les estimateurs à noyau associé, asymétrique et continu sont de la forme :

$$f_h(x) = \frac{1}{n} \sum_{i=1}^n K_{x,h}(X_i). \quad (5)$$

où h est le paramètre de lissage et $K_{x,h}$ est le noyau associé continu asymétrique.

Le noyau associé gamma a été introduit par Chen (2000) pour estimer des densités à support $\mathcal{R} = [0, \infty[$. Deux classes de noyaux ont été proposées :

$$K_{GAM(\frac{x}{h}+1,h)}(u) = \frac{u^{\frac{x}{h}} \exp(\frac{-u}{h})}{h^{\frac{x}{h}+1} \Gamma(\frac{x}{h} + 1)}, \quad (6)$$

où $\Gamma(\alpha) = \int_0^\infty \exp(-t)t^{\alpha-1}dt$. Son estimateur associé est donné par :

$$f_h^{GAM}(x) = \frac{1}{n} \sum_{i=1}^n K_{\frac{x}{h}+1,h}(X_i). \quad (7)$$

La deuxième classe est le noyau Gamma modifié qui est donné comme suit :

$$K_{GAM1(\rho_h(x),h)}(u) = \frac{u^{\rho_h(x)-1} \exp(\frac{-u}{h})}{h^{\rho_h(x)} \Gamma(\rho_h(x))}, \quad (8)$$

où

$$\rho_h(x) = \begin{cases} \frac{x}{h} & \text{si } x \geq 2h \\ \frac{1}{4}(\frac{x}{h})^2 + 1 & \text{si } 0 \leq x < 2h. \end{cases} \quad (9)$$

L'estimateur à noyau gamma modifié est donné par :

$$f_h^{GAM1}(x) = \frac{1}{n} \sum_{i=1}^n K_{\rho_h(x),h}(X_i). \quad (10)$$

Afin de faire une comparaison avec le noyau Gamma, nous utilisons le noyau gaussien-inverse-réciproque qui a été intrduit par Scaillet (2004).

4 Simulation

Afin de réaliser cette étude numérique, nous avons besoin d'un n-échantillon des montants des réclamations pour pouvoir estimer, par la méthode du noyau, la densité de probabilité f_h . Pour cela, nous prenons la loi de Cox2 et nous développons un algorithme de simulation qui contient les étapes suivantes :

- Génération d'un n-échantillon de fonction de répartition F du montant de réclamation supposé inconnue ;
- Estimation de la densité f par f_h en utilisant le noyau GAM, GAM1 et GIR ;
- Introduire le taux moyen d'arrivée des sinistres λ ;
- Détermination du montant moyen de réclamation $\mu \leftarrow \int_0^\infty x f_h(x) dx$;
- Verifier si $c > \lambda\mu$, sinon la ruine est certaine.
- Détermination du domaine de ϵ tel que $\epsilon_{min} < \epsilon < \epsilon_{max}$
 où : ϵ_{min} (respectivement ϵ_{max}) est la plus petite valeurs (respectivement la plus grande valeur) qui vérifie les deux conditions suivantes :

$$0 < \epsilon < \min\left\{\frac{1}{\mu}, \frac{c-\lambda\mu}{cu}\right\} \text{ et } u(a, a') < (1 - \rho(\epsilon))^2$$
- Détermination de l'erreur d'approximation $\Gamma = \frac{u(a, a')}{(1-\rho)((1-\rho)^2 - u(a, a'))}$

5 Conclusion

Dans ce travail, nous avons développé une approche de simulation dans une étude non paramétrique de la borne de stabilité forte des probabilités de ruines du modèle de risque classique. Pour cela, nous avons exploité trois types de noyaux, où nous avons déterminé l'erreur d'approximation avec chaque noyau.

Les résultats numériques obtenus sont significatifs par rapport à la qualité du noyau utilisé dans l'estimation de la loi des montants de réclamation.

Références

- Aissani, D. et N. V. Kartashov (1983). Ergodicity and stability of markov chains with respect to operator topology in the space of transition kernel. *Dokl.Akad.Nauk U.S.S.R, ser.A 11*, 3–5.
- Benouaret, Z. et D. Aissani (2010). Strong stability in a two-dimensional classical risk model with independent claims. *Scand. Actuar. Journal 2*, 83–93.
- Chen, S. (2000). Gamma kernel estimators for density functions. *Annals of the Institute of Statistical Mathematics 52*, 471–480.
- Enikeeva, F., V. Kalashnikov, et D. Rusaityte (2001). Continuity estimates for ruin probabilities. *Scand. Actuar. Journal 1*, 18–39.
- Kalashnikov, V. (2000). The stability concept for stochastic risk models. Working paper nr 166, Laboratory of Actuarial Mathematics, University of Copenhagen.
- Rusaityte, D. (2001). Stability bounds for ruin probabilities in a markov modulated risk model with investments. Working paper nr 178, Laboratory of Actuarial Mathematics, University of Copenhagen.
- Scaillet, O. (2004). Density estimation using inverse and reciprocal inverse gaussian kernels. *Journal of Nonparametric Statistics 16*, 217–226.

Summary

In this work, we are interested in the estimation of probability density f of the amount of claims, with modified gamma kernel estimator, in order to study the strong stability method in the classical risk model **P/G** (exponential distribution between two consecutive arrival and general distribution of claim amount). we use the simulation approach in order to evaluate numerically the approximation error between ruin probabilities of the **P/P** (exponential distribution between two consecutive claims and exponential distribution of claim amount) and **P/G** risk model.

Approche bayésienne dans l'estimation de la fonction de régression discrète par noyau binomial

Nabil Zougab^{*,**}, Smail Adjabi^{**}
Kokonendji Célestin C.^{***}

^{*}Université de Tizi-Ouzou
nabilzougab@yahoo.fr,

^{**}Laboratoire LAMOS, Université de Béjaia
adjabi@hotmail.com

^{***}Université de Franche-Comté, Laboratoire de Mathématiques de Besançon
celestin.kokonendji@univ-fcomte.fr

Résumé. L'objet de ce travail est de proposer les estimateurs bayésiens du paramètre de lissage h et de la variance σ^2 de l'erreur ϵ dans l'estimation de la fonction de régression discrète m . Une fois les lois a priori sur les paramètres h et σ^2 spécifiés, nous dérivons, à une constante près, leur lois a posteriori. Nous estimons ces paramètres par la moyenne a posteriori en utilisant les méthodes MCMC et l'échantillonnage de Gibbs. Des applications sur des données simulées et réelles pour illustrer l'approche bayésienne sont réalisées.

1 Introduction

Considérons un échantillon de n couples (x_i, y_i) , $i = 1, \dots, n$, à valeur dans $\mathbb{N} \times \mathbb{R}$. Le modèle de régression non paramétrique classique reliant les deux variables y et x est

$$y_i = m(x_i) + \epsilon_i \quad (1)$$

où $\mathbf{y} = (y_1, \dots, y_n)$ est un vecteur de réponse, $\mathbf{x} = (x_1, \dots, x_n)$ est un vecteur explicatif, $\epsilon = (\epsilon_1, \dots, \epsilon_n)$ est l'erreur de distribution gaussienne d'espérance nulle et de variance σ^2 finie ($\epsilon_i \sim \mathcal{N}(0, \sigma^2)$), et $m : \mathbb{N} \mapsto \mathbb{R}$ est la fonction discrète inconnue de régression à estimer. L'estimateur à noyau associé $\hat{m}_h$ de la fonction de régression m basé sur un échantillon i.i.d. $(X_1, Y_1), \dots, (X_n, Y_n)$ du couple (X, Y) est défini par (voir Kokonendji et al. (2009) Zougab (2013))

$$\hat{m}_h(x) = \frac{\frac{1}{n} \sum_{i=1}^n Y_i K_{x,h}(X_i)}{\frac{1}{n} \sum_{i=1}^n K_{x,h}(X_i)}, \quad x \in \mathbb{N}. \quad (2)$$

où h est le paramètre de lissage (fenêtre) et $K_{x,h}$ le noyau associé binomial donné par

$$K_{x,h} = B_{x,h}(y) = \frac{(x+1)!}{y!(x+1-y)!} \left(\frac{x+h}{x+1} \right)^y \left(\frac{1-h}{x+1} \right)^{x+1-y}, \quad y \in \mathbb{N}_x \subseteq \mathbb{N}. \quad (3)$$

où $\mathbb{N}_x = \{0, 1, \dots, x + 1\}$ et $h \in]0, 1]$. Notons qu'il existe d'autres noyaux associés discrets, voir par exemple, Kokonendji et al. (2009), Zougab (2013) et Zougab et al. (2014)). L'objectif de ce travail est double. Premièrement, nous proposons les estimateurs bayésiens du paramètre de lissage h et de la variance σ^2 de l'erreur ϵ dans l'estimation de la fonction de régression discrète m . Deuxièmement, nous comparons les performances de l'approche proposée avec la technique de validation croisée proposée par Kokonendji et al. (2009).

2 Le modèle

Nous rappelons que le modèle classique défini par (1) peut s'exprimer comme suit

$$y_i - m(x_i) \sim \mathcal{N}(0, \sigma^2),$$

Nous proposons d'estimer $m(x_i)$ en utilisant l'estimateur à noyau associé binomial de la fonction de régression et la technique de validation croisée. L'estimateur de $m(x_i)$ est donné par (Zhang et al. (2009))

$$\hat{m}_{-i}(x_i) = \frac{\frac{1}{n-1} \sum_{\substack{j=1 \\ i \neq j}}^n y_j K_{x_i, h}(x_j)}{\frac{1}{n-1} \sum_{\substack{j=1 \\ i \neq j}}^n K_{x_i, h}(x_j)}, \quad (4)$$

Nous désignons par (h, σ^2) les paramètres inconnus à estimer dans ce modèle. L'estimateur de la vraisemblance des données $\mathbf{y} = (y_i)_{i=1}^n$ sachant les paramètres h et σ^2 est donné par

$$\text{LCV}(y_1, \dots, y_n; (h, \sigma^2)) = \frac{1}{(2\pi\sigma^2)^{\frac{n}{2}}} \exp \left(-\frac{1}{2\sigma^2} \sum_{i=1}^n \{y_i - \hat{m}_{-i}(x_i)\}^2 \right). \quad (5)$$

L'estimateur de la vraisemblance donné par (5) va servir dans le calcul des lois a posteriori de h et σ^2 en partant d'une connaissance a priori sur ces paramètres.

3 Estimation bayésienne

3.1 Choix des lois a priori

Pour le paramètre de dispersion σ^2 , la loi a priori adoptée généralement dans le cadre bayésien est la loi inverse gamma, qu'on note par $\mathcal{IG}(a, b)$ dont la distribution est donnée par

$$\pi(\sigma^2) = \frac{b^a}{\Gamma(a)} \left(\frac{1}{\sigma^2} \right)^{a+1} \exp \left(-\frac{b}{\sigma^2} \right), \quad \sigma^2 \in [0, \infty[,$$

où $\Gamma(a)$ est la fonction Gamma. Pour le paramètre h , nous adoptons la loi bêta de paramètre α et β notée $\text{beta}(\alpha, \beta)$.

3.2 Calcul des lois a posteriori

Après la spécification des lois a priori, par le théorème de Bayes, la loi a posteriori de (h, σ^2) sachant les données $\mathbf{y}$ est proportionnelle à la loi a priori $\pi(h, \sigma^2) = \pi(h)\pi(\sigma^2)$ multipliée par la vraisemblance $\pi(y_1, \dots, y_n | h, \sigma^2)$. Plus explicitement la loi a posteriori des paramètres (h, σ^2) sachant les données est donnée à une constante d'intégration près par

$$\begin{aligned} \pi(h, \sigma^2 | y_1, \dots, y_n) &\propto \pi(h)\pi(\sigma^2) \frac{1}{(2\pi\sigma^2)^{\frac{n}{2}}} \exp\left(-\frac{1}{2\sigma^2} \sum_{i=1}^n \{y_i - \hat{m}_{-i}(x_i)\}^2\right) \\ &\propto \pi(h) \left(\frac{1}{\sigma^2}\right)^{\frac{n+2a}{2}+1} \exp\left(-\frac{1}{2\sigma^2} \left\{ \sum_{i=1}^n \{y_i - \hat{m}_{-i}(x_i)\}^2 + 2b \right\}\right) \end{aligned} \quad (6)$$

De (7), on retrouve une loi inverse gamma comme loi a posteriori du paramètre σ^2 sachant h et les données avec une mise à jour sur les paramètres de cette loi. La distribution de $(\sigma^2 | h, y_1, \dots, y_n)$ est donnée comme suit

$$\sigma^2 | h, y_1, \dots, y_n \sim \mathcal{IG}\left(\frac{n+2a}{2}, \frac{1}{2} \sum_{i=1}^n \{y_i - \hat{m}_{-i}(x_i)\}^2 + b\right). \quad (8)$$

Nous allons dériver la densité a posteriori de h sachant les données. En intégrant (7) par rapport à σ^2 , la loi de h sachant les données est alors

$$\pi(h | y_1, \dots, y_n) \propto \pi(h) \left(\frac{1}{2} \sum_{i=1}^n \{y_i - \hat{m}_{-i}(x_i)\}^2 + b\right)^{-\frac{n+2a}{2}}. \quad (9)$$

La loi a posteriori de $\pi(h | y_1, \dots, y_n)$ définie par (9) est trop complexe pour obtenir l'estimateur bayésien de h . Pour cette raison, nous proposons d'utiliser les méthodes MCMC pour construire une chaîne de Markov pour le paramètre h . Pour la variance σ^2 , on peut simuler directement à partir de (8) selon le principe de l'échantillonneur de Gibbs, car on connaît sa distribution conditionnelle.

4 Application sur données simulées et réelles

Dans cette partie, nous avons appliqué l'estimateur à noyau associé binomial et l'approche bayésienne pour estimer h et σ^2 sur des données simulées et sur deux jeux de données réels. Pour diagnostiquer la convergence de la méthode MCMC, nous avons utilisé le BMSE et le SIF (voir Zhang et al. (2009), Zougab (2013) et Zougab et al. (2014)). Nous avons comparé l'approche bayésienne avec la méthode classique de validation croisée en utilisant l'erreur quadratique RQ et le coefficient de détermination R^2 . Les résultats obtenus par Zougab (2013) montrent que l'approche bayésienne est performante par rapport à la technique de validation croisée, en particulier lorsque les échantillons sont de petite ou moyenne taille. De plus, la qualité graphique du lissage des données réelles est meilleure quand on utilise l'approche bayésienne (voir par exemple la figure 1).

Estimation bayésienne: fonction de régression discrète

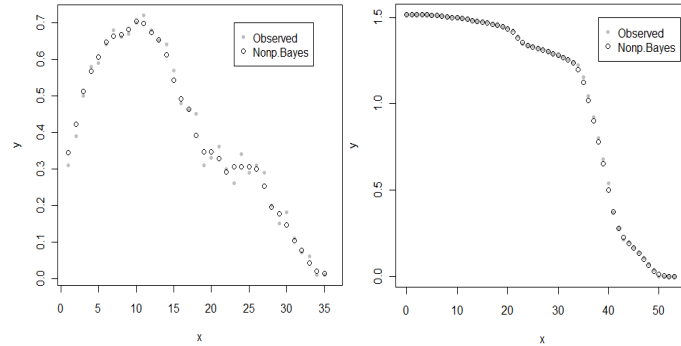


FIG. 1 – Régression nonparamétrique en utilisant le noyau binomial et l'approche bayésienne pour les données journalières de moyenne de graisse (kg/jour) (à gauche) et les données des arbres de hêtre (à droite).

Références

- Kokonendji, C. C., T. Senga Kiessé, et C. G. B. Demétrio (2009). Appropriate kernel regression on a count explanatory variable and applications. *Advances and Applications in Statistics* 12, 99–126.
- Zhang, X., R. D. Brooks, et M. L. King (2009). A bayesian approach to bandwidth selection for multivariate kernel regression with an application to state-price density estimation. *Journal of Econometrics* 153, 21–32.
- Zougab, N. (2013). *Approche bayésienne dans l'estimation non paramétrique de la densité de probabilité et la courbe de régression de la moyenne*. Thèse de doctorat, Université de Béjaia.
- Zougab, N., S. Adjabi, et C. C. Kokonendji (2014). Bayesian approach in nonparametric count regression with binomial kernel. *Communications in Statistics-Simulation and Computation* 43, 1052–1063.

Summary

In this work, we propose a Bayesian approach in the context of nonparametric count regression for estimating the bandwidth and the variance of the model error. The model error is considered as Gaussian with mean of zero and a variance of σ^2 . The Bayes estimates cannot be obtained in closed form and then, we use the well-known Markov Chain Monte Carlo (MCMC) technique to compute the Bayes estimates under the squared errors loss function. The performance of this proposed approach and the cross validation method are compared through simulation and real count data.

C. Posters

Comparaison de méthodes d'estimation pour un modèle exponentiel généralisé

Aidi khaoula*, Seddik-Ameur Nacira**

*Laboratoire de probabilités et statistiques

**Université Badji Mokhtar Annaba-Algérie

***khaoula.aidi@yahoo.fr

Résumé. Dans ce travail, on s'intéresse à un modèle exponentiel généralisé très flexible et pouvant décrire différentes durées de vie aussi bien en fiabilité qu'en analyse de survie. Comme les formes explicites des estimateurs ne peuvent être dégagées, nous nous proposons de faire une étude par simulations pour comparer les estimateurs du maximum de vraisemblance et ceux des moments en utilisant différentes méthodes numériques.

1 Introduction

Les distributions Weibull et gamma sont les distributions les plus utilisées pour modéliser les durées de vie et ceci est dû à leurs interprétations physiques et leur flexibilité. Néanmoins, la fonction de fiabilité de la distribution gamma ne peut être exprimée sous une forme explicite si le paramètre de forme n'est pas entier. Tandis que la convergence des estimateurs du maximum de vraisemblance pour la distribution Weibull est faible. Pour contourner ces inconvénients, une nouvelle distribution à deux paramètres ayant les mêmes caractéristiques que celles-ci, a été introduite par Gupta et Kundu (1999). Cette distribution appelée distribution exponentielle généralisée ou exponentielle exponentielle, a été étudiée par divers auteurs, Gupta et Kundu (1999, 2001, 2002); Raqab (2002), (2004), Raqab et Ahsanullah (2001), Zheng (2002). Ce modèle peut être considéré comme alternative au modèle gamma et au modèle Weibull dans de nombreuses situations. Dans ce travail, nous nous intéressons à une étude comparative des méthodes d'estimation pour ce modèle. Nous en déduisons l'estimateur de la fonction de fiabilité du modèle.

2 Présentation du modèle

La distribution exponentielle généralisée de paramètre de forme α et de paramètre d'échelle λ et notée $GE(\alpha, \lambda)$, a pour fonction densité

$$f(x; \alpha, \lambda) = \alpha \lambda (1 - e^{-\lambda x})^{\alpha-1} e^{-\lambda x} \quad \text{pour } x \geq 0$$

La fonction de distribution cumulative est :

$$F(x; \alpha, \lambda) = (1 - e^{-\lambda x})^\alpha \quad \text{pour } x \geq 0$$

et la fonction de survie, le taux de hasard sont donnés par :

$$S(x; \alpha, \lambda) = 1 - (1 - e^{-\lambda x})^\alpha \quad \text{pour } x \geq 0$$

$$H(x; \alpha, \lambda) = \frac{\alpha \lambda (1 - e^{-\lambda x})^{\alpha-1} e^{-\lambda x}}{1 - (1 - e^{-\lambda x})^\alpha} \quad \text{pour } x \geq 0$$

3 Estimation par la méthode du maximum de vraisemblance

Considérons $X_1, X_2, \dots, X_n$, un échantillon aléatoire de la loi exponentielle généralisée $EG(\alpha, \lambda)$,

La fonction de log-vraisemblance du modèle s'écrit :

$$L(\alpha, \lambda) = n \ln \alpha + n \ln \lambda + (\alpha - 1) \sum_{i=1}^n \ln(1 - e^{-\lambda x_i}) - \lambda \sum_{i=1}^n x_i$$

et dès lors que la log-vraisemblance est dérivable, l'estimateur du maximum de vraisemblance annule le vecteur des scores :

$$\begin{aligned} \frac{\partial L}{\partial \alpha} &= \frac{n}{\alpha} + \sum_{i=1}^n \ln(1 - e^{-\lambda x_i}) = 0 \\ \frac{\partial L}{\partial \lambda} &= \frac{n}{\lambda} + (\alpha - 1) \sum_{i=1}^n \frac{x_i e^{-\lambda x_i}}{1 - e^{-\lambda x_i}} - \sum_{i=1}^n x_i = 0 \end{aligned}$$

Nous utilisons des méthodes numériques pour résoudre ce système d'équations non linéaires.

4 L'estimation par la méthode du moments

Pour estimer les paramètres de la loi exponentielle généralisée par la méthode des moments, on a besoin des deux premiers moments car il y a deux paramètres à estimer.

L'espérance d'une variable de loi exponentielle généralisée, est donnée par :

$$E(X) = \int_0^\infty x \alpha \lambda (1 - e^{-\lambda x})^{\alpha-1} e^{-\lambda x} dx$$

$$E(X) = \frac{1}{\lambda} (\psi(\alpha + 1) - \psi(1)) = \frac{1}{n} \sum_{i=1}^n x_i \quad (6)$$

$$V(X) = -\frac{1}{\lambda^2} (\psi'(\alpha + 1) - \psi'(1)) = \frac{1}{n} \sum_{i=1}^n (x_i)^2$$

Pour les détails du calcul, on peut voir Gupta et Kundu (1999). Ici $\psi(\alpha) = \ln \Gamma(\alpha)$ représente la fonction log-gamma et $\psi'(\alpha) = \frac{d \ln \Gamma(\alpha)}{d\alpha}$ représente la dérivée de $\psi(\cdot)$.

Pour trouver les estimateurs des moments, on doit résoudre les deux égalités. Comme la forme explicite des estimateurs ne peut être obtenue, alors on utilise les méthodes numériques. Pour cela, nous adoptons la méthode de Newton Raphson et celle de l'algorithme EM et nous comparons les résultats.

4.0.1 la méthode de Newton Raphson

Cette méthode est définie par l'algorithme suivant :

$$\begin{cases} x_0 \text{ approximation initiale} \\ x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)} \end{cases}$$

et x_0 doit vérifier : $f(x_0)f''(x_0) > 0$

Dans le cas d'un modèle de durée de vie, on utilise comme fonction f la dérivée de la logvraisemblance par rapport au paramètre (le score), ce qui conduit à l'expression :

$$\lambda_{k+1} = \lambda_k - \left[\frac{\partial^2 L(\lambda_k)}{\partial \lambda \partial \alpha} \right]^{-1} \frac{\partial L(\lambda_k)}{\partial \lambda}$$

4.0.2 L'algorithme EM

L'algorithme EM (Expectation-Maximisation) est une approche générale, introduite par Dempster, Laird et Rubin (1977), que fait un calcul itératif pour trouver des estimateurs du maximum de vraisemblance lorsque les données sont incomplètes.

L'algorithme EM tire son nom du fait qu'à chaque itération, il opère deux étapes distinctes :

- la phase << Expectation >>, souvent désignée comme << l'étape E >>, procède comme son nom le laisse supposer à l'estimation des données inconnues, sachant les données observées et la valeur des paramètres déterminée à l'itération précédente .

- la phase << Maximisation >>, ou << étape M >>, procède donc à la maximisation de la vraisemblance, rendue désormais possible en utilisant l'estimation des données inconnues effectuée à l'étape précédente, et met à jour la valeur du ou des paramètres pour la prochaine itération.

En bref, l'algorithme EM procède selon un mécanisme extrêmement naturel : s'il existe un obstacle pour appliquer la méthode MV, on fait simplement sauter cet obstacle puis on applique effectivement cette méthode .

Le côté itératif de l'algorithme garantit que la vraisemblance augmente à chaque itération, ce qui conduit donc à des estimateurs de plus en plus corrects.

5 Simulations

Nous menons une étude par simulations numériques en utilisant le logiciel Maple, pour comparer les estimateurs des paramètres inconnus du modèles afin de pouvoir choisir les

plus précis. Pour cela, nous générons 10.000 échantillons de variables aléatoires de différentes tailles $n_1 = 50$, $n_2 = 150$, $n_3 = 300$, $n_4 = 500$, et avec différentes valeurs de paramètres, nous calculons ensuite les différents estimateurs et leurs erreurs quadratiques moyennes.

Si X est une variable aléatoire de loi exponentielle généralisée $EG(\alpha, \lambda)$, et U représente une variable aléatoire uniforme dans $[0, 1]$,

alors, on obtient les valeurs de X par la formule suivante :

$$X = -\frac{1}{\beta} \ln \left(1 - U^{\frac{1}{\alpha}} \right)$$

Bibliographie

- [1] Gupta, R. D. and Kundu, D. (1999a) "Generalized Exponential Distributions" , Australian and New Zealand Journal of Statistics, 41(2), 173 ± 188.
- [2] Gupta, R. D. and Kundu, D. (1999b) "Generalized Exponential Distributions : Statistical Inferences", Technical Report, The University of New Brunswick, Saint John.
- [3] Gupta, R. D. and Kundu, D. (2001a), Exponentiated exponential family ; an alternative to gamma and Weibull", Biometrical Journal, vol. 43, 117 - 130.
- [4] Gupta, R. D. and Kundu, D. (2001b), "Generalized exponential distributions : different methods of estimation", Journal of Statistical Computation and Simulation. vol. 69, 315 - 338.
- [5] Gupta, R. D. and Kundu, D. (2002), "Generalized exponential distributions : statistical inferences", Journal of Statistical Theory and Applications, vol. 1, 101 -118.

Summary

In this work, we are interested in a very flexible and generalized exponential model that can describe different lives both in reliability in survival analysis. As the explicit forms of the estimators can not be cleared, we propose to make a simulation study to compare the maximum likelihood estimators and those moments using different numerical methods.

Bayesian estimation of change point in generalized exponential distribution

Aissa Aknine*, Karima Nouali**

* a_aknine@ymail.com,

** noualikarima@yahoo.fr

Résumé. The Bayes estimates of change point and shape parameter of Generalized Exponential distribution are considered under symmetric and asymmetric loss functions.

1 Introduction

The model considered is the Generalized Exponential distribution $GE(\theta, \lambda)$ with distribution function

$$F(\theta, \lambda, x) = (1 - e^{-\lambda x})^\theta, \quad \theta > 0, \lambda > 0, x > 0.$$

and probability density of random observation x is given by

$$f(\theta, \lambda, x) = \theta \cdot \lambda (1 - e^{-\lambda x})^{\theta-1} \cdot e^{-\lambda x}, \quad \theta > 0, \lambda > 0, x > 0. \quad (1)$$

Here θ is the shape parameter and λ is the scale parameter.

Generalized Exponential $GE(\theta, \lambda)$ distribution has been introduced by the authors (Gupta and Kundu, 1999). It can be used quite effectively in analyzing many lifetime data, particularly in place of two-parameters Gamma and two parameters Weibull distribution (Gupta and Kundu, 2001). The two parameters $GE(\theta, \lambda)$ can have an increasing and decreasing failure rates depending on shape parameter.

2 Change point model

There are many studies on shifts models in a sequence of random variables in Bayesian Framework (Hinkley (1970), Broemeling Tsurum (1987), ...). In model(1), it may happen that at some unknown time m a break in the sequence of observations $x_1, x_2, \dots, x_n$ which divide on two sequences. So, we assume that the m observations $x_1, x_2, \dots, x_m$ come from $GE(\theta_1, \lambda)$ and the $(n-m)$ observations $x_{m+1}, x_{m+2}, \dots, x_n$, also follow from $GE(\theta_2, \lambda)$

Let $S_k = \sum_{i=1}^k \ln(1 - \exp(-\lambda x_i))^{-1}, \forall i = 1, \dots, n$.

$$T_1 = \sum_{m=1}^{n-1} T_1(m) \quad \text{and} \quad T_2 = \sum_{m=1}^{n-1} T_2(m)$$

Bayesian estimation of change point in generalized exponential distribution

Where $T_1(m) = \Gamma(m + a_1) \cdot \Gamma(n - m + a_2) \cdot (b_1 - S_m)^{-(m+a_1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2)}$
and $T_2(m) = \Gamma(m) \cdot \Gamma(n - m) \cdot S_m^{-m} \cdot (S_n - S_m)^{-(n-m)}$

3 Likelihood, Prior, Posterior and Marginal distribution

The likelihood function of the given sample information $\bar{x} = (x_1, \dots, x_m, x_{m+1}, \dots, x_n)$ is

$$L(\theta_1, \theta_2, m | \underline{x}) = \lambda^n \theta_1^m \cdot \theta_2^{n-m} \cdot e^{-\lambda \sum_{i=1}^n x_i} \cdot e^{S_n} \cdot e^{-\theta_1 S_m} \cdot e^{-\theta_2 (S_n - S_m)}.$$

The joint posterior density of θ_1 , θ_2 and m is

$$g(\theta_1; \theta_2, m | \underline{x}) = \frac{g(\theta_1, \theta_2) \cdot L(\theta_1, \theta_2, m | \underline{x})}{\sum_{m=1}^{n-1} \int_0^{+\infty} \int_0^{+\infty} g(\theta_1, \theta_2) \cdot L(\theta_1, \theta_2, m | \underline{x}) d\theta_1 d\theta_2}.$$

Where $g(\theta_1, \theta_2)$ is the joint prior density of θ_1 and θ_2 .

3.1 Case of conjugates priors

The distribution of θ_1 and θ_2 are taken as natural conjugate Gamma prior as

$$g_1(\theta_1) \propto \theta_1^{a_1-1} \cdot \exp(-b_1 \cdot \theta_1) \quad \text{and} \quad g_1(\theta_2) \propto \theta_2^{a_2-1} \cdot \exp(-b_2 \cdot \theta_2)$$

As in Broemeling and Tsurumi(1987), we suppose the marginal prior distribution of shift point m to be discrete uniforme over the set $\{1, 2, \dots, n\}$

$$g_1(m) = \frac{1}{n-1}$$

The joint prior density of θ_1 , θ_2 and m is

$$g_1(\theta_1, \theta_2, m) = \frac{1}{n-1} \theta_1^{a_1-1} \cdot \theta_2^{a_2-1} \cdot e^{-(b_1 \cdot \theta_1 + b_2 \cdot \theta_2)}.$$

Where the joint posterior density of θ_1 , θ_2 and m is

$$g_1(\theta_1; \theta_2, m | \underline{x}) = \frac{\theta_1^{m+a_1-1} \cdot \theta_2^{n-m+a_2-1} \cdot e^{-\theta_1 \cdot (b_1 + S_m)} \cdot e^{-\theta_2 \cdot [b_2 + (S_n - S_m)]}}{\sum_{m=1}^{n-1} K_1(m) \cdot [b_1 + S_m]^{-(m+a_1)} \cdot [b_2 + (S_n - S_m)]^{-(n-m+a_2)}}.$$

Where $K_1(m) = \Gamma(m + a_1) \cdot \Gamma(n - m + a_2)$

The marginal posterior density of θ_1 , of θ_2 and of m are respectively

$$g_1(\theta_1 | \underline{x}) = \frac{1}{T_1} \cdot \sum_{m=1}^{n-1} \Gamma(n - m + a_2) \cdot [b_2 + (S_n - S_m)]^{-(n-m+a_2)} \cdot \theta_1^{m+a_1-1} \cdot e^{-\theta_1 \cdot (b_1 + S_m)}.$$

$$g_1(\theta_2 | \underline{x}) = \frac{1}{T_1} \cdot \sum_{m=1}^{n-1} \Gamma(m + a_1) \cdot [b_1 - S_m]^{-(m+a_1)} \cdot \theta_2^{n-m+a_2-1} \cdot e^{-\theta_2 [b_2 - (S_n - S_m)]}.$$

$$g_1(m/\underline{x}) = \frac{T_1(m)}{\sum_{m=1}^{n-1} T_1(m)}$$

Where $T_1(m) = K_1(m) \cdot (b_1 - S_m)^{-(m+a_1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2)}$

and $T_1 = \sum_{m=1}^{n-1} T_1(m)$

3.2 Case of noninformative priors

Let the joint noninformative prior density of θ_1 and θ_2 be given by

$$g_2(\theta_1, \theta_2, m) = \frac{1}{(n-1)\theta_1\theta_2}$$

the joint posterior density of θ_1 , θ_2 and m be given by

$$g_2(\theta_1, \theta_2, m | \underline{x}) = \frac{\theta_1^{m-1} \cdot \theta_2^{n-m-1} \cdot e^{-\theta_1 \cdot S_m} \cdot e^{-\theta_2 (S_n - S_m)}}{\sum_{m=1}^{n-1} K_2(m) \cdot S_m^{-m} \cdot (S_n - S_m)^{-(n-m)}}$$

Where $K_2(m) = \Gamma(m) \cdot \Gamma(n-m)$

The marginal posterior density of θ_1 , of θ_2 and of m are respectively

$$g_2(\theta_1/\underline{x}) = \frac{1}{T_2} \cdot \sum_{m=1}^{n-1} \Gamma(n-m) \cdot (S_n - S_m)^{-(n-m)} \cdot \theta_1^{m-1} \cdot e^{-\theta_1 \cdot S_m}.$$

$$g_2(\theta_2/\underline{x}) = \frac{1}{T_2} \cdot \sum_{m=1}^{n-1} \Gamma(m) \cdot (S_m)^{-m} \cdot \theta_2^{n-m-1} \cdot e^{-\theta_2 \cdot (S_n - S_m)}$$

and

$$g_2(m/\underline{x}) = \frac{T_2(m)}{\sum_{m=1}^{n-1} T_2(m)}$$

Where $T_2(m) = K_2(m) \cdot S_m^{-m} \cdot (S_n - S_m)^{-(n-m)}$

and $T_2 = \sum_{m=1}^{n-1} T_2(m)$

We denote by m^* , θ_1^* and θ_2^* (respectively, m^{**} , θ_1^{**} and θ_2^{**}), the Bayes estimators of m , θ_1 and θ_2 in case of informative prior (respectively in case of noninformative prior).

4 Bayes estimation under symmetric and asymmetric loss functions

4.1 Bayes estimation under symmetric loss functions

The Bayes estimate of generic parameter α based on a Squar Error Loss (SEL) Function

$$L_1(\alpha, d) = (\alpha - d)^2$$

Bayesian estimation of change point in generalized exponential distribution

Where d is a decision rule to estimate α is posterior mean.

The Bayes estimator of m under SEL is

$$m_S^* = \frac{\sum_{m=1}^{n-1} m \cdot T_1(m)}{\sum_{m=1}^{n-1} T_1(m)} \quad \text{and} \quad m_S^{**} = \frac{\sum_{m=1}^{n-1} m \cdot T_2(m)}{\sum_{m=1}^{n-1} T_2(m)}$$

The Bayes estimator of θ_1 under SEL is

$$\theta_{1S}^* = \left[\frac{\sum_{m=1}^{n-1} I_1(m) [b_1 - S_m]^{-(m+a_1+1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2)}}{\sum_{m=1}^{n-1} T_1(m)} \right]$$

$$\theta_{1S}^{**} = \left[\frac{\sum_{m=1}^{n-1} I_2(m) S_m^{-(m+1)} \cdot (S_n - S_m)^{-(n-m)}}{\sum_{m=1}^{n-1} T_2(m)} \right]$$

Where $I_1(m) = \Gamma(m + a_1 + 1) \cdot \Gamma(n - m + a_2)$ and $I_2(m) = \Gamma(n - m) \cdot \Gamma(m + 1)$

The Bayes estimator of θ_2 under SEL is

$$\theta_{2S}^* = \left[\frac{\sum_{m=1}^{n-1} I_3(m) [b_1 - S_m]^{-(m+a_1+1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2+1)}}{\sum_{m=1}^{n-1} T_1(m)} \right]$$

$$\theta_{2S}^{**} = \left[\frac{\sum_{m=1}^{n-1} I_4(m) S_m^{-(m)} \cdot (S_n - S_m)^{-(n-m+1)}}{\sum_{m=1}^{n-1} T_2(m)} \right]$$

Where $I_3(m) = \Gamma(m + a_1) \Gamma(n - m + a_2 + 1)$ and $I_4(m) = \Gamma(m) \cdot \Gamma(n - m + 1)$

4.2 Bayes estimation under asymmetric loss functions

- Under the assumption that the minimal loss occurs at d , the Linex loss function can be expressed as

$$L_4(\alpha, d) = \exp[q_1(d - \alpha)] - q_1(d - \alpha) - I, \quad q_1 \neq 0$$

The Bayes estimate a_L^* is the value of d that minimizes $E_\alpha\{L_4(\alpha, d)\}$

$$a_L^* = -\frac{1}{q_1} \cdot \ln[E_\alpha\{\exp(-q_1\alpha)\}]$$

If $\alpha = m$ then the Bayes estimators of m are given respectively by

$$m_{1L}^* = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} e^{-q_1 m} T_1(m)}{\sum_{m=1}^{n-1} T_1(m)} \right] \quad \text{and} \quad m_{2L}^* = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} e^{-q_1 m} T_2(m)}{\sum_{m=1}^{n-1} T_2(m)} \right]$$

If $\alpha = \theta_1$ then the Bayes estimators of θ_1 are given respectively by

$$\theta_{1L}^* = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} I_5(m) [b_1 + q_1 - S_m]^{-(m+a_1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2)}}{\sum_{m=1}^{n-1} T_1(m)} \right].$$

$$\theta^{**}_{1L} = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} I_6(m) [q_1 + S_m]^{-m} [S_n - S_m]^{-(n-m)}}{\sum_{m=1}^{n-1} T_2(m)} \right].$$

If $\alpha = \theta_2$ then the Bayes estimators of θ_2 are given respectively by

$$\theta_{2L}^* = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} I_5(m) [b_1 - S_m]^{-(m+a_1)} \cdot [b_2 + q_1 - (S_n - S_m)]^{-(n-m+a_2)}}{\sum_{m=1}^{n-1} T_1(m)} \right].$$

$$\theta^{**}_{2L} = -\frac{1}{q_1} \ln \left[\frac{\sum_{m=1}^{n-1} I_6(m) S_m^{-m} [(S_n - S_m) + q_1]^{-(n-m)}}{\sum_{m=1}^{n-1} T_2(m)} \right].$$

Where $I_5(m) = \Gamma(m + a_2) \Gamma(n - m + a_2)$ and $I_6(m) = \Gamma(m) \cdot \Gamma(n - m)$

— Another loss function, called General Entropy(GE) loss function given by

$$L_5(\alpha, d) = \left(\frac{d}{\alpha} \right)^{q_2} - q_2 \ln \left(\frac{d}{\alpha} \right) - I.$$

The Bayes estimate α_E^* is the value of d that minimizes $E_\alpha \{L_5(\alpha, d)\}$

$$\alpha_E^* = [E_\alpha(\alpha^{-q_3})]^{-1/q_2}$$

If $\alpha = m$ then the Bayes estimators of m are given respectively by

$$m_E^* = \left[\frac{\sum_{m=1}^{n-1} m^{-q_2} T_1(m)}{\sum_{m=1}^{n-1} T_1(m)} \right]^{-1/q_2} \quad \text{and} \quad m_E^{**} = \left[\frac{\sum_{m=1}^{n-1} m^{-q_2} T_2(m)}{\sum_{m=1}^{n-1} T_2(m)} \right]^{-1/q_2}$$

If $\alpha = \theta_1$ then the Bayes estimators θ_1 are given respectively by

$$\theta_{1E}^* = \left[\frac{\sum_{m=1}^{n-1} I_7(m) [b_1 - S_m]^{-(m+a_1-q_2)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2)}}{\sum_{m=1}^{n-1} T_1(m)} \right]^{-1/q_2}$$

$$\theta^{**}_{1E} = \left[\frac{\sum_{m=1}^{n-1} I_8(m) S_m^{-(m-q_2+1)} \cdot (S_n - S_m)^{-(n-m)}}{\sum_{m=1}^{n-1} T_2(m)} \right]^{-1/q_2}$$

$I_7(m) = \Gamma(m + a_1 - q_2) \cdot \Gamma(n - m + a_2)$ and $I_8(m) = \Gamma(n - m) \cdot \Gamma(m - q_2 + 1)$

If $\alpha = \theta_2$ then the Bayes estimators of θ_2 are given respectively by

$$\theta_{2E}^* = \left[\frac{\sum_{m=1}^{n-1} I_9(m) [b_1 - S_m]^{-(m+a_1)} \cdot [b_2 - (S_n - S_m)]^{-(n-m+a_2-q_2)}}{\sum_{m=1}^{n-1} T_1(m)} \right]^{-1/q_2}$$

$$\theta^{**}_{2E} = \left[\frac{\sum_{m=1}^{n-1} I_{10}(m) S_m^{-(m)} \cdot (S_n - S_m)^{-(n-m-q_2)}}{\sum_{m=1}^{n-1} T_2(m)} \right]^{-1/q_2}$$

Where $I_9(m) = \Gamma(m + a_1) \Gamma(n - m + a_2 - q_2)$ and $I_{10}(m) = \Gamma(m) \cdot \Gamma(n - m - q_2)$

With simulation study, we can compare the performance of the differents Bayes estimates of m , θ_1 and θ_2 for moderate and small samples .

Références

- Broemeling, L.D, and H. Tsurumi (1987). *Econometrics and structural Change*. Marcel Dekker, New York, NY, USA.
- Gupta, R.D, and D. Kundu (1999). *Generalized Exponential Distribution*. Australian and New Zealand Journal of Statistics, 41(2), 173-188.
- Gupta, R.D, and D. Kundu (2001). *Generalized exponential distribution, an alternative to Gamma and Weibull distribution*. Biometrical journal, 43(1), 117-130.
- Hinkley, D.V (1970). *Inference about a change point in a sequence of random variables*. Biometrika, 57, 1-17.

Summary

Dans ce travail, nous abordons l'estimation Bayésienne du point de rupture et du paramètre de position de la distribution exponentielle généralisée dans le cas des fonctions de pertesymétriques et asymétriques.

Inégalités stochastiques pour le système $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle

Lala Maghnia Alem, Mohamed Boualem, Djamil Aissani

Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
alem_nanou@yahoo.fr, robertt15dz@yahoo.fr, lamos_bejaia@hotmail.com
<http://www.lamos.org>

Résumé. Le but de ce travail est d'appliquer les méthodes de comparaison stochastique, pour étudier les propriétés de monotonie du modèle $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle. pour cela nous avons montré la monotonie de l'opérateur de transition de la chaîne de Markov incluse par rapport à l'ordre stochastique et convexe. Nous avons ensuite montré que la distribution stationnaire du nombre de clients dans un système $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle est majorée (respectivement minorée) par la distribution stationnaire du nombre de clients dans un système $M_1, M_1/M_1, M_2/1$ avec rappels à communication bidirectionnelle, si la distribution des temps de service est *NBUE* (respectivement *NWUE*). Enfin, nous avons confirmé les résultats théoriques obtenus par une application numérique.

1 Introduction

Les files d'attente avec rappels ont été largement utilisés pour modéliser de nombreux problèmes dans les systèmes téléphoniques, informatiques, des réseaux locaux et des situations de la vie quotidienne. Dans la plupart des publications sur les files d'attente avec rappels, le serveur ne fournit que le service aux arrivées entrantes effectuées par les clients réguliers. Cependant, il existe des situations réelles par exemple : les centres d'appels où un opérateur non seulement sert les appels entrants, mais il effectue aussi des appels sortants vers l'extérieur lorsque le serveur est libre. Cette fonction est connue sous le nom de files d'attente avec rappels à communication bidirectionnelle.

2 Le modèle mathématique

Nous considérons un système de file d'attente à un seul serveur auquel les clients primaires entrants arrivent selon un processus de Poisson de taux λ . En outre, si le serveur est libre, alors il génère un appel sortant dans un temps exponentiellement distribué avec un taux α . Un appel entrant qui trouve le serveur occupé rejoint l'orbite et il retente d'entrer dans le serveur après une distribution du temps exponentielle avec un Taux μ , si $N(t) = j$, alors le taux de rappel

est $j\mu$. L'arrivée des flux d'appels entrants et sortants, temps de service et les inter-rappels sont supposés être mutuellement indépendants.

L'état du système au temps t peut être décrit par le processus $Y(t) = (C(t), N(t))$ où :

$$C(t) = \begin{cases} 0, & \text{si le serveur est inactif à l'instant } t; \\ 1, & \text{si le serveur est occupé par un appel entrant à l'instant } t; \\ 2, & \text{si le serveur lance un appel vers l'extérieur à l'instant } t. \end{cases}$$

et $N(t)$ représente le nombre de clients en orbite à l'instant t .
la chaîne de Markov induite associée à ce modèle est :

$$Z_n = Z_{n-1} - W_n + V_n \quad (1)$$

où V_n est le nombre d'arrivées entrantes pendant service du $n^{\text{ème}}$ client, sa distribution est donnée par :

$$k_j^l = \int_0^\infty e^{-\lambda x} \frac{(\lambda x)^j}{j!} dB_l(x), \quad l = 1, 2, \quad j \in \mathbf{Z}_+, \quad (2)$$

et $W_n = \begin{cases} 1, & \text{si le } n^{\text{ème}} \text{ client en service provient de l'orbite,} \\ 0, & \text{sinon.} \end{cases}$

3 Inégalités stochastiques pour le système $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle

Soient Σ_1 et Σ_2 deux modèles d'attente $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle de paramètres, pour $i = 1, 2$:

$\lambda^{(i)}$: taux d'arrivées entrantes dans Σ_i .

$\mu^{(i)}$: taux de rappels dans Σ_i .

$\alpha^{(i)}$: taux d'appels sortants dans Σ_i

$B_1^{(i)}(x)$: distribution du temps de service d'un appel entrant dans Σ_i .

$B_2^{(i)}(x)$: distribution du temps de service d'un appel sortant dans Σ_i .

$k_n^{(i)}$: le nombre de nouvelles arrivées durant le service du $n^{\text{ème}}$ client dans Σ_i .

$\pi_n^{(i)}$: la distribution stationnaire du nombre de clients dans le système Σ_i .

3.1 Inégalités préliminaires

Soient X et Y deux variables aléatoires non négatives de fonctions de répartition F et G , respectivement. On dit que X est inférieure à Y par rapport à :

– Ordre stochastique (noté $X \leq_{st} Y$) ssi : $F(x) \geq G(x), \quad \forall x \geq 0$.

– Ordre convexe (noté $X \leq_v Y$) ssi : $\int_x^{+\infty} \bar{F}(u) du \leq \int_x^{+\infty} \bar{G}(u) du, \quad \forall x \geq 0$.

Soit X et X_τ des variables aléatoires représentant respectivement la durée de vie et la durée de vie résiduelle d'un élément. Soient F et F_τ leurs distributions respectives. On dit que F est :

– *NBUE* (New Better than Used in Expectation), si : $E(X_\tau) \leq E(X), \quad (0 < \tau < \infty)$.

– *NWUE* (New Worse than Used in Expectation), si : $E(X) \leq E(X_\tau), \quad (0 < \tau < \infty)$.

Lemme 1 – si $\lambda^{(1)} \leq \lambda^{(2)}$ et $B_l^{(1)} \leq_{st} B_l^{(2)}$ alors $\{k_n^{(1)}\} \leq_{st} \{k_n^{(2)}\}$,
– si $\lambda^{(1)} \leq \lambda^{(2)}$ et $B_l^{(1)} \leq_v B_l^{(2)}$, $l = 1, 2$ alors $\{k_n^{(1)}\} \leq_v \{k_n^{(2)}\}$.

3.2 Monotonie de la chaîne de Markov incluse

Les probabilités de transition en un pas de la chaîne de Markov incluse pour le système sous étude sont données par la formule suivante :

$$p_{n,m} = \begin{cases} \frac{n\mu}{\lambda+\alpha+n\mu} k_0^1, & \text{si } m = n - 1 \text{ et } n \geq 1, \\ \frac{\lambda}{\lambda+\alpha+n\mu} k_{m-n}^1 + \frac{\alpha}{\lambda+\alpha+n\mu} k_{m-n}^2 + \frac{n\mu}{\lambda+\alpha+n\mu} k_{m-n+1}^1, & \text{si } 0 \leq n \leq m. \end{cases}$$

Soit l'opérateur de transition τ de la chaîne de Markov incluse. Pour chaque distribution $p = (p_n)_{n \geq 0}$, on associe une distribution $\tau_p = q = (q_m)_{m \geq 0}$ telle que

$$q_m = \sum_{n \geq 0} p_n p_{nm}.$$

Notons par $\tau^{(1)}, \tau^{(2)}$ les opérateurs de transition associés aux chaînes de Markov incluses de chaque système.

Theorem 1 si $\lambda^{(1)} \leq \lambda^{(2)}$, $\mu^{(1)} \geq \mu^{(2)}$, $\alpha^{(1)} \leq \alpha^{(2)}$, $B_1^{(1)} \leq_s B_1^{(2)}$ et $B_2^{(1)} \leq_s B_2^{(2)}$ alors $\tau^{(1)} \leq_{st} \tau^{(2)}$,
c'est-à-dire que pour une distribution quelconque p on a $\tau^{(1)}p \leq_s \tau^{(2)}p$, ($s = st$ ou v).

3.3 Inégalités stochastiques des distributions stationnaires du nombre de clients dans le système

Theorem 2 Si pour le modèle $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle, la distribution de temps de service est *NBUE* (New Better than Used in Expectation) (respectivement *NWUE* - New Worse than Used in Expectation), et si de plus $B_2^{(1)} \leq_v B_2^{(2)}$, $B_1^* = B_1^{(2)} \leq_v B_2^{(2)}$, alors la distribution stationnaire du nombre de clients dans ce système est inférieure (respectivement supérieure), par rapport à l'ordre convexe, à la distribution stationnaire du nombre de clients dans le système $M_1, M_2/M_1, M_2/1$ avec rappels à communication bidirectionnelle.

4 Application numérique

Après avoir élaboré un simulateur, sous environnement Matlab, décrivant le comportement du modèle $M_1, M_2/G_1, G_2/1$ avec rappels à communication bidirectionnelle, nous avons estimé les probabilités stationnaires d'un tel système lorsque la distribution des temps de service est *NBUE* (resp. *NWUE*). Ensuite, les comparer à celles du système $M_1, M_2/M_1, M_2/1$ avec rappels à communication bidirectionnelle afin de valider le résultat obtenu dans le Théorème 2. Pour ce faire, on a choisi deux lois de probabilités de type *NBUE* (une distribution de Weibull ($Wbl(a, b)$, avec $a > 1$) et deux autres lois de type *NWUE* (une distribution de Weibull ($Wbl(a, b)$, avec $a \leq 1$) et une distribution Gamma ($\Gamma(a, b)$, avec $0 \leq a < 1$) pour

les temps de service. Ainsi que nous avons fixé le taux d'arrivées entrantes $\lambda = 0.3$, le taux d'arrivées sortantes $\alpha = 0.2$, le taux de rappels $\mu = 1$, le temps de simulation $T_{max} = 1000$ unités de temps et $n = 100$ (le nombre de replications).

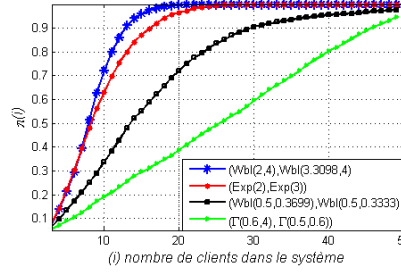


FIG. 1 – Comparaison des probabilités stationnaires.

Références

- Artalejo, J. et T. Phung-Duc. (2013). Single server retrial queues with two way communication. *Applied Mathematical Modelling* 37, 1811–1822.
- Artalejo, J. R. et A. Gómez-Corral (2008). *Retrial queueing system: A computational approach*. Berlin: Springer Edition.
- Boualem, M., N. Djellab, et D. Aïssani (2012). Stochastic approximations and monotonicity of a single server feedback retrial queue. *Mathematical Problems in Engineering* 2012, 1–13.
- M. Boualem, N. D. et D. Aïssani. (2009). Stochastic inequalities for $m/g/1$ retrial queues with vacations and constant retrial policy. *Mathematical and Computer Modelling* 50: (1-2), 207–212.
- Shaked, M. et J. Shanthikumar. (2007). *Stochastic Orders*. New York: Springer-Verlag.
- Stoyan., D. (1983). *Comparison methods for queues and other stochastic models*. New York: Wiley.

Summary

In this work we use the general theory of stochastic orderings to investigate the monotonicity properties of the system $M_1, M_2/G_1, G_2/1$ retrial queue with two way communication. We show the monotonicity of the transition operator of the embedded Markov chain relative to the strong stochastic ordering, convex ordering, and we obtain comparability conditions of the number of customers in the system. Inequalities are also obtained for the stationary distribution of the embedded Markov chain. An illustrative numerical example is presented.

Sélection de variables pour réseau bayésien

Raouf Benmakrelouf*, Wafa Karouche*, Joseph Rynkiewicz**, Mohamed Djedour*

*Faculty of mathematics PoBox 32, Al alia Bab ezzouar, USTHB-Algiers, Algeria
ismtid_raouf@hotmail.fr
karouche.w@hotmail.fr
mdjedour@yahoo.fr

**Université Paris 1, Panthéon-Sorbonne, Paris, France
joseph.rynkiewicz@univ-paris1.fr

Résumé. Notre travail permet de sélectionner des variables pour construire un réseau bayésien qui explique l'intention de vote pour un parti politique français, on classe ces variables en deux groupes de variables (v.contexte et v.opinion) explicatives. A travers ces deux groupes de variables on construit tous les modèles logistiques possibles associer au variables de contexte puis au variables d'opinion, a partir de ces deux groupes de modèles on choisit le meilleur modèle avec variables de contexte et le meilleur modèle avec variable d'opinion, au sens du critère d'information bayésien (BIC). Les variables de contexte et d'opinion obtenus a partir des modèles sélectionner ce sont les variables qui expliquent au mieux l'intention de vote. Le réseau bayésien suivra un sens imposé de causalité, variables de contexte vers variables d'opinion. La dernière étape de la construction du réseau est de tester les connexions en testons les interactions entre les variables.

1 Introduction

Le vote est un événement sociale important et ces acteurs cherches toujours a le comprendre et essayer d'anticiper les résultats du scrutin, ce qui a poussé les mathématiciens a en faire un centre d'intérêt.

Notre document présente l'étude d'une société a l'approche des élections, on veut comprendre ce qui pousse les électeurs a voté pour un certain candidat ou non, en d'autres termes on veut expliquer l'intention de vote pour notre candidat a travers la vie sociale des électeurs et leurs opinion sur les questions sociale et politique majeurs. A la fin on aimera bien construire un bon réseau bayésien pour bien comprendre ce phénomène sociale.

2 Modèle causale

Un modèle en terme générale est une représentation idéal de la réalité. Il met en évidence certains aspects de la réalité et ignore d'autres. Modèle c'est un objet mathématique qui approche un certain processus, a chaque fois qu'on donne une valeur représentative au modèle on a une sortie approché de la vraie sortie.

exemples de modèles : Fonction de répartition, Table statistique, modèle causale, ect.....

2.1 Définition : Modèle Causale

soit $M = (X, Y, F)$ un modèle causale

Avec :

- (i) "X", c'est l'ensemble de départ (variables de départ ou exogènes)
- (ii) "Y", c'est l'ensemble $\{Y_1, Y_2, \dots, Y_n\}$, de variables appelé variables endogènes, $Y_i \in X \cup Y$
se sont des variables produites à l'intérieur du modèle.
- (iii) "F", c'est l'ensemble de fonctions $\{f_1, f_2, \dots, f_n\}$ tellesque :

$$\begin{array}{ccc} F : X \mapsto Y & f_i : X \cup (Y \setminus Y_i) \mapsto Y_i \\ X \mapsto Y(x) & x_i \mapsto y_i \end{array},$$

et $Y_i = f_i(pa_i, x)$, pa_i c'est une réalisation de Y , $pa_i \in PA$.

Tout modèle causale est associer a un graphe directionnel $G(M)$.

exemple : modèle linéaire , modèle logistique, modèle loglinéaire, ect.....

3 Modèle Logistique et Sélection de Modèle

3.1 Modèle logistique

La regression logistique est une technique qui vise a construire un modèle permettant d'expliquer ou prédire les valeurs prises par une variable cible qualitative à partir d'un ensemble de variables explicatives quantitatives ou qualitatives (d'ou la nécessité d'un codage pour cette dernière).

Nous nous somme intéressé a la probabilité de vote de la variable qui représente un parti politique qui a été recodé en variable binaire : $Y=1$ pour un vote tout a fait probable ou pluto probable et $Y=0$ sinon. Lorsque la variable qualitative est dichotomique comme pour cette dernière (2 modalités, $\{0 \text{ ou } 1\}$) on parle alors de régression logistique, si elle possède plus de deux modalités, on parle de régression logistique polytomique.

Comme il y a beaucoup de modalités pour chaque variable, nous avons considéré uniquement les modèles additifs pour la régression logistique. Pour expliquer l'intention de vote du parti politique, nous avons dans un premier temps pris les variables de contexte comme explicatives et construit toutes les combinaisons de modèles possible $\{C_c^0, C_c^1, C_c^2, \dots, C_c^{c-1}\}$ avec c : nombre de variables contexte. La même procédure est appliqué cette fois si avec les variables d'opinion $\{C_o^0, C_o^1, \dots, C_o^{o-1}\}$ avec o : nombre de variables opinion.

Une fois l'énumération (dénombrement) de tout les modèles candidats établis, le problème posé fut le choix du meilleur modèle dans le cas de variables de contexte et de meme pour les variables d'opinion.

3.2 Critère d'information bayésien

Nous nous sommes confronté à de nombreux modèles possibles, correspondant aux différentes combinaisons de variables explicatives. Nous avons choisi comme critère de sélection des modèles le *BIC* qui dans le cas de grands échantillons aboutit à des modèles parcimonieux. On rappelle que pour n observations :

$$BIC = -2\log L_n + nb(param) * \log n$$

Le modèle sélectionné par ce critère est :

$$M_{BIC} = \underset{M_i}{\operatorname{argmin}} BIC_i$$

M_{BIC} représente le ième modèle dont le BIC_i est minimale.

Le modèle optimal sélectionné pour expliquer les intentions de vote pour l'extrême droite à l'aide des variables de contextes et d'opinion :

$$\operatorname{Logit}(\text{intention de vote}) = \alpha + \beta_1 \text{Diplome} + \beta_2 \text{Origine des Parents} + \beta_3 \text{Tranche d'age} + \beta_4 \text{Situation professionnelle}$$

$$\operatorname{Logit}(\text{intention de vote}) = \delta + \theta_1 \text{Opinion sur les Immigres} + \theta_2 \text{Chomage} + \theta_3 \text{Delinquance} + \theta_4 \text{Peine de mort}$$

Toutes les variables de contexte et d'opinion incluses dans les modèles sélectionnés par le BIC expliquent les intentions de vote pour le parti politique .

4 Réseaux Bayésiens

Les réseaux bayésiens constituent un langage graphique et une méthodologie simples et corrects pour exprimer et expliquer pratiquement ce de quoi on est certain ou incertain. Ils reposent sur la formule de Bayes reliant des probabilités conditionnelles avec des probabilités jointes. Nous rappelons le théorème de Bayes :

$$P(A|B) = \frac{P(A,B)}{P(B)} \quad \text{ou alors} \quad P(A|B) = \frac{P(B|A) \cdot P(A)}{P(B)}$$

L'élaboration du réseau bayésien nécessite la construction d'un modèle pour expliquer les connexions et les tests entre chacune des variables d'opinions (Opinion sur les Immigres, Mondialisation, Peine de mort) qui ont été sélectionnés à partir des modèles précédents (modèle optimal choisi par le BIC).

Notre objectif est de construire un réseau bayésien cohérent avec les données. En tenant compte du sens possible (imposé) des causalités :

$$\text{Variables de contexte} \mapsto \text{Variables d'opinion} .$$

Ce qui permet de restreindre le nombre de réseaux possibles et facilite l'analyse exhaustive de tous les réseaux possibles.

Références

Agresti, A. (2007). An introduction to categorical data analysis, Chapitre 1 2-22, Université de Floride.

Emelie lebarbier, Tristan Mary-huard ; septembre (2004), Le Critère BIC : Fondements théoriques et application.

Pearl, J. (2008). Causality, Chapitre 1 1-24, Université de Californie.

Rynkiewicz, J. (2013), Cours des Réseaux Bayésiens, Université Paris 1.

Summary

Our work allows to select variables to construct a Bayesian network that explains the intention to vote for a French political party, these variables are classified in two groups of variables (and v.contexte v.opinion) explanatory. Through these two groups of variables are constructed all possible logistic models associated with the context variables well in the opinion of variables from these two groups of models the best model is chosen with context variables and the best model with variable opinion within the meaning of the Bayesian information criterion (BIC). Context variables and opinion obtained from the models I want this are the variables that best explain the intention to vote. The Bayesian network follow an imposed sense of causality, context variables to variables of opinion. The last step in the construction of the network is to test connections to test the interactions between variables.

Parameter Estimation for time dependent drift to Stochastic Differential Equations

Meriem Bensalloua*, Kamel Boukhetala*

*University of Sciences and Technology Houari Boumediene
BP 32
El Alia 16111 Algiers
Algeria
mbensalloua@usthb.dz,
kboukhetala@usthb.dz
<http://www.usthb.dz>

Résumé. Dans ce travail, on étudie le problème de l'estimation des paramètres de dérive dépendants du temps d'un système dynamique régi par une équation différentielle stochastique linéaire unidimensionnelle excitée par des bruits brownien. D'une part, on montre que si les paramètres des bruits sont connus on peut, sous certaines conditions, construire une structure statistique. Puis on utilise le lemme de Randon-Nikodym pour approcher le problème d'estimation par la méthode du maximum de vraisemblance. On donne la forme générale de la fonction log-vraisemblance, et de sa dérivée première et seconde. D'autres part, on formule le problème d'estimation des paramètres de dérive d'une EDS analytiquement par un problème d'optimisation (maximisation de log-vraisemblance), où les paramètres de dérive sont les paramètres inconnus qu'on cherche à déterminer.

Dans la pratique, ce problème est un problème d'optimisation stochastique: sa fonction objectif (log vraisemblance) évolue au cours du temps. L'approche principalement adoptée dans cette étude consiste à adapter des algorithmes génétiques à codage réel RCGA pour l'estimation des paramètres qui sont en fonction du temps.

1 Introduction

The work we present in this study concerns a particular class of stochastic models of the dynamical system evolves in continuous time. This class is described by differential equations dimensional linear stochastic excited by Brownian noise, in the presence of control function (or control or strategy).

whether

$$\begin{cases} dX_t = (\varphi_1^{(t)} X_t + \varphi_2^{(t)} g(X_t))dt + \sigma dW_t \\ X_0 = 0, \quad t \geq 0 \end{cases} \quad (1)$$

where X_t is the state of the system at time t , $g(X_t)$ is the control applied to the system at time t , W_t is a standard Brownian motion. $\varphi_1^{(t)}$ and $\varphi_2^{(t)}$ are parameters which depend on the time drift that are unknown, σ is the diffusion parameter that we assumed constant in our study.

In practice, a variety of concrete examples which can be modeled by equation (1) is encountered. For example, in electronics, in a communication system the input signal is a linear equation and multidimensional controlled output signal is a linear equation (6).

We are primarily interested in the problem of parameter estimation of time-dependent drift in the light of the observation of a trajectory of the system state to a control selected in a class that will be defined. Various methods for estimating unknown parameters can be used; Breton (1977) and Boukhatala (1984) used the method of maximum likelihood for estimating parameters constant drift.

The estimation of model parameters in continuous time is not immediate. Many processes, such as model (1) does not admit exact discretization (6), and since its function determined likelihood is too complex calculate, we are then led to the use of heuristics methods, such as genetic algorithm.

Genetic algorithms generally offer the advantage of being better suited to continuous numerical optimization problems. They are widely used for solving optimization problems; they are recognized as very effective tools, especially when the parameters are complicated to solve and the function to optimize is not regular or its derivatives are inaccessible, poorly conditioned or complex to calculate, and the conventional optimization methods are completely inadequate.

2 RCGA for estimating parameters dependent drift time of EDS

The problem of parameter estimation is derived from an EDS is formulated mathematically as an optimization problem (maximization of log likelihood), as follows :

$$(P) \begin{cases} \text{Max } L_T(\varphi^{(t)}; \pi_t) \\ \varphi^{(t)} \in L = \{p_0, p_1, \dots, p_p\} \subseteq R^p \end{cases}$$

- $L_T(\varphi^{(t)}; \pi_t)$: log likelihood function determined by inequality 1.
- π_t : unique strong solution almost surely continuous trajectories of model 1.
- $\varphi^{(t)} = (\varphi_1^{(t)}, \varphi_2^{(t)})$: parameters derived from the model 1 that we seek to determine.
- L : The set of limited development of all the parameters $\phi^{(t)} = (\varphi_1^{(t)}, \varphi_2^{(t)})$ to order (p).

2.1 Digital Application

For $P = 10, D = 1000, T = 1, \theta = 1, N = 100, P_C = 0.9, P_M = 0.1, \gamma = 100, b = 5$
We place the RCGA Matlab 7.0 to iterate the following results were found :

Figure (1) shows a graph of the log-likelihood function after adjusting discretized. The log-likelihood function has many discrete local maxima. However, the function has global maxima, as shown in the following figure

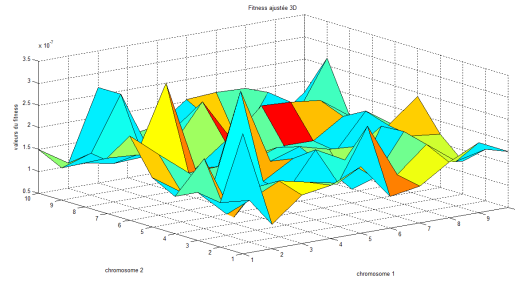


FIG. 1 – *Fitness3D*

2.1.1 Comment

RCGA method gives us a good estimate of solutions $\varphi_1^{(t)}$ and $\varphi_2^{(t)}$ to the 68th generation as shown in Figure (2). Estimators $\varphi_1^{(t)}$ and $\varphi_2^{(t)}$ provides us with a solution pt-adjusted to better fit the original solution π_t Figure (3).

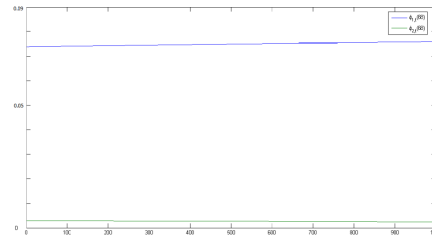


FIG. 2 –

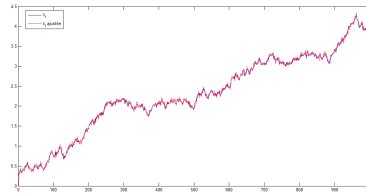


FIG. 3 –

Références

- [1] Alden H.Wright, *Genetic Algorithms for Real parameter Optimization*,Missoula, montana 1991.
- [2] Allen.E, *Modeling with Itô Stochasticc Differential Equations*, Springer 2007.
- [3] Andrea.S.Foulkes, *Applied statistical Genetics with R*,Springer 2009.
- [4] Balding.D.J, Bishop.M, Cannings.C, *Handbook of statistical genetics*, Wiley 2007.
- [5] Bischwal Jaya.P.N, *Parameter Estimation in stochastic Differential Equation*, Springer 2004.
- [6] Boukhatala Kamel, *Estimation des paramètres de dérive d'une equation différentielle stochastique linéaire, contrôlée*, thèse de Magistère 1984.
- [7] Boukhatala Kamel, *Identification and simulation of a communication system*, Maghreb Math. Rev., Vol. 4, No 2, pp.55-79, 1995.
- [8] Brodeau.F- Le Breton.A, *Identification de paramètres pour un système excité per des bruits Gaussien et Poissonien*, Annals de l'institut Henri Poincaré- Section B- Vol. XV, n° 1-1979.
- [9] Capasso.V -Bakstein.D , *An Introduction ton Continuous-Time Stochastic Processes*,Milian Research centre for industrial and applied Mathematics, pp. 127-203, 2003.
- [10] Chun-Liang Lin Ching-Huei Huang and Chin-Wei Tsai, *Structure-Specified Real Cided genetic Algorithms With Application*, TMRF e-book Advanced Knowledge based systems Model, Applications and Research (Eds.Sajja and Akerkar), pp. 160-187, 2010.
- [11] Comtet.L, *Advanced Combinatorics :The Art of Finite and Infinite Expansion*,D.Reidel Publishing,Dordrecht,The Netherlands,1974.
- [12] Doleans-Dade Cathrine, Meyer Paul-Ander, *Equations différentielles stochastiques*, Séminaire de probabilités, Strasbourg, tome 11 pp. 376-382, 1977.

Summary

In this work, we study the problem of parameter estimation of time-dependent drift of a dynamic system governed by a linear stochastic differential equation dimensional excited by a Brownian noise. On the one hand, we show that if the parameters of noises are known we can, under certain conditions, to build a statistical structure. Then we use the lemma Randon-Nikodym approach to the problem of estimating the maximum likelihood method. It gives the general form of the log-likelihood function and its first and second derivative. On the other hand, we formulate the problem of estimating parameters derived analytically by EDS of an optimization problem (maximization of log-likelihood), where the parameters are unknown parameters drift sought to determine. In practice, this problem is a stochastic optimization: the objective function (log likelihood) evolves over time. The main approach adopted in this study is to adapt Real Coded Genetic Algorithm (RCGA) for estimating parameters that are a function of time.

Estimation de la densité des retards dans un processus de type diffusion

Wahiba Benyahia* Tahar Mourid**

* ** Département de Mathématique, Université Abou Bekr Belkaid, Tlemcen

* b.stat1@yahoo.com

** t_mourid@mail.univ-tlemcen.dz

Résumé. Nous considérons l'estimation de la densité de la mesure des retards dans un processus de type diffusion en estimant ces coefficients dans un système de fonctions par la méthode du maximum de vraisemblance et de Bayes dans l'asymptotique des petites diffusions. Nous utilisons les résultats généraux de Ibragimov-Hasminski (Ibragimov et Hasminski, 1981) et (Kutoyants, 1994) pour établir la normalité asymptotique locale (LAN) (Le Cam condition), la convergence, la normalité et l'efficacité asymptotique des estimateurs. Nous illustrons cette étude par des simulations numériques.

1 Introduction

Nous considérons un processus de type diffusion $X^\varepsilon = (X_t^\varepsilon, t \in [-1, T])$ pour $\varepsilon \in]0, 1]$, défini sur un espace de probabilité complet $(\Omega, A, P, (\mathcal{F}_t))$ et solution de l'équation différentielle stochastique suivante :

$$\begin{cases} dX_t^\varepsilon = \left(\int_{\delta}^1 X_{t-s}^\varepsilon \mu(ds) \right) dt + \varepsilon dW_t, & t \in [\delta, T] \\ X_s^\varepsilon = x_0 & si \quad -1 \leq s < \delta \end{cases} \quad (1)$$

où

- $(W_t, t \in \mathbb{R}^+, (\mathcal{F}_t))$ est un processus de Wiener.
- x_0 est une constante strictement positif.

Nous étudions le problème de l'estimation de la mesure μ dans le cas où elle admet une densité par rapport à la mesure de Lebesgue admettant le développement suivant :

$$f(s) = \frac{d\mu}{ds}(s) = \sum_{i=1}^L c_i e_i(s) \quad (2)$$

où

- $(e_1, e_2, \dots, e_L)$ est un système de fonctions dans $L^2([\delta, 1])$,
- $c_1, c_2, \dots, c_L$ sont les coefficients de f ,

– L est un entier fixé.

Notre objectif est l'estimation du paramètre $\theta = (c_1, c_2, \dots, c_L)^t \in \Theta$, à partir de l'observation d'une trajectoire complète $X^\varepsilon = (X_t^\varepsilon, t \in [\delta, T])$ de (1).

Nous construisons les estimateurs du maximum de vraisemblance et de Bayes et nous étudions leurs propriétés asymptotiques quand $\varepsilon \rightarrow 0$.

Pour illustrer le comportement des estimateurs du maximum de vraisemblance étudiés, nous simulons des trajectoires de processus de diffusion par la méthode d'Euler-Maruyama. Nous utilisons le logiciel R pour l'analyse statistique des estimateurs. Nous simulons des trajectoires du mouvement brownien à l'aide de la bibliothèque "far" développée par J.Damons et S.Guillas (Modelization for Functional AutoRegressive processes Package : far Version : 0.6-0 (2005-01-10) License : LGPL-2.1 version 2.14.0 (31-10-2011) du logiciel R).

2 Notations et Hypothèses

Nous considérons l'équation différentielle stochastique (1) dans le cas où la mesure μ admet une densité par rapport à la mesure de Lebesgue de la forme (2). L'équation (1) se transforme en

$$\begin{cases} dX_t^\varepsilon = \left(\sum_{i=1}^L c_i \int_{\delta}^1 X_{t-s}^\varepsilon e_i(s) ds \right) dt + \varepsilon dW_t & \text{si } t \in [\delta, T] \\ X_s = x_0 & \text{si } -1 \leq s < \delta \end{cases} \quad (3)$$

Rappelons que l'estimateur du maximum de vraisemblance (EMV) $\hat{\theta}_\varepsilon$ est défini comme une solution de l'équation

$$\frac{dP_{\hat{\theta}_\varepsilon}^\varepsilon}{dP_{\theta_0}^\varepsilon}(X^\varepsilon) = \sup_{\theta \in \Theta} \frac{dP_\theta^\varepsilon}{dP_{\theta_0}^\varepsilon}(X^\varepsilon) \quad (4)$$

L'estimateur EMV $\hat{\theta}_\varepsilon$ vérifie le système suivant :

$$A_\varepsilon \hat{\theta}_\varepsilon = Y_\varepsilon$$

Nous montrons que la matrice A_ε est asymptotiquement inversible (p.s.).

La famille des lois $(P_\theta^\varepsilon, \theta \in \Theta)$ vérifie la condition LAN en un point $\theta \in \Theta$ si le rapport de vraisemblance

$$Z_{\varepsilon, \theta}(u) = \frac{dP_{\theta + \Phi_\varepsilon(\theta)u}}{dP_\theta}(X^\varepsilon), \quad u \in \mathbb{R}^L \quad (5)$$

admet la représentation suivante

$$Z_{\varepsilon, \theta}(u) = \exp \left[\langle u, \Delta_\varepsilon(\theta, X^\varepsilon) \rangle - \frac{1}{2} \|u\|^2 + \Psi_\varepsilon(\theta, u, X^\varepsilon) \right] \quad (6)$$

où $\Phi_\varepsilon(\theta)$ une matrice de normalisation ($L \times L$),

D'autre part, notons par $\mathbf{W}_{e,2}$ l'espace des fonctions de perte l définies sur $\mathbb{R}^L$ continues symétriques non identiquement nulles vérifiant certaines conditions.

Si la famille $(P_\theta^\varepsilon, \theta \in \Theta)$ satisfait la condition LAN en tout point $\theta \in \Theta$ avec une matrice de normalisation $\Phi_\varepsilon(\theta)$, alors pour tout estimateur $\tilde{\theta}_\varepsilon$, et $l \in \mathbf{W}_{e,2}$, nous avons l'inégalité suivante (borne de Hajek) :

$$\lim_{\delta \rightarrow 0} \liminf_{\varepsilon \rightarrow 0} \sup_{|\theta - \theta_0| < \delta} E_\theta l \left(\Phi_\varepsilon^{-1}(\theta_0)(\tilde{\theta}_\varepsilon - \theta_0) \right) \geq \frac{1}{(2\pi)^{\frac{L}{2}}} \int_{\mathbb{R}^L} l(x) \exp \left(-\frac{\|x\|^2}{2} \right) dx \quad (7)$$

Les estimateurs qui réalisent l'égalité sont dits asymptotiquement efficaces.

Posons la conditions suivante :

C 1. les fonctions $e_1, e_2, \dots, e_L$ sont continues et linéairement indépendantes et d'intégrales positives.

3 Résultats

Le théorème suivant donne la condition LAN de la famille des lois $(P_\theta^\varepsilon, \theta \in \Theta)$ solution de (3) et la borne de Hajek, ce qui implique la convergence des lois de dimensions finies du processus de vraisemblance $(Z_{\varepsilon, \theta}(u), u \in \mathbb{R}^L)$.

Théorème 3.1 Sous la condition **C 1**, la famille des lois $(P_\theta^\varepsilon, \theta \in \Theta)$ solution (3) vérifie la condition LAN (6) avec la matrice Φ_ε définie par

$$\Phi_\varepsilon(\theta) = \varepsilon I^{-\frac{1}{2}}(\theta) \quad \text{où} \quad I(\theta) = \int_0^T q_s(X^0) q_s^t(X^0) ds$$

$$q_s(X^0) = \frac{\partial S_s}{\partial \theta}(X^0, \theta) \text{ et } \Delta_\varepsilon(\theta, X^0) = \int_0^T \langle u, q_t(X^0) \rangle dW_t.$$

La fonction risque admet la minoration suivante (Inégalité de Hajek) : pour tout estimateur $\tilde{\theta}_\varepsilon, \theta_0 \in \Theta$ et $w \in \mathbf{W}_{e,2}$

$$\lim_{\eta \rightarrow 0} \liminf_{\varepsilon \rightarrow 0} \inf_{\tilde{\theta}_\varepsilon} \sup_{\|\theta - \theta_0\| < \eta} E_\theta^\varepsilon \left(l \left(\Phi_\varepsilon^{-1}(\theta_0) (\tilde{\theta}_\varepsilon - \theta) \right) \right) \geq E(l(\xi))$$

où ξ est un vecteur aléatoire gaussien centré réduit dans $\mathbb{R}^L$.

L'estimateur du maximum de vraisemblance du paramètre θ est défini par (4). Le théorème suivant donne la convergence, la normalité asymptotique et la convergence des moments de tout ordre de cet estimateur.

Théorème 3.2 . Sous la condition **C 1**, l'estimateur du maximum de vraisemblance $\hat{\theta}_\varepsilon$ vérifie sous $P_{\theta_0}^\varepsilon$ et uniformément sur tout compact K de Θ , les propriétés suivantes :

1. $\lim_{\varepsilon \rightarrow 0} \hat{\theta}_\varepsilon = \theta_0$ en probabilité.

Plus précisément, $\forall \beta > 0$,

$$\sup_{\theta_0 \in K} P_{\theta_0}^\varepsilon \left(\left\| \hat{\theta}_\varepsilon - \theta_0 \right\| > \beta \right) \leq C_1 \exp\left(\frac{-C_2 \beta}{\varepsilon^2}\right)$$

où $C_1 > 0$ et $C_2 > 0$,

2.

$$\Phi_\varepsilon^{-1}(\theta_0) \left(\hat{\theta}_\varepsilon - \theta_0 \right) \Rightarrow_{\varepsilon \rightarrow 0} \xi$$

où $\xi \hookrightarrow N(0, Id)$ dans $\mathbb{R}^L$,

3. Pour tout $p > 0$,

$$\lim_{\varepsilon \rightarrow 0} E_{\theta_0}^\varepsilon \left\| \Phi_\varepsilon^{-1}(\theta_0) \left(\hat{\theta}_\varepsilon - \theta_0 \right) \right\|^p = E \|\xi\|^p$$

4. L'estimateur $\hat{\theta}_\varepsilon$ est asymptotiquement efficace.

Références

- Benyahia, W. et T. Mourid (2011). Estimation de la densité des retards dans un processus de type diffusion. *Pub. Inst. Stat. Univ.Paris.* 2, 3, 43–64.
- Ibragimov, I. et R. Hasminski (1981). *Statistical Estimation: Asymptotic theory*. New York: Springer Verlag.
- Kuchler, U. et Y. Kutoyants (2000). Delay estimation for some stationary diffusion-type processes. *Scandinavian Journal of Statistics* 27, 3, 405–414.
- Kutoyants, Y. (1984). *Parameter Estimation For Stochastic Processus*. Berlin: Heldermann.
- Kutoyants, Y. (1994). *Identification of Dynamical Systems with small noise*. Dordrecht: Kluwer.
- Kutoyants, Y. A. (1988). Delay estimation for some stationary diffusion-type processes. *Theory Probab. Appl.* 33, 1, 175–179.
- Kutoyants, Y. A., T. Mourid, et D. Bosq (1992). Estimation paramétrique d'un processus de diffusion avec retards. *Ann. Inst. H. Poincaré Probab.Statist.* 28, 95–106.

Summary

We study the estimation of the delay measure density in type diffusion processes. We consider the maximum likelihood and Bayes estimation of a finite number of density coefficients by the Ibragimov Hasminski and Kutoyants approach. For these estimators we give the LAN condition, consistency, asymptotic normality and efficiency. Then we study the minimal distance estimation of the parameters giving consistency and asymptotic normality. We also present some numerical simulations on maximum likelihood estimators of the parameters.

Le Temps d'Occupation Pour Les $k^{ièmes}$ Extrêmes $k \geq 1$

Salah Berrouane*,

*Faculty of Mathematics, University of Sciences and Technology Houari Boumediene
Po. Box 32, El Alia, Bab Ezzouar, 16111, Algiers, Algeria
sberrouane@yahoo.fr,

Résumé. Un théorème de la limite locale pour la $k^{ième}$ ($k \geq 1$) valeur de record est prouvé au sens de l'approximation forte, et un résultat concernant le temps d'occupation est prouvé; les methode de la preuve sont différentes de celles utilisées pour les résultats des sommes partielles corespondantes. Finalement sous les conditions du lemme 2.1, nous estimons la distance entre le processus empirique inverse et le processus de quantile uniforme. Dans ce papier notre intention est le temps d'occupation.

1 Introduction

Dans la théorie des $k^{ièmes}$ valeurs extrêmes, peut de choses sont connues à propos de la qualité de la convergence dans le cas de la convergence faible des maximums normalisés d'un échantillon vers l'une des lois limites. Des tentatives pour établir des rythmes uniformes de convergences ont été sporadiques et non satisfaisantes; ainsi nous avons été appelés à introduire l'idée de l'approximation forte voir A. A. Belkema (1973); P. Deheuvels (1983); S.I. Resnick (1974).

Soit (Ω, A, P) un espace de probabilités et soit S un espace de Banach des fonctions bornées définies sur Ω par exemple $(\Omega = R)$ avec la norme $\|u(x)\| = \sup |u(x)|$; et, soit $X_1, X_2, \dots$, une suite de variables aléatoires (*v.r*) indépendantes identiquement distribuées (*i.i.d.*), de fonction de répartition continue $F(x) = P(x_1 \leq x)$. Nous dirons que X_n est une plus grande valeur de record si elle est plus grande que toutes les observations précédentes $X_1, X_2, \dots, X_{n-1}$, de la suite par convention X_1 est la plus grande valeur de record.

De la même manière si pour $n = 1, 2, \dots$, $X_{1,n} < X_{2,n} < \dots < X_{n,n}$ est la statistique d'ordre associée à la suite $X_1, X_2, \dots, X_n$;

Definition.1.1.

Pour $k \geq 1$ arbitraire, le $k^{ième}$ plus grand instant de record est défini par $\{i = 2, 3, \dots\}$

$$N^{(k)}(1) = k, N^{(k)}(i) = \min\{m > N^{(k)}(i-1), X_{m-k+1,m} > X_{N^{(k)}(i-1)-k+1, N^{(k)}(i-1)}\}. \quad (1)$$

La suite des $k^{ième}$ records (ou valeurs de record) est définie par

$$R_i^{(k)} = X_{N^{(k)}(i)-k+1, N^{(k)}(i)}, i = 1, 2, \dots$$

En d'autres termes $R_i^{(k)} = X_{N^{(k)}(i)-k+1, N^{(k)}(i)}$ est le $k^{ième}$ record s'il est plus grand que le $k^{ième}$ valeur extrême des observations

$X_1, X_2, \dots, X_n$. L'instant n quand un $k^{-ième}$ record a eu lieu est appelé un $k^{-ième}$ temps de record voir [3] P.Deheuvels (1983).

Une définition séculaire peut être donnée pour définir les plus petites valeurs de record. Par convention $X_{1,k}$ et $X_{k,k}$ sont respectivement le $k^{-ième}$ plus grand record et le $k^{-ième}$ plus petit record.

Corollaire.1.1

Soit $X_1, \dots, X_n$ une suite de variables aléatoires de fonction de répartition $F(x) = P(X_1 \leq x)$ continue, si $M(x) = -\text{Log}(1 - F(x))$.

$$P(R_n^{(k)} \leq x) = (1 - F(x))^k \sum_{i=0}^{n-1} (-kM(x))^i / i! = \int_0^{-kM(x)} t^{n-1} e^{-t} / i! dt.$$

2 Les Lois Limites

Théorème.2.1

Sous les hypothèses du corollaire 0.1 ; et ; si $F(x) = P(X_1 \leq x) = 1 - e^{-x}, x > 0$. Alors quand $n \rightarrow \infty$.

1) Si $t, t > 0$ est l'époque du premier passage à travers $+1$, nous avons :

$$P(N^{(k)}(n) > t) = \int_R (1 - e^{-(x+\text{Log}n)})^t dP(R_n^{(k)} \leq x + \text{Log}n) = (k^n / 3n^k) \int (x + \text{Log}n)^{n-1} e^{-(k-1)x} / (n-1)! \Lambda(dx)$$

2) Si t appartient à la sigma algèbre engendrée par les instants de record i.e. $t \in \sigma(N^{(k)}(n), n \geq 1)$, nous avons

$$P(N^{(k)}(n) > t) = \int_R (1 - e^{-(xn^{1/2}/k + n/k)})^t dP(R_n^{(k)} \leq xn^{1/2}/k + n/k) = \int_R N(x) dP(R_n^{(k)} \leq xn^{1/2}/k + n/k) = 2^{-1} [N^2(x)]_{-\infty}^{+\infty} = 1/2$$

3) Si $(1 - e^{-x})^{-1} = 1 + (t - x)t^{-1}$ nous avons,

$$P(N^{(k)}(n) > t) = \int_R (1 - e^{-(xn^{1/2}/k + n/k)})^t dP(R_n^{(k)} \leq xn^{1/2}/k + n/k) = (1 - e^{-n/k})^t (\text{Log}(1 - e^{-n/k}))^2.$$

On considère le processus Gaussien $\{V(t) = W(t)/\sqrt[2]{t}, 0 < t < \infty\}$, où $\{W(t), 0 < t < \infty\}$ est un processus de Wiener. Alors $EV(t) = 0$, $EV^2(t) = 1$ et $EV(t)V(s) = \sqrt[2]{s/t}$, $s < t$. La forme de cette fonction de covariance suggère immédiatement que dans le but d'obtenir un processus Gaussien stationnaire à partir de $V(t)$, nous considérons $U_\alpha(t) = V(e^{\alpha t})$, $-\infty < t < \infty$ (α fixé positif); ce dernier processus est un processus Gaussien stationnaire, pour $EU_\alpha(t)U_\alpha(s) = e^{-\alpha|t-s|/2}$, et il est appelé le processus d'Ornstein-Uhlenbeck. Nous utiliserons la notation $U(t) = U_2(t)$, et mentionnons ce résultat pour n appartenant à l'adhérence de $N = \{1, 2, \dots\}$

Théorème.2.2

$$\lim_{n \rightarrow \infty} P\{\sup U(t) \leq x/\sqrt[2]{2k} + \sqrt[2]{2n/k}\} = N(x) = N(0, 1) = (2\pi)^{-1/2} \int_{-\infty}^x e^{-s^2/2} ds.$$

$$\lim_{n \rightarrow \infty} P\{\sup |U(t)| \leq x/\sqrt[2]{2k} + \sqrt[2]{2n/k}\} = N(2x) = (2\pi)^{-1/2} \int_{-\infty}^{2x} e^{-s^2/2} ds.$$

Proof. Nous utilisons l'argumentation de Darling, D.A.-Erdős, P. voir [2] M. Csörgö. P. Révész. Theorem 1.9.1. page 56.

3 Temps d'Occupation

Les théorèmes du temps d'occupation sont fréquemment connectés avec les théorèmes de la limite locale. Voir Breiman (1968), Darling and Kac (1957). Pour le cas des $k^{-ième}$ extrêmes

nous n'avons pas trouvés une connection directe mais seulement une méthodologie analytique que nous discuterons en dessous.

Dans ce qui va suivre il convient d'avoir une représentation légèrement différente pour $F \in D(\Lambda)$ ie (F appartient au domaine d'attraction du maximum Gumbel "M.D.A.") et, nous procédons comme dans de Haan (1970), Theorem 2.4.2. Nous restreindrons notre attention au cas où $F(x) < 1$ pour tout x . En consequence de ce résultat $F \in D(\Lambda)$, i.e. (quand $n \rightarrow \infty$ $\lim F^n(x + \text{Log} n) = \Lambda(x) = \exp e^{-x}$, $-\infty < x < +\infty$) si et seulement si

$\lim_{n \rightarrow \infty} (1 - F(x)) \int_x^\infty \int_s^\infty (1 - F(u)) du ds / (\int_x^\infty (1 - F(s)) ds)^2 = 1$, et dans ce cas nous avons la représentation de Karamata (1936).

$$1 - F(x) = c(x) \exp\left\{-\int a(t)/f^*(t) dt\right\}$$

$$f^*(x) = \int_x^\infty \int_s^\infty (1 - F(u)) du ds / \int_x^\infty (1 - F(x)).$$

Lemme.3.1

Supposons $F \in D(\Lambda)$ et $a_n \rightarrow 0$ (i.e. $f^*(x) \rightarrow 0$, $x \rightarrow \infty$). Soit $M(x) = -\text{Log}(1 - F(x))$ and suppose $\lim_{t \rightarrow \infty} f^*(x+t)/f^*(t) = 1$. Alors pour tout x .

$a_n(M(x+b_n) - \text{Log} n) \rightarrow x$. ou en termes de mesure $a_n(M(\cdot + b_n) \rightarrow_v m$ où m est la mesure de Lebesgue sur $(-\infty, +\infty)$ et $\rightarrow_v$ dénote la convergence vague ;

Pour la preuve voir [7] J. Neveu.

Remarque.3.1

La condition $f^*(x+t)/f^*(t)$ est satisfaite par $1 - e^{-x^r}$, $x > 0$ for $r > 1$ puisque dans ce cas nous prenons

$f^*(x) \sim (1 - F(x))/F'(x) = r^{-1}x^{-r+1}$. Aussi dans le cas où $F(\cdot)$ est distribution normale centrée réduite i.e. $F(\cdot) = N(0, 1)$ la condition est satisfaite si $n \rightarrow \infty$

$$f^*(x) \sim (1 - F(x))/F'(x) = (1 - N(0, 1))/N'(0, 1) = e^{-x^2/2}/xe^{-x^2/2} = x^{-1}$$

Maintenant nous considérons un théorème du temps d'occupation pour la k -ième extrême et exprimons le en termes de mesures aléatoires. Définissons pour $x \in M(\cdot)$ la fonction indicatrice $1_x(A) = 1$ si $x \in A$, 0 ailleurs. La mesure du temps d'occupation pour le k -ième maximum est alors

$0^{(k)} = \sum_{i=1}^\infty 1_{X_{i-k+1,i}}(\cdot)$. Il est plus convenable d'exprimer $0^{(k)}$ en termes de k -ième plus grandes valeurs de record $\{R_n^{(k)}, n \geq 1\}$. Soit $\Delta_i^{(k)} = N^{(k)}(i) - N^{(k)}(i-1)$, $i \geq 1$. Alors

$$0^{(k)} = \sum_{i=1}^\infty \Delta_i^{(k)} 1_{R_i^{(k)}}(\cdot). \quad \text{et} \quad \hat{O}^{(k)} = \sum_{i=1}^\infty 1_{R_i^{(k)}} \quad \text{En premier lieu nous}$$

calculons la fonctionnelle de Laplace (cf, J. Neveu, 1977, page 258, pour $k = 1$) de $0^{(k)}$.

Proposition.3.1

La fonctionnelle de Laplace du processus ponctuel $\hat{O}^{(k)}$ est donnée par :

$$E \exp\{-g(x)\hat{O}^{(k)}(dx)\} = \exp\{-k \int (1 - e^{-g(x)})M(dx)\}$$

Preuve.

Cette fonctionnelle de Laplace correspond au processus ponctuel $\sum_{i=1}^\infty 1_{t_i}$ où, $\{t_i\}$ sont les points d'un processus de Poisson sur $kM(\cdot)$.

Proposition.3.2

Nous supposons $F(\cdot)$ est continue. Alors pour $g(\cdot)$ continue à support compact nous avons

$$E \exp\{-k \int g(x)0^{(k)}(dx)\} = \exp\{-k \int (1 - e^{-g(x)})(1 - e^{-g(x)}F(x))^{-1}M(dx)\}$$

Remarque.3.2

Le fait que $\hat{O}^{(1)}(t)$ est un processus de Poisson est une conséquence facile de structure additive des k -ième valeurs de record. Ce fait a été remarqué par Shorrrac (1972 b)

Théorème.3.1 ;

Nous supposons $F(\cdot) \in D(\Lambda)$ est continue et la fonction auxiliaire f^* satisfait $\lim_{t \rightarrow \infty} f^*(t+x)/f^*(t) = 1$ et $\lim_{t \rightarrow \infty} f^*(t) = 0$. Nous posons $\mu(A) = \int_A (1 - F(y))^{-1} dy$ pour A un Borelien de $M(\cdot)$. Alors

$a_n(\cdot + b_n) = a_n \sum_{i=1}^{\infty} 1_{X_{i-k+1, i-b_n}(\cdot)} k \mu$ dans le sens de la convergence faible des processus ponctuels stochastiques ;

Pour la discussion de la convergence faible des processus ponctuels voir [7], J. Neveu, 1977, page 282.

Corollaire.3.1

Sous les hypothèses du Théorème 2.1., si $k = k(n)$ i.e. (k dépend de n), et $\lim_{t \rightarrow \infty} a_n(\text{Log} k)/(k-1) \rightarrow_v \mu$, nous avons :

1) Si $F(x)$ est la fonction de distribution normale $N(0, 1)$, alors sans perte de généralités il existe sur le même espace de probabilités (Ω, A, P) une suite de (v.a) $U_1, U_2, \dots$, (i.i.d.), de fonction de distribution $G(u) = (-2 \sqrt[2]{n}(2k-1))^{-1}, (2k+1)(2k)^{-1}$, telque

$$\mu = \sqrt[2]{\pi/k} \int_0^x (u - \sqrt[2]{2k} + \sqrt[2]{2n/k}) \sqrt[2]{2k/2k+1} / G'(u) du.$$

Finalement, quand on parle de convergence vague il est intéressant de voir ce qui pourrait se passer pour le processus de quantile. Si $\text{inv}F(y)$ est la fonction réciproque de $F(x)$; on défini la fonction quantile par : $Q_n(y) = X_{i,n}$, si $(i-1)/n < y \leq i/n$, $i = 1, 2, \dots, n$, et le processus de quantile par : $q_n(y) = \sqrt[2]{n}(Q_n(y) - \text{inv}F(y))$, $0 < y \leq 1$, ce dernier processus dès fois est appelé processus empirique inverse, de la même manière on défini le quantile uniforme $U_n(y) = U_{i,n}$, si $(i-1)/n < y \leq i/n$, et le processus de quantile uniforme $u_n(y) = \sqrt[2]{n}(U_n(y) - y)$, $0 \leq y \leq 1$. Alors on a :

Corollaire.3.2 : Sous les conditions du lemme 2.1, nous avons :

$$\sup_{0 < y \leq 1} |f(\text{inv}F(y))q_n(y) - u_n(y)| = O((\text{Log} \text{Log} n) / \sqrt[2]{n})$$

Corollaire.3.3.

Sous les hypothèses du Théorème 1.1 pour tout $c > 0$, sans pertes de généralités il existe sur le même espace de probabilité (Ω, A, P) un processus de Wiener $\{W(t), t > 0\}$ telle que : $\lim_{n \rightarrow \infty} \sup_{0 \leq t \leq n - c \text{Log} n} | \frac{W(t + c \text{Log} n) - W(t)}{c} | p.s. = \lim_{n \rightarrow \infty} (x + \text{Log} n)$

Proposition.3.3.

Sous les hypothèses du Théorème 1.1 et ; supposons que $F(\cdot)$ est définie sur $(H, B) \rightarrow [0, 1]$, $H \subseteq \Omega$. Alors si H est inductif on a $0 \leq 2c - 1$, et ; la mesure de Lebesgue μ est telle que : $\mu \leq 2 \sqrt[2]{2 \text{Log} n}$.

Remark.3.3.

En considérant la proposition 2.3., en termes de mesure nous remarquons que la loi normale standard $N(0, 1)$ est telle que : $N(\cdot + \sqrt[2]{2 \text{Log} n}) \leq 1 - n^{-4}$.

4 Perspective

a) Sous les hypothèses du Théorème 1.1, en posant $S_{N^{(k)}(n)} = R_1^{(k)} + R_2^{(k)} + \dots + R_n^{(k)}$ nous avons quand $n \rightarrow \infty$ $\lim P(S_{N^{(k)}(n)} \leq x \sqrt[2]{n}) = 2\pi^{-1} \arcsin \sqrt[2]{x}$, $0 < x < 1$.

b) Si $F \in D(\Lambda)$, alors quand $n \rightarrow \infty$ on a $\lim P(f(\text{inv}F(y))q_n(y) < t) \rightarrow N(W(t))$, $t > 0$.

Références

- Salah. Berrouane (1986) : Lois Limites des k -èmes Valeurs de Record et leurs Concomitants Doctorat 3-ème cycle , Univ. Pierre et Marie Curie, Jussieu, Paris VI. FRANCE.
- M.Csörgö and P. Révész :Strong Approximation in Probability and Statistics. ACADEMIC PRESS New York San Francisco London 1981.
- Paul, Deheuvels. (1983)The Complete Characterization of an Upper and Lower Class of the Record and Inter Record Times of an I.I.D. Sequence. Z. Wahrsch. Verw. Gebiete. 62, 1-6.
- William. Feller (1966) : An Introduction to Probability Theory and its Applications. Vol II, John Wiley & Sons, Inc. New york . London . Sydney.
- Galambos Janos. The Asymptotic Theory of Extreme Order statistics. John Wiley & Sons, New York.. Chichester ; Brisbaane. Toronto.
- Haan. L. de.(1970). On Regular Variation and its Application to the Weak Convergence of Sample Extremes. Mathematicaltract 32. Amesterdam.
- Neveu. Jacques. (1977). Processus ponctuels. Ecole d'été de Probabilités de Saint-FlourVI-1976, Lecture Notes in Math. 598, Heidelberg
- Resnick, S.I., (1973c) : Extremal Processes and Record value times. J. Appl. Pro. 10, 863-868.

Summary

A limit local theorem for the k^{th} ($k \geq 1$) record valor is proved in sense of strong approximation, and a result about occupation time is proved; proof methods are quite different from those used corresponding partial sums results. In this paper, our goal is the occupation time.

Asymptotics for Distorsion risk measure under dependence

Berkoun Youcef*,

*Laboratoire de Mathématiques Pures et Appliquées [LMPA]
Université Mouloud MAMMERI de Tizi Ouzou
youberk@yahoo.com

Résumé. Using the Bahadur representation of sample quantile for strong mixing random variables, we give the asymptotic normality of the empirical estimator of a distorsion risk measure.

1 Introduction

Let $(\Omega, \mathcal{A}, P)$ be some probability space and $X : \Omega \rightarrow R$ be the loss or gain variable (risk). X can be seen as the outcome of a financial position. Denote by $\mathcal{X}$ the set of all risks and by F the cumulative distribution of X .

A risk measure is a function $\rho : \mathcal{X} \rightarrow R$.

Risk measures are usually used by actuaries in premium setting see Wang(1996). There is a considerable literature on the concept of risk measure, see for example Azner et al (1999), Dhaene et al(2006), and references therein.

Let $g : [0,1] \rightarrow [0,1]$ be a distorsion function, namely, a nondecreasing function with $g(0) = 0$ and $g(1) = 1$. A distorsion risk measure (DRM) associated with $g(x)$ noted by $\rho_g(F)$ is defined as

$$\rho_g(F) = \int_0^1 F^{-1}(x)dg(x) = \int_R xdg(F(x)) = - \int_{-\infty}^0 g(F(x))dx + \int_0^{+\infty} (1 - g(F(x)))dx \quad (1)$$

provided that $E(|X|) < \infty$, and F^{-1} is the quantile function corresponding to F , that is $F^{-1}(t) = \inf\{x : F(x) \geq t\}$.

The risk measure can be viewed as a distorted expectation of X in the sense of Choquet integral, see Denneberg(1994).

Value at Risk (VAR) is a distorsion risk measure with associated distorsion function

$$g(x) = \begin{cases} 0 & \text{if } 0 \leq x < 1 - \alpha \\ 1 & \text{if } 1 - \alpha \leq x \leq 1 \end{cases}$$

The class of DRM introduced first by Wang (1996) has received much attention in the finance and actuarial science. Any DRM possesses the properties of additivity, positive homogeneity, translation invariance and monotonicity (see Balbas(2009)). We will assume thorough this paper that g is concave. The concavity of g makes the corresponding distorsion measure coherent see Hardy et al (2001). The aim of this paper is to estimate the distorsion risk measure $\rho_g(F)$ when the observations $X_1, X_2, \dots, X_n, \dots$ are weakly dependent. Estimating the risk

measure from a finite sample has been studied in the literature, see for instance Jones et al (2003). Using the relation between the estimators of DRM and the class of L -estimators, Jones and Zitikis(2007), Brakauskas et al(2008), have studied the behavior of those estimators. For dependent observations, recently D. Belomestny et al (2012) and Tsukahara.H (2014), established some asymptotic properties for the estimators of DRM based on L -statistics. In this paper, using a different approach than used by Tsukahara, we give the asymptotic behavior of L statistics. Our approach is based on the Bahadur representation for sample quantile for strong mixing random variables.

Let $F_n(x) = \frac{1}{n} \sum_{i=1}^n I(X_i \leq x)$ be the empirical distribution function belonging to the sample $X_1, X_2, \dots, X_n$. The estimate F_n can then be plugged in (1) in order to obtain an estimate $\tilde{\rho}_g(F_n)$ of $\rho_g(F)$. A natural estimator of the DRM is then given by

$$\tilde{\rho}_g(F_n) = \int_0^1 F_n^{-1}(x) dg(x) = - \int_{-\infty}^0 g(F_n(x)) dx + \int_0^{+\infty} (1 - g(F_n(x))) dx \quad (2)$$

The last expression can be written as

$$\tilde{\rho}_g(F_n) = \sum_{i=1}^n X_{(i)} \left[g\left(\frac{n-i+1}{n-i}\right) - g\left(\frac{n-i}{n}\right) \right]$$

Let $C_{n,i} = g\left(\frac{n-i+1}{n}\right) - g\left(\frac{n-i}{n}\right) = \frac{1}{n} J\left(\frac{i}{n+1}\right)$ where J is a weight function defined on $[0,1]$

Denote $\tilde{\rho}_g(F) = T_n = \sum_{i=1}^n J\left(\frac{i}{n+1}\right) X_{(i)}$.

This shows that the natural estimator of DRM is an L -estimator. In the following, we give some properties of the estimator considered.

Proposition 1. $\tilde{\rho}_g(F_n)$ is biased estimator, that is:

$$E(\tilde{\rho}_g(F_n)) \leq \rho_g(F)$$

Proof. Use the fact that for any concave function g , we have $E(g(X)) \leq g(E(X))$

Proposition 2. [Strong consistency]

Assume that $(X_t)_t$ is a strict stationary processes, and $g \in L^p(0,1)$, $F^{-1} \in L^q(0,1)$, $\frac{1}{p} + \frac{1}{q} = 1$, $1 \leq p \leq \infty$ then

$$\tilde{\rho}_g(F_n) \longrightarrow \rho_g(F) \quad a.s$$

Proof. see Gilat et al (1997) or Baklanov, E.A(2006). The proof is based only the SLLN and the Glivenko-cantelli for α -mixing sequence.

Assume the following conditions:

- **Condition (A)**

$$- J(t) = 0 \quad \forall \delta \notin [\delta, 1 - \delta], \quad 0 < \delta < \frac{1}{2}.$$

- J and F have no common discontinuity point.
- F^{-1} is a function of bounded variation on compact sets of $[0,1]$.

- Condition (B)

Let $F(x_p) = p$. Assume that $F' = f$ is bounded in some neighborhood V_p of x_p , $0 < x_p < \infty$ and $0 < f(x_p) < \infty$ and f' is bounded in V_p .

- Condition (C)

- The processes $(X_t)_t$ is strong mixing with mixing coefficients $(\alpha(n))_n$ satisfying $\alpha(n) \leq Cn^{-\beta}$ for some positive constants C and β where $\beta > 3$.
- $0 < d_1 = \inf(f(x) : x \in V_p) \leq d_2 = \sup(f(x) : x \in V_p) < \infty$

Let us introduce the concept of α -mixing processes.

Définition Let $(X_t)_t$ be a strict stationary process and $\mathcal{F}_a^b$ the σ -algebra generated by $X_a, X_{a+1}, \dots, X_{a+b}$, $-\infty \leq a \leq b \leq +\infty$. $(X_t)_t$ is said strong mixing or α -mixing if

$$|P(A \cap B) - P(A)P(B)| \leq \alpha(k), \quad \alpha(k) \rightarrow 0 \quad \text{as} \quad k \rightarrow \infty$$

Where $A \in \mathcal{F}_{-\infty}^m$, $B \in \mathcal{F}_{-\infty}^{m+k}$

Qinchi Zhang et al (2012) obtained the following result on Bahadur representation for sample quantile under α -mixing sequence.

Théorème 1. Assume that $(X_t)_t$ is a strict stationary process α -mixing satisfying conditions (B) and (C), then, wp1,

$$\hat{X}_{np} = x_p - \frac{F_n(x_p) - p}{f(x_p)} + O(n^{-\frac{1}{2}} + (\log \log n \cdot \log n)^{\frac{1}{2}}), \quad n \rightarrow \infty$$

Using this representation, we obtain

Théorème 2. If conditions (A), (B), and (C) hold, then

$$\sqrt{n}(\tilde{\rho}_g(F) - \rho_g(F)) \rightarrow N(0, \sigma^2)$$

Where

$$\sigma^2 = \sum_{k=-\infty}^{+\infty} \int_0^1 \int_0^1 (F_k(x, y) - F(x)F(y)) \frac{J(x)J(y)}{f(F^{-1}(x))f(F^{-1}(y))} dx dy$$

$$F_k(x, y) = P(X_0 < x, X_k < y)$$

Bibliography

- Artzner P, et al, *Coherent measures of risk*, Math. Finan(1999) 9 (3), 203-228.

- Baklanov E.A, *The strong law for large numbers for L-statistics with dependent data* (2006), Siberian Mathematical Journal, V47, 975-979.
- Balbas. A et al, *Properties of distorsion risk measures*, (2009), Methodol.comput Appl.Probab,11, 385-399.
- Bruce L Jones and Ricardis Zitikis, *Empirical estimation of risk measures and related quantities*, (2003), North American Actuarial Journal 7 (october), 44-54.
- Bruce L.Jones and Ricardis Zitikis, *Risk measures, distorsion parameters, and their empirical estimation*, Insurance: Mathematics and Econometrics (2007), V41, Issue 2, 279-297.
- David.Gilat and Roelof Helmers, *On strong law for generalized L- statistics with dependent data*, Commentationes Mathematicae Universitatis Carolinae (1997), V37, N° 1, 187-192.
- Denis Belomestny and Walker Kratschmer, *Central limit theorems for law invariant coherent risk measures*, J. Appl. Prob 49, 1-21(2012).
- Dhaene.J et al, *Risk measures and comotonicity, a review*, (2006), Stoch.Models, 22, 573-606.
- Gorodotskii, V.V, *On the strong mixing properties for linear sequences*, Theory.proba. appli, 22, pp 411-413
- Qinchu Zhang et al, *On Bahadur representation for sample quantiles under α -mixing sequence*, Stat Papers (2012), 165, 579-596.
- Tsukahara, Hideatsu, *Estimation of distorsion risk measures*, Journal of financial econometrics (2014), V12, Issue 1, 213-235.
- Van Zwet W.R, *A strong law for linear functions of order statistics*, Anna. Probab 8(1980),986-990.

Summary

Using the Bahadur representation of sample quantile for strong mixing random variables, we give the asymptotic normality of the empirical estimator of a distorsion risk measure.

Some asymptotic normality result of k -Nearest Neighbour estimator of the conditional mode function for independent functional data.

Wahiba Bouabça, Mohammed Kadi Attouch

Univ. Djillali Liabès, Sidi Bel Abbès
wahiba_bouab@yahoo.fr, attou_kadi@yahoo.fr.

Résumé. Dans ce papier, nous étudions l'estimation nonparamétrique du mode conditionnel par la méthode des k -Plus Proches Voisins pour une variable de réponse scalaire donnée une variable aléatoire prise des valeurs dans un espace semi-métrique. Nous établissons la normalité asymptotique pour des données fonctionnelles indépendantes et proposons des bandes de confiance pour la fonction de mode conditionnelle. Quelques simulations ont été conduites pour montrer comment notre méthodologie peut être mise en oeuvre.

1 Introduction

The first result of the asymptotic normality on the kernel estimator comes from Masry (2005), he considers the case of α -mixing data but he did not give the explicit expression of the dominant asymptotic terms of bias and variance. After Ferraty et al.(2007) gives the explicit expression of the asymptotic law (that is the dominant terms of bias and variance) In the case of set of independent data . The results exposed in the papers of Delsol (2007a, 2008b) make the link between these articles, they generalize the results of Ferraty et al.(2007) in the case of data α -mixing, Attouch and Benchikh (2012) established the asymptotic normality of robust nonparametric regression function.

For the k -NN conditional mode estimator Attouch and Bouabça (2013) obtained the almost complete convergence with rates in independent and identically distributed (i.i.d.) functional data case.

Let $(X_i, Y_i), i = 1, \dots, n$ be n copies of random vector identically distributed as (X, Y) where X is valued in infinite dimensional semimetric vector space $(\mathcal{F}, d)$ and Y 's are valued in $\mathbb{R}$, In most practical applications, $\mathcal{S}$ is a normed space which can be of infinite dimension (e.g. Hilbert or Banach space) with norm $\|\cdot\|$ so that

$$d(x, x') = \|x - x'\|.$$

We will denote the conditional distribution function of Y by :

$$\forall y \in \mathbb{R} \quad F^x(y) = \mathbb{P}(Y \leq y | X = x).$$

Then we will denote by f^x (resp $f^{x(j)}$) the conditional density (resp its j^{th} order derivative) of Y . Furthermore we will give almost complete convergence results (with rates) for nonparametric estimates of $f^{x(j)}$. Since $f^x = f^{x(0)}$, we will deduce immediately the convergence of

Some asymptotic normality result of k -Nearest Neighbour estimator of the conditional mode

conditional density estimate from the general results concerning $f^{x(j)}$.

The conditional distribution in Ferraty et al.(2006b) is defined as follows :

$$\hat{F}^x(y) = \frac{\sum_{i=1}^n K(h_K^{-1}d(x, X_i))G(h_G^{-1}(y - Y_i))}{\sum_{i=1}^n K(h_K^{-1}d(x, X_i))} \quad \forall y \in \mathbb{R}. \quad (1)$$

K is a kernel function, G is a distribution function (df) and $h = h_K := h_{K,n}$ and $h_G = h_{G,n}$ are sequence of positive real numbers which goes to zero as n goes to infinity. Under a differentiability assumption of $G(\cdot)$, we can obtain the conditional density estimation function as :

$$\hat{f}^x(y) = \frac{h_G^{-1} \sum_{i=1}^n K(h_K^{-1}d(x, X_i))G^{(1)}(h_G^{-1}(y - Y_i))}{\sum_{i=1}^n K(h_K^{-1}d(x, X_i))} \quad (2)$$

and its j^{th} partial derivative with respect to y as :

$$\hat{f}^{x(j)}(y) = \frac{h_G^{-j-1} \sum_{i=1}^n K(h_K^{-1}d(x, X_i))G^{(j)}(h_G^{-1}(y - Y_i))}{\sum_{i=1}^n K(h_K^{-1}d(x, X_i))} \quad (3)$$

$G^{(j)}$ is the j^{th} order derivative function of G . We assume that $f^x(\cdot)$ has a unique mode, denoted by $\theta(x)$ which is defined by

$$f^x(\theta(x)) = \sup_{y \in \mathcal{S}} f^x(y). \quad (4)$$

A kernel estimator of the conditional mode $\theta(x)$ is defined as the random variable $\hat{\theta}(x)$ which maximizes the kernel estimator $\hat{f}^x(\cdot)$ of $f^x(\cdot)$

$$\hat{f}^x(\hat{\theta}(x)) = \sup_{y \in \mathcal{S}} \hat{f}^x(y). \quad (5)$$

Let consider quantities

$$\hat{f}^x(y) = \frac{\hat{f}_N^x(y, h)}{\hat{F}_D^x(h)}. \quad (6)$$

Where

$$\hat{f}_N^x(y, h) = \frac{1}{n h_G \phi(h)} \sum_{i=1}^n K(h_K^{-1}d(x, X_i))G^{(1)}(h_G^{-1}(y - Y_i)), \quad (7)$$

$$\hat{F}_D^x(h) = \frac{1}{n \phi(h)} \sum_{i=1}^n K(h_K^{-1} d(x, X_i)), \quad (8)$$

where $\phi(\cdot)$ is a function supposed to be strictly positive and it will be described later.

An analogous estimator to (7) was already given in Ferraty et al. (2006a) in the general setting, remark that (7) and (8) are unbiased estimators of $a_1 f^x(y) F_D^x(h) = a_1 f_N^x(y, h)$ and $a_1 F_D^x$ respectively where a_1 and F_D^x will be described later.

The k -Nearest Neighbour estimator is a weighted average of response variables in the neighbourhood of x , this estimation method take into account the local structure of the data. The k -NN kernel estimate has a significant advantage over the classical kernel estimate. The main advantage of the k -NN method is in the nature of the smoothing parameter. Indeed, in the classical kernel method, the smoothing parameter is the bandwidth h_n , which is a real positive number and in k -NN method, the smoothing parameter k_n takes its values in discrete set.

Now we focuses on the estimation of the j^{th} order derivative of the conditional density f on x by the k nearest neighbors (k -NN), where $\hat{f}_{kNN}^{x(j)}(y)$ is defined by :

$$\hat{f}_{kNN}^{x(j)}(y) = \frac{h_G^{-j-1} \sum_{i=1}^n K(H_{n,k}^{-1} d(x, X_i)) G^{(j+1)}(h_G^{-1}(y - Y_i))}{\sum_{i=1}^n K(H_{n,k}^{-1} d(x, X_i))}. \quad (9)$$

$H_{n,k}(\cdot)$ is defined by :

$$H_{n,k}(x) = \min \left\{ h \in \mathbb{R}^+ / \sum_{i=1}^n \mathbb{I}_{B(x,h)}(X_i) = k \right\}.$$

$H_{n,k}$ is a positive random variable (r.v) which depends $(X_1, \dots, X_n)$.

From now on, when we refer to the bandwidth of the k -NN conditional mode estimation, we mean the number of neighbours k we are considering.

A natural $\hat{\theta}_{kNN}(x)$ is given by :

$$\hat{f}^x(\hat{\theta}_{kNN}) = \sup_{y \in \mathbb{R}} \hat{f}_{kNN}^x(y). \quad (10)$$

so that

$$\hat{f}_{kNN}^x(y) = \frac{\hat{f}_N^x(y, H_{n,k})}{\hat{F}_D^x(H_{n,k})}. \quad (11)$$

Note that this estimate $\hat{\theta}_{kNN}(x)$ is not necessarily unique, so the remainder of the paper concerns any value $\hat{\theta}_{kNN}(x)$ satisfying (10).

We consider two types of non-parametric models. The first one is called continuity-type, which means that the density function f is continuous and the second one, the Hölder-type, which means that the density function f is Hölder continuous with constant $b_1, b_2 > 0$.

Our main goal is to establish the asymptotic normality of the estimator in (10) when suitably normalized. As a consequence we get the asymptotic normality of a predictor and propose confidence bands for the conditional mode function with the k -NN method.

Some asymptotic normality result of k -Nearest Neighbour estimator of the conditional mode

Assume that (H1) and (H6) hold, then for any $x \in \Upsilon$, we have

$$\left(\frac{h_G^3 k (a_1 f^{x(2)}(\theta(x)))^2}{\sigma^2(x, \theta(x))} \right)^{1/2} \left(\hat{\theta}_{kNN}(x) - \theta(x) \right) \xrightarrow{\mathcal{D}} \mathcal{N}(0, 1) \text{ as } n \rightarrow \infty. \quad (12)$$

$\xrightarrow{\mathcal{D}}$ means the convergence in distribution, $\Upsilon = \{x : x \in \mathcal{S}, f^x(\theta(x)) \neq 0\}$

$$\sigma^2(x, \theta(x)) = a_2 f^x(\theta(x)) \int_{\mathbb{R}} \left(G^{(2)}(t) \right)^2 dt \quad (\text{with } a_j = K^j(1) - \int_0^1 (K^j(s))' \zeta_0(s) ds, \text{ for } j = 1, 2). \quad (13)$$

Références

- [1] Attouch, M., A.Laksaci, E.Ould-Saïd (2009). Asymptotic distribution of robust estimator for functional nonparametric models. *Comm. Statist. Theory and Methods* 38 (8), 1317-1335.
- [2] Attouch, M., t.Benchikh (2012). Asymptotic distribution of robust k -nearest neighbour estimator for functional nonparametric models. *Mathematic Vesnic* 64, No. 4, pp. 275-285.
- [3] Delsol, L. (2007a). CLT and L^q errors in nonparametric functional regression. *C. R. Math. Acad. Sci* 345 (7), 411-414.
- [4] Delsol, L. (2008b). Advances on asymptotic normality in nonparametric functional Time Series Analysis Statistics.
- [5] Ferraty, F., A.Mas, P.Vieu (2007). Nonparametric regression on functional data : Inference and practical aspects, *Australian and New Zealand Journal of Statistics* 49, 267-286.
- [6] Ferraty, F., A.Laksaci, P.Vieu (2006a). Estimating some characteristics of the conditional distribution in nonparametric functional models. *Statist. Inf. for Stoch. Proc* 9, No.2, 47-76.
- [7] Ferraty, F. and P.Vieu (2006b). Nonparametric functional data analysis. *Springer-Verlag*, New- York.
- [8] Masry, E. (2005). Nonparametric regression estimation for dependent functional data : asymptotic normality. *Stochastic Process. Appl* 115, No.1, 155-177.

Summary

In this paper, we study the nonparametric estimator of the conditional mode using the k -Nearest Neighbors (k -NN) estimation method for a scalar response variable given a random variable taking values in a semi-metric space. We establish the asymptotic normality for independent functional data, and propose confidence bands for the conditional mode function. Some simulations have been driven to show how our methodology can be implemented.

Modélisation stochastique de risque de dégradation par processus de diffusion

Kamal Boukhetala*, Nawel Khellouf**

* Faculté de Mathématiques

Bp. 32, El-Alia, Bab-Ezzouar, USTHB Alger, Algeria

* Conseil National des Assurances, Ministère des Finances

1 rue chahid Aïssa Azzi -16302- Daly Ibrahim, Alger

kboukhetala@usthb.dz

** Faculté de Mathématiques

Bp. 32, El-Alia, Bab-Ezzouar, USTHB Alger, Algeria

** Bureau Spécialisé en Tarification, Ministère des Finances

1 rue chahid Aïssa Azzi -16302- Daly Ibrahim, Alger

nawelkhellouf@gmail.com

Résumé. Ce travail a pour but d'étudier le modèle de dégradation structurelle, qui représente l'évolution temporelle de la taille d'un défaut de fissuration dans des structures physique. La description détaillée du phénomène et la définition de ses paramètres physiques sont représentées. Des solutions analytiques du processus de dégradation sont examinées. La loi de probabilité du premier instant d'affranchissement de la taille de la fissure d'un seuil de dégradation indésirable est estimée.

1 Introduction

La problématique de la fatigue des structures est issue de la révolution industrielle du XIXe siècle. Plus précisément, un certain nombre de graves accidents (ferroviaires pour la plupart) ont motivé les ingénieurs pour sur sujet. La rupture brutale d'éléments mécaniques soumis à des charges cycliques, telles que les essieux sur le matériel roulant par exemple, a été un sujet d'étude de Rankine en Angleterre dès 1843 ou Andrean en Allemagne dès 1847. A cet effet, les problèmes de dégradation constitue un souci majeur d'une entreprise d'assurance en matière de couverture de risque dynamique structurel. Dans cet article, nous proposons une étude d'un modèle de dégradation structurelle, see Sobczyk, K. (1986), Peter, E., Loeden Eckhard Platen (1992), qui représente l'évolution temporelle de la taille d'un défaut de fissuration dans certains métaux, a savoir le plastique et l'aluminium . Nous montrons, dans certains cas, que ce modèle admet une solution unique et que cette solution est explicitement analytique. Dans d'autres cas la solution exacte est difficile à déterminer , les packages Sim.Diffproc et Sim.diffprocGUI, Boukhetala, K. et Guidoum, A. (2011, 2012), sont utilisés pour effectuer de l'échantillonnage trajectoires, solutions du modèle. Une étude statistique est réalisée, permettant d'approximer la solution et certaines grandeurs d'intérêt pratique. La loi de probabilité

du processus de dégradation, à tout instant, ainsi que la loi de la variable aléatoire, premier instant d'affranchissement de la taille de la fissure d'un seuil de dégradation indésirable sont approximées par des approches non paramétriques.

2 Modèle de diffusion pour la propagation des fissures

Sobczyk, K. (1986) a proposé un modèle de dégradation qui s'appuie sur l'approximation asymptotique de processus de taille de fissure L_t à diffusion, gouvernée par une équation différentielle stochastique, interprétée au sens de Stratonovich de la manière suivante :

Pour $p > 0$; $t \geq 0$:

$$dL_t = mf(L_t)^p dt + f(L_t)^p \circ dw_t \quad (1)$$

où $(w_t)_{t \geq 0}$ est le mouvement brownien standard, $f = mcg(R)S^{2p} = \text{constant}$ et S : Facteur d'intensité de contraintes de fissuration, telle que définie par :

$$S = \frac{1}{2}(S_{max} - S_{min}),$$

m : la moyenne des influences extérieures, comme la température,

$g(R)$: fonction de ratio des contraintes, $R = \frac{S_{min}}{S_{max}}$,

c, p sont des paramètres de matériaux (c , et p sont supposés constant pour un matériel donné).

Lorsque $p = 1$ l'équation (1) est linéaire et ses solutions existent pour tout $t \geq 0$.

Si $p > 1$, l'équation (1) est une équation différentielle stochastique non linéaire, et les résultats de cette équation sont définis pour un temps infini.

L'équation (1) interprétée comme équation différentielle de Stratonovich, qui est équivalente à l'Équation Différentielle Stochastique d'Itô, (EDS), suivante :

$$dL_t = (mf(L_t)^p + \frac{p}{2}f^2(L_t)^{2p-1})dt + f(L_t)^p dw_t \quad (2)$$

3 Étude de l'équation (1) dans le cas de $p = 1$

Dans le cas de $p = 1$, l'équation (1) est linéaire, signifiant que la relation entre les charges appliquées et le taux de croissance des fissures sont presque linéaires.

Physiquement, c'est le cas de la taille d'une fissure courte d'un matériel plastique, définie par la relation :

$$dL_t = mfL_t dt + fL_t \circ dw_t \quad (3)$$

qui est équivalente à l'équation différentielle stochastique d'Itô suivante :

$$dL_t = \left(mf + \frac{1}{2}f^2 \right) L_t dt + fL_t dw_t \quad (4)$$

3.1 La solution associée à la condition initiale $L_0 > 0$

En posant $\tilde{L}(t) = \log L(t)$ et en appliquant la formule d'Itô à la fonction $\tilde{L}(t)$:

$$d \ln(L_t) = \left[\left(mf + \frac{1}{2} f^2 \right) L_t \frac{1}{L_t} + \frac{1}{2} L_t^2 \frac{-1}{dL_t^2} \right] dt + \left[f L_t \frac{1}{L_t} \right] dw_t$$

L'équation différentielle stochastique est obtenue comme suit :

$$d \ln(L_t) = m f dt + f dw_t \quad (5)$$

Ce qui résulte que $\tilde{L}_t = \ln L_t$ est le processus de taux d'accroissement de fissure, de drift $m f$ et de diffusion f^2 ; donné sous la forme intégrale suivante :

$$\begin{aligned} \ln(L_t) &= \ln(L_0) + m f \int_0^t ds + f \int_0^t dw_s \\ &= \ln(L_0) + m f t + f w_t \end{aligned}$$

D'où la solution de l'équation (3), pour $t \geq 0$ et $L_0 > 0$, a la forme la suivante :

$$L_t = L_0 \exp(m f t + f w_t) \quad (6)$$

La solution (6) est un processus de diffusion de Markov, formulée par l'équation de Fokker-Planck régissant la densité de transition $p = p(t, L, t', L')$ de L_t , où L' étant l'état initial du processus (3) :

$$\frac{\partial p}{\partial t'} = \frac{1}{2} \frac{\partial^2 [f^2 L'^2 p]}{\partial^2 L'^2} - \frac{\partial \left[\left(mf + \frac{1}{2} f^2 \right) L' p \right]}{\partial L'} \quad (t, L) \text{ fixe} \quad (7)$$

Sous l'hypothèse de la taille de fissure de fatigue L_t , considérée comme processus de diffusion, la solution de l'équation (7) est donné par :

$$p(t, L, t', L') = \frac{1}{f L' \sqrt{2\pi (t - t')}} \exp \left[\frac{\left[-\log \left(\frac{L}{L'} \right) - m f (t - t') \right]^2}{2 f^2 (t - t')} \right] \quad (8)$$

L_t suit une loi Log-Normale, dont la fonction de répartition est donnée par :

$$F(L_t) = \phi \left(\frac{\left[-\log \left(\frac{L}{L'} \right) - m f (t - t') \right]}{\sqrt{2 f^2 (t - t')}} \right), \quad (9)$$

, où ϕ est la fonction de répartition de la loi Normale.

3.2 L'instant de Premier Passage

En pratique, Les instants de premier passage jouent un rôle très important. L'objectif de cette étude est de déterminer la densité de probabilité de la variable T , l'instant de premier passage du processus de dégradation L_t .

Les modèles de l'instant de premier passage se basent essentiellement sur :

- Une fonction de dégradation stochastique L_t qui représente l'évolution d'une taille de la fissure au cours du temps ;
- Un seuil c à partir duquel la dégradation est considérée dans un état critique.

A partir de ces deux critères, le premier instant d'atteinte T_c du seuil de dégradation critique est :

$$T_c = \inf \{t \geq 0, L_t \in \delta\Delta, L_t = c\} \quad (10)$$

$\{L_t, t \geq 0\}$ processus de diffusion solution de l'équation différentielle stochastique (3), Δ une partie de $\mathbb{R}$ supposé ouverte et bornée et c un élément de Δ .

L'équation (3) est utilisée pour déterminer explicitement la loi de probabilité de T_c .

La probabilité de transition $P(c/a)$ représente la densité de L_t avec $L_0 = a$, qui est solution de l'équation (7). Ce qui donne :

$$p(c, t_c/a, t_0) = \frac{1}{f c \sqrt{2\pi (t_c - t_0)}} \exp \frac{\left[-\log \left(\frac{c}{a} \right) - m f (t_c - t_0) \right]^2}{2 f^2 (t_c - t_0)} \quad (11)$$

Où la variable aléatoire T_c est définie par :

$$T_c = \inf \{t \geq 0, L_t \leq c, L_0 = a\} \quad (12)$$

Pour des petites valeurs de c , on a :

$$\begin{aligned} R(t) &= P(T_c > t) \\ &= P(L_t \leq c) \\ &= 1 - F_{T_c}(t) \end{aligned}$$

Donc :

$$\begin{aligned} P(T_c \leq t) &= 1 - P(L_t \leq c) \\ &= 1 - \phi \left(\frac{\left[-\log \left(\frac{c}{a} \right) - m f t \right]}{2 f^2 (t_c - t_0)} \right) \end{aligned}$$

La distribution de probabilité de l'instant de premier de passage est Gaussienne Inverse, et représentée par la fonction de densité suivante :

$$f_{T_c}(t) = \frac{\log \left(\frac{c}{a} \right)}{f \sqrt{2\pi t^{\frac{3}{2}}}} \exp \left(\frac{\left[-\log \left(\frac{c}{a} \right) - m f t \right]^2}{2 f^2 t} \right). \quad (13)$$

En pratique, il est connu que la loi de la durée de vie d'une structure, dont le comportement réel est observé par des données, suit approximativement l'une de ces lois : Exponentielle, Lognormale ou Weibull.

4 Conclusion

Dans ce travail, nous avons présenté une famille de modèles de processus de dégradation, construite par des équations différentielles stochastiques de type Stratonovich. Nous avons également développé quelques éléments théoriques fondamentaux du modèle, utiles pour la modélisation du processus de dégradation et l'étude de la fiabilité des structures physiques en question. Ainsi, à travers quelques exemples pratiques, nous avons montré l'importance et l'intérêt de la simulation, ainsi que l'adéquation de ce modèle avec les phénomènes réels de dégradation. L'échantillonnage trajectoires généré par le modèle nous a permis d'estimer statistiquement certaines grandeurs d'intérêt physique, liées à la fiabilité de la structure, en matière de durée de vie probable et de résistance aux effets internes et externes sur la structure.

Référence

- Akira Tsurui and Hiroshi Ishikawa(1986), Application of the Fokker-Planck equation to a stochastic fatigue crack growth model, Structural Safety.
- Boukhetala, K. et Guidoum A (2012), Sim.DiffProc : Simulation of Diffusion Processes, R package version 2.5, <http://CRAN.R-project.org/package=Sim.DiffProc>.
- Douglas, H. and Peter, P. (2006), Stochastic Differential Equations in Science and Engineering, World Scientific Publishing.
- Peter, E and Eckhard, P. (1995), Numerical Solution of Stochastic Differential Equations (Second Edition), Imperial College Press.
- Risken, H. (2001), The Fokker-Planck Equation : Methods of Solutions and Applications, 2nd edition, Springer Series in Synergetics.
- Saito, Y. and Mitsui, T. (1993), Simulation of Stochastic Differential Equations, The Annals of the Institute of Statistical Mathematics.
- Sobczyk, K. (1986), Modelling of random fatigue crack growth, *Engineering Fracture Mechanics*.

Summary

This work is to study a model of structural degradation, which represents the temporal evolution of structures cracking. A detailed description of phenomenon and the definition of its physical parameters are presented. Analytical solutions of the process $L(t)$, the length of dominant crack, are discussed. We put the associated reliability framework by considering the failure of the structure once the degradation process reaches a critical threshold.

Analyse bayésienne de la file d'attente M/M/1

Hayette Braham*, Louiza Berdjoudj**
Mohamed Boualem**

*Département de Mathématiques et Informatique, Université de Bejaia, 06000 Algérie
braham_hayet@hotmail.fr

**Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes)
Université de Bejaia, 06000 Algérie
l_berdjoudj@yahoo.fr, robertt15dz@yahoo.fr
<http://www.lamos.org>

Résumé. L'objectif de ce travail est de faire une analyse bayésienne de la file d'attente M/M/1 en considérant plusieurs cas dans l'estimation des mesures de performances de ce système, à savoir le nombre moyen de clients dans le système et le temps moyen de séjour d'un client dans le système, et ce pour voir quel paramètre influe le plus sur le système et sur l'approche bayésienne. Finalement un exemple numérique, en utilisant la simulation est présenté pour illustrer les résultats et une généralisation à la file d'attente M/G/1 est en cours de réalisation.

1 Introduction

La théorie des files d'attente est une technique de la recherche opérationnelle qui permet de modéliser un système admettant un phénomène d'attente, de calculer ses performances et de déterminer ses caractéristiques pour aider les gestionnaires dans leurs prises de décisions. Pour l'application de la théorie des files d'attente à l'évaluation des performances, il est apparu que même les modèles de file d'attente les plus simples fournissaient des résultats qui correspondaient de près aux observations réelles. Dans ce papier on se focalise dans un premier temps sur le système d'attente le plus simple M/M/1 qui est caractérisé par les deux processus d'entrée principaux qui sont le processus des arrivées et le processus de service (Armero et Bayarri (1994)). Lorsque la file est dans son état stationnaire les paramètres les plus intéressants ne sont pas les flots d'arrivées et de service mais plutôt les mesures de performances à savoir le nombre moyen de clients dans le système, le temps moyen de séjour d'un client dans le système et dans la file (Boualem et al. (2013)). Notons que toutes ces quantités sont observables et on peut les estimer en utilisant différentes méthodes.

L'analyse bayésienne des systèmes de file d'attente remonte aux travaux de Carly 1970, le lecteur peut se référer à Bagchi et Cunningham (1972), Muddapur (1972), Reynolds (1973), Armero (1985) et McGrath et Singpurwalla (1987). Récemment le nombre de travaux sur les techniques bayésiennes appliquées aux systèmes de files d'attente a remarquablement crû, citons Armero et Bayarri (1994, 1995, 1997), Armero et Conesa (1998), Wiper (1998) et Rios et al. (1998).

Notre objectif dans ce travail est d'utiliser les méthodes bayésiennes pour estimer les caractéristiques de la file M/M/1 en considérant plusieurs cas sur les paramètres du système et de les comparer à celles obtenues avec les méthodes classiques. Ces résultats seront généralisés sur la file d'attente M/G/1 par la suite.

2 Le Système M/M/1

Considérons le système de file d'attente M/M/1, le flot des arrivées est poissonien de taux λ et les durées de service sont indépendantes et exponentiellement distribuées de paramètre μ . On note par T "le temps d'inter-arrivées", donc $T \rightsquigarrow \mathcal{E}(\lambda)$ et par S "le temps de service", donc $S \rightsquigarrow \mathcal{E}(\mu)$. L'intensité de trafic ou le coefficient d'utilisation du système est noté par : $\rho = \frac{\lambda}{\mu}$. La variable aléatoire X ="nombre de clients dans le système au régime stationnaire" suit une loi géométrique de paramètre $(1 - \frac{\lambda}{\mu})$ et on a (Sanku (2008)) :

$$\mathbb{P}(X = x) = (1 - \frac{\lambda}{\mu})(\frac{\lambda}{\mu})^x, \quad x = 0, 1, 2, \dots \quad (1)$$

Les expressions des caractéristiques de la file M/M/1 que nous analyserons par la suite sont données par :

- Nombre moyen de clients dans le système : $L = \frac{\lambda}{\mu - \lambda} = \frac{\rho}{1 - \rho}$.
- Temps moyen de séjour d'un client dans le système : $W = \frac{1}{\mu - \lambda}$.

3 Estimation bayésienne des caractéristiques du système M/M/1

On observe n_α temps d'inter-arrivées $t = \{t_1, \dots, t_{n_\alpha}\}$ et n_β temps de service $s = \{s_1, \dots, s_{n_\beta}\}$ et on suppose que les paramètres λ , μ et ρ sont inconnus et on leur associe les lois a priori suivantes : $\lambda \rightsquigarrow \gamma(a_1, b_1)$, $\mu \rightsquigarrow \gamma(a_2, b_2)$ et $\rho \rightsquigarrow \beta(a_3, b_3)$, où $a_i > 0$, $b_i > 0$, $i = 1, 2, 3$.

Après le calcul des lois a posteriori on obtient les estimateurs de la moyenne a posteriori suivants :

$$\hat{\lambda}_B = \frac{a_1 + n_\alpha}{b_1 + \sum_{j=1}^{n_\alpha} T_j} \quad (2)$$

$$\hat{\mu}_B = \frac{a_2 + n_\beta}{b_2 + \sum_{j=1}^{n_\beta} S_j} \quad (3)$$

$$\hat{\rho}_B = \frac{a_3 + \sum_{j=1}^n X_j}{n + a_3 + b_3 + \sum_{j=1}^n X_j} \quad (4)$$

Pour estimer L et W on considère les 4 cas suivants :

1. Le paramètre inconnu est λ ($L = \frac{\hat{\lambda}_B}{\mu - \hat{\lambda}_B}$ et $W = \frac{1}{\mu - \hat{\lambda}_B}$),
2. Le paramètre inconnu est μ ($L = \frac{\lambda}{\hat{\mu}_B - \lambda}$ et $W = \frac{1}{\hat{\mu}_B - \lambda}$),
3. Les paramètres inconnus sont λ et μ ($L = \frac{\hat{\lambda}_B}{\hat{\mu}_B - \hat{\lambda}_B}$ et $W = \frac{1}{\hat{\mu}_B - \hat{\lambda}_B}$),
4. Le paramètre inconnu est ρ ($L = \frac{\hat{\rho}_B}{1 - \hat{\rho}_B}$).

4 Simulation

Afin de voir la différence entre les 4 cas considérés et de faire une comparaison de l'estimation bayésienne avec celle de l'inférence classique et les valeurs réelles des paramètres, on a élaboré un simulateur sous environnement MatLab, qui simule des vecteurs aléatoires suivants les différentes lois considérées précédemment et en faisant varier leur taille pour voir son influence à son tour, et calcule les valeurs des paramètres estimés. Les résultats obtenus nous permettent de décider lequel des quatre cas considérés est meilleur et aussi de voir que le choix des paramètres des lois a priori est très crucial.

5 La file d'attente M/G/1

Dans le but de généraliser les résultats obtenus pour la file M/M/1 on considère la file M/G/1 où le service n'est plus exponentiel mais d'une loi générale inconnue, sa distribution est caractérisée par ses moments ou bien sa transformée de Laplace (Mohamedi et al. (2013)). On considère des cas particuliers de lois générales (Cox d'ordre 2,...). Une simulation pour ce système est en cours de réalisation en prenant en considération les mêmes cas que ceux pris pour le système M/M/1.

Références

- Armero, C. et M. J. Bayarri (1994). Bayesian prediction in M/M/1 queue. *Queueing Systems* 15, 401–417.
- Boualem, M., M. Cherfaoui, N. Djellab, et D. Aïssani (2013). Analyse des performances du système M/G/1 avec rappels et feedback. *Journal Européen des Systèmes Automatisés* 1-2-3, 181–193.
- Mohamedi, A., M. R. Salehi-Rad, et E. C. Wit (2013). Using mixture of gamma distributions for bayesian analysis in an M/G/1 queue with optional second service. *Comput Stat* 28, 683–700.
- Sanku, D. (2008). A single server feedback retrial queue with collisions. *Data Sciences Journal* 7.

Summary

This paper deals with the bayesian analysis of the M/M/1 queue by considering several cases in estimating of the performance measures of the system, namely the average number of customers in the system and the mean waiting time in the system, in order to see what parameter affected more the system and the bayesian approach. Finally, a numerical example using simulation is given to illustrate the results and the generalization for an M/G/1 queue is in progress.

Tests d'hypothèses dans un processus de diffusion non linéaire à retard

Malika DALI-KORSO FECIANE*,

*Laboratoire de Statistiques et Modélisations Aléatoires LASMA,
Département de Mathématiques, Université de Tlemcen,
B.P 119, 13000, Algérie
korsom@yahoo.fr

Résumé. Le problème de tests d'hypothèses est l'un des thèmes centraux de la théorie des statistiques et ses applications. Il est important de vérifier le degré d'adéquation d'un modèle statistique aux observations. Nous présentons un certain nombre de résultats sur la construction de tests d'hypothèses simples et composites pour des processus à retard et à temps continu. Comme modèle nous considérons une équation différentielle stochastique avec petit bruit. Nous proposons des tests de niveau α donné, des tests de niveau asymptotique α et nous discutons l'asymptotique de la fonction puissance sous les alternatives locales.

1 Introduction

Les tests d'hypothèses permettent de vérifier si le modèle mathématique proposé est bien adapté aux observations. Notre but est de présenter certains tests pour un système dynamique avec petit bruit ($\varepsilon \rightarrow 0$). Supposons que le processus observé est la solution de l'équation différentielle stochastique

$$dX_t = S(X_{t-\theta})dt + \varepsilon dW_t, \quad 0 \leq t \leq T, \quad X_s = x_0, \quad s \leq 0,$$

où le paramètre inconnu θ peut prendre différentes valeurs dans $\Theta = (\alpha, \beta)$, $0 < \alpha < \beta < T$. $\{W_t, 0 \leq t \leq T\}$ est un processus de Wiener. L'inférence statistique concerne le coefficient de la dérive uniquement. Nous considérons le problème avec une hypothèse simple contre une alternative simple. Nous présentons certains tests classiques ainsi que le calcul de leur puissance. Dans la seconde section le problème abordé est celui d'alternatives composées unilatérales. Il existe des voisinages Θ_ε de $\theta_0 = 0$ tels que, si $\theta_\varepsilon \in \Theta_\varepsilon$, la fonction puissance $\beta_\varepsilon(\phi_\varepsilon, \theta_\varepsilon)$ admet une limite non dégénérée (i.e. $\neq \alpha, \neq 1$). L'approche envisagée dans ce cas est due à Pitman.

2 Hypothèse simple contre alternative simple

Nous considérons la diffusion

$$dX_t = S(X_{t-\theta})dt + \varepsilon dW_t, \quad 0 \leq t \leq T, \quad X_s = x_0, \quad s \leq 0, \quad (1)$$

Tests d'hypothèses

et nous proposons des tests pour une hypothèse simple $\mathcal{H}_0 : \theta = 0$ (pas de retard) contre une alternative également simple $\mathcal{H}_1 : \theta = \theta_1$ (il y a un retard) sur la base de l'observation d'une trajectoire complète $X = \{X_t, 0 \leq t \leq T\}$.

La dérive $S(\cdot)$ est connue au paramètre θ près et vérifie l'hypothèse suivante

H : S est une fonction positive deux fois continûment dérivable de dérivée non nulle sur un intervalle $[a, b] \subset [x_0, x_T]$ et bornée sur $\mathbb{R}$. Ici $x_\cdot$ est solution de l'équation déterministe

$$\frac{dx_t}{dt} = S(x_{t-\theta}), \quad 0 \leq t \leq T, \quad x_s = x_0, \quad s \leq 0. \quad (2)$$

Nous précisons tout d'abord quelques définitions et notations :

$P_0^{(\varepsilon)}$ (respectivement $P_1^{(\varepsilon)}$) est la mesure de probabilité induite par le processus sous l'hypothèse de base (respectivement l'alternative). $\mathbb{E}_0$ (respectivement $\mathbb{E}_1$) est l'espérance mathématique sous $P_0^{(\varepsilon)}$ (respectivement $P_1^{(\varepsilon)}$).

Notons x^0 la solution de l'équation différentielle ordinaire sans retard ($\theta = 0$), et x^1 celle de la même équation mais avec un retard ($\theta = \theta_1$). Enfin posons :

$$I_0 = \int_0^T (S(x_t^0) - S(x_{t-\theta_1}^0))^2 dt, \quad I_1 = \int_0^T (S(x_t^1) - S(x_{t-\theta_1}^1))^2 dt.$$

Pour α fixé dans l'intervalle $(0, 1)$, $\phi_\varepsilon(X)$ est la probabilité d'accepter l'alternative $\mathcal{H}_1$ à partir de l'observation : $X = \{X_t, 0 \leq t \leq T\}$ et $\mathcal{K}_\alpha$ (respectivement $\mathcal{K}'_\alpha$) est la classe des tests au seuil α (respectivement au seuil asymptotique α) :

$$\mathcal{K}_\alpha = \{\phi_\varepsilon : \mathbb{E}_0 \phi_\varepsilon(X) \leq \alpha\} \quad \mathcal{K}'_\alpha = \{\phi_\varepsilon : \limsup_{\varepsilon \rightarrow 0} \mathbb{E}_0 \phi_\varepsilon(X) \leq \alpha\}.$$

La puissance du test ϕ_ε est notée $\beta_\varepsilon(\phi_\varepsilon)$ et représente la probabilité de la bonne décision sous l'hypothèse $\mathcal{H}_1$: $\beta_\varepsilon(\phi_\varepsilon) = \mathbb{E}_{\theta_1} \phi_\varepsilon(X)$. Un test ϕ_ε est le plus puissant dans la classe $\mathcal{K}_\alpha$ si pour tout test $\psi_\varepsilon \in \mathcal{K}_\alpha$ nous avons $\beta_\varepsilon(\phi_\varepsilon) \geq \beta_\varepsilon(\psi_\varepsilon)$.

Notre premier problème est de tester $\mathcal{H}_0 : \theta = 0$ (pas de retard) contre $\mathcal{H}_1 : \theta = \theta_1$ moyennant les tests classiques qui suivent.

Test de niveau asymptotique α

Considérons la statistique $\Delta_\varepsilon(X) = \int_0^T \frac{\Delta S^\varepsilon(t)}{\varepsilon} [dX_t - S(X_t)dt]$ et le test $\phi(X) = 1_{\{\Delta_\varepsilon(X) \geq c_{\alpha, \varepsilon}\}}$. Nous établissons la consistance de ce test contre toute alternative fixée dans la classe des tests de niveau asymptotique α avec $c_\alpha = \sqrt{I_0} z_\alpha$.

Test de niveau α

Nous proposons un test de niveau α pour chaque ε fixé dans $(0, 1)$ à la différence du test précédent. Pour $t \in [T, T+1]$ posons $\Delta S^\varepsilon(t) := \sqrt{I_1}$, et remarquons toutefois que les observations X_t , sont en fait définies sur $[0, T]$ uniquement. Comme Kutoyants (2004) introduisons le temps d'arrêt

$$\tau_\varepsilon = \inf \left\{ t \in [0, T+1] : \int_0^t (\Delta S^\varepsilon(t))^2 dt \geq I_1 \right\}$$

et considérons alors la statistique

$$\Delta_\tau(\theta_1, 0, X) := \int_0^{\tau_\varepsilon} \frac{\Delta S^\varepsilon(t)}{\varepsilon} [dX_t - S(X_t)dt]$$

avec

$$dX_t - S(X_t)dt = \varepsilon dW_t, \text{ pour } t \in [T, T+1].$$

Ainsi, si $\tau_\varepsilon > T$ et sous $\mathcal{H}_0$, la statistique Δ_τ est définie par

$$\Delta_\tau(\theta_1, 0, X) = \int_0^{\tau_\varepsilon} \frac{\Delta S^\varepsilon(t)}{\varepsilon} [dX_t - S(X_t)dt] = \int_0^{\tau_\varepsilon} \Delta S^\varepsilon(t) dW_t.$$

Si z_α est le quantile de la loi $\mathcal{N}(0, 1)$. et pour tout ε fixé dans $(0, 1)$, le test $\phi_\varepsilon^*(X)$ est défini moyennant la statistique Δ_τ par $\phi_\varepsilon^*(X) = \chi_{\{\Delta_\tau(\theta_1, 0, X) \geq z_\alpha \sqrt{T_1}\}}$. Nous obtenons alors que Le test ϕ_ε^* ainsi défini est dans la classe $\mathcal{K}_\alpha$, et il est consistant contre toute alternative fixée.

3 Hypothèse simple contre alternative composée

Nous considérons le même système dynamique que dans la section précédente et nous nous proposons de tester l'hypothèse de base simple $\mathcal{H}_0 : \theta = 0$ (pas de retard) contre une alternative composée unilatérale $\mathcal{H}_1 : \theta > 0$ sur la base de l'observation d'une trajectoire $X = \{X_t, 0 \leq t \leq T\}$.

Les alternatives contigües correspondent à $\theta = \varepsilon u$. En effet remarquons que le rapport de vraisemblance $Z_\varepsilon(u) = L(\varepsilon u, 0, X)$ a une limite non dégénérée car, la condition **H** étant satisfaite, les mesures induites $P_0^{(\varepsilon)}$ et $P_1^{(\varepsilon)}$, respectivement sous $\mathcal{H}_0$ et $\mathcal{H}_1$, par le processus solution sont équivalentes (voir Lipster et Shiriyayev (2000)) donc contigües. Nous pouvons réécrire le problème de test d'hypothèses :

$\mathcal{H}_0 : u = 0$ (pas de retard) contre $\mathcal{H}_1 : u > 0$,

et nous adoptons une approche due à Pitman pour les alternatives contigües.

Nous introduisons la quantité

$$I(\theta) = \int_0^T S^2(x_{t-2\theta}) S'^2(x_{t-\theta}) dt$$

qui représente l'information de Fisher pour ce type de problème. Enfin nous notons ξ une variable aléatoire de loi $\mathcal{N}(0, 1)$ et z_α son quantile d'ordre α . Un test ϕ_ε est dit localement asymptotiquement uniformément le plus puissant dans la classe $\mathcal{K}'_\alpha$ si pour tout autre test $\psi_\varepsilon \in \mathcal{K}'_\alpha$ nous avons

$$\lim_{\varepsilon \rightarrow 0} \inf_{0 \leq u \leq K} [\beta_\varepsilon(u, \phi_\varepsilon) - \beta_\varepsilon(u, \psi_\varepsilon)] \geq 0.$$

Test de niveau α

Sous $\mathcal{H}_0$, pour tout $t \in [T, T+1]$ posons

$$dX_t - S(X_t)dt = \varepsilon dW_t \quad \text{et} \quad S'(X_t)S(X_t) = S'(x_t)S(x_t)$$

Tests d'hypothèses

où $\{W_t\}$ est un Wiener indépendant du processus $\{X_t, 0 \leq t \leq T\}$.
Introduisons le temps d'arrêt

$$\tau_\varepsilon := \inf \left\{ t \in [0, T+1] : \int_0^t (S'(X_t))^2 (S(X_t))^2 dt \geq I(0) \right\}$$

$\tau_\varepsilon \leq T+1$ p.s.

Considérons alors la statistique $\Delta_\tau(\varepsilon, 0, X) = \int_0^{\tau_\varepsilon} \frac{S'(X_t)S(X_t)}{\varepsilon} [dX_t - S(X_t)dt]$

et le test $\phi_\varepsilon^* = \chi_{\{\Delta_\tau(\varepsilon, 0, X) \geq c_{\alpha, \varepsilon}\}}$.

Le test ϕ_ε^* est, avec $c_\alpha = \sqrt{I(0)}z_\alpha$ de niveau α et uniformément asymptotiquement le plus puissant ; sa puissance asymptotique est

$$\beta_\alpha^* = P\{\xi \geq z_\alpha - \sqrt{I(0)}u\}.$$

Test du rapport de vraisemblance

Nous considérons le processus du rapport de vraisemblance pour les deux valeurs du paramètre 0 et εv

$$Z_\varepsilon(v) = L(\varepsilon v, 0, X) \quad v \in U_\varepsilon = \{u \in \mathbb{R}_+, 0 \leq \varepsilon v \leq B\}$$

où B est un réel positif. Ce processus est défini par : Nous introduisons alors la statistique $\delta_\varepsilon(X) = \sup_{v>0} Z_\varepsilon(v)$ et le test associé $\Psi_\varepsilon(X) = \chi_{\{\delta_\varepsilon > c_\alpha\}}$.. le test Ψ_ε est de niveau asymptotique α . Pour le calcul de la fonction puissance nous introduisons le rapport de vraisemblance aux points $\varepsilon u, \varepsilon u + \varepsilon v$: $Z_\varepsilon^*(v)$ et remarquons alors que le rapport de vraisemblance $L(\varepsilon v, 0, X)$ admet, asymptotiquement la représentation

$$L(\varepsilon v, 0, X) = Z_\varepsilon(u)Z_\varepsilon^*(v)$$

avec

$$\begin{aligned} Z_\varepsilon(u) &= \exp\{u\Delta_\varepsilon(0, X) + \frac{1}{2}u^2I(0) + \Psi_\varepsilon(u, 0, X)\} \\ Z_\varepsilon^*(v) &= \exp\{v\Delta_\varepsilon(0, X) - \frac{1}{2}v^2I(0) + \Psi_\varepsilon^*(v, 0, X)\} \end{aligned}$$

Nous pouvons alors établir que le test Ψ_ε est, avec $c_\alpha = \exp(z_\alpha^2/2)$ de niveau asymptotique α et localement asymptotiquement uniformément le plus puissant et

$$\beta_\alpha = P\{\xi \geq z_\alpha - \sqrt{I(0)}v\}.$$

Test de l'estimateur du maximum de vraisemblance

L'estimateur du maximum de vraisemblance $\hat{\theta}_\varepsilon$ est défini pour ce problème par l'équation suivante

$$L(\hat{\theta}_\varepsilon, \theta_0; X^\varepsilon) = \sup_{\theta \in \Theta} L(\theta, \theta_0; X^\varepsilon) = \sup_{\theta \in \Theta} \frac{dP_\theta^{(\varepsilon)}}{dP_{\theta_0}^{(\varepsilon)}}(X)$$

Introduisons la statistique $\eta_\varepsilon(X) := \frac{\hat{\theta}_\varepsilon}{\varepsilon}$ et le test $\hat{\Psi}_\varepsilon := \chi_{\{\eta_\varepsilon(X) \geq c_{\alpha, \varepsilon}\}}$ et rappelons (voir thèse Korsó Féciàne) que uniformément en θ sur tout compact de Θ l'estimateur du maximum

de vraisemblance est asymptotiquement normal. Moyennant en particulier cette convergence une fois sous l'hypothèse de base ($u = 0$) puis sous l'alternative ($u > 0$), nous pouvons établir l'optimalité locale du test $\hat{\Psi}_\varepsilon$, avec $c_\alpha = z_\alpha I(0)^{-1/2}$ et évaluer sa puissance asymptotique est

$$\beta_\alpha = P \left\{ \xi \geq I(0)^{1/2} c_\alpha - I(0)^{1/2} u \right\}.$$

Références

Durbin, J.(1973). *Distribution Theory for Tests Based on the Sample Distribution Function*. Conference Board of the Mathematical sciences Regional conference series in Applied mathematics, No.9 Society for Industrial and Applied Mathematics, Philadelphia, Pa.,.

Kutoyants, Yu. A.(2004). *Statistical Inference for Ergodic Diffusion Processes*. London : Springer-Verlag.

Kutoyants, YU.A(1994). *Identification Of Dynamical Systems With Small Noise*. Kluwer :Dordrecht

Lehmann, E.L. and Romano,(2005). *Testing Statistical Hypotheses*. (3rd ed.) New York :Springer.

Rabhi, A.(2008). *On the goodness-of-fit testing of composite hypothesis for dynamical systems with small noise*. prepublication 08–6, Université du Maine, ([http ://www.univ-lemans.fr/sciences/statist/download/Rabhi/GoFpaper.pdf](http://www.univ-lemans.fr/sciences/statist/download/Rabhi/GoFpaper.pdf)).

Summary

The problem of testing is one of the central themes of statistical theory and practice. It is important to verify the degree of correspondence between observed outcomes and expected outcomes that is the foundation of classical statistics. We present a number of results on the construction of classical simple and composite hypothesis testing for time-delay and continuous processes. A stochastic differential equation with small noise is considered as model. We consider the case of hypotheses testing with the basic simple hypothesis: there is no delay. We propose tests with a given level or asymptotic level α . The alternatives are either simple or composite. For the composite ones we are interesting with one-sided parametric and one-sided non-parametric composite alternatives. And so we construct locally asymptotically uniformly most powerful tests.

Asymptotic normality of the kernel estimator of the conditional density under association for censored data

Samra Dhiabi*, Ourida Sadki**

*Univ. Mohamed Khider, Biskra
samradhiabi@yahoo.fr,
<http://www.univ-biskra.dz/>

**USTHB, Alger, <http://www.usthb.dz/>
osadki@usthb.dz

** Univ. Larbi Ben M'hidi Oum El Bouaghi <http://www.univ-oeb.dz/>

Résumé. In this paper we study a smooth estimator of the conditional density function in the censorship model, when the data exhibit some dependence structure. We show that, under some regularity conditions, that the kernel estimator of the conditional density function suitably normalized is asymptotically normally distributed

1 Introduction

Let $T_1, \dots, T_n$ be a sequence of the survival times of n individuals in a life table. These random variables (r.v.s) are not assumed to be mutually independent but are positively-associated (P.A) and strictly stationary with common unknown absolutely distribution function (df) F . In many situations, we observe only censored lifetimes of items under study. That is, assuming that $\{C_i; i \geq 1\}$ is a sequence of independent censoring times with common unknown df G , we observe only the n pairs $\{(Y_i, \delta_i), i = 1, 2, \dots, n\}$, with $Y_i = T_i \wedge C_i$ and $\delta_i = I_{\{T_i \leq C_i\}}$ where $\wedge$ denotes minimum and $I_{\{\cdot\}}$ is the indicator r.v. of the specified event. We will suppose that T and C are independent to ensures the identifiability of the model.

Let $X_1, \dots, X_n$ be a stationary sequence of real-valued r.v.s, and $F(\cdot | \cdot)$ be the conditional df of T given $X = x$, that is,

$$F(t | x) = E[1_{\{T \leq t\}} | X = x]$$

wich can be written as $F(t | x) =: \frac{F_1(t, x)}{l(x)}$, where l is the marginal density of X with respect to Lebesgue measure and

$$f(t | x) =: \frac{F'_1(t, x)}{l(x)}$$

where $F'_1(\cdot, \cdot)$ is the joint probability density function of (T, X) .

The definition of the underlying dependence considered here is as follows

The random variables $(X_1, X_2, \dots, X_n)$ is said to be *positively-associated (PA)*, if for every $g_1, g_2 : R^n \rightarrow R$, which are coordinatewise nondecreasing, and for which $Eg_1^2(X_j, 1 \leq j \leq n) < \infty, Eg_2^2(X_j, 1 \leq j \leq n) < \infty$, it hold that

$$\text{cov}(g_1(X), g_2(X)) \geq 0.$$

The above r.v.s. are said to be *negatively-associated (NA)*, if for every nonempty proper subset A of $\{1, 2, \dots, n\}$ and for every $g_1, g_2 : R^{\#A} \rightarrow R$, which are nondecreasing as above, and for which $Eg_1^2(X_j, j \in A) < \infty, Eg_2^2(X_j, j \in A^c) < \infty$, it holds that :

$$\text{cov}(g_1(X_i, i \in A), g_2(X_j, j \in A^c)) \leq 0;$$

here, $\#A$ denotes the cardinality of A .

Association has found extensive applications in systems reliability and various problems in statistical mechanics (see, for example, Esary et al. (1967)). Roussas (2001) studied the normality of the kernel estimator of the conditional density for associated complete data. Properties of associated random sequences are given in Bulinski et al.(2007).

In this paper, we study the large sample properties of the conditional density estimator of $f(t|x)$ for the case in which the underlying failure times and the covariate are assumed to be positively or negatively associated and under the condition that C and (X, T) are independent.

2 Definitions and results

It is well known that a linear estimator of the conditional df $F(t|x)$ is given by (see Carbonez, Györfi and Van der Meulin (1995), Kohler, Mathé and Pinter (2002) and Ould-Said (2006))

$$\tilde{F}_n(t|x) = \sum_{i=1}^n W_{i,n}(x) \frac{\delta_i I_{\{Y_i \leq t\}}}{\bar{G}(Y_i)} \quad (1)$$

where $W_{i,n}(\cdot)$ are the well-known Nadaraya-Watson weights which are measurable functions of x and depending on the sample $X_1, X_2, \dots, X_n$ and $\bar{G}(\cdot) := 1 - G(\cdot)$.

In practice G is usually be unknown, hence it impossible to use the estimator (1). Then, we replace G by the Kaplan-Meier estimator G_n given by

$$1 - G_n(t) = \begin{cases} \prod_{i=1}^n \left(1 - \frac{1 - \delta_{(i)}}{n - i + 1}\right)^{I_{\{Y_{(i)} \leq t\}}} & \text{if } t < Y_{(n)} \\ 0 & \text{if } t \geq Y_{(n)} \end{cases}$$

Therefore the feasible estimator of the conditional df $F(t|x)$ is given by

$$F_n(t/x) = \frac{\sum_{i=1}^n \frac{\delta_i}{\bar{G}_n(Y_i)} K\left(\frac{x - X_i}{h}\right) H\left(\frac{t - Y_i}{h}\right)}{\sum_{i=1}^n K\left(\frac{x - X_i}{h}\right)} =: \frac{F_{1,n}(t, x)}{l_n(x)},$$

with the convention $0/0=0$ and where K is a probability density function (so called kernel function), H is a distribution function, $h_n =: h$ is a sequence (so called bandwidth) of positive real

numbers which goes to zero as n goes to infinity and $l_n(\cdot)$ is the well-known kernel estimator of $l(\cdot)$.

The conditional density estimator is given by

$$f_n(t|x) = \frac{F'_{1,n}(t, x)}{l_n(x)},$$

where

$$F'_{1,n}(x, t) = \frac{1}{nh} \sum_{i=1}^n \frac{\delta_i}{\bar{G}_n(Y_i)} K\left(\frac{x - X_i}{h}\right) H'\left(\frac{t - Y_i}{h}\right).$$

Somme properties of $f_n(t|x)$ are studied by Ould Saïd & Sadki (2008) and (2011) in the i.i.d. case and the α -mixing case respectively.

Under some classical regularity assumptions on K , H , and h , and a assumption on the dependence structure of the data, more precisely

$$u(n) = \sup_{n \geq 1} \sum_{i=1}^n \sum_{j=1}^n \theta_{|i-j|}^\gamma < \infty,$$

where $0 < \gamma < 1$, and $\theta_{|i-j|} := |\text{cov}(X_i, X_j)|$, we establish the following results

PROPOSITION 1 *for n large enough, we have*

$$(nh^2)^{\frac{1}{2}} (f_n(t|x) - f(t|x)) \longrightarrow^D N(0, \sigma^2(t, x))$$

where $\longrightarrow^D$ denotes the convergence in distribution, and

$$\sigma^2(t, x) = \frac{1}{l^2(x)} \left[\frac{\kappa \sigma F'_1(t, x)}{\bar{G}(t)} \right]$$

where κ and σ are specified constants.

Références

- [1] BULINSKI, A. and SHASHKIN, A. (2007). Limit theorems for associated random fields and related systems. Vol. 10. Advanced series on statistical science & applied probability.
- [2] CARBONEZ, A., GYÖRFI, L. and VAN DER MEULIN, E.C. (1995). Partitioning estimates of a regression function under random censoring, *Statist. Decisions* **13** 21–37
- [3] ESARY, J., PROSCHAIN, F., and WALKUP, D. (1967). Association of random variables with applications Ann Maths. Statist. **38**, 1466-1474.

- [4] KAPLAN, E. M. and MEIER, P. Non parametric estimation from incomplete observations, *J. Amer. Statistic. Assoc.* **53** (1958), 457-481.
- [5] KOHLER, M., MÁTHÉ, P. and PINTÉR, M. (2002). Prediction from randomly right censored data, *J. Multivariate Anal.* **80** 73–100.
- [6] OULD SAÏD, E. and SADKI, O. (2011). Asymptotic normality for a smooth kernel estimator of the conditional quantile for censored time series, *South African Statistical Journal.* **45** 65–98.
- [7] OULD SAÏD, E. and SADKI, O. (2008). Prediction via the conditional quantile for right censorship model, *Far East J. Theoretical Statist.* **25** 145–179.
- [8] OULD-SAÏD, E. (2006). A strong uniform convergence rate of kernel conditional quantile estimator under random censorship, *Statist. Probab. Lett.* **76**, 579-586.
- [9] ROUSSAS, G. G. (2000). Asymptotic normality of the kernel estimate of a probability density function under association. *Statis. Probab. Lett.*, **50**(1), 1-12.

Summary

In this paper we study a smooth estimator of the conditional density function in the censorship model, when the data exhibit some dependence structure. We show that, under some regularity conditions, that the kernel estimator of the conditional density function suitably normalized is asymptotically normally distributed

Simulation de l'estimateur à noyau de la fonction de régression dans un modèle de troncature à gauche

Farida Hamrani*, Zohra Guessoum**

*Faculté Mathématique USTHB
BP 32, El-Alia 16111, Algiers, Algeria,
fhamrani@usthb.dz

**Laboratoire MSTD, Faculté Mathématique, USTHB
BP 32, El-Alia 16111, Algiers, Algeria,
zguessoum@usthb.dz

Résumé. Nous nous intéressons dans ce travail de simulation à calculer l'estimateur à noyau de la fonction régression pour un modèle tronqué à gauche (RLT), lorsque les données sont indépendantes ainsi lorsqu'elles sont α -mélangeantes et associées, afin de tester la performance des résultats théoriques retrouvés pour ces modèles.

1 Introduction

Soit $\{(X_i, Y_i); i = 1, \dots, N\}$ une suite strictement stationnaire de vecteurs aléatoires définie dans le même espace de probabilité $(\Omega, F, \mathbb{P})$ à valeurs dans $R^d \times R$, ayant la même loi que (X, Y) . Si X et Y ne sont pas indépendants et X est observable, il est raisonnable de prédire Y à partir des informations apportées par X . Cette relation dite de régression est modélisée par $Y_i = m(X_i) + \epsilon_i; i = 1, \dots, N$, où $m(\cdot)$ désigne la fonction de régression et $\{\epsilon_i; i = 1, \dots, N\}$ est une suite d'erreurs indépendante de $\{X_i; i = 1, \dots, N\}$. Le meilleur prédicteur conditionnel de Y sachant X est donné par $m(x) = \mathbb{E}[Y/X = x], x \in R^d$, qui peut être écrit sous la forme suivante

$$m(x) = \frac{\psi(x)}{v(x)}$$

avec $\psi(x) = \int_R y f(x, dy)$, $f(\cdot, \cdot)$ est la densité conjointe de (X, Y) et $v(\cdot)$ est la densité marginale de X .

Il existe plusieurs estimateurs de la fonction de régression et nous nous focalisons sur l'estimateur à noyau.

Dans certains échantillons de survie, il arrive que la variable d'intérêt Y ne soit pas complètement observée, par conséquent, nous nous sommes intéressés au modèle de troncature à gauche qui, à l'origine, est apparu en astronomie et ensuite a été étendu à plusieurs domaines. Dans ce modèle, nous n'observons que $\{(X_i, Y_i, T_i); i = 1, \dots, N\}$ pour lesquelles $Y_i \geq T_i$, $\{T_i, i = 1, \dots, N\}$ désigne une suite de variables de troncature de même loi que T . Il est clair que nous disposons donc d'un échantillon observé de taille $n \leq N$.

En pratique, Le fait de supposer que les données sont toujours indépendantes est peu réaliste, c'est pour cela depuis quelques années plusieurs auteurs ont concentré leurs études sur un autre type de données, qui sont les données dépendantes.

Dans notre travail, nous allons présenter premièrement l'estimateur à noyau de la fonction de régression pour un modèle tronqué à gauche, ainsi que des résultats sur sa convergence dans le cas des données indépendantes, α -mélangeantes et associées. Puis on effectue des simulations de l'estimateur dans les différents modèles afin de vérifier sa qualité et de confronter les résultats pratiques à ceux attendus par la théorie.

2 Définition de l'estimateur et présentations des résultats

Ould Saïd et Lemdani (2006) ont défini un estimateur de la fonction régression pour un modèle tronqué à gauche par

$$\hat{m}_n(x) = \frac{\hat{\psi}_n(x)}{\hat{v}_n(x)}$$

$$\text{avec } \hat{\psi}_n(x) = \frac{\alpha_n}{nh_n^d} \sum_{i=1}^n \frac{Y_i}{G_n(Y_i)} K_d\left(\frac{x - X_i}{h_n}\right) \text{ et } \hat{v}_n(x) = \frac{\alpha_n}{nh_n^d} \sum_{i=1}^n \frac{1}{G_n(Y_i)} K_d\left(\frac{x - X_i}{h_n}\right)$$

- $K_d : R^d \rightarrow R$ est la fonction à noyau,
- h_n est appelée fenêtre qui tend vers 0 quand $n \rightarrow \infty$,
- G_n est l'estimateur de Lynden-Bell (1971) de la f.r G de la v.a de troncature T ,
- α_n est l'estimateur de $\alpha := \mathbb{P}(Y \geq T)$ proposé par He et Yang (1998).

Dans le cas indépendant et $d = 1$, Ould Saïd et Lemdani (2006) ont établi sa convergence uniforme ainsi que sa normalité asymptotique.

Quand les données sont α -mélangeantes, Liang et al. (2009) ont donné le taux de convergence uniforme de cet estimateur et Liang (2011) a étudié son comportement asymptotique.

Dans le cas où les $(X_i, Y_i)_{i=1, \dots, N}$ sont associés, nous avons établi la convergence uniforme sur un compact du même estimateur sous certaines conditions sur la fenêtre, le noyau et des conditions de régularité sur les densités conjointes et marginales. Notre résultat est énoncé dans le théorème suivant.

Théorème :

$$\sup_{x \in D} |\hat{m}_n(x) - m_n(x)| = O \left\{ \sqrt{\frac{\log n}{nh_n^d}} \vee \left(\frac{\log \log n}{n} \right)^\theta \vee h_n \right\} \mathbf{P} - \text{p.s quand } n \rightarrow \infty$$

avec $0 < \theta < \frac{2\gamma}{4\gamma+18+3\kappa}$, κ, γ sont des réels positifs.

3 Simulation

Dans cette partie, nous réalisons des simulations pour étudier la performance de l'estimateur $\hat{m}_n(x)$ de $m(x)$ dans le cas où $d = 1$ sur des échantillons de taille finie. L'estimateur $\hat{m}_n(x)$ dépend du choix du noyau et de la fenêtre h_n . Le choix du noyau influe peu sur la performance de cet estimateur ce qui nous a conduit à privilégier la solution classique du noyau

gaussien dans tous les modèles. Le choix de la fenêtre h_n est, en revanche, cruciale. Nous choisissons dans les différents cas les fenêtres optimales suivantes :

- cas indépendant : $h_n = Cn^{-1/5}$ (obtenue par Silverman (1986) pour les données complètes),
- cas α -mélange et associé : $h_n = C \left(\frac{\log n}{n} \right)^{1/3}$ (obtenue par Liebscher (2002) pour les données complètes α -mélangeantes).

avec C une constante à adapter pour chaque modèle.

On simule dans un premier temps, N valeurs $(X_i, Y_i, T_i)_{i=1, \dots, N}$ du triplet de variables aléatoires (X, Y, T) avec T indépendante de (X, Y) , et $T \sim \exp(\lambda)$, λ est adapté de manière à obtenir les différentes valeurs de α . Ensuite on retient les observations $(X_i, Y_i, T_i)_{i=1, \dots, n}$ vérifiant $Y_i \geq T_i$.

La qualité de l'estimateur est affecté beaucoup plus par n la taille de l'échantillon que par le taux de troncature α .

Cas indépendant : Modèle $\begin{cases} X_i \sim N(0, 1), \\ Y_i = 2X_i + 1 + \epsilon_i, \epsilon_i \sim N(0.2, 1). \end{cases}$

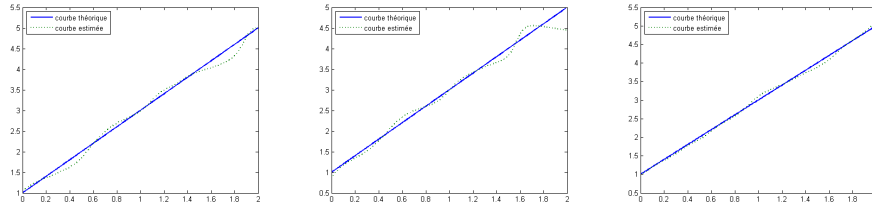


FIG. 1 – $m(x) = 2x + 1$ avec $\alpha \approx 70\%$ et $n = 50, 100, 300$, respectivement.

Cas α -mélange : Modèle $\begin{cases} X_i = 0.1X_{i-1} + \epsilon 1_i, \epsilon 1_i \sim N(0, 1), \\ Y_i = 2X_i + 1 + \epsilon 2_i, \epsilon 2_i \sim N(0.2, 1). \end{cases}$

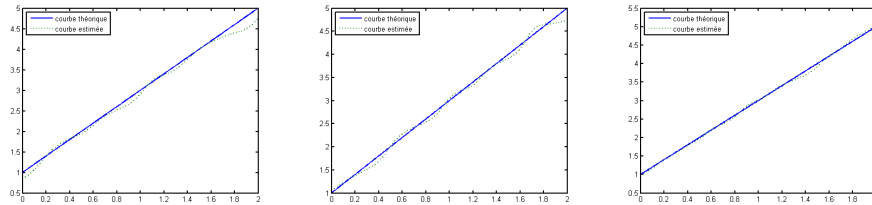


FIG. 2 – $m(x) = 2x + 1$ avec $\alpha \approx 70\%$ et $n = 50, 100, 300$, respectivement.

Simulation de l'estimateur à noyau de la fonction de régression dans le modèle RLT

Cas associé : Modèle $\begin{cases} X_i = \exp(W_{i-1}/2) \exp(W_{i-2}/2), W_i \text{ sont } N+1 \text{ v.a iid } \sim N(0, 1), \\ Y_i = 2X_i + 1 + \epsilon_i, \epsilon_i \sim N(0.2, 1). \end{cases}$

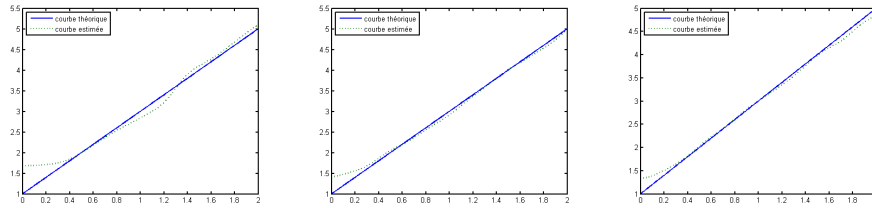


FIG. 3 – $m(x) = 2x + 1$ avec $\alpha \approx 70\%$ et $n = 50, 100, 300$, respectivement.

Références

- He, S. et G. Yang (1998). Estimation of the truncation probability in the random truncation model. *The Annals of Statistics* 26, 1011–1027.
- Liang, H. (2011). Asymptotic normality for regression function estimate under truncation and α -mixing conditions. *C. Statistics-Theory and Methods* 40, 1999–2021.
- Liang, H., D. Li, et Y. Qi (2009). Strong convergence in nonparametric regression with truncated dependent data. *Journal of J.Multivariate Anal* 100, 162–174.
- Liebscher, E. (2002). Kernel density and hazard rate estimation for censored data under α -mixing condition. *Ann.Inst.Statist.Math* 34, 81–106.
- Lynden-Bell, D. (1971). A method of allowing for known observational selection in small samples applied to 3CR quasars. *Monthly Notices Royal Astronomy Society* 155, 95–118.
- Ould Saïd, E. et M. Lemdani (2006). Asymptotic properties of a nonparametric regression function estimator with randomly truncated data. *Ann.Inst.Statist.Math* 58, 357–378.
- Silverman, B. (1986). *Density estimation for statistics and data analysis*. London: Chapman and Hall.

Summary

we have interested in this work of simulation to calculate the kernel estimator of the function of regression in the random left truncated (RLT) model, when the data are independent, α -mixing and associated in order to test the performance of the theoretical results founded for this models.

Impact d'une modélisation à volatilité stochastique sur la prime d'un Call européen

Souad Kadem*, Kamal Boukhetala*

*Faculté de Mathématiques-USTHB, PB 32, Bab-Ezzouar, Alger, Algérie
skadem@usthb.dz
kboukhetala@usthb.dz

Résumé. Les premières modélisations d'évaluation d'options européennes se fondent sur un paramètre de volatilité constant, et par souci de mieux prendre en compte la réalité et de mieux appréhender l'incertitude inhérente au marché financier, nous essayerons de voir si le fait de permettre à la volatilité d'évoluer au cours du temps dans le processus de prix de l'actif sous jacent va influencer les primes données par un modèle où la volatilité est supposée constante et quel est l'impact de ses différents paramètres sur ces primes. Nous pouvons ainsi conclure sur l'amélioration de la qualité de l'évaluation avec une telle modélisation.

1 Introduction

Les actifs financiers se comportent de façon aléatoire traduisant la complexité et la diversité du monde financier et économique. Mesurer et gérer des risques est ainsi devenu un enjeu majeur pour les opérateurs des marchés financiers, et donc le gain en précision d'un modèle doit être non seulement statistiquement, mais économiquement significatif. Le modèle le plus utilisé sur le parquet de la bourse suppose que la volatilité des rendements est constante (Black et Sholes, 1973) ce qui est loin d'être la réalité car des études empiriques (Cont, 2001) et (Ammraoui, 2008) ont constaté la non-conformité de cette hypothèse. Une nouvelle modélisation certes plus complexe mais qui adhère mieux à la réalité des marchés boursiers permettra-t-elle une meilleure évaluation ? Dans ce contexte, nous essayerons de répondre aux questions naturelles qui nous viennent en esprit :

- quel est l'influence des différents paramètres de ce modèle sur les primes des Calls ?
- la variabilité de la maturité a-t-elle un quelconque effet sur cette influence ?
- la variabilité de la maturité joue-t-elle un rôle important sur les primes des Calls ?

Nous étudierons donc l'influence des paramètres du modèle à volatilité stochastique SV (Heston, 1993) sur les primes de Calls européen.

Nous essayerons de répondre à ces questions par voie numérique, pour ce faire, nous avons généré des courbes et des surfaces représentant la différence des primes de Calls européens de deux modèles, le premier est à volatilité stochastique (Heston, 1993) et le deuxième est à volatilité constante (Black et Sholes, 1973).

2 Présentation des modèles

Avant d'entamer notre analyse sur la sensibilité des prix d'options européennes due à la variation des paramètres du modèle à volatilité stochastique (Heston, 1993), il est nécessaire de présenter les deux modèles.

2.1 Le modèle à volatilité constante de Black & Scholes :

le cours de l'actif S_t , obéit au processus d'évolution stochastique continu défini par l'équation différentielle stochastique :

$$dS_t = \mu S_t dt + \sigma S_t dW_t, \quad S_0 = s_0 \quad (1)$$

où μ est l'espérance des rendements de l'actif, σ est sa volatilité et W_t est un mouvement brownien standard.

2.2 Le modèle à volatilité stochastique :

Le modèle SV (Heston, 1993) fait intervenir un processus stochastique pour la volatilité qui est bien distinct de celui du cours de l'actif sous-jacent. Le modèle étant régi par la dynamique suivante :

$$\begin{cases} dS_t = \mu S_t dt + \sqrt{V_t} S_t dW_t^s \\ dV_t = \kappa(\theta - V_t)dt + \sigma\sqrt{V_t}dW_t^v \\ dW_t^s dW_t^v = \rho dt \end{cases} \quad (2)$$

où μ est l'espérance des rendements de l'actif sous-jacent, κ est la vitesse de retour au niveau moyen, θ est le niveau moyen à long terme de la variance, σ est la volatilité de la variance, ρ est le coefficient de corrélation entre le prix et la variance de l'actif sous-jacent, et W^s et W^v sont deux mouvements browniens standards corrélés.

3 Influence des paramètres du modèle SV sur les primes des Calls

3.1 Influence de la vitesse de retour à la moyenne κ sur les calls

En se basant sur l'allure de la courbe donnée par la figure 1, il en ressort les observations suivantes :

1. les deux modèles SV et BS convergent en primes lorsque les Calls sont très en dehors ($M < 0.55$) ou très en dedans de la monnaie ($M > 1.8$), c'est-à-dire que ces primes ont la même valeur quelle que soit la modélisation utilisée.
2. les Calls en dehors de la monnaie et à la monnaie $0.55 \leq M \leq 1.05$ sont sous-évalués par rapport au modèle BS. Cette sous-évaluation est fonction décroissante du paramètre κ .
3. les Calls en dedans de la monnaie $1.05 < M \leq 1.8$ sont surévalués par rapport au modèle BS, cette surévaluation est fonction décroissante du paramètre κ .

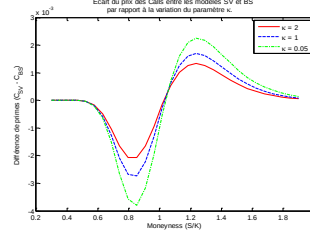


FIG. 1 – différence de primes des Calls entre les modèles SV et BS par rapport à la variation de κ .

¹La figure 1 correspond au plan de coupe $T = 1$ dans les figures suivantes :

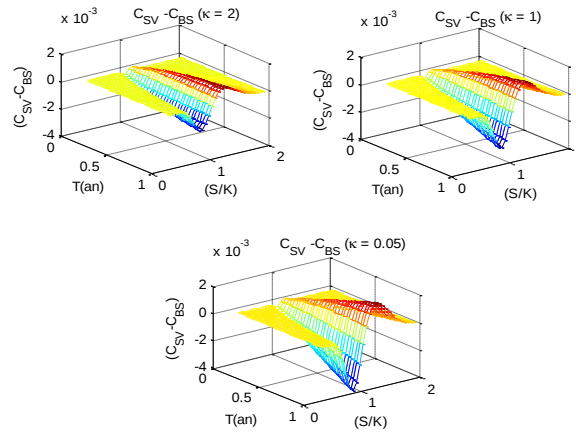


FIG. 2 – différence de primes de Calls entre les modèles SV et BS en fonction de la moneyness et de la maturité pour $\kappa = 2, 1$ et 0.05 .

Les surfaces illustrées dans²la figure 2 permettent de capter l'évolution de l'influence de la vitesse de retour de la variance à son niveau moyen dans le temps :

1. Toutes les observations précédentes sont d'autant plus vérifiées que la maturité est grande et l'influence du paramètre κ est d'autant plus significative que la maturité est plus élevée.
2. L'écart entre les primes des Calls évaluées par le modèle SV et le modèle BS semble s'accroître de façon non linéaire avec la maturité.

1. Les paramètres utilisés dans cette simulation sont : $\theta = 0.04, \sigma = 0.1, \rho = -0.5, V_0 = 0.04, r = 0.01, T = 1, S_0 = 1, \kappa = 0.05, 1, 2$. La volatilité dans le modèle de BS est égale à 0.2 .

2. Les paramètres utilisés dans cette simulation sont les mêmes que ceux de la figure 1 pour $T = [1/12, 1]$

L'effet d'un modèle SV sur les primes d'options

Nous pouvons conclure donc que plus la valeur de la vitesse de retour à la moyenne du processus de la variance augmente, moins est l'écart entre les primes des Calls évaluées par le modèle SV et celles évaluées par le modèle BS, cette constatation concerne toutes les options d'achat à l'exception de celles qui sont très en dedans ou très en dehors de la monnaie car elles ne sont pas influencées par une variation de κ ou le sont "très faiblement", c'est le cas aussi des options ayant une courte maturité. ($SV \approx BS$).

3.2 Influence de la volatilité du processus de la variance σ sur les calls

La courbe représentant la différence des primes ($C_{SV} - C_{BS}$) des Calls entre les modèles SV et BS en fonction de la Moneyness $M = \frac{S}{K}$ pour différentes valeurs de σ est donnée par la figure 3.

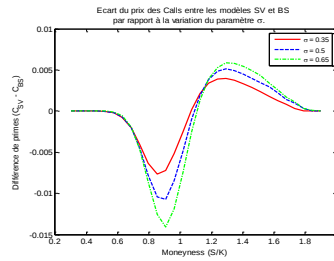


FIG. 3 – différence de primes des Calls entre les modèles SV et BS par rapport à la variation de σ

1. les Calls très en dehors de la monnaie ($M < 0,55$) ou très en dedans de la monnaie ($M > 1,85$) convergent en primes c'est-à-dire que quelle que soit la modèle utilisé, la variation de la volatilité de la variance σ n'a aucun effet sur les prix des Calls puisque la différence $C_{SV} - C_{BS}$ est nulle.
2. les Calls en dehors de la monnaie et à la monnaie $0,55 \leq M \leq 1,1$ sont sous-évalués avec le modèle SV par rapport au modèle BS.
3. les Calls en dedans de la monnaie $1,1 < M < 1,85$ sont surévalués avec le modèle SV par rapport à leurs valeurs avec le modèle BS.
4. l'écart des prix entre les deux modèles SV et BS pour tous les Calls excepté ceux qui sont très en dehors ou très en dedans de la monnaie est une fonction croissante de la volatilité de la variance σ . Cet écart est le plus élevé lorsque les Calls sont proches de la monnaie ($M = 0,9$).

Les courbes bleu et rouge représentées dans la figure 3 correspondent à l'intersection entre le plan $T = 1$ an et les surfaces décrites par la figure 4 pour les valeurs $\sigma = 0,35$ et $\sigma = 0,5$. En effet, en traçant la différence des primes des Calls entre les modèles SV et BS ($C_{SV} - C_{BS}$) en fonction de la Moneyness $M = \frac{S}{K}$ et de la maturité, on obtient :

3. Les paramètres utilisés dans cette simulation sont : $\theta = 0.04$, $\rho = -0.5$, $V_0 = 0.04$, $r = 0.01$, $T = 1$, $S_0 = 1$, $\kappa = 2$, $\sigma = 0.35, 0.5, 0.65$. La volatilité dans le modèle de BS est égale à 0.2.
4. Les paramètres utilisés dans cette simulation sont les mêmes que ceux de la figure 3 sauf pour $\sigma = 0.05, 0.35, 0.5$ et $T = [1/12, 1]$

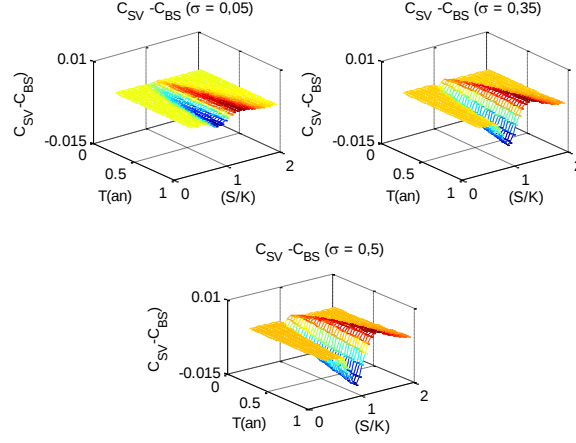


FIG. 4 – différences de primes de Calls entre les modèles SV et BS en fonction de la Moneyness et de la maturité pour $\sigma = 0.05, 0.35$ et 0.5 .

À partir des surfaces représentées dans la figure 4 nous pouvons conclure que :

1. l'effet de la volatilité de la variance σ sur les Calls est d'autant plus significatif que la maturité est grande, lorsque celle-ci tend vers zéro, l'écart des primes $C_{SV} - C_{BS}$ tend à s'estomper et lorsqu'elle augmente, l'écart augmente.
2. pour la première surface où $\sigma = 0.05$, que le Call soit en dehors, proche ou dans la monnaie l'écart est nul ou l'est presque. De ce fait, nous déduisons que lorsque la volatilité du processus de la variance tend vers 0, cette variance stochastique tend à devenir déterministe et les rendements du cours de l'actif sous jacent tendent à décrire une distribution log-normale.

3.3 Influence de la corrélation ρ du cours du sous jacent et sa variance sur les calls :

Le coefficient de corrélation entre le cours de l'actif sous jacent et sa volatilité affecte de façon significative l'asymétrie de la distribution des rendements. Ainsi, une corrélation négative (positive) crée une queue de distribution plus épaisse à gauche (à droite) et une queue plus mince à droite (à gauche) que dans le cas d'une distribution gaussienne.

⁵La figure 5 montre l'allure de la différence de primes $C_{SV} - C_{BS}$ des modèles SV et BS pour $\rho = 0.5$ et $\rho = -0.5$.

5. Les paramètres utilisés dans cette simulation sont : $\kappa = 2, \theta = 0.04, \sigma = 0.1, V_0 = 0.04, r = 0.01, T = 1, S_0 = 1, \rho = -0.5, 0.5$. La volatilité dans le modèle de BS est égale à 0.2.

L'effet d'un modèle SV sur les primes d'options

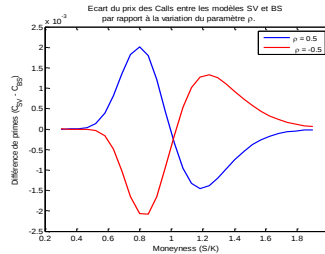


FIG. 5 – différence de primes des Calls entre les modèles SV et BS par rapport à la variation de ρ .

Nous pouvons constater qu'une corrélation négative dans le modèle SV a pour effet de diminuer (augmenter) les primes de Calls en dehors de la monnaie (dans la monnaie) et inversement pour une corrélation positive.

4 Conclusion

Nous avons étudié dans cet article l'influence de chaque paramètre du modèle SV (Heston, 1993) sur les primes des Calls européens par rapport à leurs valeurs avec le modèle BS. Nous pouvons prétendre que les paramètres κ , σ et ρ permettent l'ajustement de la sous-évaluation ou la surévaluation des primes calculées avec le modèle BS. Nous pouvons donc dire que :

- les Calls très en dehors et très en dedans de la monnaie ne sont influencés par aucune variation des trois paramètres κ , σ et ρ .
- la vitesse du retour à la moyenne permet d'ajuster la sous-évaluation des Calls en dehors et proche de la monnaie par rapport modèle BS, et permet aussi d'ajuster la surévaluation des Calls dans la monnaie.
- la volatilité de la variance stochastique σ permet l'ajustement de la sous-évaluation des Calls en dehors et proche de la monnaie par rapport au modèle BS.
- une corrélation non nulle influe fortement sur les Calls en dehors de la monnaie et dans la monnaie.
- l'influence de la vitesse de retour à la moyenne et de la volatilité de la variance croît de façon non linéaire avec la maturité.
- une corrélation négative permet d'ajuster la sous-évaluation des Calls en dehors de la monnaie et la surévaluation des Calls dans la monnaie par rapport au modèle BS et inversement pour une corrélation positive.
- l'écart des prix des Calls ($C_{SV} - C_{BS}$) entre les modèles SV et BS est au maximum pour les Calls proches de la monnaie quelque soit les valeurs de κ , σ et ρ .

Références

Ammraoui, M. (2008). *Le marché financier sous la dynamique de la volatilité stochastique*. Thèse de doctorat, Université PANTHEON-ASSAS (PARIS II).

- Black, F. et M. Sholes (1973). The pricing of options and corporate liabilities. *Journal of Political Economy* 81, 637–659.
- Cont, R. (2001). Empirical properties of assets returns: Stylized facts and statistical issues. *Quantitative Finance* 1, 1–14.
- Heston, S. (1993). A closed-form solution for options with stochastic volatility with applications to bond and currency options. *Review of Financial Studies* 6, 327–343.

Summary

The first modeling of the valuation of European options are based on a constant volatility parameter, and for the sake of better reflect reality and to better understand the inherent uncertainty of the financial market, we will try to see if permitting volatility to evolve over time in the price process of the underlying asset will influence the premium given by a model where the volatility is assumed constant, and what is the impact of its various parameters on these premiums. We can then conclude on improving the quality of the evaluation with such modeling.

Analyse des données catégorielles

Wafa Karouche*, Raouf Benmakrelouf*, Joseph Rynkiewicz**, Mohamed Djedour*

*Faculty of mathematics PoBox 32, Al alia Bab ezzouar, USTHB-Algiers, Algeria
karouche.w@hotmail.fr,
ismtid_raouf@hotmail.fr
mdjedour@yahoo.fr

**Université Paris 1, Panthéon-Sorbonne, Paris, France
joseph.rynkiewicz@univ-paris1.fr

Résumé. Les données d'enquête, sont bien souvent composées de données qualitatives. L'analyse de données catégorielles (qualitatives) permet d'évaluer les facteurs qui affectent les opinions et les attitudes sur divers questions controversées. Dans ce travail, nous discuterons l'utilisation des méthodes statistiques pour les données catégorielles en présentant quelques exemples d'applications.

1 Introduction

L'analyse de données catégorielles est une étape primaire lorsqu'on souhaite interroger une base de données quelconque. Le document que nous présentons constitue le premier jalon d'une étude que nous menons dans le domaine des mathématiques appliquées aux sciences humaines. Il s'agit donc pour nous de faire une première étude d'analyse des observations et de déterminer les variables significatives qui permettent d'expliquer la manière dont les opinions, jugements et comportements politiques se structurent.

2 Nature de variables

Les modèles statistique nécessitent la distinction entre les variables qualitatives et quantitatives. Elles s'expriment en modalités ou en catégories pour les variables qualitatives quant aux variables quantitatives elles sont représenté par des quantités et donc en valeurs.

2.1 Variables catégorielles

Une variable catégorielle à une échelle de mesure constituant un ensemble de catégories. Elles sont de types assez variés et nous pouvons les organiser selon les catégories suivantes :

- Variables réponse et variables explicatives
- Variables nominales et ordinales
- Variables continues et discrètes

3 Distribution pour données catégorielles

3.1 Distribution binomiale

Généralement, les données catégorielles résultent d'essais n indépendants et identiques avec deux résultats nommé Succé et Echec.

Essais identiques : la probabilités de succès p est la même pour chaque essai.

Essais indépendants : les variables $\{Y_i\}$ sont indépendants.

La variable $Y = \sum_{i=1}^n Y_i$ est distribuée selon une loi Binomiale $B(n, p)$ telle que :

$$\mathbb{P}(y) = C_n^y p^y (1-p)^{n-y}, y = 0, 1, 2, \dots, n.$$

Avec :

$$\mathbb{E}(Y_i) = p \text{ et } \mathbb{V}(Y_i) = p(1-p)$$

$$\mathbb{E}(Y) = np \text{ et } \mathbb{V}(Y) = np(1-p)$$

Dans la pratique, il n'est pas toujours garanti que des observations binaires successives soient indépendantes. On utilise dans ce cas d'autres distributions.

3.2 Distribution multinomiale

La variable réponse Y dans ce cas contient plus de deux modalités. Notons c le nombre de résultats possible. Soit $y_{ij} = 1$ si le résultat de l'essai i appartient à la catégorie j , sinon $y_{ij} = 0$.

Alors $y_i = (y_{i1}, y_{i2}, \dots, y_{ic})$ est un essai multinomial, et le vecteur $(n_1, n_2, \dots, n_c)$ a une distribution multinomiale.

Soit $p_j = \mathbb{P}(Y_{ij} = 1)$. La loi multinomiale est :

$$p(n_1, n_2, \dots, n_c) = \left(\frac{n!}{n_1! n_2! \dots n_c!} \right) p_1^{n_1} p_2^{n_2} \dots p_c^{n_c}$$

Avec :

$$\mathbb{E}(n_j) = np_j \text{ et } \mathbb{V}(n_j) = np_j(1-p_j)$$

De plus :

$$\text{cov}(n_j, n_l) = -np_j p_l$$

Lorsque $c = 2$, nous nous retrouvons dans le cas d'une distribution binomiale.

3.3 Distribution de poisson

La loi de poisson est utilisée pour compter les événements survenant aléatoirement dans le temps ou dans l'espace, sa loi s'écrit sous la forme :

$$p(y) = \frac{e^{-\lambda} \lambda^y}{y!}, y = 0, 1, 2, \dots$$

C'est aussi une approximation de la loi binomial pour une taille d'échantillon élevé et une probabilité p petite.

De plus, $\mathbb{E}(Y) = \mathbb{V}(Y) = \lambda$.

4 Structure probabiliste des tables de contingences

4.1 Table de contingence

Prenons X et Y deux variables à réponses catégorielles, possédant respectivement I et J catégories. La distribution (X, Y) est représentée par une table à I entrées pour la catégorie X , et J entrées pour la catégorie Y ou chaque cellule est associé à une combinaison de $I \times J$ de (X, Y) . Cette table est appelée tableau de contingence ou tableau croisé à deux dimension.

Le tableau ci-dessous donne un exemple de table de contingence, prise d'une enquête sur les intentions de vote d'un certain candidat. Nous nous intéressons à l'effet de la variable sexe sur l'intention de vote pour un candidat précis.

	Homme	Femme
Défavorable	2427	2633
Favorable	269	307

TAB. 1 – Tableau croisé entre la variable intention de vote et la variable sexe

On note $p_{ij} = \mathbb{P}(X = i, Y = j)$.

p_{ij} est appelée distribution jointe de X et Y . Les distributions marginales représentent les totaux par ligne et colonne.

On note : $\mathbb{P}(X = i) = p_{i.} = \sum_j p_{ij}$ et $\mathbb{P}(Y = j) = p_{.j} = \sum_i p_{ij}$

La distribution conditionnelle de Y sachant $X = i$ est donnée par le vecteur $(p_{1|i}, p_{2|i}, \dots, p_{J|i})$

Avec : $p_{j|i} = \mathbb{P}(Y = j | X = i) = \frac{p_{ij}}{p_{i.}}$

La distribution marginale de X est estimé par :

$$f_{i.} = \sum_{j=1}^J f_{ij} \quad \forall 1 \leq i \leq I$$

Avec : $f_{i.} = \frac{n_{i.}}{n}$ et $f_{.j} = \frac{n_{.j}}{n}$

4.2 Indépendance des variables catégorielles

On parle d'indépendance entre deux variables catégorielles si la distribution jointe est le produit des distributions marginales

$$p_{ij} = p_{i.} p_{.j} \quad \forall i = \overline{1, I} \text{ et } \forall j = \overline{1, J}$$

4.3 Distribution pour les tables de contingences

Les distributions introduites précédemment peuvent s'étendre sur les tables de contingences. Lorsque la variable Y est dichotomique (binaire), la distribution binomiale s'applique pour chaque distribution conditionnelle, et si les catégories sont supérieurs à deux, la distribution multinomiale le sera aussi.

Références

Agresti, A. (2007). An introduction to gategorical data analysis. Université de Floride.

Pearl, J. (2008). Causality. Université de Californie.

Cristensen, R. (1997). Log linear models and logistic regression. Université du nouveau Mexique.

Summary

Survey data are often composed of qualitative data. The analysis of categorical data, assesses factors affecting the attitudes and opinions on a variety of controversial issues. In this work , we discuss the use of statistical methods for categorical data by presenting some examples of applications.

Sur l'estimation des processus $SARFIMA - S\alpha S$

Radia Lessak*

*Département de Mathématiques, Université Constantine 1, Constantine, Algérie
lessak_radia@yahoo.fr

Résumé. On considère dans cette communication l'estimation d'un processus $SARFIMA - S\alpha S$. Ce processus englobe trois autres classes de processus, à savoir: les processus à mémoire longue, les processus saisonniers et les processus à variance infinie. La méthode utilisée est l'une des méthodes classiques les plus simples malgré qu'elle possède pas mal d'avantages. C'est la méthode de la distance minimale.

1 Introduction

Afin de prendre en compte la notion de mémoire longue, Granger et Joyeux (1980) ainsi que Hosking (1981) ont introduit les processus $ARFIMA$ en supposant que les résidus sont des bruits blancs gaussiens. Par ailleurs, dans tous les champs d'application des statistiques, de nombreuses séries possèdent un comportement cyclique qui ne pourra malheureusement pas être modélisé par un tel modèle. Dans ce contexte, Reisen et al. (2006) ont proposé les processus $SARFIMA$ qui prennent en compte les phénomènes de mémoire longue et de saisonnalité. Ils ont proposé des estimateurs obtenus à partir du log-périodogramme et à titre comparatif ils ont examiné les estimateurs du maximum de vraisemblance et ceux obtenus par la méthode semi-paramétrique présentée dans Geweke et Porter-Hudak (1983) et Porter-Hudak (1990). Comme la classe de ces modèles est limitée face à des données présentant une variance infinie, on va considérer dans cette communication la classe des processus $SARFIMA$ symétriques $\alpha - stables$ qui englobe ces trois éléments : mémoire longue, saisonnalité et variance infinie.

A de nombreux égards, l'estimation par une méthode de distance minimale appropriée est l'une des idées les plus naturelles dans les statistiques et c'est d'ailleurs l'objectif du présent travail. Autrement dit, nous allons introduire dans la deuxième section le modèle $SARFIMA - S\alpha S$. Dans la section 3, nous présentons le principe de la méthode de la distance minimale utilisée. La dernière section est consacrée à l'étude par simulation des propriétés asymptotiques des estimateurs obtenus à travers cette méthode.

2 Le modèle

Soit $(\varepsilon_t)_{t \in \mathbb{Z}}$ une suite de variables aléatoires indépendantes et identiquement distribuées $S\alpha S$ ($0 < \alpha < 2$) de moyenne nulle et de paramètre de dispersion égal à 1. Soit s la période

Sur l'estimation des processus $SARFIMA - S\alpha S$

de saisonnalité. Les polynômes d'ordres p, q, P, Q définis respectivement par :

$$\begin{aligned}\phi(B) &:= 1 - \sum_{j=1}^p \phi_j B^j, \theta(B) := 1 + \sum_{j=1}^q \theta_j B^j, \\ \Phi(B^s) &:= 1 - \sum_{j=1}^P \Phi_j B^{js}, \Theta(B^s) := 1 + \sum_{j=1}^Q \Theta_j B^{js},\end{aligned}$$

sont les opérateurs non-saisonniers et saisonniers d'un processus $ARMA$. Ainsi, un processus $(X_t)_{t \in \mathbf{Z}}$ est dit processus $ARFIMA$ saisonnier $S\alpha S$, noté $SARFIMA(p, d, q) \times (P, D, Q)_s - S\alpha S$, s'il satisfait l'équation suivante :

$$\phi(B)\Phi(B^s)X_t = (1-B)^{-d}(1-B^s)^{-D}\theta(B)\Theta(B^s)\varepsilon_t, \quad (1)$$

où $d, D \in \mathbf{R}$ sont les paramètres de mémoire longue. Par développement en série, le filtre $(1-B)^{-d}(1-B^s)^{-D}$ peut s'écrire sous la forme

$$(1-B)^{-d}(1-B^s)^{-D} = \sum_{j=0}^{\infty} \zeta_j B^j,$$

où le calcul des coefficients $(\zeta_j)_{j \geq 0}$ se fait en utilisant les polynômes de Gegenbaur.

Supposons que les polynômes $\phi(z)$, $\Phi(z^s)$, $\theta(z)$ et $\Theta(z^s)$ n'ont aucun zéro commun et soit la condition suivante :

$$[C.1] \quad |d+D| < 1 - \frac{1}{\alpha} \text{ et } |D| < 1 - \frac{1}{\alpha} \text{ avec } 0 < \alpha \leq 2.$$

Sous [C.1] on a :

1. Si les zéros des polynômes $\phi(z)$ et $\Phi(z^s)$ sont à l'extérieur du cercle unité, alors le processus (??) est strictement stationnaire, causal et admet une unique représentation moyenne mobile infinie donnée par :

$$X_t = \sum_{j=0}^{\infty} \psi_j \varepsilon_{t-j}, \quad (2)$$

où la formule explicite des coefficients $(\psi_j)_{j \geq 0}$ est obtenue en développant l'équation

$$\Theta(z^s)\theta(z) \sum_{j=0}^{\infty} \zeta_j z^j = \Phi(z^s)\phi(z) \sum_{j=0}^{\infty} \psi_j z^j, |z| < 1.$$

2. Si les polynômes $\theta(z)$ et $\Theta(z^s)$ n'ont pas de zéros dans le cercle unité, alors le processus (??) est inversible et admet une représentation autorégressive infinie définie par

$$\varepsilon_t = \sum_{j=0}^{\infty} \pi_j X_{t-j},$$

où les coefficients $(\pi_j)_{j \geq 0}$ sont définis par

$$\frac{\phi(z)\Phi(z^s)}{\theta(z)\Theta(z^s)} \prod_{j=0}^{\lfloor \frac{s}{2} \rfloor} (1 - 2(\cos \lambda_j)z + z^2)^{-d_j} = \sum_{j=0}^{\infty} \pi_j z^j,$$

où $\lambda_j = \frac{2\pi j}{s}$, $j = 0, \dots, \lfloor \frac{s}{2} \rfloor$ ($[x]$ est la partie entière de x).

Sur la base de ces résultats, on va passer à la partie estimation.

3 Principe de la méthode de la distance minimale

La distance minimale est une méthode d'estimation classique dans la littérature économétrique. Elle est fondée sur la minimisation d'une fonction critère de la forme : $Q_N(\underline{\beta}) = H_N(\underline{\beta})' \mathbf{W} H_N(\underline{\beta})$, où $H_N(\underline{\beta})$ est une fonction des observations X_t , $t = 1, \dots, N$, et du vecteur des paramètres d'intérêt $\underline{\beta}$ qui appartient à un ensemble compact Θ ; $\mathbf{W}$ est une matrice poids symétrique définie positive. Si $\underline{\beta}_0$ désigne la vraie valeur de $\underline{\beta}$, alors cette valeur est caractérisée par la propriété suivante : $p \lim_{N \rightarrow \infty} H_N(\underline{\beta}_0) = 0$. L'estimateur de la distance minimale compte sur le choix de la matrice poids optimale¹ $\mathbf{W}_{opt}$, et la question qui se pose est la suivante : Quelle matrice poids choisissons-nous afin d'obtenir un estimateur convergent et efficace du vecteur $\underline{\beta}$?

Il est connu que si $p \lim_{N \rightarrow \infty} \text{var}(\sqrt{N} H_N(\underline{\beta}_0)) = \mathbf{\Gamma}$, alors $\mathbf{W}_{opt} = \mathbf{\Gamma}^{-1}$. Dans ce cas, la matrice asymptotique de variance-covariance de l'estimateur de $\underline{\beta}$ se simplifie à $(\mathbf{J}'_{\underline{\beta}_0} \mathbf{\Gamma}^{-1} \mathbf{J}_{\underline{\beta}_0})^{-1}$, où $\mathbf{J}_{\underline{\beta}_0}$ est la limite de la matrice Jacobienne de H_N .

Notons que les différents choix de la fonction H_N génèrent différents estimateurs. Si, par exemple, $H_N(\underline{\beta}) = \frac{1}{N} \sum_{t=1}^N h(X_t, \underline{\beta})$, où $E\{h(X_t, \underline{\beta})\} = 0$, la minimisation de Q_N engendre un estimateur *GMM* (méthode des moments généralisée).

La méthode de la distance minimale présente de nombreux avantages tels que : la simplicité, elle permet d'estimer tous les paramètres simultanément, l'estimateur obtenu à partir de cette méthode admet de très bonnes propriétés asymptotiques, etc. Ainsi, cette méthode a été l'objet de nombreux travaux, citons par exemple Mayoral (2007) et Kouamé et Hili (2008). Ces derniers ont considéré l'estimation par la distance minimale dans la classe des processus *GARMA* à k facteurs.

4 Etude par simulation

Nous présentons dans cette section un ensemble de simulations pour illustrer le comportement asymptotique des estimateurs d'un processus *SARFIMA* $(p, d, q) \times (P, D, Q)_s - S\alpha S$. Pour simplifier, nous nous limitons sans perte de généralité au modèle avec $p = q = P = Q =$

1. Optimale au sens que la variance asymptotique de l'estimateur soit minimale.

Sur l'estimation des processus $SARFIMA - S\alpha S$

0. Plus précisément, nous simulons, en utilisant la méthode proposée dans la section 3, le processus donné par

$$X_t = (1 - B)^{-d} (1 - B^s)^{-D} \varepsilon_t.$$

Notre objectif est de générer une trajectoire X_t , $t = 1, \dots, N$, issu d'un tel processus en utilisant sa représentation moyenne mobile infinie donnée par l'équation (??). En d'autres termes, nous pouvons approximer cette trajectoire de taille N par sa version tronquée de la manière suivante :

$$X_t = \sum_{j=1}^M \psi_j \varepsilon_{t-j},$$

où M est le paramètre de troncature (on garde les mêmes notations de la deuxième section).

Notons aussi qu'à la différence du cas gaussien, certaines modifications doivent être apportées. Par exemple, l'outil de base utilisé dans notre méthode d'estimation qui est la fonction d'autocorrélation n'existe pas. Néanmoins, même si cette quantité n'a pas d'existence théorique, on peut tout de même calculer sa valeur empirique et l'utiliser comme dans le cas classique.

Références

- Geweke J. and S. Porter-Hudak (1983). The estimation and application of long memory time series models. *Journal of Time Series Analysis* 4, 221–238.
- Granger C. W. J. and R. Joyeux (1980). An introduction to long-memory time series models and fractional differencing. *Journal of Time Series Analysis* 1, 15-29.
- Hosking J. R. M. (1981). Fractional differencing. *Biometrika* 68, 165-176.
- Kouamé E. F. and O. Hili (2008). Minimum distance estimation of k-factors GARMA processes. *Statistics and Probability Letters* 78, 3254-3261.
- Mayoral L. (2007). Minimum distance estimation of stationary and non-stationary ARFIMA processes. *The Econometrics Journal* 10, 124-148.
- Porter-Hudak S. (1990). An application of the seasonal fractionally differenced model to the monetary aggregates. *Journal of the American Statistical Association* 85, 338-344.
- Reisen V. A., A. L. Rodrigues, and W. Palma (2006). Estimation of seasonal fractionally integrated processes. *Computational Statistics and Data Analysis* 50, 568-582.

Summary

We consider in this communication the estimation of a $SARFIMA - S\alpha S$ process. This process includes three other classes of processes, namely: long memory processes, seasonal processes and processes with infinite variance. The method used here is one of the simplest classical methods although it has a lot of advantages. This is the method of the minimum distance.

Estimation du Maximum du Taux de Hasard dans un Modèle Associé-Tronqué

Dalila Madani*, Abdelkader Tatachak*

* Laboratoire MSTD, Département de Probabilités et Statistique, USTHB, Alger, Algeria
dmadani@usthb.dz & atatachak@usthb.dz

1 Introduction

L'estimation non paramétrique du taux de hasard fut étudiée par les précurseurs Watson et Leadbetter (64) pour des données i.i.d. La généralisation à des données incomplètes (censurées ou tronquées) a suscité une grande littérature, tels les travaux de Gurler et Lang (93) et Sun et Zhou (01) dans le cas de données tronquées à gauche. Ces derniers ont étudié le comportement asymptotique de l'estimateur du taux de hasard pour des données i.i.d. et α -mélangeantes respectivement. Il arrive parfois que les données à traiter ne soient ni indépendantes ni α -mélangeantes, mais de type associées. Cette notion de dépendance a été introduite par Esary et al. (67) dans le cas de données complètes. Récemment Doukhan et Neumann (07) ont proposé une inégalité exponentielle pour le traitement de données faiblement dépendantes et Bulinski et Shashkin (07) ont présenté des travaux sur le sujet en étudiant les processus associés. Lorsque les données sont incomplètes, l'article de Guessoum et al. (12) étend les travaux de Woodroffe (85) à des données associées alors que l'article de Cai et Roussas (98) traite le cas de données censurées (Kaplan-Meier) associées.

2 Estimation non paramétrique du taux de hasard

Soit Y and T deux variables aléatoires (v.a.) réelles indépendantes de fonctions de répartition inconnues F and G respectivement. Notons par f la densité de probabilité (p.d.f) de Y . Soit $\{(Y_i, T_i); i = 1, \dots, N\}$ une suite de N copies de (Y, T) où la taille de l'échantillon N est déterministe mais inconnue. Dans le modèle de troncature aléatoire gauche (TAG), (Y_i, T_i) n'est observé que lorsque $Y_i > T_i$ sinon rien n'est observé. Notons par $\{(Y_i, T_i); i = 1, \dots, n; (n \leq N)\}$ l'échantillon réellement observé et par $\alpha = P(Y > T)$ la probabilité d'observer au moins une paire de (Y, T) . Nous supposons que $\alpha > 0$ sinon aucune donnée n'est observée. Rappelons que la taille de l'échantillon observé n est une v.a. binomiale $B(N, \alpha)$. Notons par $a_F = \inf\{u : F(u) > 0\}$ et $b_F = \sup\{u : F(u) < 1\}$ (de même pour G). Nous imposons aussi à $\{Y_i, 1 \leq i \leq N\}$ et $\{T_i, 1 \leq i \leq N\}$ d'être strictement stationnaires et associés, par conséquent, les observations demeurent aussi associées Esary et al. (67). Pour assurer l'identifiabilité du modèle, nous supposons l'indépendance de Y et T et les conditions $a_G < a_F$ et $b_G < b_F$ et $\int_{a_F}^{\infty} \frac{dF}{G} < \infty$.

Estimation du maximum du taux de hasard

La fonction de hasard à laquelle nous nous sommes intéressées dans ce travail est définie par :

$$\lambda(y) = \frac{f(y)}{1 - F(y)}$$

et son maximum

$$\lambda(\theta) = \max_{y \in \mathcal{D}} \lambda(y).$$

Leurs estimateurs que nous proposons sont respectivement

$$\hat{\lambda}_n(y) = \frac{\hat{f}_n(y)}{(1 - F_n(y))}$$

et

$$\hat{\lambda}(\hat{\theta}_n) = \max_{y \in \mathcal{D}} \hat{\lambda}_n(y),$$

où

$$\hat{f}_n(y) = \frac{\alpha}{nh} \sum_{i=1}^n \frac{1}{G_n(Y_i)} K\left(\frac{y - Y_i}{h}\right)$$

et

$$F_n(y) = 1 - \prod_{i: Y_i \leq y} \left[\frac{nC_n(Y_i) - 1}{nC_n(Y_i)} \right].$$

$F_n(y)$ étant l'estimateur de Lynden-Bell.

3 Résultats principaux

Sous les hypothèses de régularités classiques et les conditions exploitant la structure de dépendance de type associée TAG nous avons le résultat 1) établi par Guessoum et al. (12), cependant les deux autres 2) et 3) font l'objet de ce présent travail.

- 1) $\sup_{y \in D} |\hat{f}_n(y) - f(y)| = \mathcal{O} \left(\left(\frac{\log \log n}{n} \right)^\delta \vee h^2 \vee \left(\frac{\log n}{nh} \right)^{1/2} \right)$ a.s as $n \rightarrow +\infty$.
 - 2) $\sup_{y \in D} |\hat{\lambda}_n(y) - \lambda(y)| = \mathcal{O} \left(\left(\frac{\log \log n}{n} \right)^\delta \vee h^2 \right)$ a.s as $n \rightarrow +\infty$.
 - 3) $\sup_{y \in D} |\hat{\theta}_n - \theta| = \mathcal{O} \left(\left(\frac{\log \log n}{n} \right)^\delta \right)$ a.s as $n \rightarrow +\infty$.
- où $\delta > 0$ est le même que dans les hypothèses.

4 Simulations

- Modèles i.i.d. tronqués : On choisit la variable d'intérêt Y suivant une weibull (1.5,2), la variable de troncature T une exp(1.5), le noyau K est gaussien, n la taille de l'échantillon observé, α = taux de troncature et la fenêtre h varie en fonction de n . Les graphes obtenus représentent $f(y)$ et son estimateur $\hat{f}_n(y)$ et la fonction de hasard $\lambda(y)$ et son estimateur $\hat{\lambda}_n(y)$.

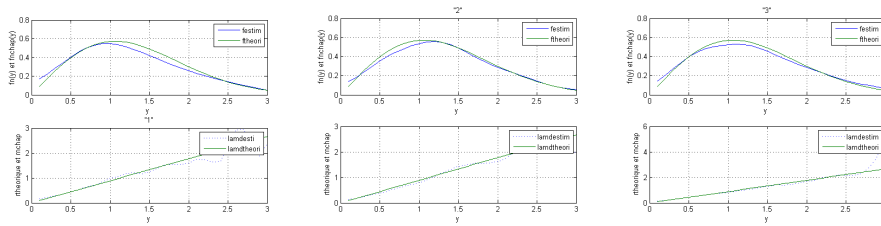


FIG. 1 – *Modèle i.i.d., $f(y)$, $\hat{f}_n(y)$ et $\lambda(y)$, $\hat{\lambda}_n(y)$*

- (1) $n=100$, $\alpha=0.7092$, $h = 0.3675$ weibull(1.5,2)
- (2) $n=200$, $\alpha=0.6993$, $h = 0.3055$ weibull(1.5,2)
- (3) $n=300$, $\alpha=0.6977$, $h = 0.2811$ weibull(1.5,2)

- Modèles TAG associés : Pour tous les modèles que nous avons choisi dans cette partie de simulation sont en cohérence avec les hypothèses de la théorie menée en particulier vérifient $a_G < a_F$ et $b_G < b_F$;

A partir d'un vecteur Z de dimension $(n+2)$ issu d'une exp(1) on construit le vecteur v telque $v_i = (z_{i-1} + z_{i-2})/2$ et en posant $Y(i) = v(i+2)$ ainsi Y vérifie la dépendance d'associé, et on introduit une variable de troncature T type exp(0.5)

on a obtenu les graphes suivants

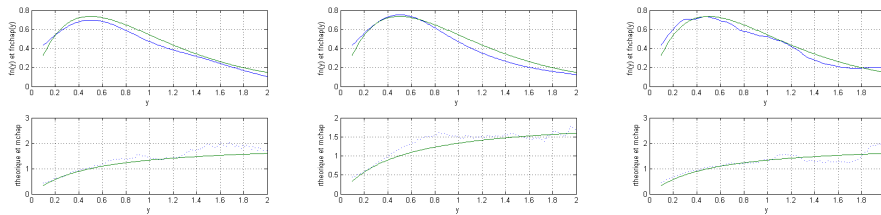


FIG. 2 – *Modèle dépendant (associé) $f(y)$, $\hat{f}_n(y)$ et $\lambda(y)$, $\hat{\lambda}_n(y)$*

- (1) $n=100$, $h=0.2891$, $\alpha=0.7092$, noyau gaussien
- (2) $n=200$, $h=0.28121$, $\alpha=0.6969$, noyau gaussien
- (3) $n=200$, $h=0.2736$, $\alpha=0.7067$, noyau d'Epanchnikov

Références

- Bulinski, A. et A. Shashkin (07). *Limit theorems for associated random fields and related systems*. Berlin: Advanced series on statistical science & applied probability.
- Cai, Z. et G. Roussas (98). Kaplan-meier estimator under association. *J. Multivariate Anal* 67, 318–348.
- Doukhan, P. et M. Neumann (07). Probability and moment inequalities for sums of weakly dependent random variables, with applications. *Stochastic Processes and their Applications* 117, 818–903.
- Esary, J., F. Proschan, et D. Walkup (67). Association of random variables with applications. *Ann. Math. Stat* 38, 1466–1476.
- Guessoum, Z., E. OuldSaid, O. Sadki, et A. Tatachak (12). A note on the lynden-bell estimator under association. *Statist. Probab. lett* 82, 1994–2000.
- Gurler, G. et J. Lang (93). Nonparametric estimation of hasard functions and their derivatives under truncation model. *Ann. Inst. Statist. Math* 45, 249–264.
- Sun, L. et X. Zhou (01). Survival function and density estimation for truncated dependent data. *Statistics and Probability Letters* 52, 47–57.
- Watson, G. et M. Leadbetter (64). Hasard analysis i. *Biometrika* 51, 175–184.
- Woodroffe, M. (85). Estimating a distribution function with truncated data. *Ann. Statist* 123, 163–177.

Summary

In this work we propose an estimator of the maximum of the hazard function under left-truncated and associated model. We establish the strong consistency of the proposed estimator as well as its rate. A simulation study is presented to illustrate the behavior of the estimator.

Estimateur à noyau de la fonction de régression dans un modèle de censure à droite:applications

Nassira Menni*, Abdelkader Tatachak**

*Faculté de Mathématiques, USTHB
BP 32, El-Alia 16111, Algiers, Algeria,
nmenni@usthb.dz

**Laboratoire MSTD, Faculté de Mathématiques, USTHB
BP 32, El-Alia 16111, Algiers, Algeria,
atatchak@usthb.dz

Résumé. Dans ce travail, nous nous sommes intéressés à la construction et à l'étude des propriétés d'un estimateur à noyau de la fonction de régression pour des données censurées à droite, sous l'hypothèse d'indépendance d'abord, puis, en supposant deux structures de dépendances, α -mélange et association. Dans cette présentation, nous illustrons par des simulations les résultats théoriques obtenus.

1 Introduction

Soit $Y_1, Y_2, \dots, Y_n$ une suite strictement stationnaire de variables aléatoires (v.a) admettant fonction de répartition (f.r) commune F . Souvent ces v.a sont considérées comme des durées de vie qui ne peuvent pas être complètement observables. Soit $(C_i)_{1 \leq i \leq n}$ une suite de v.a de censure à droite, indépendantes et identiquement distribuées (i.i.d) de f.r G . Dans le modèle de censure à droite, on observe seulement le couple $\{(T_i, \delta_i), i = 1, 2, \dots, n\}$ où $T_i = Y_i \wedge C_i$ et $\delta_i = \mathbf{1}_{\{Y_i \leq C_i\}}$, (l'indicatrice de non censure). Nous supposons dans toute la suite que les variables $(Y_i)_{i \geq 1}$ et $(C_i)_{i \geq 1}$ sont mutuellement indépendantes afin d'assurer l'identifiabilité du modèle. Soit X un vecteur de covariables à valeurs dans R^d , de densité v . Afin d'étudier la relation entre X et Y , on utilise le modèle de régression qui s'écrit sous la forme $Y_i = r(X_i) + \varepsilon_i, i = 1, \dots, n$ où $r(\cdot)$ dénote la fonction de régression inconnue et ε_i sont des erreurs. Il est connu que la fonction de régression qui réalise le minimum de l'erreur quadratique moyenne est donnée par l'espérance conditionnelle $r(x) = E(Y/X = x)$. qui peut s'écrire sous la forme :

$$r(x) = \frac{\int_R y f_{X,Y}(x, y) dy}{v(x)} =: \frac{r_1(x)}{v(x)},$$

où $f_{X,Y}(\cdot, \cdot)$ et $v(x)$ désignent la densité conjointe de (X, Y) et la densité de la covariable X , respectivement. Il existe plusieurs estimateurs non paramétriques de $r(\cdot)$. Dans notre travail, nous nous intéressons à l'estimateur à noyau défini par Nadaraya (1964).

2 Modèles et principaux résultats

Dans un premier temps, nous rappelons certains résultats sur le sujet, obtenus par Guessoum et OuldSaïd (2008). Nous présentons les résultats concernant la consistance de l'estimateur étudié pour les différents modèles considérés, sous certaines hypothèses sur le noyau, la fenêtre et des conditions de régularité sur les lois conjointes et marginales.

L'expression de l'estimateur est donnée par :

$$r_n(x) =: \frac{r_{1,n}(x)}{v_n(x)}, \quad (1)$$

où

$$r_{1,n}(x) = \frac{1}{nh_n^d} \sum_{i=1}^n \frac{\delta_i T_i}{\bar{G}_n(T_i)} K_d \left(\frac{x - X_i}{h_n} \right)$$

et

$$v_n(x) = \frac{1}{nh_n^d} \sum_{i=1}^n K_d \left(\frac{x - X_i}{h_n} \right).$$

- $K_d : R^d \rightarrow R$ est une fonction, appelée noyau,
- h_n est une suite de réels qui tend vers 0 quand $n \rightarrow \infty$, appelée fenêtre.
- $\bar{G}_n$ est l'estimateur de Kaplan (1958) de la loi de censure.

2.1 Modèle 1 : cas indépendant :

Théorème 1 : Guessoum et OuldSaïd (2008)

$$\sup_{x \in D} |r_n(x) - r(x)| = O \left(\max \left\{ \left(\sqrt{\frac{\log n}{nh_n^d}} \right), h_n^2 \right\} \right) \text{ p.s quand } n \rightarrow \infty.$$

2.2 Modèle 2 : cas α -mélange

Théorème 2 : Guessoum et OuldSaïd (2010)

$$\sup_{x \in D} |r_n(x) - r(x)| = O \left(\sqrt{\frac{\log n}{nh_n^d}} + \sqrt{h_n^{d(\nu-2)} \log n} \right) + O(h_n) \text{ p.s quand } n \rightarrow \infty.$$

2.3 Modèle 3 : cas associé

Supposons que $(X_i, Y_i)_{i \geq 1}$ sont associés. Alors le résultat suivant est établi :

Théorème 3 :

$$\sup_{x \in D} |r_n(x) - r(x)| = O \left((h_n) \vee \left(\sqrt{\frac{\log n}{nh_n^d}} \right) \right) \text{ p.s quand } n \rightarrow \infty.$$

3 Simulation

Le but des simulations suivantes est d'examiner la performance de l'estimateur défini dans l'expression (1), pour les différents modèles. Pour toutes les simulations, le taux de non censure est fixé à 0.8.

Modèle 1. On simule un n-échantillon des v.a X_i et Y_i comme suit :

$$\begin{cases} X_i \sim N(0, 1), \\ Y_i = X_i^2 + \varepsilon_i, \varepsilon_i \sim N(0.2, 1). \end{cases}$$

On simule aussi n v.a i.i.d de la censure $C_i \sim \exp(\lambda)$. Pour réaliser notre simulation, nous avons choisi le noyau gaussien et la fenêtre $h_n = 0.2$. Les graphes montrent que plus la taille n augmente plus l'estimateur $r_n(x)$ est meilleur.

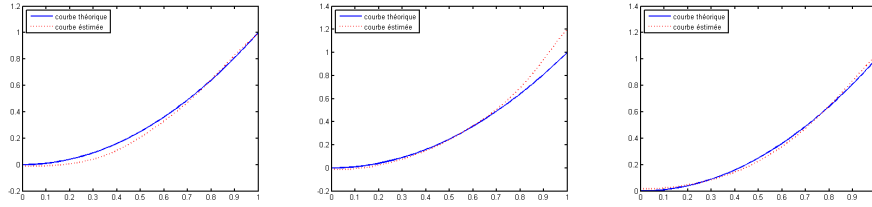


FIG. 1 – $r(x) = x^2$, $n = 100, 300, 500$.

Modèle 2 $\begin{cases} X_i = \rho X_{i-1} + \sqrt{1 - \rho^2} \varepsilon_i, \varepsilon_i \sim N(0, 1), \text{ et } \rho = 0.2 \\ Y_i = X_i^2 + \varepsilon_{1i}, \varepsilon_{1i} \sim N(0.2, 1). \text{ et } C_i \sim \exp(1.5) \end{cases}$

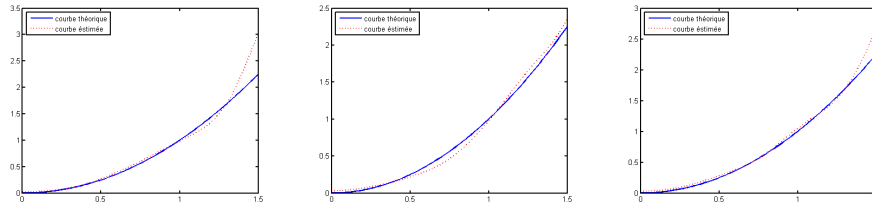


FIG. 2 – $r(x) = x^2$, $n = 100, 300, 500$.

Modèle 3 $\begin{cases} X_i = (z_{i-1} + z_{i-2})/2, z_i \text{ sont } n+1 \text{ v.a iid } \sim \exp(1), \\ Y_i = X_i^2 + \varepsilon_i, \varepsilon_i \sim N(0, 1). \text{ et } C_i \sim \exp(10) \end{cases}$

Estimateur à noyau de la fonction de régression

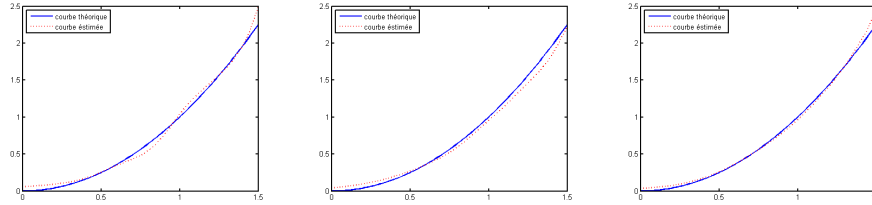


FIG. 3 – $r(x) = x^2$, $n = 100, 300, 500$.

Références

- Guessoum, Z. et E. OuldSaïd (2008). On the nonparametric estimation of the regression function under censorship model. *Statist. and Decisions* 26, 159–177.
- Guessoum, Z. et E. OuldSaïd (2010). Kernel regression uniform rate estimation for censored data under α -mixing condition. *Statist. and Decisions* 4, 117–132.
- Kaplan, E. L., M. P. (1958). Nonparametric estimation from incomplete observations. *Statist. Assoc* 53, 159–177.
- Nadaraya, E. A. (1964). On estimating regression. *Statist. and Decisions* 9, 141–142.

Summary

In this work, we studied the asymptotic behavior of a kernel regression function estimator, when the data are subject to right censoring under three models: independence, α -mixing and association cases. We presented a simulation study to illustrate and reinforce our theoretical results.

Estimation réursive à noyau de la densité des variables faiblement dépendantes

Kenza Assia Mezhoud*, Zaher Mohdeb**,
Sana Louhichi*,**

*Université Constantine 1, Département des Mathématiques, Laboratoire LMSD
m.kenzassia@yahoo.fr

**Université Constantine 1, Département des Mathématiques, Laboratoire LMSD
zaher.mohdeb@umc.edu.dz

***Laboratoire Jean Kuntzmann, UMR 5224, Tour IRMA IMG,
51 rue des Mathématiques, BP. 53 F-38041 Grenoble Cedex 9, France
sana.louhichi@imag.fr

Résumé. On établit des résultats de convergence en moyenne quadratique et de normalité asymptotique de l'estimateur réursif à noyau de la densité des variables η -faiblement dépendantes au sens généralisé de P.Doukhan et S. Louhichi.

1 Introduction

Le comportement asymptotique de l'estimateur à noyau a été largement étudié dans de nombreux travaux, Parzen et Rosenblatt (1956, 1962) ont étudié l'estimateur de la densité, Paul Doukhan et Sana Louhichi(1999) ont travaillé sur cet estimateur sous une hypothèse de dépendance faible généralisant le mélange. En revanche, lorsqu'on a à investiguer des séries temporelles, la mise à jour de l'estimation nécessite une version réursive, pour cela, on propose d'étudier des propriétés de l'estimateur réursif de la densité inconnue proposé par A. Amiri (2010) d'une suite de variables aléatoires faiblement dépendantes défini par

$$\hat{f}_n^\ell(x) = \frac{1}{\sum_{i=1}^n h_i^{d(1-\ell)}} \sum_{i=1}^n \frac{1}{h_i^{d\ell}} K\left(\frac{x - X_i}{h_i}\right), \quad (1)$$

h_n est une suite de positifs décroissant vers 0, K est appelée noyau de l'estimateur, $\ell \in [0, 1]$. Cet estimateur généralise d'autres estimateurs réursifs, tels que l'estimateur de Deheuvels(1974), l'estimateur de Wegman et Davies(1979), et celui de Wolverton et Wagner (1969).

2 Résultats

2.1 Convergence en moyenne quadratique

Proposition

Soit $(X_t)_{t \in \mathbb{Z}}$ une suite i.i.d. η -faiblement dépendantes. Pour $\ell \in [0, 1]$, on pose

$$U_{n,\ell}(N) = \frac{1}{n} \sum_{k=N+1}^{n-1} \sum_{p=1}^n \left(\frac{h_{k+p}}{h_n} \right)^{-2\ell-1} \eta_k.$$

– Si $\lim_{N \rightarrow \infty} \overline{\lim}_{n \rightarrow \infty} N^{2+\nu} U_{n,\ell}(N) = 0$, pour $\nu > 0$, alors

$$\text{Var} f_n^\ell(x) = \frac{\beta_{1-2\ell}}{nh_n \beta_{1-\ell}^2} f(x) \int K^2(x) dx + o\left(\frac{1}{nh_n}\right),$$

quand $n \rightarrow \infty$.

– Si $\lim_{N \rightarrow \infty} \overline{\lim}_{n \rightarrow \infty} N^2 U_{n,\ell}(N) < \infty$, alors

$$\text{Var} f_n^\ell(x) = \frac{\beta_{1-2\ell}}{nh_n \beta_{1-\ell}^2} f(x) \int K^2(x) dx + O\left(\frac{1}{nh_n}\right),$$

quand $n \rightarrow \infty$.

Comme la structure de la dépendance n'influe pas sur l'expression du biais obtenu par A. Amiri sous l'hypothèse de l'indépendance, la convergence en moyenne quadratique se déduit alors ainsi.

Corollaire

$$MSE(n, h_n) \sim \frac{h_n^4 \beta_{3-\ell}^2}{4\beta_{1-\ell}^2} \left(f^{(2)}(x) \int x^2 K(x) dx \right)^2 + \frac{\beta_{1-2\ell}}{nh_n \beta_{1-\ell}^2} f(x) \int K^2(x) dx.$$

2.2 Normalité asymptotique

En se basant sur la technique des blocs de Bernstein, on obtient la normalité asymptotique de notre estimateur récursif pour des processus faiblement dépendants.

Théorème

Soit $(X_t)_{t \in \mathbb{Z}}$ une suite i.i.d. η -faiblement dépendantes avec $\eta_k = O(\exp(-\lambda k))$, pour $\lambda > 0$.

On a

$$\left(\sqrt{nh_n} \left(f_n^\ell(x) - E f_n^\ell(x) \right) \right)_n$$

converges en distribution quand $n \rightarrow \infty$, vers la loi normale $\mathcal{N}(0, \sigma^2)$ avec $\sigma^2 = \frac{\beta_{1-2\ell}}{\beta_{1-\ell}^2} f(x) \int K^2(x) dx$.

Références

- Amiri, A.(2010). Estimateurs fonctionnels récursifs et leurs applications à la prévision. Thèse de Doctorat. Académie d'Aix-Marseille, Université d'Avignon et des pays de Vaucluse 1.
- Doukhan, P., Louhichi, S. (1999). A new weak dependence condition and applications to moment inequalities. *Stochastic Process. Appl.*, 84, 313-342.

Summary

We establish results of mean square error convergence and asymptotic normality of the recursive kernel density estimator of the weak dependent variables.

Sur l'évolution des performances du protocole DHCP/IP

Meriem Rachedi*, Natalia Djellab**

* Département d'informatique, Université Badji Mokhtar Annaba
meryannaba@hotmail.com

** Département de Mathématique, Université Badji Mokhtar Annaba
djellab@yahoo.fr

Résumé. Dynamic Host Configuration Protocol (DHCP) est un protocole réseau dont le rôle est d'assurer la configuration automatique des paramètres IP d'une station, notamment en lui affectant automatiquement une adresse IP et un masque de sous-réseau. DHCP peut aussi configurer l'adresse de la passerelle par défaut, des serveurs DNS (Domain Name System) et des serveurs NBNS (NetBios Name System), connus sous le nom de serveurs WINS sur les réseaux de la société Microsoft. Le propos de nos investigations consiste à évaluer la performance du mécanisme d'allocation dynamique d'adresses IP du protocole DHCP à l'aide de l'approche basée sur la modélisation mathématique ainsi que sur la simulation. A cet effet, nous présentons un modèle analytique (dont la construction s'appuie sur le formalisme de la théorie des files d'attente avec rappels) suivis par celui de simulation (dont la réalisation est assurée par le logiciel GNS3). Nous terminons notre travail par la présentation des résultats des expérimentations numériques et de simulation.

1 Introduction

Le transfert des données sur Internet se réalise à l'aide d'un ensemble de protocoles appelé TCP/IP : TCP (Transmission Control Protocol) et IP (Internet Protocol). Le protocole IP permet aux ordinateurs reliés à l'Internet de dialoguer entre eux, c'est-à-dire il se charge de l'acheminement des données sur le réseau. Comme pour un courrier postal, sur l'Internet, chaque message (paquet de données) s'accompagne de différentes informations : adresse de l'expéditeur, adresse du destinataire et, éventuellement, des données supplémentaires. Pour que les paquets parviennent à leur destination, une adresse IP unique est attribuée à chaque ordinateur connecté au réseau. Les paquets de données transitent par dizaines d'ordinateurs jusqu'à atteindre leur destinataire. Le protocole TCP se charge de la communication entre les applications, c'est-à-dire entre les logiciels utilisés par les ordinateurs. A cet effet, il vérifie que le destinataire est prêt à recevoir les données, fractionne les messages en paquets et numérote ces derniers. A la réception, le protocole en question vérifie que tous les paquets sont bien arrivés, réassemble les paquets avant de les transmettre aux logiciels et envoie des accusés de réception. Le protocole TCP/IP permet donc à des logiciels situés sur des ordinateurs différents de communiquer de façon fiable.

Évaluation des performances du protocole DHCP

La configuration des paramètres TCP/IP d'un ordinateur (notamment l'attribution d'une adresse IP, d'un masque de sous-réseau, d'un routeur, d'un domaine et d'un serveur DNS (Domain Name System)) est assurée par un protocole réseau qui est DHCP (Dynamic Host Configuration Protocol). Ce dernier permet aussi de gérer depuis un seul ordinateur toute la configuration réseau des ordinateurs d'un parc informatique. Il y a trois mécanismes d'assignation d'adresses IP : allocation automatique (attribution d'une adresse permanente à un client), allocation manuelle (adresse IP du client est attribuée par administrateur) et allocation dynamique (DHCP attribue une adresse IP au client pour une période limitée appelée bail).

2 Modèle analytique

2.1 analyse mathématique

L'état du système peut être décrit par le processus stochastique $\{I(t), J(t), t \geq 0\}$, qui est celui de Markov, d'espace d'états $S = \{0, 1, \dots, C\} \times N$. Les probabilités d'états sont :

$$p_{ij}(t) = p(I(t) = i, J(t) = j), (i, j) \in S.$$

Les taux de transition en régime stationnaire $\rho = \lambda \div c \times \mu < 1$ sont donnés par :

- pour $0 \leq i \leq c - 1$

$$q_{ij}(k, l) = \begin{cases} \lambda si(k, l) = (i + 1, l) \\ i\mu si(k, l) = (i - 1, j); \\ j\nu si(k, l) = (i + 1, j - 1) \end{cases}$$

$$q_{ij}(k, l) = -(\lambda + i\mu + j\nu) si(k, l) = (i, j)$$

0 si ailleurs

- pour $i = c$

$$q_{cj}(k, l) = \lambda si(k, l) = (c, j + 1)$$

$$c\mu si(k, l) = (c - 1, j)$$

$$-(\lambda + c\mu) si(k, l) = (c, j)$$

0 si ailleurs

Les caractéristiques importantes de la qualité de service sont :

- probabilité que toutes les adresses IP sont allouées $p_c = \lim_{t \rightarrow \infty} p(C(t) = c)$;
- nombre moyen de demandes d'adresses IP refusées n_o ;
- nombre moyen d'adresses IP allouées $\bar{c} = \lim_{t \rightarrow \infty} E[C(t)]$.

Due à l'absence des formules analytiques explicites pour les caractéristiques probabilistes principales des modèles de files d'attente avec appels à multi-serveurs, la seule méthode pour

obtenir des données numériques précises consiste à résoudre les équations de Kolmogorov numériquement . Mais dès que ce système d'équations est infini, il ne peut pas être résolu directement même sur ordinateur. Les transformations qui réduisent ces équations à une solution d'un petit problème fini dans le cas général ne sont pas disponibles. Donc, nous avons besoin d'une méthode pour approximer la solution numérique de ce système. La méthode en question est celle de troncation . A cet effet, on remplace le système de files d'attente décrit ci-dessus (où le nombre de clients en orbite est illimité) par un système similaire mais avec le nombre limité M de clients en orbite. M est le niveau de troncation. Ce dernier système a un ensemble d'équations de Kolmogorov fini. Par conséquent, il peut être résolu numériquement à l'aide d'un algorithme.

A présent, nous pouvons donner les formules des caractéristiques de la qualité de service.

Mesures de performance

- Probabilité que toutes les adresses IP sont allouées .

$$p_c^{(M)} = \sum_{j=0}^M r_{cj}^{(M)} * p_{0M}^{(M)}$$

- Nombre moyen d'adresses IP allouées

$$\bar{c}^{(M)} = \sum_{i=0}^c \sum_{j=0}^M i r_{ij}^{(M)} * p_{0M}^{(M)}$$

- Nombre moyen de demandes d'adresse IP refusées

$$\bar{n}_0^{(M)} = \sum_{i=0}^c \sum_{j=0}^M j r_{ij}^{(M)} * p_{0M}^{(M)}$$

3 Application

A présent, nous étudions l'influence de l'intensité du trafic ρ ainsi que du taux des rappels (du taux de répétition de demande pour le serveur DHCP) ν sur les valeurs numériques des

mesures de performance citées ci-dessus. A travers de ce paragraphe, nous supposons que le taux de service $\mu = 1$, le taux des rappels $\nu = 0.05$, le nombre d'adresses allouables

$c \in \{5, 15, 20\}$, Le taux d'arrivée des messages DHCPDISCOVERY , λ , est choisi de manière que le système se trouve dans un régime stationnaire (un modèle reflète le phénomène réel lorsqu'il est dans son régime stationnaire), c'est-à-dire il faut s'assurer que $\rho = \lambda \div c\mu < 1$. En outre, le niveau de troncation adopté est $M = 100$, les valeurs numériques de p_{ij} et, par conséquent, des mesures de performance subissent de très faibles changements ¹).

1. Note : les valeurs numériques sur l'axe doivent être multipliées par 10⁻¹.

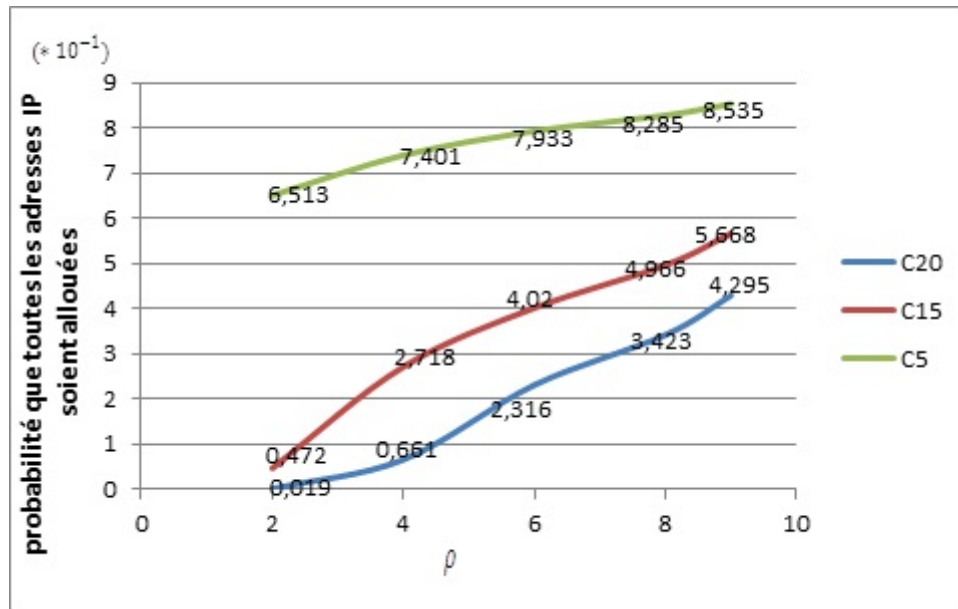


FIG. 1 – influence de l'intensité du trafic sur la probabilité que toutes les adresses ip soient allouées pour $c=15, c=5$ et $c=20$

Références

- [1] G. I. Falin, J. G. C. Templeton, Retrial Queues, Chapman Hall, London, 1997,
- [2] J. R. Artalejo, A. Gomez-Corral, Retrial Queueing Systems, Springer, 2008,
- [3] Tien Van Do, An Efficient Solution to a Retrial Queue for the Performance Evaluation of DHCP, 2007,

Summary

Dynamic Host Configuration Protocol (DHCP) is a network protocol whose role is to ensure the automatic configuration of IP station, including settings it automatically assigning an IP address and subnet mask. DHCP can also configure the address of the default gateway, DNS (Domain Name System) servers and NBNS (NetBIOS Name System) servers, known as WINS servers on networks of Microsoft Corporation. The purpose of our investigation is to evaluate the performance of the mechanism of dynamic allocation of IP addresses DHCP using the approach based on mathematical modeling and simulation. To this end, we present two analytical models (whose construction is based on the formalism of the theory of waiting lines with reminders), followed by the simulation.

Procédures d'inférence Bayésienne séquentielle appliquées aux essais cliniques

Amel Zerari *, Hayet Merabet **

*Centre Universitaire de Mila, Institut des Sciences et de la Technologie
zerari.amel@yahoo.fr

**Laboratoire de mathématiques appliquées et modélisation,
Département de Mathématiques, Université Constantine 1
merabethammadi@outlook.com

Résumé. Dans ce travail on étudie les plans d'expérience séquentiels appliqués aux essais cliniques. Dans le développement d'un nouveau médicament, c'est au cours de la phase 2 que la relation dose-réponse est évaluée et que les doses les plus prometteuses sont sélectionnées pour la phase 3. A coté du dispositif en groupes parallèles de doses, il existe des dispositifs adaptatifs visant à réduire le nombre de patients soumis aux doses les moins efficaces. En particulier lorsque la réponse est binaire (succès / échec) plusieurs dispositifs adaptatifs ont pu être proposés, qui de plus permettent que les réponses soient différées, notamment les modèles d'urne de Freedman généralisés, d'allocation linéaire, ou encore les dispositifs dits "Drop-the-Loser". Dans le cas de deux traitements, l'étude détaillée des performances fréquentistes de procédures bayésiennes non informatives pour différents dispositifs adaptatifs conduit à une conclusion favorable à ces méthodes. Nous étendrons cette étude au cas de deux traitements, situation plus usuelle dans les études de phase 2, lorsque loi a priori est uniforme, pour trouver les intervalles de crédibilité dans le même cas.

1 Introduction

Nous considérerons la situation où l'efficacité des traitements est appréciée par un critère dichotomique, la réponse étant typiquement le succès ou l'échec du traitement. La règle expérimentale communément utilisée est d'attribuer les traitements de manière aléatoire, équiprobable, aux sujets. L'objectif est de réduire le plus possible les biais expérimentaux et réduire le nombre de patients recevant le traitement le moins efficace et d'augmenter au contraire le nombre de patients recevant le traitement le plus efficace. Différentes procédures adaptatives ont été proposées, le plus souvent dans le cadre de la détermination de la dose, c'est-à-dire de la détermination de la quantité optimale de médicament à donner à un patient.

Les différentes règles sont traditionnellement modélisées comme des modèles d'urne de Friedman, 1965. Le principe est le suivant. Typiquement, on considère une urne avec T types de boules représentant les T -traitements. Au départ l'urne contient respectivement $(Y_0^1, \dots, Y_0^T)$ boules de chacun des types. Quand un nouveau sujet n est inclus, on tire avec remise une boule

dans l'urne, qui contient alors $(Y_{n-1}^1, \dots, Y_{n-1}^T)$ boules, et on attribue à ce sujet le traitement correspondant. Quand l'issue du traitement est connue, on ajoute des boules dans l'urne en fonction de la réponse observée. Par exemple, pour deux traitements, si le traitement t a été attribué on ajoute $u + v$ boules de la manière suivante : en cas de succès on ajoute u boules du type t et v boules de l'autre type, et en cas d'échec v boules du type t et u boules de l'autre type (on prend le plus souvent $u = 1$ et $v = 0$). La nouvelle composition de l'urne sera alors $(Y_n^1, \dots, Y_n^T)$, avec $\sum_{i=1}^T Y_i = n(u + v)$ (le nombre de boules dans l'urne à l'étape n est le même, quels que soient les événements passés). Différentes généralisations peuvent être envisagées, en particulier les "nombres de boules" peuvent ne pas être des entiers.

On peut donner une autre modélisation, comme un modèle d'adaptation à états. Pour chaque sujet n , on a deux événements observables, le traitement t_n qui lui est attribué et l'issue correspondante r_n . L'état de l'expérimentateur est représenté par un vecteur $z_n = (z_n^1, \dots, z_n^T)$ avec $0 \leq z_n^i \leq 1$ et $\sum z_n^i = 1$.

Différentes procédures adaptatives ont été proposées, le plus souvent dans le cadre de la détermination de la dose, c'est-à-dire de la détermination de la quantité optimale de médicament à donner à un patient.

Pour la simplification, nous présentons la situation la plus simple le cas de deux traitements et des réponses immédiates (l'issue $r_n - 1$ est connue au moment de l'attribution du traitement t_n). Pour chaque sujet n , il y a deux événements observés : (1) le traitement t_n auquel ce sujet est assigné, $t_n = t^i$ avec la probabilité z_{n-1}^i , et (2) les résultats correspondants r_n , $r_n = 1$ (succès) avec la probabilité φ_1 pour t^1 et la probabilité φ_2 pour t^2 (telle que $\varphi_t = P(r_n = 1 | t_n = t)$). Nous composons un premier état $z_0 = (z_0^1, z_0^2)$.

2 Description de quelques plans à deux traitements

2.1 La règle 'Play-The-Winner' (PTW)

La règle «Play-The-Winner» a été introduite par Zelen (1969) et ensuite étendue par Wei et Durham (1978). Conçue à l'origine pour la comparaison de deux traitements, cette méthode repose sur un processus d'adaptation « en tout ou rien » : si le sujet $n-1$ a reçu le traitement t , et si la réponse est un succès on attribue au sujet n le même traitement t , si au contraire l'issue est un échec, on attribue au sujet n l'autre traitement. Par la suite différents plans ont été proposés pour répondre à des objectifs plus générales. Il s'agit notamment de pouvoir prendre en compte le cas de réponses différées et d'étendre la méthode à plus de deux traitements (e.g : Hoel et Wei et Durham (1978), Andersen et Tamura (2003), Bai et Shen (2002), Biswas (2003)).

2.2 La règle 'Play The Winner' randomisée (PTWR)

La règle «Play The Winner» randomisée (PTWR) a été introduite par Wei et Durham (1978), comme une extension de la règle « Play The Winner» de Zelen. Ils ont modifié le modèle pour prendre explicitement en compte les réponses différées. Toujours pour deux traitements.

Quand un nouveau sujet n est inclus, l'urne contient, pour chaque traitement t , un certain nombre de boules. Une boule est tirée au hasard et est remise dans l'urne. Le sujet reçoit le traitement correspondant. Quand le résultat est connu des boules sont ajoutées à l'urne. Par

exemple, pour deux traitements, $u+v$ boules sont ajoutées : u boules pour le traitement attribué et v boules pour l'autre traitement en cas de succès ; les nombres inverses en cas d'échec. Bai, Hu et Shen (2002) ont considéré une classe générale de plans sur ce modèle. A l'examen ces modèles sont relativement peu performants, en ce sens que le processus peut converger très lentement vers la limite théorique des proportions d'attribution des traitements (Lecoutre & ElQasyr, 2008). Ceci est dû au fait que le nombre de boules dans l'urne reste constant.

2.3 Règle 'Drop-the-Loser' (DTL)

La règle « Drop-the-Loser » a été proposée par Ivanova (2003). Son principe est le suivant : si une boule de type "traitement" est tirée et si le résultat est un échec la boule est remise dans l'urne, dont la composition reste par conséquent inchangée ; si au contraire c'est un succès la boule n'est pas remise et par conséquent le nombre de boules diminue d'une unité. En outre, un nombre constant de boules de type « pas de traitement » (boules d'immigration) sont incluses dans l'urne : si une telle boule est tirée, aucun sujet n'est traité et la boule est remise avec des boules supplémentaires, une de chaque type de traitement. Par rapport aux modèles de Bai, Hu et Shen, elle a propriété intéressante de réduire la variabilité des proportions d'attribution.

2.4 Comparaison des règles

La règle 'Play The Winner' (PTW) est un procédé stochastique, depuis qu'il dépend des probabilités de succès sur chaque traitement (Lecoutre & ElQasyr, 2008).

La règle 'Play The Winner' randomisée (PTWR) converge plus lentement que la règle 'Play The Winner' (PTW) (Lecoutre & ElQasyr, 2008).

La présentation de la règle 'Drop-the-Loser' (DTL) est moins variable que la règle 'Play The Winner' (PTW), avec amélioration du pouvoir Ivanova (2003), mais la règle 'Drop-the-Loser' (DTL) est convergente plus lentement que la règle 'Play The Winner' (PTW) (Lecoutre, Derzko & ElQasyr, 2009)

3 Calcul de l'intervalle de crédibilité

Soit n_{t1} et n_{t0} les nombres respectifs de succès et d'échecs pour chacun des T traitements t . Pour tous les dispositifs adaptatifs considérés, la fonction de vraisemblance est proportionnelle à

$$\prod_{t=1}^T \varphi_t^{n_{t1}} (1 - \varphi_t)^{n_{t0}}$$

C'est-à-dire proportionnelle à la fonction de vraisemblance associée à la comparaison de proportions binomiales indépendantes. Une solution bayésienne simple et usuelle suppose des distributions *a priori* indépendantes $Bêta(v_{t1}, v_{t0})$ pour chacun des paramètres φ_t .

C'est une distribution conjuguée et les distributions *a posteriori* marginales sont encore des $Bêta$ indépendantes : $Bêta(v_{t1} + n_{t1}, v_{t0} + n_{t0})$. L'approche bayésienne permet d'obtenir, à partir de la distribution *a posteriori* conjointe, la distribution de tout paramètre dérivé.

Celle-ci peut être facilement approximée en simulant des distributions $Bêta$ indépendantes.

Une solution simple raisonnable pour une distribution *a priori* objective est de prendre pour chaque traitement une $B\hat{e}ta(1, 1)$. Dans le cas particulier de deux traitements, on peut ainsi obtenir des Intervalles de crédibilité Bayésienne sur le rapport, la différence ou l'odds-ratio.

3.1 Loi a posteriori de la différence de deux lois Bêta de paramètres entiers

Pham-Gia et Turkkan (1993) ont donné la loi a posteriori de $\varphi = \varphi_1 - \varphi_2$ conditionnelle au données observées $D = (v_{11}, v_{10}, v_{21}, v_{20})$

$$P(\varphi|D) = \begin{cases} \frac{B(p_2, q_1)}{A} \varphi^{q_1+q_2-1} (1-\varphi)^{p_2+q_1-1} F_1(q_1, p_1+p_2+q_1+q_2-2, 1-p_1; \\ p_2+q_1; 1-\varphi, 1-\varphi^2) \text{ pour } 0 \leq \varphi \leq 1 \\ \frac{B(p_2, q_1)}{A} (-\varphi)^{q_1+q_2-1} (1+\varphi)^{p_1+q_2-1} F_1(q_2, p_1+p_2+q_1+q_2, 1-p_2; \\ p_1+q_2; 1+\varphi, 1+\varphi^2) \text{ pour } -1 \leq \varphi \leq 0 \end{cases}$$

Où

$$\begin{cases} p_1 = v_{11} + n_{11} , q_1 = v_{10} + n_{10} \\ p_2 = v_{21} + n_{21} , q_2 = v_{20} + n_{20} \\ A = B(p_1, q_1)B(p_2, q_2) , B(p, q) = \frac{\Gamma(p)\Gamma(q)}{\Gamma(p+q)} \end{cases}$$

et F_1 fonction hypergéométrique.

$$F_1 = (a, b_1, b_2; c; x_1, x_2) = \sum_{i \geq 0} \sum_{j \geq 0} \frac{a^{[i+j]} b_1^{[i]} b_2^{[j]} x_1^i x_2^j}{c^{[i+j]} i! j!} \text{ et } a^{[b]} = a(a+1)\dots(a+b-1)$$

Un intervalle de crédibilité $100(1-\gamma)\%$ pour $[C_1^{(\varphi)}, C_2^{(\varphi)}]$ peut être obtenu en déterminant $C_1^{(\varphi)}$ et $C_2^{(\varphi)}$ telles que :

$$\int_{C_2^{(\varphi)}}^{C_1^{(\varphi)}} P(\varphi|D) d\varphi = 1 - \gamma$$

3.2 Loi a posteriori du rapport de deux lois Bêta de paramètres entiers

En utilisant la transformation de Mellin, Steece (1976) a donné la densité du rapport $W = \varphi_1/\varphi_2$

$$F_W(w|D) = \frac{\Gamma(v_{11} + v_{10} + n_1)}{\Gamma(v_{11} + n_{11})\Gamma(v_{10} + n_{10})} \frac{\Gamma(v_{21} + v_{20} + n_2)}{\Gamma(v_{21} + n_{21})\Gamma(v_{20} + n_{20})} \times \sum_{i=0}^{v_{10}+n_{10}-1} \sum_{j=0}^{v_{20}+n_{20}-1} (-1)^{i+j} \binom{v_{10}+n_{10}-1}{i} \binom{v_{20}+n_{20}-1}{j} \times J_w(v_{11} + n_{11} + i - 1, v_{21} + n_{21} + j - 1)$$

Où

$$J_w(c_1, c_2) = \iint_{\frac{\varphi_1}{\varphi_2} \leq w} \varphi_1^{c_1} \varphi_2^{c_2} d\varphi_1 d\varphi_2$$

Un intervalle de crédibilité $100(1-\gamma)\%$ pour $[C_1^{(w)}, C_2^{(w)}]$ peut être obtenu en déterminant $C_1^{(w)}$ et $C_2^{(w)}$ telles que :

$$\int_{C_2^{(w)}}^{C_1^{(w)}} F_W(w|D) d\varphi = 1 - \gamma$$

3.3 Loi a posteriori de l'odds ratio de deux lois Bêta de paramètres entiers

En utilisant les formules précédentes, et on pose :

Un intervalle de crédibilité $100(1 - \gamma)\%$ pour X/Y peut être obtenus comme pour le rapport φ_1/φ_2 et la différence $\varphi_1 - \varphi_2$ précédentes.

Cadre de simulations Considérons un nombre fixé des sujets par exemple $N = 150$ sujets

($N = 150$, simulation de 100000 échantillons). Les nombres des échecs (ici 20 et 21). Une solution simple et raisonnable pour un objectif *a priori* est de prendre deux marginales indépendants $Bêta(1, 1)$ (c.-à-d $\varphi_1 \sim Bêta(75, 21)$ et $\varphi_2 \sim Bêta(36, 22)$). Les pourcentages d'erreurs pour la limite à 95%. Taux de couverture de l'intervalle de crédibilité Bayésienne en utilisant le logiciel *LEPAC*.

La transition de probabilité pour le modèle avec deux traitements décrits ci-dessus est donnée dans les courbes suivants :

$$[0.013, 0.310] \text{ pour } \varphi_1 - \varphi_2$$

$$[1.018, 1.613] \text{ pour } \varphi_1/\varphi_2$$

$$[1.068, 4.561] \text{ pour } \varphi_1(1 - \varphi_2)/\varphi_2(1 - \varphi_1)$$

3.4 Discussion

Une application importante du calcul de l'intervalle de crédibilité dans les essais cliniques est qu'elle peut être employée pour distributions initiales Bêta indépendantes. Par exemple, nous pouvons nous appliquer distributions uniforme sur $[0, 1]$ pour deux groupes indépendants ($Bêta(1, 1)$), et distribution initiale de Jeffreys pour deux groupes indépendants ($Bêta(0.5, 0.5)$) (ElQasyr, 2008). De plus il est facile de considérer les principaux critères classiques pour comparer proportions (différence, rapport et odds ratio) et la différence entre la distribution uniforme et loi de Jeffreys, (les intervalles de confiance 95%).

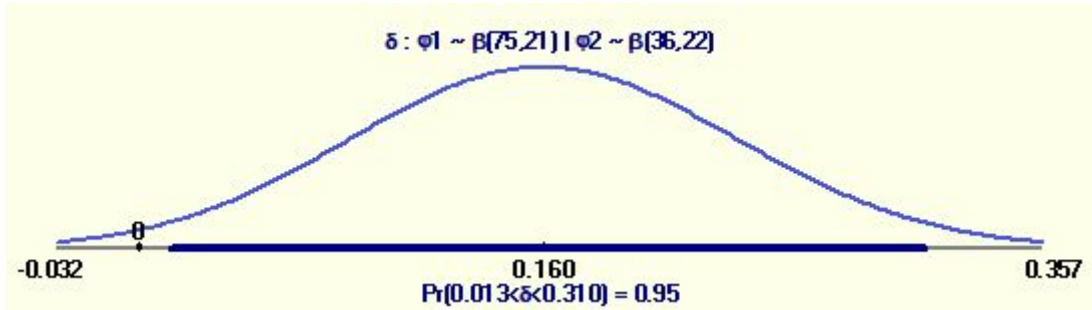


FIG. 1 – Taux de couverture de l'intervalle des crédibilité bayésien : Pourcentages d'erreurs pour le limite a 95 ($N = 150$, simulation de 100000 échantillons) pour $\varphi_1 - \varphi_2$

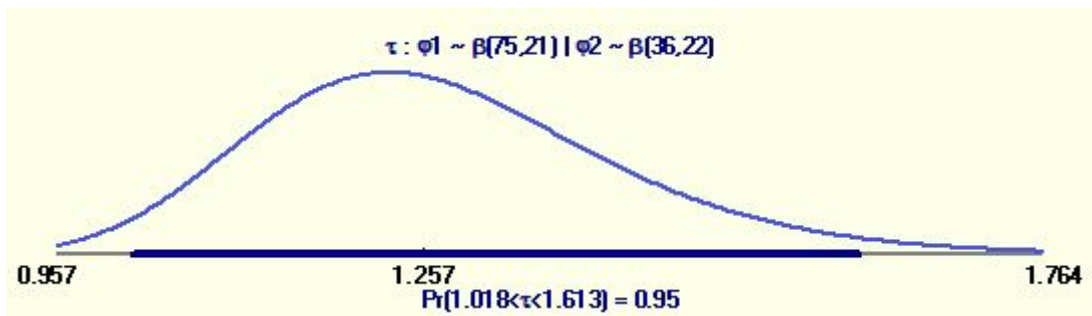


FIG. 2 – Taux de couverture de l'intervalle des crédibilité bayésien : Pourcentages d'erreurs pour le limite a 95 ($N = 150$, simulation de 100000 échantillons) pour φ_1 / φ_2

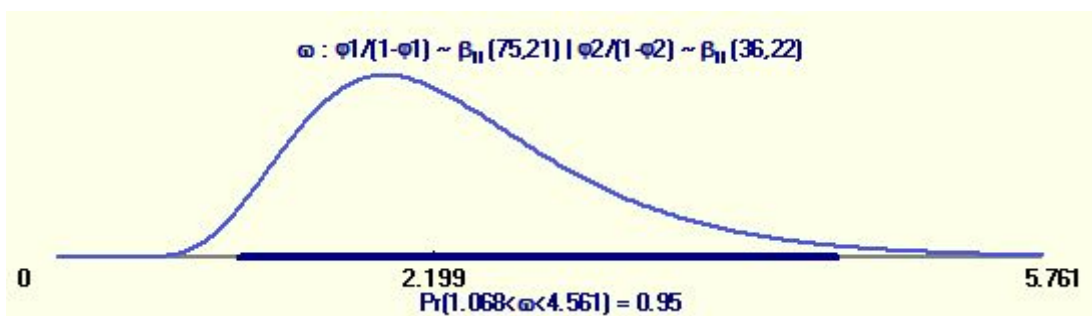


FIG. 3 – Taux de couverture de l'intervalle des crédibilité bayésien : Pourcentages d'erreurs pour le limite a 95 ($N = 150$, simulation de 100000 échantillons) pour $\varphi_1(1 - \varphi_2) / \varphi_2$

Pour la différence $\varphi_1 - \varphi_2$ les intervalles de crédibilité très proche et très étroite, ne pas trouver des intervalles symétriques : $[0.013, 0.310]$ et $[0.013, 0.312]$ (pour la loi uniforme respective la loi de Jeffreys). Pour le rapport φ_1/φ_2 les intervalles de confiance plus petits les deux et très proche : $[1.018, 1.613]$ et $[1.02, 1.62]$ respectivement. Si on a la curiosité de calculer ces mêmes intervalles pour le odds ratio $\varphi_1(1 - \varphi_2)/\varphi_2(1 - \varphi_1)$, on a la surprise de ne pas trouver des intervalles symétriques, mais au contraire des intervalles largement différents : respectivement $[1.068, 4.561]$ et $[1.07, 4.64]$.

4 Conclusion

Les plans d'expériences séquentiels sont explicitement intéressants puisqu'ils permettent, pour des raisons éthiques, d'exposer le moins de sujets possibles aux traitements les moins efficaces. Nous avons développé des méthodes d'inférence bayésienne pour analyser les données de différents plans. Dans le cas de deux traitements, nous avons déterminé les intervalles de crédibilité pour une a priori $B\hat{e}ta(1, 1)$ et comparé les résultats avec les résultats antérieurs.

Références

- B. Lecoutre, G. Derzko, J. G. (1995). Bayesian predictive approach for inference about proportions. *Statistics in Medecine* 14, 1057–1063.
- B. Lecoutre, K. E. (2008). Adaptative designs for multi-arm clinical trials : The playthe- winner rule revisited. *Communications in Statistics - Simulation and Computation* 37, 590–601.
- B. Lecoutre, G. Derzko, K. E. (2009). Frequentist performance of bayesian inference with response-adaptive designs. *Soumis pour publication*.
- B. Lecoutre, G. Derzko, K. E. (2010). Frequentist performance of bayesian inference with response-adaptive designs. *Soumis pour publication*.
- Biswas, A. (2003). Generalized delayed response in randomized play-the-winner rule communications in statistics. *Simulation and Computation* 32(1).
- ELQasyr, K. (2008). *Modélisation et analyse statistique des plans d'expérience séquentiels*. Thèse de doctorat, Université de Rouen.
- Freedman, D. (1965). Bernard friedman's urn. *Annals of Mathematical Statistics* 36, 956–970.
- Hölldobler, B. et E.-O. Wilson (1990). *The Ants*. Berlin: Springer Verlag.
- Ivanova, A. (2003). A play-the-winner-type urn design with reduced variability. *Metrika* 58, 1–13.
- L. J. Wei, S. D. (1978). The randomized play-the-winner rule in medical trial. *Journal of the American Statistical Association* 73, 840–843.
- R. Sun, S. H. Cheung, L. Z. (2007). A generalized drop-the-loser rule for multitreatment clinical trials. *Journal of Statistical Planing and inference* 137, 2011–2023.
- Z.D. Bai, F. H. (2005). Asymptotics in randomized urn models. *The Annals of Applied Probability* 15, 914–940.

- Z.D. Bai, F. Hu, L. S. (2002a). An adaptive design for multi-arm clinical trials. *Journal of Multivariate Analysis* 81, 1–18.
- Z.D. Bai, F. Hu, R. W. (2002b). Asymptotic properties of adaptive designs for clinical trials with delayed response. *The Annals of Statistics* 30, 122–139.
- Zelen, M. (1969). Play the winner rule and the controlled clinical trial. *Journal of the American Statistical Association* 64, 131–146.

Summary

On with this work one studies the sequential experimental designs applied to the clinical trials. In the development of a new drug, it is during phase 2 that the relation amount-answer is evaluated and that the most promising amounts are selected for phase 3. With dimensions of the device in parallel groups of amounts, there are adaptive devices aiming at reducing the number of patients subjected to the least effective amounts. In particular when the answer is binary (success/failure) several adaptive devices could be proposed, which moreover allow that the answers are differed, in particular the models of ballot box of Freedman generalized, of linear allowance, or devices known as "Drop-the-Loser". In the case of two treatments, the detailed study of the performances frequentists of noninformative procedures bayésiennes for various adaptive devices leads to a conclusion favorable to these methods. We will extend this study to the case of more than two treatments, situation more usual in the studies of phase 2, when law apriori is uniform, to find the intervals of credibility in the same case.

Sur l'étude de l'effet d'allongement des vecteurs dans le cas linéaire

Dalel ZERDAZI*, Ahmed CHIBAT**

*Rue Ain Elbay, Constantine, Algérie
sky_stats@hotmail.com

**Rue Ain Elbay, Constantine, Algérie
achibat@gmail.com

Résumé. Lorsqu'il existe une forte multi colinéarité entre les variables explicatives, l'estimateur des moindres carrés ordinaires MCO, bien qu'il soit le meilleur parmi ceux sans biais, souffre d'une difficulté énorme qui est celle de l'instabilité. C'est cet écueil qui a conduit à l'élaboration d'autres estimateurs, nécessairement biaisés, mais dont l'erreur quadratique moyenne est inférieure à celle de l'estimateur des MCO. L'intérêt est d'aboutir à une bonne qualité de généralisation même s'il faut pour cela concéder un petit biais.

1 Introduction

L'une parmi les questions les plus importantes qui continuent de se poser avec acuité dans le domaine de la régression et de la modélisation, linéaire ou non linéaire, est celle concernant la qualité de la généralisation.

Cette question reste d'une grande actualité, tout aussi bien dans le cadre des techniques de la statistique traditionnelle, que dans le cadre des techniques récentes telles que les réseaux de neurones.

La qualité des estimateurs construits est tributaire des propriétés des données qui conduisent à leur construction.

Nous entreprenons actuellement la démonstration de l'équivalence formelle (ou du moins fonctionnelle) entre les méthodes conçus pour l'amélioration de la généralisation par les réseaux de neurones (telles que la régularisation, l'arrêt prématuré, l'OBD) d'une part, et d'autre part les estimateurs raccourcis dans le cadre du modèle linéaire (connus comme étant meilleurs que l'estimateur des moindres carrés en terme d'erreur quadratique moyenne).

La question qui se pose est de voir s'il existe des liens formels, même restreints au modèle linéaire, entre ces techniques.

Les deux problématiques semblent être différentes. Cependant les solutions préconisées pour les deux cas présentent certaines similitudes dignes d'intérêt et qui demandent d'être analysées en profondeur.

Les estimateurs conçus dans le cadre des techniques de la statistique traditionnelle, à savoir :

Allongement des vecteurs dans le cas linéaire

L'estimateur Ridge de Hoerl et Kennard

L'estimateur par inverse généralisée de Marquardt

L'estimateur de James- Stein

sont tous des estimateurs raccourcis, dans le sens où la norme du vecteur des paramètres est plus courte que celle des MCO.

De leur côté également, les techniques mises à contribution dans le cadre des réseaux de neurones, à savoir :

La régularisation

Et l'arrêt prématuré

Conduisent aussi à un raccourcissement du vecteur des paramètres.

La recherche actuelle a touché la question de l'influence de la norme des vecteurs en cas linéaire, et le problème reste tout à fait ouvert. Notre avancée dans son traitement est satisfaisante pour l'heure et nous a produit des résultats originaux.

2 Méthodologie de travail

La première hypothèse à vérifier est : l'influence de l'étendue de l'erreur sur le coefficient de la variable. On veut dire que, puisque les variables les plus fortes ont leurs coefficients quasiment exacts, alors on va augmenter l'erreur jusqu'à voir les coefficients des plus fortes variables devenir perturbés eux aussi.

On va ensuite voir s'il y a une relation entre la norme de chaque variable et le niveau de perturbation.

Ensuite, On a divisé les variables (une à une ou toutes en même temps) par leurs normes et voir si on gagne en précision.

On a entraîné le réseaux avec

- 10 aléa
- 100 aléa
- 1000 aléa
- 10.000 aléa

Nous traiterons plus en détail le second processus de réduction de la dimension du vecteur des entrées : l'élimination des entrées dont l'influence peut légitimement être négligée, car elle est inférieure à l'influence du bruit ou à celle des perturbations non mesurées.

2.1 Résumé de toute l'expérimentation

Il y a trois choses à surveiller :

- La performance
- Les valeurs des paramètres
- La norme du vecteur des paramètres

L'influence sur ces trois choses est provoquée par :

- Le nombre d'observations
- L'étendue de l'erreur

-La platitude de la fonction (qui traduit la prépondérance de certaines variables par rapport à d'autres). Cette platitude est réglée par le domaine de variation de la variable x .

2.2 Synthèse de ce qui s'est produit

Après notre étude de simulation affectée sur 180 fichiers et 1000 observations .. On a ces résultats :

- Il y a l'influence de la taille de l'échantillon.
- Il y a l'influence de la taille de l'erreur
- Il y a l'influence du domaine de variation

Lorsque le nombre d'observations est 1000, il y a comme une solidité du modèle : les variables prépondérantes ne sont touchées que lorsque la taille de l'erreur est excessive. La performance bien qu'inférieure à la variance de l'échantillon elle lui est légèrement différente. Le nombre d'itérations est raisonnable. La norme du vecteur des paramètres reste petite (inférieure ou très voisine du vecteur vrai) même à 100 fois aléa.

Il faut signaler que le surajustement existe toujours, mais qu'il ne s'accompagne pas nécessairement d'un allongement du vecteur des paramètres : Il faut que cela s'accompagne d'un certains nombre de conditions.

Il faudrait donc être prudent dans l'application de l'algorithme de régularisation.

Nous connaissons la taille de l'échantillon.

Nous pouvons vérifier aussi la prépondérance des variables.

Mais nous n'avons pas une idée très claire de la taille de l'erreur.

Peut être qu'en supprimant les variables les moins prépondérantes.. cela aurait comme conséquence un raccourcissement substantiel du vecteur des paramètres.

Ceci est à mettre en parallèle avec la fonction sinusoïdale, lorsque nous devons réduire le domaine de variation, et supprimer la partie où le ratio signal-bruit est mauvais.

Donc Lorsque le nombre d'observations est conséquent (1000 ou 10.000), le réseau a tendance à estimer très correctement. Mais lorsque le nombre d'observations est petit (100), le réseau a tendance à mesestimer, surtout dans le sens du surajustement.

Références

- [1] Engelbrecht, A.P. A (2001). *A new pruning heuristic based on variance analysis of sensitivity information*. IEEE Transactions on Neural Networks,
- [2] Hagiwara, K (2002). *Regularization learning, early stopping and biased estimator*. Neurocomputing, 48, 937.
- [3] Hoerl, A. E., and R. W. Kennard. *Ridge Regression : Applications to Nonorthogonal Problems*. Technometrics. Vol. 12, No. 1, 1970, pp. 69-82.
- [4] Hoerl, A. E., and R. W. Kennard (1970). *Ridge Regression : Biased Estimation for Nonorthogonal Problems*. Technometrics. Vol. 12, No. 1, pp. 55-67

Allongement des vecteurs dans le cas linéaire

- [5] James, W. ; Stein, C. (1961). *Estimation with quadratic loss*. Proc. Fourth Berkeley Symp. Math. Statist. Prob., 1, pp. 361-379
- [6] Le Cun, Y., Denker, J.S., and Solla, S.A (1990). *Optimal brain damage, Advances in Neural Information Processing*. Vol. 2, D.S. Touretzky, ed., Morgan Kaufmann Publishers, Los Altos, CA, 598.
- [7] Marquardt, D.W.(1970). *Generalized Inverses, Ridge Regression, Biased Linear Estimation, and Nonlinear Estimation*. Technometrics. Vol. 12, No. 3, , pp. 591-612.

Annexe

On a eu des résultats originaux de cette étude à l'aide du logiciel "Matlab"

Summary

When there exists a strong multi colinearity between the explanatory variables, the estimator of ordinary least squares MCO, although it is the best among those without biais, suffers from an enormous difficulty which is that of instability. This is the cause which led to the development of other estimators, necessarily biased, but whose average quadratic error is lower than that of the estimator of the MCO. The interest is to lead to a good quality of generalization even if it is necessary for that to concede a small biais.

Index des Auteurs

Aïssani, D., 34, 42, 50, 54, 58, 73, 119, 139
Abbas, K., 77
Abi-ayad, I., 30
Adjabi, S., 58, 81, 119, 123
Aidi, K., 129
Ait Amokhtar, S., 86
Aknine, A., 133
Alem, L.M., 139
Ameraoui, A., 90
Attouch, M.K., 16, 164

Bareche, A., 34
Bedouhene, F., 38
Belabed, F.Z., 16
Belhadj, S.A., 19
Benmakrelouf, R., 143, 196
Benmoussat, N., 19
Benouaret, Z., 42, 119
Bensalloua, M., 147
Benyahia, A., 46
Benyahia, W., 151
Berdjoudj, L., 50, 173
Berkoun, Y., 160
Berrouane, S., 155
Bouabça, W., 164
Boualem, M., 54, 58, 139, 173
Bouderbala, I., 96
Boukhetala, K., 90, 108, 147, 168, 189
Braham, H., 173

Célestin C, K., 123
Challali, N., 38
Cherfaoui, M., 54, 58
Chibat, A., 227
Coquet, F., 5

Dali-Korso, M., 46, 176
Dassamiour, M., 100
Della Krachai, M., 19
Dhiabi, S., 181
Diop, A., 7
Djedour, M., 143, 196
Djellab, N., 54, 215
Djeridi, Z., 105

Dupuy, J.F., 7, 90

EL Saadi, N., 86

Gamboa, F., 10
Graiche, F., 24
Guerbyenne, H., 62, 66
Guessoum, Z., 185
Guidoum, A.C., 108

Hamadouche, D., 24
Hamrani, F., 185
Hanafi, M., 12

Ikhlef, L., 73

Kadem, S., 189
Kara Terki, N., 69
Karouche, W., 143, 196
Kessira, A., 66
Khellouf, N., 168
Kherchi, H., 96

Lekadir, O., 73
Lessak, R., 200
Louhichi, S., 212

Madani, D., 204
Mellah, O., 38
Menni, N., 208
Merabet, D., 24
Merabet, H., 105, 219
Messahli, L., 62
Mezghache, H., 100
Mezhoud, K.A., 212
Mohdeb, Z., 115, 212
Mourid, T., 30, 69, 151

Ndao, P., 7
Nouali, K., 133

Rachedi, M., 215
Raynaud De Fitte, P., 38
Rynkiewicz, J., 143, 196

Sadki, O., [181](#)
Saidi, G., [96](#)
Seddik-Ameur, N., [129](#)

Takhedmit, B., [77](#)
Tatachak, A., [204](#), [208](#)
Touazi, A., [42](#), [119](#)

Zerari, A., [219](#)
Zerdazi, D., [227](#)
Ziane, Y., [81](#)
Zougab, N., [81](#), [123](#)

Sponsorisé par:

