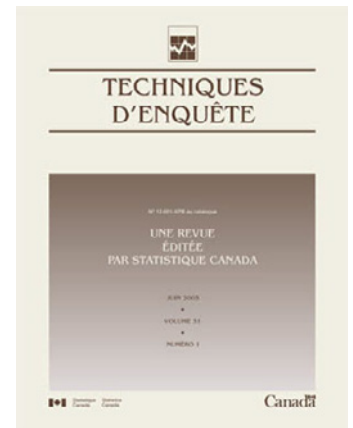


Techniques d'enquête

Réconciliation bayésienne dans le modèle de Fay-Herriot par suppression aléatoire

par Balgobin Nandram, Andreea L. Erciulescu et Nathan B. Cruze

Date de diffusion : le 27 juin 2019



Statistique
Canada

Statistics
Canada

Canada

Comment obtenir d'autres renseignements

Pour toute demande de renseignements au sujet de ce produit ou sur l'ensemble des données et des services de Statistique Canada, visiter notre site Web à www.statcan.gc.ca.

Vous pouvez également communiquer avec nous par :

Courriel à STATCAN.infostats-infostats.STATCAN@canada.ca

Téléphone entre 8 h 30 et 16 h 30 du lundi au vendredi aux numéros suivants :

- | | |
|---|----------------|
| • Service de renseignements statistiques | 1-800-263-1136 |
| • Service national d'appareils de télécommunications pour les malentendants | 1-800-363-7629 |
| • Télécopieur | 1-514-283-9350 |

Programme des services de dépôt

- | | |
|-----------------------------|----------------|
| • Service de renseignements | 1-800-635-7943 |
| • Télécopieur | 1-800-565-7757 |

Normes de service à la clientèle

Statistique Canada s'engage à fournir à ses clients des services rapides, fiables et courtois. À cet égard, notre organisme s'est doté de normes de service à la clientèle que les employés observent. Pour obtenir une copie de ces normes de service, veuillez communiquer avec Statistique Canada au numéro sans frais 1-800-263-1136. Les normes de service sont aussi publiées sur le site www.statcan.gc.ca sous « Contactez-nous » > « [Normes de service à la clientèle](#) ».

Note de reconnaissance

Le succès du système statistique du Canada repose sur un partenariat bien établi entre Statistique Canada et la population du Canada, les entreprises, les administrations et les autres organismes. Sans cette collaboration et cette bonne volonté, il serait impossible de produire des statistiques exactes et actuelles.

Publication autorisée par le ministre responsable de Statistique Canada

© Sa Majesté la Reine du chef du Canada, représentée par le ministre de l'Industrie 2019

Tous droits réservés. L'utilisation de la présente publication est assujettie aux modalités de l'[entente de licence ouverte](#) de Statistique Canada.

Une [version HTML](#) est aussi disponible.

This publication is also available in English.

Réconciliation bayésienne dans le modèle de Fay-Herriot par suppression aléatoire

Balgobin Nandram, Andreea L. Erciulescu et Nathan B. Cruze¹

Résumé

La réconciliation d'estimations de niveau inférieur à des estimations de niveau supérieur est une activité importante au *National Agricultural Statistics Service* (NASS) du département de l'Agriculture des États-Unis (par exemple, réconcilier les estimations de superficie d'ensemencement en maïs des comtés aux estimations au niveau des États). Nous posons qu'un comté est un petit domaine et employons le modèle initial de Fay-Herriot pour obtenir une méthode bayésienne générale pour réconcilier les estimations des comtés aux estimations des États (constituant la cible). Dans ce cas, nous supposons que les estimations cibles sont connues et dégageons les estimations des comtés avec pour contrainte que leur addition donne la valeur cible. C'est là une réconciliation externe qui a de l'importance pour la statistique officielle, et non seulement pour les données du NASS, et on le rencontre plus généralement dans les estimations sur petits domaines. Il est possible de réconcilier de telles estimations en « supprimant » un des comtés (habituellement le dernier) de manière à intégrer la contrainte de réconciliation au modèle. Il est tout aussi vrai cependant que les estimations peuvent changer selon le comté qui est supprimé au moment d'inclure la contrainte dans le modèle. Dans la présente étude, nous accordons à chaque petit domaine une chance de suppression et parlons pour toute cette procédure de méthode de réconciliation par suppression aléatoire. Nous démontrons empiriquement que les estimations accusent des différences selon le comté supprimé et qu'il existe des différences entre ces estimations et celles obtenues par suppression aléatoire. Ces différences peuvent être jugées petites, mais il est hautement logique de procéder par suppression aléatoire; aucun comté n'a alors droit à un traitement préférentiel et nous observons également une modeste hausse de la précision par rapport à une réconciliation avec suppression du dernier petit domaine.

Mots-clés : Contrainte; estimations directes; modèle de Fay-Herriot; densité normale multivariée; statistique officielle; estimation sur petits domaines.

1 Introduction

Dans la statistique officielle, il importe que l'addition des estimations de niveau inférieur donne les estimations de niveau supérieur. Ainsi, le *National Agricultural Statistics Service* (NASS) va fréquemment de haut en bas dans la diffusion de ses estimations officielles, car il publie les estimations du pays et des États (estimations de superficie totale ensemencée en maïs, par exemple) avant de mener à bien la collecte de données supplémentaires et de calculer les estimations correspondantes des comtés (Cruze, Erciulescu, Nandram, Barboza et Young, 2019). Dans ces petits domaines administratifs, les données d'enquête sont souvent clairsemées. Plusieurs techniques prisées de modélisation permettent d'obtenir des estimations sur petits domaines plus fiables, mais celles-ci peuvent ne pas concorder automatiquement avec les estimations produites à des niveaux supérieurs d'agrégation, et on peut appliquer des procédures de réconciliation pour imposer cette concordance.

Les techniques de réconciliation qui ont servi à forcer une concordance entre niveaux multiples et à protéger contre toute mauvaise spécification des modèles ont déjà une histoire considérable. Elles appartiennent à deux catégories : réconciliation interne où une cible est tirée des données courantes

1. Balgobin Nandram, Worcester Polytechnic Institute et USDA National Agricultural Statistics Service, Department of Mathematical Sciences, Stratton Hall, 100 Institute Road, Worcester, MA 01609. Courriel : balnan@wpi.edu; Andreea L. Erciulescu, Westat, 1600 Research Boulevard, Rockville, MD 20850. Courriel : alerciulescu@gmail.com; Nathan B. Cruze, USDA National Agricultural Statistics Service, 1400 Independence Avenue, SW, Room 6412 A, Washington, D.C. 20250-2054. Courriel : nathan.cruze@nass.usda.gov.

d'enquête; réconciliation externe où la cible peut être puisée à d'autres sources comme des données administratives ou des estimations déjà établies. Il sera ici question de réconciliation externe selon la procédure « de haut en bas » du NASS avec le modèle de Fay-Herriot (FH) (Fay et Herriot, 1979).

Le plus récent examen consacré à l'estimation sur petits domaines est de Pfeiffermann (2013); le manuel le plus à jour sur cette estimation est de Rao et Molina (2015). Plus tôt, Jiang et Lahiri (2006) avaient présenté un examen poussé de la démarche inférentielle classique pour les modèles mixtes linéaires et linéaires généralisés servant à l'estimation sur petits domaines. Ces documents traitent de réconciliation aussi, mais le dernier exposé ne portait pas sur le système hiérarchisé de Bayes qui nous intéresse avant tout ici.

Avec ce cadre hiérarchique, You, Rao et Dick (2004) ont étudié les estimateurs réconciliés sur petits domaines en se fondant sur des modèles d'échantillonnage et de couplage sans appariement qui avaient déjà été proposés par You et Rao (2002). Ils ont appliqué ce cadre à l'estimation de sous-couverture pour les 10 provinces canadiennes dans le cadre du recensement du Canada de 1991. Wang, Fuller et Qu (2008) ont livré une caractérisation du meilleur prédicteur linéaire sans biais (BLUP) des moyennes de petits domaines dans un modèle au niveau des domaines qui satisfait à une contrainte de réconciliation et minimise le critère de fonction de perte auquel se conforment tous les prédicteurs linéaires sans biais. Ils ont également proposé une autre façon d'imposer la contrainte de réconciliation de sorte que l'estimateur BLUP offre une propriété d'autocalage (voir You et Rao, 2002). Wang et coll. (2008) ont caractérisé une catégorie d'estimateurs réconciliés comme étant les prédicteurs qui minimisent une fonction de perte quadratique sous une restriction de réconciliation. Le modèle augmenté d'autocalage qu'ils avancent vient réduire le biais tant dans l'ensemble qu'au niveau des petits domaines. On peut trouver d'autres procédures de réconciliation dans Bell, Datta et Ghosh (2013), Ghosh et Steorts (2013), Pfeiffermann, Sikov et Tiller (2014) et Pfeiffermann et Tiller (2006).

Que l'on ajuste des modèles au niveau des unités ou des domaines, intégrer une cible externe fixe se ramène à imposer la contrainte générale $\sum_{i=1}^{\ell} w_i \theta_i = a$, où a est une constante connue et où θ_i désigne les estimations sur petits domaines à réconcilier; pour les totaux, les poids w_i sont tous égaux à 1. Une façon de le faire est d'utiliser la transformation $\phi = a - \sum_{i=1}^{\ell} \theta_i$, en gardant θ_i , $i = 1, \dots, \ell - 1$ inchangé et en « supprimant » le dernier petit domaine pour le remplacer par $\theta_{\ell} = \phi - \left(a - \sum_{i=1}^{\ell-1} \theta_i\right)$. Janicki et Vesper (2017) ont conçu une transformation légèrement différente, $\phi_i = \theta_i$, $i = 1, \dots, \ell - 1$, $\phi_{\ell} = \sum_{i=1}^{\ell} \theta_i$, qui consiste pour l'essentiel en une réconciliation interne qui préserve la somme de toutes les estimations ℓ . Si cette somme (de toutes les estimations ℓ non réconciliées) est prescrite comme cible externe, $\theta_{\ell} = \phi_{\ell} - \sum_{i=1}^{\ell-1} \theta_i$, et la transformation de Janicki et Vesper équivaut alors à supprimer l'estimation du dernier petit domaine.

Nandram et Sayit (2011) et Nandram, Berg et Barboza (2014) ont étudié les procédures de réconciliation externe avec suppression de l'estimation du dernier petit domaine à des fins de réconciliation de prévisions et de probabilités binomiales de rendements cultureux respectivement. (Dans ce double contexte, la contrainte visait en réalité la somme pondérée des estimations sur petits domaines.) Erciulescu, Cruze et Nandram (2019) se sont attachés à diverses techniques de réconciliation externe : réconciliation avec suppression, réconciliation par les différences ou les rapports, etc., dans le contexte des modèles

hiérarchiques bayésiens de petit domaine. Collectivement, la contrainte de réconciliation externe a été introduite dans la fonction de vraisemblance (voir Toto et Nandram (2010), par exemple), la densité conjointe des effets de domaine (voir Nandram et Sayit (2011), par exemple) ou la densité a posteriori de ces mêmes effets (Janicki et Vesper, 2017). Précisons que, dans ce dernier cas, on utilise les connaissances antérieures ou les exigences qu'incarne la contrainte a posteriori plutôt que les distributions antérieures elles-mêmes.

Datta, Ghosh, Steorts et Maples (2011) ou DGSM ont proposé une catégorie générale d'estimateurs de Bayes avec contrainte pour la production d'estimations réconciliées. À propos de la méthode de Toto et Nandram (2010) en particulier pour les modèles au niveau des unités, DGSM ont écrit : « *Un inconvénient avec cette méthode est que les résultats peuvent varier selon l'unité qui est retranchée* [TRADUCTION]. » Cette remarque vaut aussi pour Nandram et Toto (2010), Nandram, Toto et Choi (2011), Janicki et Vesper (2017) et d'autres, tout comme pour un modèle au niveau des domaines avec contrainte externe. Les procédures de DGSM dépendent d'un important paramètre à spécificité de domaine (voir la section 4) et qui reçoit plusieurs spécifications distinctes. On pourrait soutenir que les estimations résultantes pourraient subir l'effet du choix de la spécification. De plus, les procédures de DGSM ne présentent ni erreurs-types a posteriori ni intervalles de crédibilité.

En réaction à cette observation de DGSM au sujet de la réconciliation avec suppression du dernier domaine, nous présentons une réconciliation à suppression aléatoire où chaque domaine a des chances d'être retranché, et non le dernier seulement. La justification mathématique de cette méthode de réconciliation à suppression aléatoire figure à l'annexe A. Les résultats empiriques font voir de légères différences entre une réconciliation à suppression du dernier domaine et une réconciliation à suppression aléatoire.

Il sera question de cette dernière méthode de réconciliation dans le contexte d'un modèle bayésien de Fay-Herriot (BFH). À la section 2, nous décrivons le modèle BFH sans contrainte. À la section 3, nous exposons la méthode d'imposition d'une cible externe au modèle BFH par voie de suppression aléatoire. À la section 4, nous décrivons les études empiriques permettant d'évaluer les caractéristiques des estimations tirées de la réconciliation à suppression aléatoire et parlons aussi des mesures d'incertitude liées. À la section 5, nous livrons nos observations en conclusion. Les détails techniques figurent dans plusieurs annexes.

2 Modèle de Fay-Herriot bayésien

Supposons que les données observées sont $(\hat{\theta}_i, s_i)$, $i = 1, \dots, \ell$, où $\hat{\theta}_i$ et s_i sont respectivement une estimation et son erreur-type (que nous supposons connue pour simplifier) de la quantité à l'étude, c'est-à-dire la θ_i , totale du i^e domaine. Le modèle BFH est

$$\begin{aligned}\hat{\theta}_i &| \theta_i \stackrel{\text{ind}}{\sim} \text{Normale}(\theta_i, s_i^2), \\ \theta_i &| \boldsymbol{\beta}, \sigma^2 \stackrel{\text{ind}}{\sim} \text{Normale}(\mathbf{x}_i' \boldsymbol{\beta}, \sigma^2), \\ \pi &(\boldsymbol{\beta}, \sigma^2),\end{aligned}\tag{2.1}$$

où $i = 1, \dots, \ell$, $\mathbf{x}_i$ est un jeu de covariables à p composantes (dont la valeur à l'origine) et $\pi(\boldsymbol{\beta}, \sigma^2)$, la codistribution antérieure pour $(\boldsymbol{\beta}, \sigma^2)$. Nous posons a priori que $\pi(\boldsymbol{\beta}, \sigma^2) = \pi(\boldsymbol{\beta})\pi(\sigma^2)$, c'est-à-dire que $\boldsymbol{\beta}$ et σ^2 sont indépendants avec

$$\pi(\boldsymbol{\beta}) \propto 1 \quad \text{et} \quad \pi(\sigma^2) = 1/(1 + \sigma^2)^2. \quad (2.2)$$

Par le théorème de Bayes, la densité a posteriori conjointe des $\ell + p + 1$ paramètres du modèle est

$$\pi(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 | \hat{\boldsymbol{\theta}}) \propto \frac{1}{(1 + \sigma^2)^2} \left(\frac{1}{\sigma^2} \right)^{\ell/2} \prod_{i=1}^{\ell} \left\{ \exp \left[-\frac{1}{2} \left\{ \frac{1}{s_i^2} (\hat{\theta}_i - \theta_i)^2 + \frac{1}{\sigma^2} (\theta_i - \mathbf{x}_i' \boldsymbol{\beta})^2 \right\} \right] \right\}. \quad (2.3)$$

Dans (2.2), $\pi(\sigma^2)$ est une distribution antérieure propre, plus plate $\pi(\sigma^2) \propto 1/\sigma^2$ près de zéro et sans moments. En fait, toute distribution antérieure propre pour σ^2 convient et une distribution antérieure impropre peut mener dans ce cas à une densité a posteriori impropre. Comme $\pi(\boldsymbol{\beta})$ est impropre, le produit $\pi(\boldsymbol{\beta}, \sigma^2) = \pi(\boldsymbol{\beta})\pi(\sigma^2)$ est *impropre* lui aussi, d'où la possibilité que la densité a posteriori conjointe de $(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2)$ le soit tout autant, scénario peu souhaitable. Le théorème 1 qui suit établit la propriété de la densité a posteriori conjointe (2.3).

L'annexe B donne plus de détails sur le modèle BFH. Plus précisément, avec $\lambda_i = \frac{\sigma^2}{s_i^2 + \sigma^2}$, $i = 1, \dots, \ell$, nous avons démontré que

$$\theta_i | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}} \stackrel{\text{ind}}{\sim} \text{Normale} \left\{ \lambda_i \hat{\theta}_i + (1 - \lambda_i) \mathbf{x}_i' \boldsymbol{\beta}, (1 - \lambda_i) \sigma^2 \right\}, i = 1, \dots, \ell, \quad (2.4)$$

$$\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}} \sim \text{Normale}(\hat{\boldsymbol{\beta}}, \hat{\Sigma}), \quad (2.5)$$

$$\pi_3(\sigma^2 | \hat{\boldsymbol{\theta}}) \propto Q(\sigma^2) \frac{1}{(1 + \sigma^2)^2}, \quad (2.6)$$

où

$$\hat{\boldsymbol{\beta}} = \hat{\Sigma} \sum_{i=1}^{\ell} \frac{\hat{\theta}_i \mathbf{x}_i}{s_i^2 + \sigma^2}, \quad \hat{\Sigma}^{-1} = \sum_{i=1}^{\ell} \frac{\mathbf{x}_i \mathbf{x}_i'}{s_i^2 + \sigma^2},$$

$$Q(\sigma^2) = |\hat{\Sigma}|^{1/2} \prod_{i=1}^{\ell} \frac{1}{(s_i^2 + \sigma^2)^{1/2}} \exp \left\{ -\frac{1}{2} \sum_{i=1}^{\ell} \frac{1}{s_i^2 + \sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \hat{\boldsymbol{\beta}})^2 \right\}.$$

Théorème 1

La densité a posteriori conjointe (2.3) est propre à condition que la matrice de plan soit de plein rang.

Preuve du théorème 1

Comme la matrice de plan est de plein rang, $\hat{\boldsymbol{\beta}}$ et $\hat{\Sigma}$ sont bien définis pour tous les σ^2 d'où l'implication que $Q(\sigma^2)$ est borné dans σ^2 . Ainsi, comme $\pi(\sigma^2) = \frac{1}{(1 + \sigma^2)^2}$ est propre, la densité a posteriori $\pi_3(\sigma^2 | \hat{\boldsymbol{\theta}})$ l'est aussi. Si on applique la règle de multiplication en probabilité, il s'ensuit que la densité a posteriori conjointe, $\pi(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 | \hat{\boldsymbol{\theta}})$, est propre.

La propriété étant acquise, nous pouvons échantillonner la densité a posteriori conjointe et faire notre inférence sur θ_i par une procédure simple de Monte Carlo. Nous prenons la règle de multiplication et tirons des échantillons de (2.6), (2.5) et (2.4). Voici la procédure qui s'applique. D'abord, nous tirons un échantillon de $\pi_3(\sigma^2 | \hat{\theta})$ en (2.6). Les tirages sur cette distribution peuvent s'opérer par la méthode de la grille (voir l'annexe B). Il est ensuite facile de tirer des échantillons de la densité a posteriori conditionnelle de β en (2.5). Enfin, nous pouvons tirer les échantillons de θ_i indépendamment de (2.4). Échantillonner de la sorte à partir de la densité a posteriori conjointe dans le modèle BFH sans contrainte n'exige pas de surveillance des données (à la différence de la méthode de Monte Carlo à chaîne de Markov). Le modèle BFH sans contrainte donnera une base de comparaison des estimations et des mesures liées d'incertitude obtenues par la méthode proposée de réconciliation à suppression aléatoire.

3 Méthode de suppression aléatoire

Comme nous l'avons fait remarquer, la méthode de suppression aléatoire exige qu'on introduise une nouvelle variable prenant les valeurs 1, ..., ℓ en équiprobabilité (poids) et avec peut-être des différences. À la section 3.1, nous indiquons comment construire la densité a posteriori conjointe des paramètres du modèle BFH à suppression aléatoire et, à la section 3.2, comment échantillonner à partir de cette densité conjointe.

3.1 Construction de la densité a posteriori conjointe

Le théorème de base de la réconciliation s'explique par les opérations examinées à la section 1. Le but qui suit est d'élaborer la densité a priori conjointe de $\theta_1, \dots, \theta_\ell$ sous la contrainte de réconciliation $\sum_{i=1}^{\ell} \theta_i = a$, où a est une cible externe connue. Nous nous servons de la densité a priori conjointe de $\theta_1, \dots, \theta_\ell$ pour compléter le modèle Fay-Herriot avec contrainte pour une réconciliation à suppression du dernier domaine (voir la section 1).

Théorème 2

Soit $\theta_i \stackrel{\text{ind}}{\sim} \text{Normale}(\mathbf{u}_i' \beta, \delta^2)$, $i = 1, \dots, \ell$. Avec la contrainte $\sum_{i=1}^{\ell} \theta_i = a$, où a est constant, soit $\theta_{(\ell)}$ le vecteur de tous les θ_i sauf le dernier. La densité conjointe de θ_i , $i = 1, \dots, \ell$, est alors

$$\theta_{(\ell)} \sim \text{Normale} \left\{ \left(I - \frac{1}{\ell} J \right) \mathbf{c}, \delta^2 \left(I - \frac{1}{\ell} J \right) \right\}, \theta_\ell = a - \sum_{i=1}^{\ell-1} \theta_i, \quad (3.1)$$

$\mathbf{c}' = a\mathbf{j}' + \left((\mathbf{u}_1 - \mathbf{u}_\ell)' \beta, \dots, (\mathbf{u}_{\ell-1} - \mathbf{u}_\ell)' \beta \right)$, J et $\mathbf{j}$ sont respectivement une matrice $(\ell - 1) \times (\ell - 1)$ et un vecteur $(\ell - 1)$ de uns.

Preuve du théorème 2

Voir l'annexe C.

La preuve du théorème 2 fait appel à la distribution normale multivariée et servira à démontrer le théorème plus général avec correction de la distribution antérieure pour la suppression de n'importe quel domaine.

Le modèle BFH avec contrainte est

$$\begin{aligned}\hat{\theta}_i \mid \theta_i &\stackrel{\text{ind}}{\sim} \text{Normale}(\theta_i, s_i^2), i = 1, \dots, \ell, \\ \theta_i \mid \boldsymbol{\beta}, \sigma^2 &\stackrel{\text{ind}}{\sim} \text{Normale}(\mathbf{x}_i' \boldsymbol{\beta}, \sigma^2), \sum_{i=1}^{\ell} \theta_i = a, \\ \pi(\boldsymbol{\beta}, \sigma^2) &.\end{aligned}\tag{3.2}$$

La contrainte de la distribution antérieure sur $\theta_1, \dots, \theta_\ell$ corrige pour l'essentiel la densité a priori conjointe, mais pour intégrer cette contrainte, nous emploierons le théorème 2 et procéderons à une correction de la densité a posteriori dans le modèle sans contrainte.

Pour $i = 1, \dots, \ell$, soit $\lambda_i = \frac{\sigma^2}{\sigma^2 + s_i^2}$ et $\hat{\theta}_i = \lambda_i \hat{\theta}_i + (1 - \lambda_i) \mathbf{x}_i' \boldsymbol{\beta}$. Soit également $\hat{\sigma}_i^2 = \sigma^2 (1 - \lambda_i)$, $i = 1, \dots, \ell$. Dans le modèle sans contrainte, la densité a posteriori conjointe est

$$\pi_{nb}(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 \mid \hat{\boldsymbol{\theta}}) \propto \pi(\boldsymbol{\beta}, \sigma^2) \prod_{i=1}^{\ell} \left[\text{Normale}_{\hat{\theta}_i} \left(\mathbf{x}_i' \boldsymbol{\beta}, \frac{\sigma^2}{\lambda_i} \right) \text{Normale}_{\theta_i} \left(\hat{\theta}_i, \hat{\sigma}_i^2 \right) \right].$$

Nous élargissons maintenant le résultat du théorème 2, car nous nous intéressons au modèle BFH avec contrainte. Le théorème 3 qui suit servira à élaborer la méthode de réconciliation à suppression aléatoire.

Théorème 3

Dans une notation générale, soit $y_i \stackrel{\text{ind}}{\sim} \text{Normale}(\mu_i, \sigma_i^2)$, $i = 1, \dots, \ell$. Soit $\mathbf{y}_{(i)}$ le vecteur de tous les y_i sauf le i^{e} . Soit $v_i = \sigma_i^2 / \sqrt{\sum_{j=1}^{\ell} \sigma_j^2}$, $i = 1, \dots, \ell$, et $v_i^* = \sigma_i^2 / \sum_{j=1}^{\ell} \sigma_j^2$, $i = 1, \dots, \ell$. Sous la contrainte $\sum_{i=1}^{\ell} y_i = a$,

$$\mathbf{y}_{(i)} \sim \text{Normale} \left\{ \mathbf{u}_{(i)} - \left(\sum_{j=1}^{\ell} u_j - a \right) \mathbf{v}_{(i)}^*, \text{diagonale}(\boldsymbol{\sigma}_{(i)}^2) - \mathbf{v}_{(i)} \mathbf{v}_{(i)}' \right\} \tag{3.3}$$

avec $y_i = a - \sum_{j \neq i} y_j$. Dans ce cas,

$$p(\mathbf{y}, z = r \mid \phi = 0) = \delta_{y_r} \left(a - \sum_{j \neq r} y_j \right) \text{Normale}_{\mathbf{y}_{(r)}} \left\{ \boldsymbol{\mu}_{(r)} - \left(\sum_{j=1}^{\ell} \mu_j - a \right) \mathbf{v}_{(r)}^*, \text{diagonale}(\boldsymbol{\sigma}_{(r)}^2) - \mathbf{v}_{(r)} \mathbf{v}_{(r)}' \right\}, \tag{3.4}$$

pour $r = 1, \dots, \ell$ et $\psi = a - \sum_{i=1}^{\ell} y_i$.

Preuve du théorème 3

La preuve du théorème 3 ressemble à celle du théorème 2. Voir l'annexe C.

Dans ce qui suit, un des paramètres de domaine ℓ sera supprimé au hasard (probabilité $1/\ell$). Soit $z = 1, \dots, \ell$ représentant le comté en suppression. Ainsi, $P(z = r) = 1/\ell$, $r = 1, \dots, \ell$. Dans le modèle BFH avec contrainte et avec le théorème 3, la densité a posteriori conjointe est donc

$$\begin{aligned}\pi_b(\boldsymbol{\theta}, z = r, \boldsymbol{\beta}, \sigma^2 \mid \hat{\boldsymbol{\theta}}) &\propto \pi(\boldsymbol{\beta}, \sigma^2) \prod_{i=1}^{\ell} \left[\text{Normale}_{\hat{\theta}_i} \left(\mathbf{x}_i' \boldsymbol{\beta}, \frac{\sigma^2}{\lambda_i} \right) \right] \times \delta_{\theta_r} \left(a - \sum_{j \neq r} \theta_j \right) \\ &\times \text{Normale}_{\boldsymbol{\theta}_{(r)}} \left\{ \hat{\boldsymbol{\theta}}_{(r)} - \left(\sum_{j=1}^{\ell} \hat{\theta}_j - a \right) \mathbf{v}_{(r)}^*, \text{diagonale}(\hat{\boldsymbol{\sigma}}_{(r)}^2) - \mathbf{v}_{(r)} \mathbf{v}_{(r)}' \right\}, \end{aligned} \tag{3.5}$$

où $r = 1, \dots, \ell$, et pour $i = 1, \dots, \ell$, $v_i^* = \hat{\sigma}_i^2 / \sum_{j=1}^{\ell} \hat{\sigma}_j^2$ et $v_i = \hat{\sigma}_i^2 / \sqrt{\sum_{j=1}^{\ell} \hat{\sigma}_j^2}$.

3.2 Échantillonnage de la densité a posteriori conjointe

À la différence du modèle BFH, le modèle avec contrainte en (3.2) ne peut être ajusté par des tirages au hasard; nous devons utiliser un échantillonneur de Gibbs. La densité a posteriori conjointe conditionnelle (dpc) de $(\boldsymbol{\theta}, z)$ est

$$\begin{aligned} \pi(\boldsymbol{\theta}, z = r | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}}) &\propto \delta_{\theta_r} \left(a - \sum_{j \neq r} \theta_j \right) \\ &\times \text{Normale}_{\boldsymbol{\theta}_{(r)}} \left\{ \hat{\boldsymbol{\theta}}_{(r)} - \left(\sum_{j=1}^{\ell} \hat{\theta}_j - a \right) \mathbf{v}_{(r)}^*, \text{diagonale}(\hat{\boldsymbol{\sigma}}_{(r)}^2) - \mathbf{v}_{(r)} \mathbf{v}_{(r)}' \right\}, \end{aligned} \quad (3.6)$$

où $r = 1, \dots, \ell$, $\boldsymbol{\theta}_{(r)}$ désigne le vecteur de θ_i avec suppression de la r^{e} composante et où $\mathbf{v}$ et $\mathbf{v}^*$ sont définis au théorème 3. Ainsi, la densité a posteriori conjointe conditionnelle de $(\boldsymbol{\beta}, \sigma^2)$ est

$$\begin{aligned} \pi(\boldsymbol{\beta}, \sigma^2 | \boldsymbol{\theta}, \hat{\boldsymbol{\theta}}, z = r) &\propto \pi(\boldsymbol{\beta}, \sigma^2) \prod_{i=1}^{\ell} \text{Normale}_{\hat{\theta}_i} \left(\mathbf{x}_i' \boldsymbol{\beta}, \frac{\sigma^2}{\lambda_i} \right) \\ &\times \text{Normale}_{\boldsymbol{\theta}_{(r)}} \left\{ \hat{\boldsymbol{\theta}}_{(r)} - \left(\sum_{j=1}^{\ell} \hat{\theta}_j - a \right) \mathbf{v}_{(r)}^*, \text{diagonale}(\hat{\boldsymbol{\sigma}}_{(r)}^2) - \mathbf{v}_{(r)} \mathbf{v}_{(r)}' \right\}. \end{aligned} \quad (3.7)$$

Il est simple d'échantillonner les dpc de $\boldsymbol{\theta}$ et z , mais la chose ne sera pas si simple pour $\boldsymbol{\beta}$ et σ^2 dans ce qui va suivre.

Premièrement, pour obtenir les dpc de $\boldsymbol{\beta}$ et σ^2 , nous définissons

$$\mathbf{g}_1 = \frac{1}{\sigma^2} \sum_{i=1}^{\ell} \lambda_i \hat{\theta}_i \mathbf{x}_i \quad \text{et} \quad A_1 = \frac{1}{\sigma^2} \sum_{i=1}^{\ell} \lambda_i \mathbf{x}_i \mathbf{x}_i'.$$

Deuxièmement, pour $i = 1, \dots, \ell$, soit $d_i = \lambda_i \hat{\theta}_i - \left(\sum_{j=1}^{\ell} \lambda_j \hat{\theta}_j - a \right) v_i^*$ et $\mathbf{x}_i = (1 - \lambda_i) \mathbf{x}_i - v_i^* \sum_{j=1}^{\ell} (1 - \lambda_j) \mathbf{x}_j$. Soit $\mathbf{d}_{(r)}$ et $\tilde{X}_{(r)}$ désignant respectivement le vecteur à entrées d_i sans la r^{e} composante et la matrice où parmi les colonnes $\mathbf{x}_i$ on retranche $\mathbf{x}_r$. Soit

$$\mathbf{g}_2 = (\boldsymbol{\theta}_{(r)} - \mathbf{d}_{(r)})' \Sigma_{(r)} \tilde{X}_{(r)}' \quad \text{et} \quad A_2 = \tilde{X}_{(r)} \Sigma_{(r)} \tilde{X}_{(r)}',$$

où $\Sigma_{(r)}$ est la matrice des covariances sans la ligne et la colonne r^{e} . Troisièmement, avec une distribution antérieure normale multivariée sur $\boldsymbol{\beta}$ de la forme $\boldsymbol{\beta} \sim \text{Normale}(\boldsymbol{\beta}_0, \Sigma_0)$, où $\boldsymbol{\beta}_0$ et Σ_0 sont spécifiés, soit

$$\mathbf{g}_0 = \boldsymbol{\beta}_0 A_0 \quad \text{et} \quad A_0 = \Sigma_0^{-1};$$

ce qui assure une certaine protection contre l'impropriété de la distribution a posteriori. Il s'ensuit que

$$\boldsymbol{\beta} \mid \boldsymbol{\theta}, z = r, \sigma^2, \hat{\boldsymbol{\theta}} \sim \text{Normale} \left\{ \left(\sum_{s=0}^2 A_s \right)^{-1} \sum_{s=0}^2 A_s \mathbf{g}_s, \left(\sum_{s=0}^2 A_s \right)^{-1} \right\}.$$

Nous pouvons éliminer la distribution antérieure de $\boldsymbol{\beta}$ en posant $\Sigma_0 \rightarrow \infty$ (distribution antérieure non informative) pour obtenir $\pi(\boldsymbol{\beta}) = 1$ comme dans le modèle BFH.

Enfin, nous considérons la dpc de σ^2 . Soit

$$U_i = \lambda_i \hat{\theta}_i + (1 - \lambda_i) \mathbf{x}_i' \boldsymbol{\beta} - \left\{ \sum_{j=1}^{\ell} \left\{ \lambda_j \hat{\theta}_j + (1 - \lambda_j) \mathbf{x}_j' \boldsymbol{\beta} \right\} - a \right\} v_i^*, i = 1, \dots, \ell,$$

et

$$\Sigma = \sigma^2 \left\{ \text{diagonale}(1 - \lambda_1, \dots, 1 - \lambda_{\ell}) - [(1 - \lambda_i)(1 - \lambda_{i'})] / \sum_{j=1}^{\ell} (1 - \lambda_j) \right\}.$$

La dpc de σ^2 est alors

$$\pi(\sigma^2 \mid \boldsymbol{\theta}, z = r, \boldsymbol{\beta}, \hat{\boldsymbol{\theta}}) \propto \pi(\sigma^2) \left[\prod_{i=1}^{\ell} \text{Normale}_{\hat{\theta}_i} \{ \mathbf{x}_i' \boldsymbol{\beta}, \sigma^2 + s_i^2 \} \right] \text{Normale}_{\theta_{(r)}} \{ \mathbf{U}_{(r)}, \Sigma_{(r)} \},$$

où $\mathbf{U}_{(r)}$ désigne le vecteur des U_i sans la r^{e} composante et où $\pi(\sigma^2)$ est une distribution antérieure sur σ^2 . Comme dans le modèle BFH (voir la section 2), nous attribuons la densité antérieure $\pi(\sigma^2) = 1/(1 + \sigma^2)^2$, $\sigma^2 > 0$ à σ^2 . Comme la densité a posteriori de base du modèle BFH est propre, la densité a posteriori du modèle BFH avec contrainte le sera aussi.

4 Études empiriques

Le but de ces études empiriques est double : d'abord, nous démontrons que le modèle BFH peut être ajusté comme on l'indique à la section 2 et que les méthodes de réconciliation à suppression du dernier domaine et à suppression aléatoire peuvent s'appliquer; ensuite, nous comparons les méthodes de réconciliation dans une étude par simulation qui exploite un ensemble de données bien utilisé dans les études consacrées à l'estimation sur petits domaines.

Dans le processus de génération des données, nous prenons les données sur les superficies en maïs et en soya dans Battese, Harter et Fuller (1988), lesquelles sont disponibles pour 12 comtés (domaines) en Iowa. Nous dressons le tableau résultant des superficies en maïs et en soya au niveau des comtés à l'aide d'un certain nombre de segments échantillonnés dans la population (nombre connu de segments). Nous disposons également de données satellitaires Landsat sur le nombre de pixels de maïs et de soya dans les segments échantillonnés (deux covariables). Nous pouvons ajouter des moyennes de population finie pour le nombre de pixels caractérisés comme maïs et soya dans chaque comté. Avec cet ensemble de données, nous construisons de nouveaux jeux de données avec tout nombre de domaines.

La génération des données se fait en deux étapes. D'abord, nous ajustons le modèle au niveau des unités $y_{ij} = \mathbf{x}'_{ij}\boldsymbol{\beta} + e_{ij}$, $i = 1, \dots, \ell$, $j = 1, \dots, n_i$, où $e_{ij} \stackrel{\text{iid}}{\sim} (0, \sigma^2)$, par rapport aux données disponibles pour les $\ell = 12$ comtés de l'Iowa. Les tailles d'échantillon de domaines sont $n_1 = n_2 = n_3 = 1$, $n_4 = 2$, $n_5 = n_6 = n_7 = n_8 = 3$, $n_9 = 4$, $n_{10} = n_{11} = 5$, et $n_{12} = 6$. Par la méthode des moindres carrés, nous estimons $\boldsymbol{\beta}$ et σ^2 par $\hat{\boldsymbol{\beta}}$ et $\hat{\sigma}^2$, respectivement. Pour les domaines dont la taille d'échantillon est supérieure à l'unité, nous posons que s_i^2 est égal à la variance estimée de la moyenne d'échantillon $\bar{y}_i$ ($\bar{y}_i = \sum_{j=1}^{n_i} y_{ij} / n_i$) et nous prenons S^2 comme moyenne géométrique. Pour les domaines dont la taille d'échantillon correspond à l'unité, nous posons s_i^2 égal à S^2 . Le vecteur des covariables $\bar{\mathbf{X}}_i$ est à trois éléments, à savoir l'entier un (comme valeur à l'origine) et ensuite les moyennes de population des pixels caractérisés comme maïs et soya.

Comme deuxième étape, nous illustrons le processus de génération de données pour tout nombre désiré ℓ de petits domaines. Nous échantillonnons les covariables $\mathbf{x}_i$, $i = 1, \dots, \ell$, avec remise à partir de $\bar{\mathbf{X}}_i$, $i = 1, \dots, 12$. Nous tirons alors les moyennes au niveau des domaines à l'aide de

$$\theta_i \stackrel{\text{iid}}{\sim} \text{Normale}(\mathbf{x}'_i \hat{\boldsymbol{\beta}}, \hat{\sigma}^2), i = 1, \dots, \ell,$$

où $\hat{\boldsymbol{\beta}}$ et $\hat{\sigma}^2$ sont les estimations déjà définies par les moindres carrés. Nous dégagons les variances d'échantillon s_i^2 en deux étapes. Nous tirons d'abord les tailles d'échantillon d'une distribution uniforme, $n_i \stackrel{\text{iid}}{\sim} \text{Uniforme}(5, 25)$, $i = 1, \dots, \ell$. En deuxième lieu, nous posons $s_i^2 = S^2 V_i / (n_i - 1)$, où $V_i \stackrel{\text{iid}}{\sim} \chi^2_{n_i-1}$ et S^2 sont définis comme ci-dessus. Enfin, nous tirons les estimations d'enquête sur petits domaines par $\hat{\theta}_i \stackrel{\text{iid}}{\sim} \text{Normale}(\theta_i, s_i^2)$, $i = 1, \dots, \ell$. Nous posons la cible de réconciliation comme égale à la somme des $\hat{\theta}_i$ et des variantes de cette valeur $\sum_{i=1}^{\ell} \hat{\theta}_i$ majorées ou minorées de 50 %. Dans la pratique du NASS pour les estimations culturelles des comtés, cette cible est une valeur déjà établie au niveau de l'État. Pour évaluer les méthodes de réconciliation dans des cas extrêmes, nous considérons d'autres scénarios de simulation où une taille d'échantillon de domaines est fixée à 2 ou 50 ou où le facteur S^2 est multiplié par 10.

Dans ce qui suit, nous présentons les résultats empiriques surtout d'un scénario de simulation avec 12 domaines. Nous examinons brièvement des exemples avec un plus grand nombre de domaines. Ainsi, l'Iowa compte 99 comtés et le NASS a pour intérêt notamment de réconcilier les estimations de superficie ensemencée et récoltée et de production (en boisseaux) par rapport à un total préétabli au niveau de l'État. Lorsque le nombre de domaines est si petit, aucune correction n'est à apporter aux procédures de réconciliation avec suppression du dernier domaine ou suppression aléatoire dont nous avons parlé dans les sections précédentes. Toutefois, le calcul peut être inacceptable lorsque le nombre de domaines est extrêmement grand (un million, disons), auquel cas il faudrait apporter certaines modifications aux procédures actuellement appliquées.

Il est utile d'examiner les calculs dans un scénario de simulation avec 12 domaines. Pour l'inférence a posteriori dans le modèle BFH, nous avons recouru à 1 000 tirages aléatoires exécutés en seulement

quelques secondes. En revanche, il est plus difficile d'exécuter un échantillonneur de Gibbs dans une réconciliation à suppression d'un domaine à la fois ou à suppression aléatoire. Nous pouvons toutefois proposer un échantillonneur de Gibbs efficace. Nous avons effectué un long passage de 20 000 itérations avec les 10 000 premières en « rodage » et avec choix de chaque 10^e itération par la suite. Nous avons contrôlé ce processus par tâtonnement en prenant les autocorrélations, le test de stationnarité de Geweke et les tailles effectives d'échantillon. Pour les 1 000 itérations sélectionnées, les valeurs d'autocorrélation sont toutes négligeables. Pour la réconciliation à suppression aléatoire, les valeurs p du test de Geweke pour les trois coefficients de régression et δ^2 sont respectivement de 0,651; 0,087; 0,828 et 0,699 (la stationnarité n'est pas rejetée). Les tailles effectives d'échantillon sont toutes de 1 000. Ajoutons que les tracés temporels ne montrent aucun signe de non-stationnarité. Ainsi, l'échantillonneur de Gibbs est efficace et ne prend que quelques secondes malgré l'abondance des itérations.

Nous évaluons le rendement des méthodes de réconciliation avec un ensemble de mesures : moyennes a posteriori (MP), écarts-types a posteriori (ETP) et, si la chose est pratique, coefficients de variation a posteriori (CVP), erreurs-types numériques (ETN) des estimations et intervalles de plus haute densité a posteriori (HDP) à 95 %. Les résultats numériques figurent aux tableaux 4.1 à 4.8.

Nous présentons au tableau 4.1 une version condensée des résultats de base qui sert à comparer la moyenne, l'erreur-type et le coefficient de variation des données observées aux MP, ETP et CVP du modèle BFH et des modèles de réconciliation à suppression du dernier domaine (DD) et à suppression aléatoire (SA). Les résultats au tableau 4.1 s'appliquent à deux scénarios de simulation où $S^2 = 163$ pour une légère variation des données observées et $S^2 = 1\,630$ pour une variation relativement supérieure de ces mêmes données. Quand $S^2 = 163$, il n'y a à peu près pas de différences entre les données observées et les quantités a posteriori pour les modèles BFH, DD et SA. Vu les petits coefficients de variation des estimations d'enquête, il est difficile de réduire encore plus la variabilité pour tout modèle. Ainsi, les CVP sont comparables aux CV des estimations d'enquête. Par ailleurs, deux points intéressants ressortent du scénario où $S^2 = 1\,630$. D'abord, les MP du modèle BFH peuvent être très différentes de celles des modèles DD et SA et ces deux dernières MP sont très proches l'une de l'autre. Ensuite, les ETP sont bien moindres que les erreurs-types des données observées; on relève des gains appréciables de précision avec le modèle BFH. Il reste que les ETP sont de quatre à cinq fois moindres que ceux des données observées et que les ETP des modèles DD et SA sont environ le double de ceux du modèle BFH. Il en va de même des CVP. Enfin, les modèles DD et SA sont très proches dans les trois mesures (MP, ETP et CVP); le modèle SA présente des ETP qui sont seulement un peu moindres. Comme on pouvait s'y attendre, le modèle DD s'écarter légèrement du modèle SA, mais on doit aussi remarquer que la réconciliation du modèle BFH est importante, puisque nous obtenons des réponses différentes de celles de ce modèle du moins pour ce qui est des écarts-types et des coefficients de variation a posteriori. La réconciliation est une procédure à bruit aléatoire qui aide à protéger le modèle contre les défauts de spécification et qui, par conséquent, doit créer une plus ample variabilité des estimations sur petits domaines.

Tableau 4.1

Comparaison des modèles BFH sans réconciliation, à suppression du dernier domaine et à suppression aléatoire sous l'angle des moyennes a posteriori (MP) des écarts-types a posteriori (ETP) et des coefficients de variation a posteriori (CVP) pour deux valeurs de S^2

	A	MP				ETP				CVP			
		OB	BFH	DD	SA	OB	BFH	DD	SA	OB	BFH	DD	SA
a. $S^2 = 163$; $a = 1\,435$	1	135,6	134,0	133,8	133,5	6,03	5,62	5,47	5,41	0,044	0,042	0,041	0,041
	2	102,0	103,5	103,1	103,0	7,10	6,50	6,11	5,82	0,070	0,063	0,059	0,057
	3	117,7	121,0	120,7	120,5	7,31	6,72	6,55	6,25	0,062	0,056	0,054	0,052
	4	77,0	81,5	81,4	81,0	5,88	6,00	5,46	5,53	0,076	0,074	0,067	0,068
	5	126,9	127,8	127,5	127,5	5,63	5,25	5,25	5,06	0,044	0,041	0,041	0,040
	6	113,1	113,4	112,9	113,1	8,06	7,15	6,82	6,74	0,071	0,063	0,060	0,060
	7	137,2	133,7	133,5	133,9	6,74	6,38	5,93	6,02	0,049	0,048	0,044	0,045
	8	124,8	124,7	124,7	124,7	4,03	3,91	3,83	3,76	0,032	0,031	0,031	0,030
	9	118,3	116,5	115,8	116,6	7,54	6,79	6,29	6,65	0,064	0,058	0,054	0,057
	10	156,5	153,4	153,3	153,3	4,37	4,45	4,12	4,18	0,028	0,029	0,027	0,027
	11	109,5	110,3	110,3	110,2	4,88	4,64	4,70	4,70	0,045	0,042	0,043	0,043
	12	116,3	118,1	117,9	117,7	7,23	6,62	6,26	6,00	0,062	0,056	0,053	0,051
b. $S^2 = 1\,630$; $a = 1\,482$	1	129,1	129,8	127,2	126,5	19,07	4,64	10,71	10,45	0,148	0,036	0,084	0,083
	2	117,3	126,3	122,1	122,1	22,46	5,08	12,73	12,51	0,191	0,040	0,104	0,102
	3	120,0	145,5	137,3	136,9	23,11	5,93	12,91	12,68	0,193	0,041	0,094	0,093
	4	68,8	107,3	94,0	93,6	18,60	7,47	12,04	11,86	0,270	0,070	0,128	0,127
	5	142,4	146,4	142,3	142,2	17,80	4,52	11,98	11,15	0,125	0,031	0,084	0,078
	6	108,8	120,2	115,2	115,4	25,49	5,43	11,75	11,66	0,234	0,045	0,102	0,101
	7	136,8	116,2	118,2	119,0	21,31	5,37	11,32	11,90	0,156	0,046	0,096	0,100
	8	124,5	132,5	127,3	127,3	12,76	4,39	9,00	8,91	0,102	0,033	0,071	0,070
	9	144,2	127,5	128,0	129,5	23,86	5,33	12,74	14,00	0,165	0,042	0,100	0,108
	10	172,9	129,2	145,5	145,3	13,81	9,23	10,28	10,37	0,080	0,071	0,071	0,071
	11	109,1	114,7	110,6	110,2	15,42	4,31	10,53	10,43	0,141	0,038	0,095	0,095
	12	108,4	120,3	114,6	114,2	22,87	5,10	12,42	12,01	0,211	0,042	0,108	0,105

Note : OB : données observées; BFH : modèle de Fay-Herriot bayésien; DD : réconciliation à suppression du dernier domaine; SA : modèle à suppression aléatoire; a : cible. Pour les données observées, nous présentons l'estimation directe, l'erreur-type et le coefficient de variation pour MP, ETP et CVP respectivement. Dans la procédure de réconciliation de DGSM pour $S^2 = 163$, les valeurs de réconciliation sont de 133,7; 103,3; 120,8; 81,3; 127,6; 113,2; 133,4; 124,5; 116,3; 153,1; 110,1; 117,9 et, pour $S^2 = 1\,630$, de 126,9; 123,5; 142,3; 105,0; 143,2; 117,5; 113,6; 129,5; 124,7; 126,4; 112,1; 117,6.

Dans le scénario de simulation de base, nous comparons les méthodes de réconciliation à suppression à une des méthodes de DGSM qui donne des estimations a posteriori réconciliées sans suppression. Pour reprendre la notation dans DGSM, nous devons reformuler l'équation de réconciliation sous la forme suivante :

$$\sum_{i=1}^{\ell} \omega_i \theta_i = \frac{a}{\ell} = t,$$

où $\omega_i = 1/\ell$, $\sum_{i=1}^{\ell} \omega_i = 1$. Soit $\hat{\theta}_i^{(B)}$ la moyenne a posteriori du modèle BFH. Définissons $\bar{\hat{\theta}}_B = \sum_{i=1}^{\ell} \omega_i \hat{\theta}_i^{(B)}$,

$$\phi_i = \frac{\omega_i}{\hat{\theta}_i^{(B)}}, r_i = \frac{\omega_i}{\phi_i}, i = 1, \dots, \ell,$$

et $S^* = \sum_{i=1}^{\ell} \omega_i^2 / \phi_i$. À noter que, parmi les spécifications de DGSM, nous avons choisi ϕ_i au hasard (sans traitement préférentiel). Dans ce cas, les estimateurs de Bayes réconciliés de DGSM sont

$$\hat{\theta}_i^{(BM)} = \hat{\theta}_i^{(B)} + (t - \bar{\theta}_B) r_i / S^*, i = 1, \dots, \ell.$$

Nous présentons les résultats empiriques avec l'estimateur $\hat{\theta}_i^{(BM)}$ dans la note du tableau 4.1. La plus grande différence entre les estimations réconciliées selon les méthodes de réconciliation est celle du domaine 10 (données observées : 172,9; BFH : 129,2; DD : 145,5; SA : 145,3; DGSM : 126,4). En général, les MP des modèles DD et SA sont plus proches de celles des données observées. Dans les autres cas, les estimations se comparent relativement bien à celles des modèles DD et SA, mais avec quelques légères différences; le modèle DGSM ne présente pas d'écarts-types a posteriori ni d'intervalles de crédibilité.

Nous présentons des résultats plus détaillés pour $S^2 = 163$ aux tableaux 4.2 à 4.8 et aux figures 4.1 à 4.4. Notre intérêt est surtout de comparer les modèles à suppression d'un domaine à la fois (le dernier) et à suppression aléatoire.

Par les résultats au tableau 4.2, nous concluons que les MP du modèle BFH (sans réconciliation) sont légèrement différentes de celles des estimations directes et, comme on pouvait s'y attendre, supérieures et inférieures respectivement aux valeurs des estimations directes (en moins et en plus). Sauf pour deux domaines et comme il était à prévoir, les ETP sont inférieurs aux écarts-types des estimations directes. Ainsi, l'estimation directe la plus petite (76,997) est celle qui rétrécit le plus et présente un écart-type plus grand (5,881 contre 5,995); les résultats concordent avec le rétrécissement caractéristique d'une estimation sur petits domaines. Il convient de noter que les CVP sont tous petits et que les ETN le sont raisonnablement aussi.

Tableau 4.2

Comparaison de l'estimateur direct avec l'inférence a posteriori par le modèle de Fay-Herriot bayésien pour les paramètres de domaine

Domaine	<i>n</i>	$\hat{\theta}$	<i>s</i>	MP	ETP	CVP	ETN	HDP à 95 %
1	5	135,575	6,031	133,985	5,617	0,042	0,057	(123,422; 145,402)
2	7	101,980	7,101	103,461	6,498	0,063	0,065	(90,598; 116,134)
3	24	117,655	7,309	121,006	6,716	0,056	0,066	(107,730; 134,124)
4	23	76,997	5,881	81,473	5,995	0,074	0,058	(69,046; 92,578)
5	21	126,917	5,629	127,832	5,248	0,041	0,052	(117,850; 138,406)
6	9	113,132	8,061	113,393	7,147	0,063	0,068	(99,441; 127,451)
7	5	137,236	6,739	133,661	6,378	0,048	0,064	(121,771; 146,662)
8	20	124,839	4,034	124,732	3,906	0,031	0,039	(117,233; 132,309)
9	16	118,306	7,544	116,479	6,785	0,058	0,071	(103,225; 130,003)
10	9	156,503	4,368	153,355	4,449	0,029	0,045	(144,785; 162,031)
11	23	109,546	4,877	110,348	4,637	0,042	0,047	(101,179; 119,294)
12	9	116,314	7,232	118,098	6,623	0,056	0,068	(105,135; 131,186)

Note : *n* est la taille d'échantillon de domaines, $\hat{\theta}$ est l'estimateur direct et *s* son erreur-type. MP est la moyenne a posteriori, ETP l'écart-type a posteriori et HDP l'intervalle de plus haute densité a posteriori. ETN est l'erreur-type numérique de la moyenne a posteriori. La valeur de réconciliation est de 1 435 et la somme des moyennes a posteriori, de 1 437,823 (sans réconciliation).

Nous présentons aux tableaux 4.3 et 4.4 les estimations du modèle BFH à suppression du dernier domaine et à suppression aléatoire dans le cas d'une distribution antérieure uniforme (en équipondération).

Les poids a posteriori diffèrent à peine de 0,083; le plus grand (0,097) est celui du dernier domaine et le plus petit (0,056), du 8^e domaine. La suppression aléatoire comme la suppression du dernier domaine améliore la précision et les ETP des estimations réconciliées sont tous inférieurs aux erreurs-types des données observées pour les deux méthodes de réconciliation. Les ETN sont supérieures à celles du modèle sans réconciliation, mais la chose importe peu, puisqu'il s'agit d'erreurs des moyennes a posteriori (la caractéristique de la MP est à trois chiffres).

Tableau 4.3

Comparaison de l'estimateur direct avec l'inférence a posteriori par le modèle Fay-Herriot bayésien pour les paramètres de domaine avec réconciliation à suppression aléatoire

Domaine	<i>n</i>	$\hat{\theta}$	<i>s</i>	MP	ETP	CVP	ETN	HDP à 95 %
1	5	135,575	6,031	133,516	5,431	0,041	0,171	(123,414; 143,541)
2	7	101,980	7,101	102,903	5,793	0,056	0,199	(92,378; 114,250)
3	24	117,655	7,309	120,671	6,237	0,052	0,194	(107,744; 132,190)
4	23	76,997	5,881	81,170	5,597	0,069	0,202	(69,781; 91,177)
5	21	126,917	5,629	127,652	5,036	0,039	0,170	(118,293; 137,228)
6	9	113,132	8,061	112,805	6,707	0,059	0,223	(100,926; 126,074)
7	5	137,236	6,739	133,908	6,007	0,045	0,177	(122,135; 145,344)
8	20	124,839	4,034	124,703	3,757	0,030	0,120	(117,962; 132,304)
9	16	118,306	7,544	116,451	6,650	0,057	0,249	(103,400; 129,316)
10	9	156,503	4,368	153,222	4,216	0,028	0,134	(144,392; 160,854)
11	23	109,546	4,877	110,221	4,694	0,043	0,150	(101,038; 119,570)
12	9	116,314	7,232	117,780	5,997	0,051	0,208	(104,619; 128,158)

Note : *n* est la taille d'échantillon de domaines, $\hat{\theta}$ l'estimateur direct et *s* son erreur-type. MP est la moyenne a posteriori, ETP est l'écart-type a posteriori et HDP l'intervalle de plus haute densité a posteriori. ETN est l'erreur-type numérique de la moyenne a posteriori. La valeur de réconciliation est 1 435. Pour une distribution antérieure uniforme (en équipondération), les probabilités a posteriori que les domaines 1; 2; 3; 4; 5; 6; 7; 8; 9; 10; 11; 12 soient supprimés sont respectivement de 0,090; 0,084; 0,095; 0,077; 0,066; 0,093; 0,097; 0,056; 0,098; 0,068; 0,079; 0,097.

Tableau 4.4

Comparaison de l'estimateur direct avec l'inférence a posteriori par le modèle Fay-Herriot bayésien pour les paramètres de domaine avec suppression du dernier domaine

Domaine	<i>n</i>	$\hat{\theta}$	<i>s</i>	MP	ETP	CVP	ETN	HDP à 95 %
1	5	135,575	6,031	133,772	5,519	0,041	0,151	(122,213; 143,991)
2	7	101,980	7,101	103,026	6,319	0,061	0,171	(89,424; 113,857)
3	24	117,655	7,309	120,470	6,458	0,054	0,209	(108,783; 134,261)
4	23	76,997	5,881	81,391	5,906	0,073	0,171	(69,636; 92,634)
5	21	126,917	5,629	127,883	5,158	0,040	0,142	(117,282; 137,305)
6	9	113,132	8,061	112,895	6,270	0,056	0,216	(100,664; 124,320)
7	5	137,236	6,739	133,298	5,948	0,045	0,178	(121,831; 144,727)
8	20	124,839	4,034	124,664	3,810	0,031	0,124	(117,321; 131,941)
9	16	118,306	7,544	116,542	6,531	0,056	0,203	(104,238; 129,622)
10	9	156,503	4,368	153,229	4,353	0,028	0,132	(144,443; 161,593)
11	23	109,546	4,877	109,997	4,563	0,041	0,168	(101,428; 118,953)
12	9	116,314	7,232	117,835	6,344	0,054	0,215	(106,421; 131,483)

Note : *n* est la taille d'échantillon de domaines, $\hat{\theta}$ l'estimateur direct et *s* son erreur-type. MP est la moyenne a posteriori, ETP l'écart-type a posteriori et HDP l'intervalle de plus haute densité a posteriori. ETN est l'erreur-type numérique de la moyenne a posteriori. La valeur de réconciliation est 1 435.

Nous comparons les trois méthodes (BFH, DD, SA) à l'aide des résultats au tableau 4.5. Les MP sont comparables et la réconciliation (SA et DD) ne déforme (rétrécit) pas les estimations bien au-delà du rétrécissement caractéristique du modèle BFH. Ajoutons que, dans les modèles DD et SA, les ETP sont presque toujours inférieurs à ceux du modèle BFH. Pour 8 des 12 domaines, les ETP sont moindres dans le modèle SA que dans le modèle DD. Pour ces domaines, les ETP baissent en gros de 1 % dans le modèle SA par rapport au modèle DD et environ de 4 % par rapport au modèle BFH.

Pour étudier à quel point les ETP sont sensibles aux différentes cibles de réconciliation, nous présentons les résultats pour trois choix de cibles au tableau 4.6. Les ETP varient seulement un peu selon les cibles et l'emportent toujours sur les erreurs-types des estimations directes.

Dans la conception d'un jeu complexe de simulations, nous envisageons de recourir à des probabilités (poids) inégales pour la réconciliation à suppression aléatoire. Nous mentionnons les résultats au tableau 4.7. Nous comparons les poids uniformes (EW) aux poids inversement proportionnels (IW) aux tailles d'échantillon, ainsi qu'aux poids directement proportionnels (DW) à ces mêmes tailles. Là encore, nous relevons de légères différences entre les trois MP et les trois ETP. Les ETP demeurent inférieurs à ceux des estimations directes.

Au moyen des résultats au tableau 4.8, nous étudions dans quelle mesure les tailles d'échantillon extrêmes pour le dernier comté (à supprimer) influent sur l'inférence a posteriori. Nous posons à cette fin que la taille d'échantillon du dernier comté se situe à l'extérieur de la fourchette de simulation qui va de 5 à 25 en prenant les valeurs 2 et 50. Nous considérons d'abord le cas où la taille d'échantillon du dernier comté est de 2. Comme pour les résultats précédents, les MP présentent de légères différences par rapport à l'absence de réconciliation, à la suppression du dernier domaine et à la suppression aléatoire pour tous les comtés. Les ETP des modèles DD et SA sont inférieurs à ceux du modèle BFH et neuf de ces ETP sont moindres dans le modèle SA que dans le modèle DD. Pour le dernier comté, nous observons cependant des écarts-types a posteriori relativement importants (10,00; 8,771; 8,525) avec une baisse qui est en gros de 15 % de l'ETP du modèle SA par rapport au modèle sans réconciliation. Ensuite, nous considérons le cas où la taille d'échantillon du dernier comté est de 50. La configuration est semblable sauf que les ETP du dernier comté sont comparables à ceux des modèles BFH, DD et SA et que, là encore, la baisse est environ de 10 % (6,282; 5,958; 5,702) des ETP du modèle SA par rapport à l'absence de réconciliation. Il semblerait que, si on prend délibérément le comté avec la taille d'échantillon la plus extrême (en moins ou en plus) comme dernier comté, il peut y avoir une incidence sur la procédure de réconciliation. En revanche, nous relevons des variations légères lorsque les domaines dont la taille d'échantillon est extrême ne sont pas systématiquement supprimés. Quand la taille d'échantillon est de 2, les nouvelles valeurs MP et ETP sont les suivantes : BFH : 124,307; 9,993; DD : 123,371; 9,000; SA : 123,540; 8,887. Quand la taille d'échantillon est de 50, les nouvelles valeurs MP et ETP sont les suivantes : BFH : 118,167; 6,284; DD : 117,802; 6,094; SA : 117,716; 5,948.

Tableau 4.5

Résumé de la comparaison de l'inférence avec l'estimateur direct, le modèle de Fay-Herriot bayésien (BFH), la réconciliation à suppression aléatoire (SA) et la réconciliation à suppression du dernier domaine (DD)

Domaine	n	$\hat{\theta}$	s	BFH		SA		DD	
				MP	ETP	MP	ETP	MP	ETP
1	5	135,575	6,031	133,985	5,617	133,516	5,431	133,772	5,519
2	7	101,980	7,101	103,461	6,498	102,903	5,793	103,026	6,319
3	24	117,655	7,309	121,006	6,716	120,671	6,237	120,470	6,458
4	23	76,997	5,881	81,473	5,995	81,170	5,597	81,391	5,906
5	21	126,917	5,629	127,832	5,248	127,652	5,036	127,883	5,158
6	9	113,132	8,061	113,393	7,147	112,805	6,707	112,895	6,270
7	5	137,236	6,739	133,661	6,378	133,908	6,007	133,298	5,948
8	20	124,839	4,034	124,732	3,906	124,703	3,757	124,664	3,810
9	16	118,306	7,544	116,479	6,785	116,451	6,650	116,542	6,531
10	9	156,503	4,368	153,355	4,449	153,222	4,216	153,229	4,353
11	23	109,546	4,877	110,348	4,637	110,221	4,694	109,997	4,563
12	9	116,314	7,232	118,098	6,623	117,780	5,997	117,835	6,344

Note : n est la taille d'échantillon de domaines, $\hat{\theta}$ l'estimateur direct et s son erreur-type. MP est la moyenne a posteriori et ETP l'écart-type a posteriori. La valeur de réconciliation est 1 435. Dans le cas d'une distribution antérieure uniforme, les probabilités a posteriori que les domaines 1; 2; 3; 4; 5; 6; 7; 8; 9; 10; 11; 12 soient supprimés sont respectivement de 0,090; 0,084; 0,095; 0,077; 0,066; 0,093; 0,097; 0,056; 0,098; 0,068; 0,079; 0,097.

Tableau 4.6

Comparaison de l'inférence a posteriori des paramètres de domaine avec la réconciliation à suppression aléatoire et différentes cibles ($a = 1\ 435$)

Domaine	n	$\hat{\theta}$	s	a		1,5a		0,5a	
				MP	ETP	MP	ETP	MP	ETP
1	5	135,575	6,031	133,516	5,431	189,249	5,385	77,769	5,561
2	7	101,980	7,101	102,903	5,793	175,963	5,794	29,847	5,899
3	24	117,655	7,309	120,671	6,237	197,219	6,099	44,145	6,461
4	23	76,997	5,881	81,170	5,597	134,628	5,871	27,771	5,460
5	21	126,917	5,629	127,652	5,036	177,209	5,165	78,125	5,053
6	9	113,132	8,061	112,805	6,707	201,949	7,145	23,614	6,995
7	5	137,236	6,739	133,908	6,007	200,989	6,018	66,781	6,024
8	20	124,839	4,034	124,703	3,757	151,951	3,952	97,484	3,924
9	16	118,306	7,544	116,451	6,650	196,849	6,990	35,990	6,607
10	9	156,503	4,368	153,222	4,216	184,720	4,019	121,708	4,706
11	23	109,546	4,877	110,221	4,694	148,724	4,966	71,752	4,760
12	9	116,314	7,232	117,780	5,997	193,050	5,954	42,514	6,081

Note : n est la taille d'échantillon de domaines, $\hat{\theta}$ est l'estimateur direct et s son erreur-type. MP est la moyenne a posteriori et ETP l'écart-type a posteriori. La valeur de réconciliation est 1 435. Dans le cas d'une distribution antérieure uniforme, les probabilités a posteriori que les domaines 1; 2; 3; 4; 5; 6; 7; 8; 9; 10; 11; 12 soient supprimés sont respectivement de 0,090; 0,084; 0,095; 0,077; 0,066; 0,093; 0,097; 0,056; 0,098; 0,068; 0,079; 0,097. Quand la valeur de réconciliation est majorée de 50 %, les probabilités sont de 0,090; 0,084; 0,095; 0,079; 0,064; 0,093; 0,097; 0,056; 0,098; 0,068; 0,079; 0,097. Quand la valeur de réconciliation est minorée de 50 %, les probabilités sont de 0,090; 0,084; 0,095; 0,077; 0,066; 0,093; 0,097; 0,057; 0,097; 0,068; 0,079; 0,097.

Tableau 4.7

Comparaison de l'inférence a posteriori des paramètres de domaine avec la réconciliation à suppression aléatoire en équipondération (EW), en pondération inversement proportionnelle aux tailles d'échantillon (IW) et en pondération directement proportionnelle aux tailles d'échantillon (DW)

Domaine	n	$\hat{\theta}$	s	EW		IW		DW	
				MP	ETP	MP	ETP	MP	ETP
1	5	135,575	6,031	133,516	5,431	133,508	5,518	133,436	5,404
2	7	101,980	7,101	102,903	5,793	103,042	5,737	103,049	5,809
3	24	117,655	7,309	120,671	6,237	120,529	6,176	120,634	6,247
4	23	76,997	5,881	81,170	5,597	81,167	5,571	81,111	5,567
5	21	126,917	5,629	127,652	5,036	127,669	5,079	127,541	5,055
6	9	113,132	8,061	112,805	6,707	112,762	6,704	113,074	6,716
7	5	137,236	6,739	133,908	6,007	133,965	5,968	133,798	6,027
8	20	124,839	4,034	124,703	3,757	124,829	3,734	124,719	3,757
9	16	118,306	7,544	116,451	6,650	116,300	6,707	116,502	6,640
10	9	156,503	4,368	153,222	4,216	153,238	4,198	153,204	4,220
11	23	109,546	4,877	110,221	4,694	110,190	4,697	110,208	4,690
12	9	116,314	7,232	117,780	5,997	117,802	6,010	117,726	5,989

Note : n est la taille d'échantillon de domaines, $\hat{\theta}$ est l'estimateur direct et s son erreur-type. MP est la moyenne a posteriori et ETP l'écart-type a posteriori. La valeur de réconciliation est 1 435. Dans le cas d'une distribution antérieure uniforme, les probabilités a posteriori que les domaines 1; 2; 3; 4; 5; 6; 7; 8; 9; 10; 11; 12 soient supprimés sont respectivement de 0,090; 0,084; 0,095; 0,077; 0,066; 0,093; 0,097; 0,056; 0,098; 0,068; 0,079; 0,097. Quand la réconciliation se fait en pondération inversement proportionnelle aux tailles d'échantillon, les probabilités sont de 0,167; 0,127; 0,039; 0,037; 0,026; 0,105; 0,184; 0,030; 0,061; 0,078; 0,039; 0,107. Quand la réconciliation se fait en pondération directement proportionnelle aux tailles d'échantillon, les probabilités sont de 0,032; 0,048; 0,168; 0,124; 0,103; 0,061; 0,036; 0,083; 0,112; 0,044; 0,123; 0,066.

À des fins de comparaison, nous présentons les différentes densités a posteriori aux figures 4.1 à 4.4. Il s'agit aux figures 4.1 et 4.2 des densités a posteriori des 12 paramètres de domaine quand chaque domaine est retranché à son tour. Nous observons que la densité a posteriori varie légèrement selon les modes, mais sans rien de remarquable. Aux figures 4.3 et 4.4, nous présentons les densités a posteriori des 12 paramètres de domaine pour le modèle FH (sans contrainte) et les deux réconciliations à suppression aléatoire et à suppression du dernier domaine. Des différences existent entre les trois densités, mais sans rien d'alarmant.

Nous livrons enfin les résultats empiriques d'un scénario de simulation avec 99 domaines correspondant aux 99 comtés de l'Iowa. Nous générons les données comme nous l'avons décrit et ajustons les modèles BFH sans réconciliation et avec réconciliation à suppression aléatoire et à suppression du dernier domaine au moyen de 20 000 itérations de l'échantillonneur de Gibbs. Dans l'ajustement de chaque modèle, les 10 000 premières itérations servent au « rodage » et, par la suite, nous gardons chaque 10^e itération. L'ajustement du modèle BFH prend 15 secondes et celui des modèles à réconciliation par suppression, moins de trois minutes chacun. Pour les paramètres du modèle de réconciliation à suppression aléatoire et avec les coefficients de régression β et la variance σ^2 , les valeurs p du test de Geweke sont respectivement de 0,822; 0,128; 0,752 et 0,219 et les tailles effectives d'échantillon sont toutes de 1 000 pour les

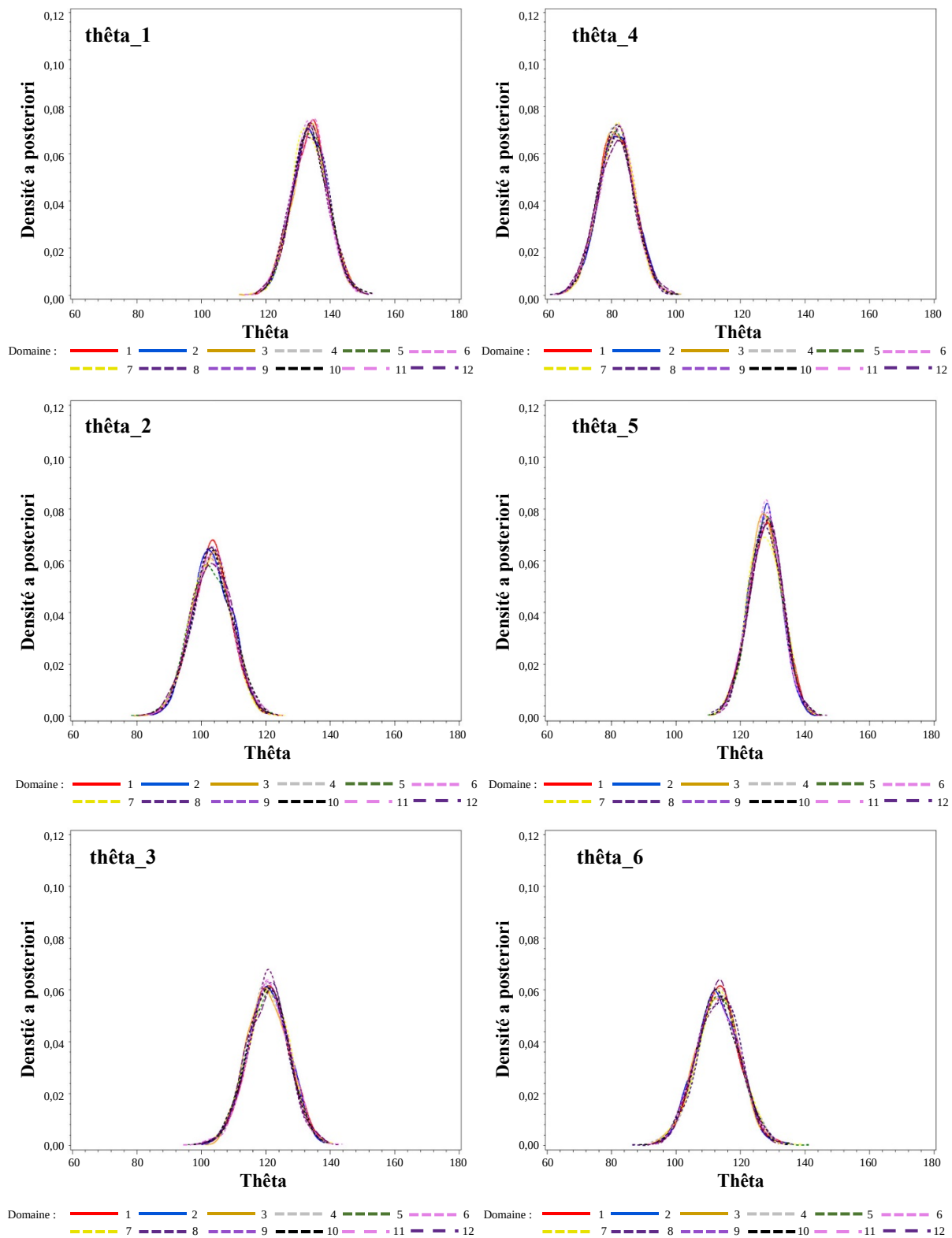
1 000 itérations sélectionnées (échantillonneur de Gibbs efficace). À noter que la cible est 12 162,93 et la somme des MP du modèle BFH, de 12 168,49, une différence de 5,56. À la figure 4.5, nous présentons le tracé des coefficients de variation pour la réconciliation à suppression aléatoire, la réconciliation à suppression du dernier domaine et le modèle BFH par rapport aux estimations directes par domaine. Les différences entre ces modèles ne sont pas remarquables. La plupart des points dont les CV des estimations directes sont supérieurs à 0,04 environ se situent sous la droite à 45 °. Toutefois, certains points (en losange) pour le modèle BFH se trouvent au-dessus de cette même droite et quatre d'entre eux sont à remarquer, peut-être parce qu'ils rétrécissent trop. Nous en concluons qu'il est logique d'opter pour la réconciliation à suppression aléatoire.

Tableau 4.8

Résumé de la comparaison d'inférence avec l'estimateur direct, le modèle Fay-Herriot bayésien (BFH), la réconciliation à suppression du dernier domaine (DD) et la réconciliation à suppression aléatoire (SA) dans les cas où le dernier comté est extrême

	Domaine	n	$\hat{\theta}$	s	BFH		DD		SA	
					MP	ETP	MP	ETP	MP	ETP
a. La taille du dernier comté est de 2.	1	5	135,575	6,031	134,116	5,607	133,772	5,473	133,510	5,409
	2	7	101,980	7,101	103,205	6,482	102,818	6,118	102,745	5,837
	3	24	117,655	7,309	121,110	6,730	120,911	6,577	120,666	6,260
	4	23	76,997	5,881	81,586	6,021	81,741	5,544	81,196	5,631
	5	21	126,917	5,629	127,901	5,252	127,552	5,264	127,619	5,041
	6	9	113,132	8,061	113,454	7,147	112,889	6,818	113,074	6,815
	7	5	137,236	6,739	133,938	6,339	133,479	5,968	133,947	5,994
	8	20	124,839	4,034	124,753	3,906	124,699	3,824	124,738	3,735
	9	16	118,306	7,544	116,199	6,806	115,329	6,327	116,065	6,785
	10	9	156,503	4,368	153,419	4,434	153,148	4,174	153,240	4,213
	11	23	109,546	4,877	110,512	4,645	110,473	4,696	110,324	4,686
	12	2	121,881	12,75	124,243	10,00	123,755	8,771	123,444	8,525
b. La taille du dernier comté est de 50.	1	5	135,575	6,031	133,984	5,618	133,745	5,461	133,452	5,385
	2	7	101,980	7,101	103,462	6,499	103,136	6,086	103,044	5,780
	3	24	117,655	7,309	121,006	6,716	120,832	6,536	120,698	6,232
	4	23	76,997	5,881	81,473	5,995	81,596	5,512	81,162	5,728
	5	21	126,917	5,629	127,832	5,248	127,519	5,238	127,661	5,001
	6	9	113,132	8,061	113,393	7,146	112,929	6,777	112,899	6,675
	7	5	137,236	6,739	133,659	6,380	133,351	5,947	133,851	5,941
	8	20	124,839	4,034	124,732	3,906	124,713	3,821	124,726	3,825
	9	16	118,306	7,544	116,480	6,785	115,766	6,269	116,319	6,601
	10	9	156,503	4,368	153,355	4,449	153,225	4,173	153,306	4,230
	11	23	109,546	4,877	110,347	4,637	110,378	4,692	110,155	4,689
	12	50	116,538	6,791	118,117	6,282	118,035	5,958	117,952	5,702

Note : n est la taille d'échantillon de domaines, $\hat{\theta}$ est l'estimateur direct et s sont erreur-type. MP est la moyenne a posteriori et ETP l'écart-type a posteriori. Quand la taille d'échantillon du dernier comté est de 50 (2), la valeur de réconciliation est de 1 435 (1 441). La distribution antérieure uniforme intervient dans la réconciliation à suppression aléatoire.



Figures 4.1 Comparaison des densités a posteriori de θ_1 à θ_6 lorsque chaque domaine est retranché à son tour (le premier domaine est supprimé dans le premier panneau, etc.).

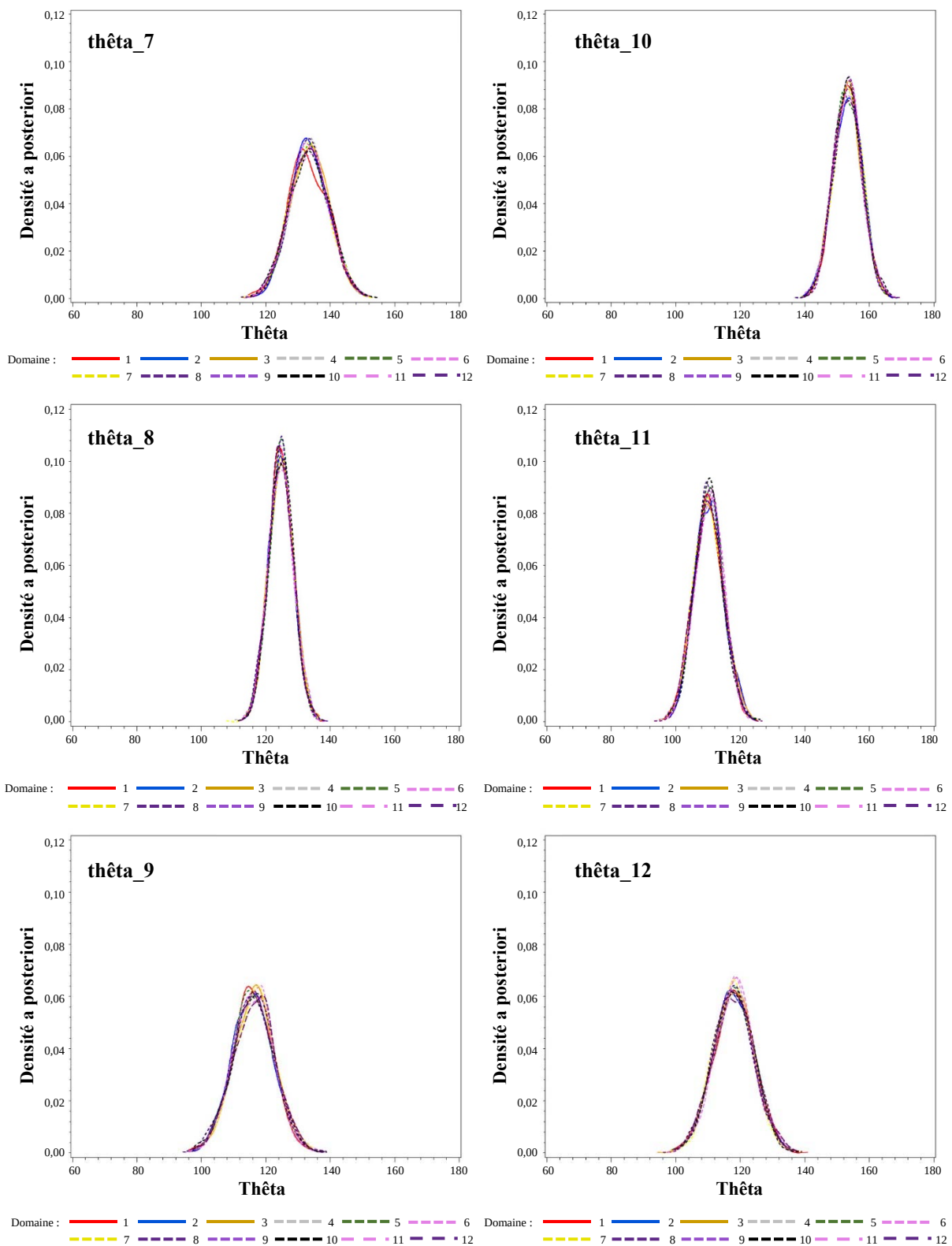


Figure 4.2 Comparaison des densités a posteriori de θ_7 à θ_{12} lorsque chaque domaine est supprimé à son tour.

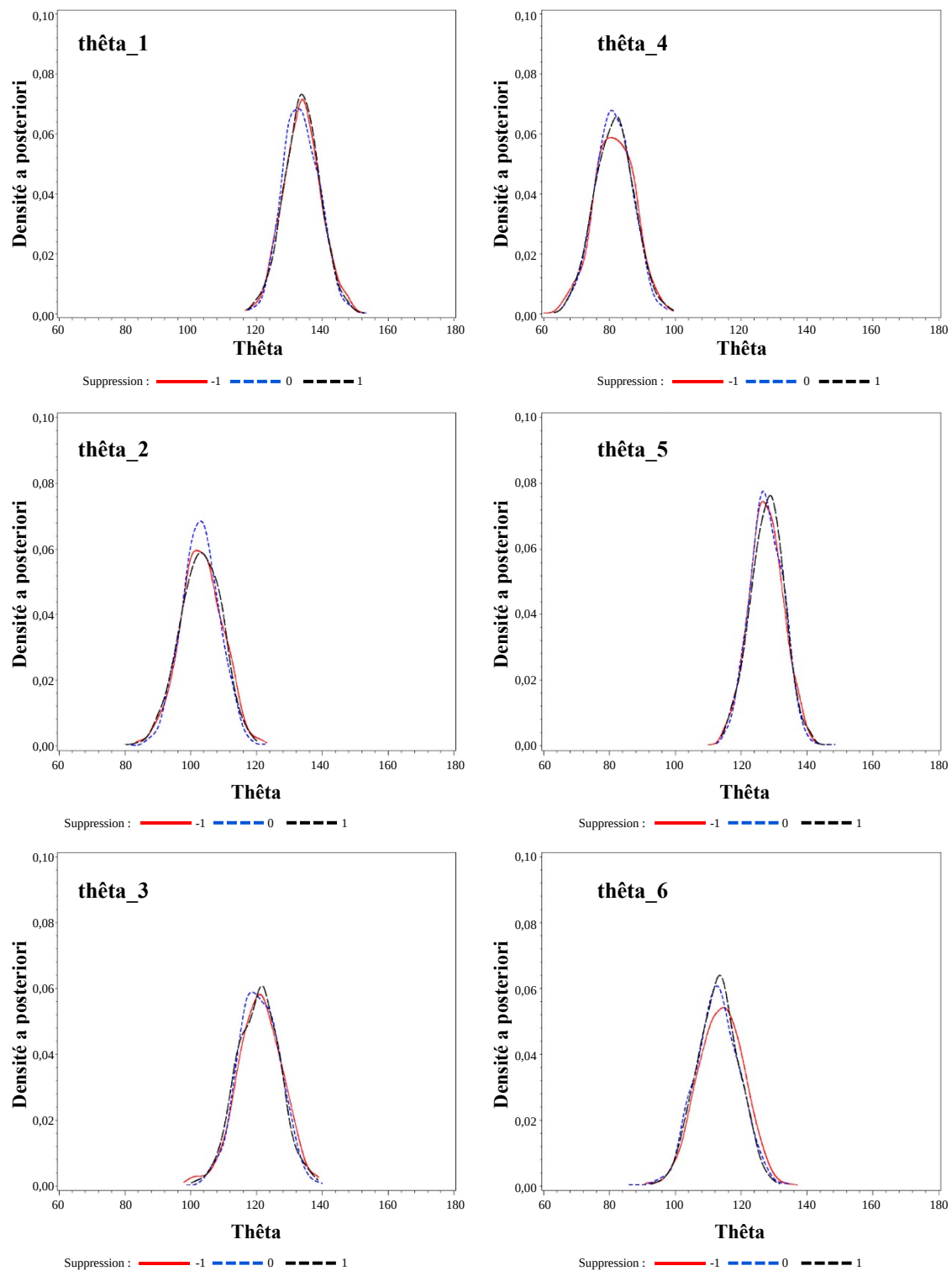


Figure 4.3 Comparaison des densités a posteriori de θ_1 à θ_6 avec le modèle de Fay-Herriot (-1), la réconciliation à suppression aléatoire (0) et la réconciliation à suppression du 12^e domaine.

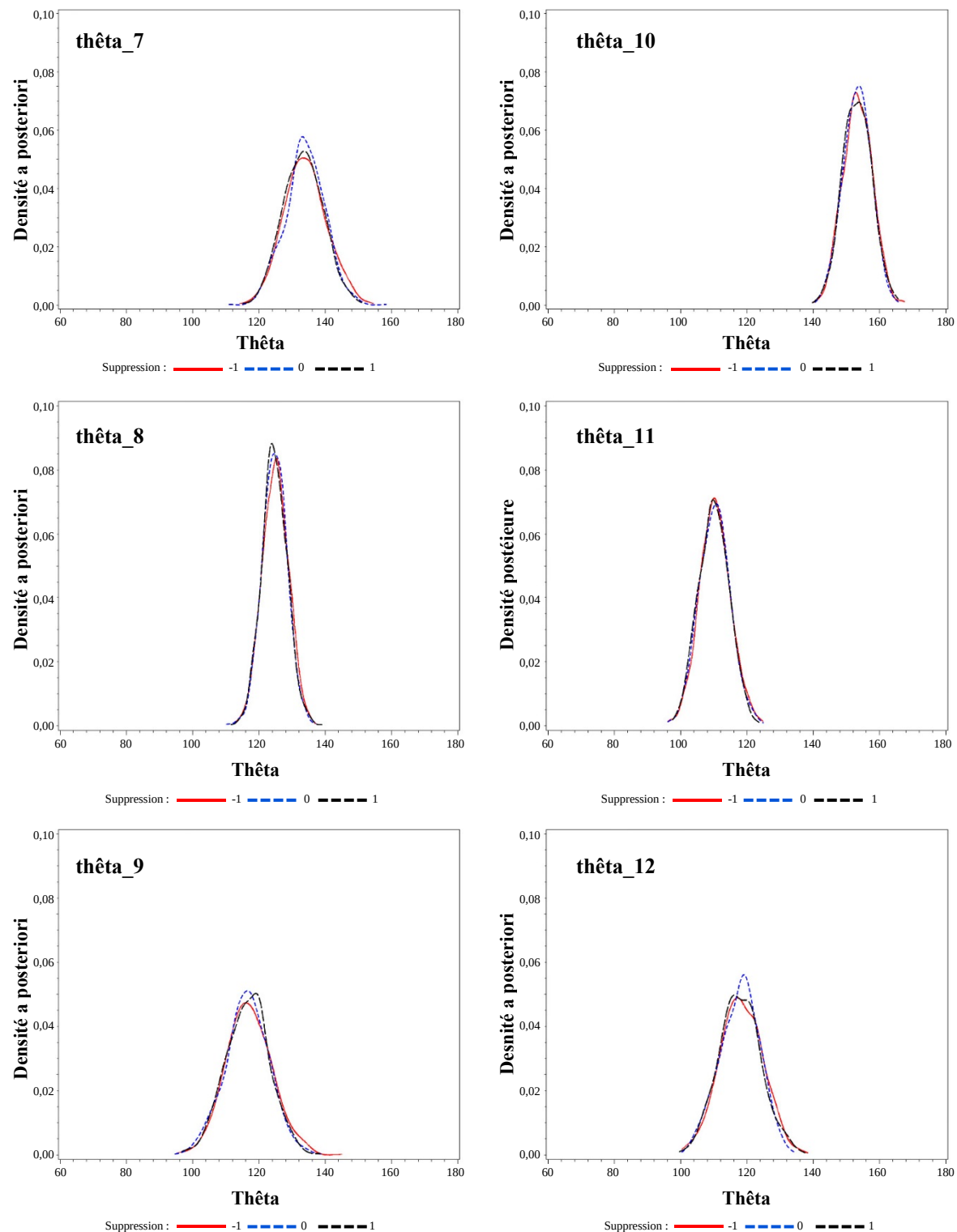


Figure 4.4 Comparaison des densités a posteriori de θ_7 à θ_{12} avec le modèle de Fay-Herriot (-1), la réconciliation à suppression aléatoire (0) et la réconciliation à suppression du 12^e domaine.

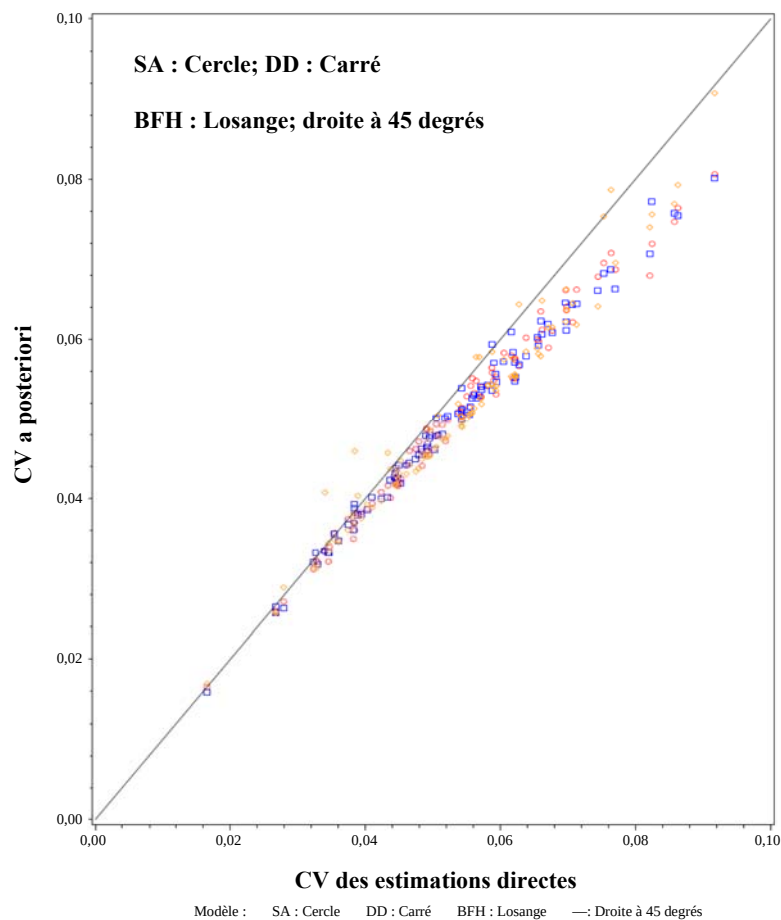


Figure 4.5 Tracé des coefficients de variation avec la réconciliation à suppression aléatoire, la réconciliation à suppression du dernier domaine et le modèle de Fay-Herriot bayésien pour 99 domaines.

5 Observations en conclusion

Nous examinons ici en détail le modèle de Fay-Herriot bayésien (BFH). Nous démontrons que ce modèle peut faire l'objet d'un ajustement à l'aide d'échantillons aléatoires plutôt qu'avec un échantillonneur de Monte Carlo à chaîne de Markov. Comme les échantillons aléatoires n'exigent aucune surveillance, la méthode est avantageuse, car le NASS ne dispose guère de temps entre la réception des données sommaires d'enquête au niveau des comtés et la présentation des estimations finales. Dans le sens du modèle BFH, nous montrons que la densité a posteriori sous le modèle BFH est propre, pouvant servir de base à une réconciliation. Nous étudions les effets de la réconciliation dans une étude par simulation où nous comparons les modèles BFH sans réconciliation et aux modèles BHF avec deux méthodes de réconciliation.

Dans cette étude, nous supposons que la contrainte de réconciliation est de la forme $\sum_{i=1}^{\ell} \theta_i = a$. On peut concevoir une généralisation simple des méthodes de réconciliation pour une contrainte de la forme

$\sum_{i=1}^{\ell} w_i \theta_i = a$, où les w_i sont les poids. Ce dernier cas se présente, par exemple, pour la réconciliation des rendements, des rapports de production et des superficies de récolte.

Notre grande contribution consiste à élargir le modèle BFH en fonction d'une réconciliation. Les méthodes du passé venaient supprimer le dernier domaine, d'où la question de savoir si le choix du domaine à supprimer a de l'importance. Dans cet article, nous concevons et illustrons une méthode donnant une chance de suppression à chaque domaine. Nous indiquons comment ajuster ce modèle BFH élargi au moyen de l'échantillonneur de Gibbs. Une méthode à base d'échantillonnage sans chaînes de Markov ne peut être employée ici en raison de la complexité de la densité a posteriori conjointe. Nous démontrons par des études empiriques que les différences de moyennes a posteriori sont très petites entre les modèles sans réconciliation, avec réconciliation à suppression du dernier comté et avec réconciliation à suppression aléatoire.

Nous avons étudié dans une analyse de sensibilité les effets d'un changement de cible de réconciliation. Comme on pouvait le prévoir, un tel changement mène à des estimations différentes, mais ce qui est inattendu, c'est que les différences des écarts-types a posteriori sont petites. Nous relevons pour les méthodes de réconciliation une faible variation des estimations selon les différentes probabilités de suppression.

On s'attend à ce que, à la suite d'une réconciliation à suppression du dernier domaine et à suppression aléatoire, les écarts-types a posteriori soient plus élevés qu'avec le modèle BFH à cause du bruit aléatoire que crée la réconciliation. Il reste que, dans les études empiriques que nous présentons, les réconciliations à suppression du dernier domaine et à suppression aléatoire donnent en gros les mêmes écarts-types a posteriori avec une modeste baisse dans le cas de la suppression aléatoire. Le grand atout avec la méthode de réconciliation à suppression aléatoire est qu'aucun domaine ou comté ne reçoit de traitement préférentiel.

Avertissement et remerciements

Le département de l'Agriculture des États-Unis n'a pas diffusé officiellement les constatations et conclusions de cette publication préliminaire et celle-ci ne doit pas être interprétée comme représentant une décision ou une politique de l'organisme. La présente étude a été soutenue en partie par le programme de recherche interne du *National Agricultural Statistics Service* du département de l'Agriculture.

Les travaux du D^r Nandram ont été financés par une subvention de la Simons Foundation (353953, Balgobin Nandram). Les auteurs remercient le rédacteur en chef adjoint et les examinateurs de leurs observations et leurs suggestions. Les travaux de Erciulescu ont été réalisés en tant que chercheur associé pour les projets du NASS au *National Institute of Statistical Sciences* (NISS).

Annexe A

Illustration de la sensibilité de la suppression

Soit $y_i \stackrel{\text{ind}}{\sim} \text{Normale}(\mu_i, \sigma_i^2)$, $i = 1, 2$, de sorte que $y_1 + y_2 = a$, et $\lambda = \sigma_2^2 / (\sigma_1^2 + \sigma_2^2)$. Si nous commençons par supprimer y_2 , la densité conjointe de (y_1, y_2) est

$$f(y_1, y_2 | \phi = 0) = \delta_{y_2}(a - y_1) \text{ Normale} \{ \lambda \mu_1 + (1 - \lambda)(a - \mu_2), (1 - \lambda) \sigma_2^2 \},$$

où $\delta_a(b) = 1$ si $a = b$ et $\delta_a(b) = 0$ si $a \neq b$. Toutefois, si nous supprimons d'abord y_1 , la densité conjointe de (y_1, y_2) est

$$g(y_1, y_2 | \phi = 0) = \delta_{y_1}(a - y_2) \text{ Normale} \{ \lambda(a - \mu_1) + (1 - \lambda)\mu_2, (1 - \lambda)\sigma_2^2 \}.$$

Dans la procédure d'estimation, la variable qui est supprimée en particulier a de l'importance, parce que les deux codistributions sont différentes. À noter que les deux distributions se confondent si et seulement si

$$\lambda \mu_1 + (1 - \lambda)(a - \mu_2) = \lambda(a - \mu_1) + (1 - \lambda)\mu_2,$$

ce qui donne

$$\lambda(\mu_1 - a/2) = (1 - \lambda)(\mu_2 - a/2).$$

Même si nous supposons que $\sigma_1^2 = \sigma_2^2$, les deux restent différentes. Toutefois, $\lambda = 1/2$, avec cette hypothèse et la condition pour que les deux distributions se confondent est que $\mu_1 = \mu_2$. En d'autres termes, la condition à respecter dans l'ensemble pour cette identité des codistributions est que $\mu_1 = \mu_2$ et $\sigma_1 = \sigma_2$, ce qui rend y_1 et y_2 interchangeables. C'est néanmoins là un cas très restreint.

Pour éviter la difficulté, on peut en réalité supprimer tant y_1 que y_2 d'une certaine manière. Soit $z = 1$ si y_1 est supprimé et $z = 0$ si y_2 l'est. Dans ce cas,

$$p(y_1, y_2, z | \phi = 0) = [pg(y_1, y_2 | \phi = 0)]^z [(1 - p)f(y_1, y_2 | \phi = 0)]^{1-z},$$

où nous avons pris $z \sim \text{Bernoulli}(p)$ et, comme z n'est pas réellement identifiable, nous allons prendre $p = 1/2$ (suppression aléatoire de l'un ou de l'autre). À noter cependant que

$$z | y_1, y_2, \phi = 0 \sim \text{Bernoulli} \left\{ \frac{pg(y_1, y_2 | \phi = 0)}{[pg(y_1, y_2 | \phi = 0)] + [(1 - p)f(y_1, y_2 | \phi = 0)]} \right\}.$$

Annexe B

Ajustement du modèle de Fay-Herriot bayésien

Le modèle de Fay-Herriot bayésien (BFH) est donné en (2.1) et la densité a posteriori conjointe avec ce modèle, en (2.3). Nous l'exprimons ainsi par commodité ici :

$$\pi(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 | \hat{\boldsymbol{\theta}}) \propto \frac{1}{(1 + \sigma^2)^2} \left(\frac{1}{\sigma^2} \right)^{\ell/2} \prod_{i=1}^{\ell} \left\{ \exp \left[-\frac{1}{2} \left\{ \frac{1}{s_i^2} (\hat{\theta}_i - \theta_i)^2 + \frac{1}{\sigma^2} (\theta_i - \mathbf{x}_i' \boldsymbol{\beta})^2 \right\} \right] \right\}. \quad (\text{B.1})$$

Nous montrons comment procéder à l'ajustement de la densité a posteriori conjointe des paramètres à l'aide d'échantillons aléatoires (pas même avec un échantillonneur de Gibbs), ce qui permet d'éviter toute surveillance des données. Nous emploierons la règle de multiplication pour écrire

$$\pi(\boldsymbol{\theta}, \boldsymbol{\beta}, \sigma^2 | \hat{\boldsymbol{\theta}}) = \pi_1(\boldsymbol{\theta} | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}}) \pi_2(\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}}) \pi_3(\sigma^2 | \hat{\boldsymbol{\theta}}),$$

où $\pi_1(\boldsymbol{\theta} | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}})$ et $\pi_2(\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}})$ ont des formes types et où $\pi_3(\sigma^2 | \hat{\boldsymbol{\theta}})$ est non standard, mais étant la densité d'un paramètre unique.

Pour le moment, nous oublierons le terme $\frac{1}{(1+\sigma^2)^2} \left(\frac{1}{\sigma^2}\right)^{\ell/2}$, parce qu'il touche seulement la densité a posteriori de σ^2 . En d'autres termes,

$$\pi_1(\boldsymbol{\theta} | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}}) \propto \prod_{i=1}^{\ell} \left\{ \exp \left[-\frac{1}{2} \left\{ \frac{1}{s_i^2} (\hat{\theta}_i - \theta_i)^2 + \frac{1}{\sigma^2} (\theta_i - \mathbf{x}_i' \boldsymbol{\beta})^2 \right\} \right] \right\}.$$

Les calculs standard réduisent l'argument (sans $-1/2$) du terme exponentiel à

$$\frac{1}{(1-\lambda_i)\sigma^2} \left\{ \theta_i - (\lambda_i \hat{\theta}_i + (1-\lambda_i) \mathbf{x}_i' \boldsymbol{\beta}) \right\}^2 + \frac{\lambda_i}{\sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \boldsymbol{\beta})^2, \quad \lambda_i = \frac{\sigma^2}{s_i^2 + \sigma^2}, \quad i = 1, \dots, \ell.$$

Ainsi, pour $\pi_1(\boldsymbol{\theta} | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}})$,

$$\theta_i | \boldsymbol{\beta}, \sigma^2, \hat{\boldsymbol{\theta}} \stackrel{\text{ind}}{\sim} \text{Normale} \left\{ \lambda_i \hat{\theta}_i + (1-\lambda_i) \mathbf{x}_i' \boldsymbol{\beta}, (1-\lambda_i) \sigma^2 \right\}, \quad i = 1, \dots, \ell. \quad (\text{B.2})$$

Pour le moment aussi, nous oublierons le terme $\prod_{i=1}^{\ell} [(1-\lambda_i) \sigma^2]^{1/2}$. Si nous marginalisons les θ_i , nous obtenons

$$\pi_2(\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}}) \propto \exp \left\{ -\frac{1}{2} \sum_{i=1}^{\ell} \frac{\lambda_i}{\sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \boldsymbol{\beta})^2 \right\}.$$

Ainsi, l'exposant (sans $-1/2$) peut se formuler comme

$$\sum_{i=1}^{\ell} \frac{\lambda_i}{\sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \hat{\boldsymbol{\beta}})^2 + (\boldsymbol{\beta} - \hat{\boldsymbol{\beta}})' \hat{\Sigma}^{-1} (\boldsymbol{\beta} - \hat{\boldsymbol{\beta}}),$$

où

$$\hat{\boldsymbol{\beta}} = \hat{\Sigma} \sum_{i=1}^{\ell} \frac{\hat{\theta}_i \mathbf{x}_i}{s_i^2 + \sigma^2} \quad \text{et} \quad \hat{\Sigma}^{-1} = \sum_{i=1}^{\ell} \frac{\mathbf{x}_i \mathbf{x}_i'}{s_i^2 + \sigma^2}.$$

Il convient de noter que $\hat{\boldsymbol{\beta}}$ et $\hat{\Sigma}$ sont bien définis pour tous les σ^2 à condition que la matrice de plan, X , avec $X' = (\mathbf{x}_1, \dots, \mathbf{x}_{\ell})$, soit de plein rang. Dans ce cas,

$$\pi_2(\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}}) \propto \exp \left\{ -\frac{1}{2} \sum_{i=1}^{\ell} \frac{\lambda_i}{\sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \hat{\boldsymbol{\beta}})^2 - \frac{1}{2} (\boldsymbol{\beta} - \hat{\boldsymbol{\beta}})' \hat{\Sigma}^{-1} (\boldsymbol{\beta} - \hat{\boldsymbol{\beta}}) \right\}.$$

En d'autres termes,

$$\boldsymbol{\beta} | \sigma^2, \hat{\boldsymbol{\theta}} \sim \text{Normale}(\hat{\boldsymbol{\beta}}, \hat{\Sigma}). \quad (\text{B.3})$$

Si nous marginalisons $\boldsymbol{\beta}$ et incorporons dans σ^2 les termes que nous avons négligés, nous obtenons

$$\pi_3(\sigma^2 | \hat{\theta}) \propto Q(\sigma^2) \frac{1}{(1 + \sigma^2)^2}, \quad (\text{B.4})$$

où

$$Q(\sigma^2) = |\hat{\Sigma}|^{1/2} \prod_{i=1}^{\ell} \frac{1}{(s_i^2 + \sigma^2)^{1/2}} \exp \left\{ -\frac{1}{2} \sum_{i=1}^{\ell} \frac{1}{s_i^2 + \sigma^2} (\hat{\theta}_i - \mathbf{x}_i' \hat{\boldsymbol{\beta}})^2 \right\}.$$

Pour tirer un échantillon aléatoire de (B.1), nous échantillonnons σ^2 par (B.4), $\boldsymbol{\beta}$ par (B.3) et les θ_i indépendamment par (B.2). La densité a posteriori conditionnelle en (B.4) est non standard et, pour en tirer un échantillon, nous employons une méthode de grille (voir Nandram et Yin (2016), par exemple). D'abord, nous transformons σ^2 en $\phi = \sigma^2 / (1 + \sigma^2)$ de sorte que $0 < \phi < 1$. Nous divisons ensuite $(0, 1)$ en 100 grilles. En réalité, nous avons situé la fourchette de ϕ dans $(0, 1)$ et divisé cet intervalle en 100 grilles, ce qui nous donne une fonction de masse de probabilité que nous échantillonnons. Nous procédons par bruit aléatoire dans la grille sélectionnée pour obtenir des écarts qui seront différents avec une probabilité un (pour plus de détails, voir Nandram et Yin, 2016).

Annexe C

Preuve du théorème 2

Il est pratique d'opérer les transformations suivantes,

$$\theta_i = \theta_i, i = 1, \dots, \ell - 1, \phi = \sum_{i=1}^{\ell} \theta_i - a.$$

Là, ϕ est une variable nominale correspondant à la contrainte de réconciliation, ce qui garantit que la transformation sera non singulière. Le jacobien est l'unité et la transformation inverse est

$$\theta_i = \theta_i, i = 1, \dots, \ell - 1, \theta_{\ell} = \phi + a - \sum_{i=1}^{\ell-1} \theta_i.$$

La densité transformée est

$$\tilde{\pi}(\theta_1, \dots, \theta_{\ell-1}, \phi).$$

Ainsi, la densité qui correspond exactement à la contrainte de réconciliation est

$$\tilde{\pi}(\theta_1, \dots, \theta_{\ell-1} | \phi = 0), \theta_{\ell} = a - \sum_{i=1}^{\ell-1} \theta_i.$$

Ainsi,

$$\pi(\theta_1, \dots, \theta_{\ell-1} | \phi = 0) \propto \exp \left\{ -\frac{1}{2\sigma^2} \left[\sum_{i=1}^{\ell-1} (\theta_i - \mathbf{u}_i' \boldsymbol{\beta})^2 + \left\{ \sum_{i=1}^{\ell-1} \theta_i - (a - \mathbf{u}_i' \boldsymbol{\beta}) \right\}^2 \right] \right\}.$$

Si nous oublions les termes sans $\boldsymbol{\theta}_{(\ell)} = (\theta_1, \dots, \theta_{\ell-1})'$, il est facile de démontrer que l'exposant est

$$\frac{1}{2\delta^2} \left\{ \boldsymbol{\theta}_{(\ell)}' (I + J) \boldsymbol{\theta}_{(\ell)} - 2 \left[(\mathbf{u}_1 - \mathbf{u}_\ell)' \boldsymbol{\beta}, \dots, (\mathbf{u}_{\ell-1} - \mathbf{u}_\ell)' \boldsymbol{\beta} + a\mathbf{j}' \right] \boldsymbol{\theta}_{(\ell)} \right\}.$$

Si nous employons les propriétés d'une densité normale multivariée, nous obtenons

$$\boldsymbol{\theta}_{(\ell)} | \phi = 0 \sim \text{Normale} \left((I + J)^{-1} \left(a\mathbf{j}' + (\mathbf{u}_1 - \mathbf{u}_\ell)' \boldsymbol{\beta}, \dots, (\mathbf{u}_{\ell-1} - \mathbf{u}_\ell)' \boldsymbol{\beta} \right)', \delta^2 (I + J)^{-1} \right).$$

Par la formule de Sherman-Morrison, $(I + J)^{-1} = I - \frac{1}{\ell} J$, nous dégageons enfin la formule

$$\boldsymbol{\theta}_{(\ell)} | \phi = 0 \sim \text{Normale} \left\{ \left(I - \frac{1}{\ell} J \right) \left(a\mathbf{j}' + (\mathbf{u}_1 - \mathbf{u}_\ell)' \boldsymbol{\beta}, \dots, (\mathbf{u}_{\ell-1} - \mathbf{u}_\ell)' \boldsymbol{\beta} \right)', \delta^2 \left(I - \frac{1}{\ell} J \right) \right\}.$$

Il convient de noter que le lemme du déterminant de la matrice donne $\det(I + J) = \ell$ et donc $\det(\delta^2 (I - \frac{1}{\ell} J)) = \frac{1}{\ell} (\delta^2)^{\ell-1}$.

Bibliographie

- Battese, G., Harter, R. et Fuller, W. (1988). An error-components model for prediction of county crop areas using survey and satellite data. *Journal of the American Statistical Association*, 83(401), 28-36. doi:10.2307/2288915.
- Bell, W.R., Datta, G.S. et Ghosh, M. (2013). Benchmarking small area estimators. *Biometrika*, 100(1), 189-202.
- Cruze, N.B., Erciulescu, A.L., Nandram, B., Barboza, W. et Young, L.J. (2019). Producing official county-level agricultural estimates in the United States: Needs and challenges. *Statistical Science*. À paraître.
- Datta, G.S., Ghosh, M., Steorts, R. et Maples, J. (2011). Bayesian benchmarking with applications to small area estimation. *Test*, 20(3), 574-588.
- Erciulescu, A.L., Cruze, N.B. et Nandram, B. (2019). Model-based county level crop estimates incorporating auxiliary sources of information. *Journal of Royal Statistical Society, Series A*, 182, 283-303. doi:10.1111/rssa.12390.
- Fay, R.E., et Herriot, R.A. (1979). Estimates of income for small places: An application of James-Stein procedures to census data. *Journal of the American Statistical Association*, 74(366), 269-277.
- Ghosh, M., et Steorts, R. (2013). Two stage Bayesian benchmarking as applied to small area estimation. *Test*, 22(4), 670-687.
- Janicki, R., et Vesper, A. (2017). Benchmarking techniques for reconciling Bayesian small area models at distinct geographic levels. *Statistical Methods and Applications*, 26, 4, 557-581.
- Jiang, J., et Lahiri, P. (2006). Mixed model prediction and small area estimation. *Test*, 15(1), 1-96.

- Nandram, B., et Sayit, H. (2011). Une analyse bayésienne des probabilités de réponse dans les petits domaines sous une contrainte. *Techniques d'enquête*, 37, 2, 147-162. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2011002/article/11603-fra.pdf>.
- Nandram, B., et Toto, M.C.S. (2010). Bayesian predictive inference for benchmarking crop production for Iowa counties. *Journal of the Indian Society of Agricultural Statistics*, 64 (2), 191-207.
- Nandram, B., et Yin, J. (2016). A nonparametric Bayesian prediction interval for a finite population mean. *Journal of Statistical Computation and Simulation*, 86 (16), 3141-3157.
- Nandram, B., Berg, E. et Barboza, W. (2014). A hierarchical Bayesian model for forecasting state-level corn yield. *Journal of Environmental and Ecological Statistics*, 21, 507-530.
- Nandram, B., Toto, M.C.S. et Choi, J.W. (2011). A Bayesian benchmarking of the Scott-Smith model for small areas. *Journal of Statistical Computation and Simulation*, 81 (11), 1593-1608.
- Pfeffermann, D. (2013). New important developments in small area estimation. *Statistical Science*, 28(1), 40-68.
- Pfeffermann, D., Sikov, A. et Tiller, R. (2014). Single-and two-stage cross-sectional and time series benchmarking procedures for small area estimation. *Test*, 23(4), 631-666.
- Pfeffermann, D., et Tiller, R. (2006). Small area estimation with state-space models subject to benchmark constraints. *Journal of the American Statistical Association*, 101 (476), 1387-1397.
- Rao, J.N.K., et Molina, I. (2015). *Small Area Estimation*, Wiley Series in Survey Methodology.
- Toto, M.C.S., et Nandram, B. (2010). A Bayesian predictive inference for small area means incorporating covariates and sampling weights. *Journal of Statistical Planning and Inference*, 140, 2963-2979.
- Wang, J., Fuller, W.A. et Qu, Y. (2008). Estimation pour petits domaines sous une contrainte. *Techniques d'enquête*, 34, 1, 33-40. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2008001/article/10619-fra.pdf>.
- You, Y., et Rao, J.N.K. (2002). Small area estimation using unmatched sampling and linking models. *The Canadian Journal of Statistics*, 30, 3-15.
- You, Y., Rao, J.N.K. et Dick, J.P. (2004). Benchmarking hierarchical Bayes small area estimators in the Canadian census under coverage estimation. *Statistics in Transition*, 6, 631-640.