

Techniques d'enquête

Estimation du niveau et de la variation du chômage au moyen de modèles de séries chronologiques structurels

par Harm Jan Boonstra et Jan A. van den Brakel

Date de diffusion : le 17 décembre 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.

Estimation du niveau et de la variation du chômage au moyen de modèles de séries chronologiques structurels

Harm Jan Boonstra et Jan A. van den Brakel

Résumé

On obtient les estimations mensuelles du chômage provincial fondées sur l'Enquête sur la population active (EPA) des Pays-Bas au moyen de modèles de séries chronologiques. Les modèles tiennent compte du biais de renouvellement et de la corrélation sérielle causée par le plan d'échantillonnage à panel rotatif de l'EPA. L'article compare deux méthodes d'estimation de modèles de séries chronologiques structurels (MSCS). Dans la première méthode, les MSCS sont exprimés sous forme de modèles espace-état, auxquels sont appliqués le filtre et le lisseur de Kalman dans un cadre fréquentiste. L'autre solution consiste à exprimer ces MSCS sous forme de modèles multiniveaux de séries chronologiques dans un cadre bayésien hiérarchique et à les estimer à l'aide de l'échantillonneur de Gibbs. Nous comparons ici les estimations mensuelles du chômage et les erreurs-types fondées sur ces modèles pour les 12 provinces des Pays-Bas. Nous discutons ensuite des avantages et des inconvénients de la méthode multiniveau et de la méthode espace-état.

Les MSCS multivariés conviennent pour l'emprunt d'information dans le temps et l'espace. La modélisation de la matrice de corrélation complète entre les composantes des séries chronologiques accroît rapidement le nombre d'hyperparamètres qu'il faut estimer. La modélisation de facteur commun est une des façons possibles d'obtenir des modèles plus parcimonieux qui continuent de tenir compte de la corrélation transversale. L'article propose une méthode encore plus parcimonieuse, dans laquelle les domaines ont en commun une tendance globale et leurs propres tendances indépendantes pour les écarts propres au domaine par rapport à la tendance globale. L'approche par modélisation de séries chronologiques est particulièrement adaptée à l'estimation de la variation mensuelle du chômage.

Mots-clés : Estimation sur petits domaines; modèles de séries chronologiques structurels; modèles multiniveaux de séries chronologiques; estimation du chômage.

1 Introduction

Le bureau central de la statistique des Pays-Bas utilise les données de l'Enquête néerlandaise sur la population active (EPA) pour estimer la situation d'emploi à divers niveaux d'agrégation. Les estimations nationales sont produites mensuellement, les estimations provinciales trimestriellement et les estimations municipales annuellement. Traditionnellement, les publications mensuelles sur la population active étaient fondées sur des chiffres trimestriels mobiles compilés au moyen d'une estimation par la régression généralisée directe (GREG), voir par exemple Särndal, Swensson et Wretman (1992). La nature continue de l'EPA permet d'emprunter de l'information non seulement à d'autres domaines, mais aussi dans le temps. Un modèle de séries chronologiques structurel (MSCS) permettant d'estimer la situation d'emploi mensuelle nationale pour six catégories d'âge selon le sexe est utilisé depuis 2010 (van den Brakel et Krieg, 2009, 2015).

Jusqu'à maintenant, les estimations provinciales sont produites tous les trimestres au moyen de l'estimateur GREG. Afin de produire des chiffres mensuels, il faut une stratégie d'estimation fondée sur un modèle pour surmonter le problème de la taille trop petite des échantillons provinciaux mensuels. Nous

1. Harm Jan Boonstra, bureau central de la statistique des Pays-Bas, Département des méthodes statistiques. Courriel : h.jh.boonstra@cbs.nl; Jan A. van den Brakel, bureau central de la statistique des Pays-Bas, Département des méthodes statistiques et Université de Maastricht, Département d'économie quantitative.

proposons dans l'article un modèle qui combine, d'une part, une approche de modélisation de séries chronologiques servant à emprunter de l'information dans le temps et, d'autre part, des modèles transversaux d'estimation sur petits domaines servant à emprunter de l'information dans l'espace dans le but de produire des estimations mensuelles fiables du chômage provincial. En raison du plan d'échantillonnage par panel de l'EPA, les estimations mensuelles GREG sont autocorrélées et les estimations fondées sur les vagues de suivi sont biaisées par rapport aux estimations de la première vague. On désigne souvent le dernier phénomène par l'expression biais de renouvellement (Bailar, 1975). Les deux caractéristiques doivent être prises en compte dans le modèle (Pfeffermann, 1991). D'autres comptes rendus de travaux d'estimation sur petits domaines du chômage régional avec emprunt d'information dans le temps et l'espace ont été publiés, citons notamment Rao et Yu (1994); Datta, Lahiri, Maiti et Lu (1999); You, Rao et Gambino (2003); You (2008); Pfeffermann et Burck (1990); Pfeffermann et Tiller (2006); van den Brakel et Krieg (2016); voir aussi Rao et Molina (2015), section 4.4 pour une vue d'ensemble.

Dans le présent article, on a élaboré les MSCS multivariés appliqués aux données provinciales mensuelles sur la population active sous forme d'estimation sur petits domaines pour emprunter de l'information dans le temps et l'espace, afin de tenir compte du biais de renouvellement et de la corrélation sérielle induite par le plan d'échantillonnage à panel rotatif. Dans un MSCS, une série observée est décomposée en plusieurs composantes non observées comme une tendance, une composante saisonnière, des composantes de régression, d'autres composantes cycliques et un terme de bruit blanc pour la variation inexpliquée restante. Ces composantes sont basées sur des modèles stochastiques, ce qui leur permet de varier au fil du temps. De manière classique, pour ajuster les MSCS, on les exprime sous forme de modèle espace-état et l'on applique le filtre et le lisseur de Kalman afin d'obtenir des estimations optimales pour les variables d'état et les signaux. Les hyperparamètres inconnus des modèles pour les variables d'état sont estimés au moyen du maximum de vraisemblance (MV) (Harvey, chapitre 3). On peut aussi ajuster les modèles espace-état dans un cadre bayésien à l'aide d'un filtre à particules (Andrieu, Poucet et Holenstein (2010); Durbin et Koopman (2012), chapitre 9). Les MSCS peuvent également être exprimés sous forme de modèles multiniveaux de séries chronologiques et peuvent être considérés comme un prolongement du modèle classique de Fay-Herriot (Fay et Herriot, 1979). Les liens entre les modèles de séries chronologiques structurels et les modèles multiniveaux ont été examinés de plusieurs points de vue, notamment dans Knorr-Held et Rue (2002); Chan et Jeliaskov (2009); McCausland, Miller et Pelletier (2011); Ruiz-Cárdenas, Krainski et Rue (2012); Piepho et Ogutu (2014); Bollineni-Balabay, van den Brakel, Palm et Boonstra (2016). Ces articles décrivent de façon plus explicite l'équivalence entre les composantes du modèle espace-état et les composantes du modèle multiniveau. Les modèles multiniveaux peuvent être ajustés dans un cadre fréquentiste ou un cadre hiérarchique bayésien, voir Rao et Molina (2015), sections 8.3 et 10.9, respectivement.

Le présent article contribue à la littérature sur l'estimation sur petits domaines en comparant les différences entre les MSCS pour les plans d'échantillonnage à panel rotatif exprimés sous forme de modèles espace-état et de modèles multiniveaux de séries chronologiques. Les modèles espace-état sont ajustés au

moyen d'un filtre et d'un lisseur de Kalman dans un cadre fréquentiste où les hyperparamètres sont estimés au moyen du MV. Dans ce cas, les modèles sont comparés au moyen du critère d'information d'Akaike (AIC) et du critère d'information bayésien (BIC). Les modèles multiniveaux de séries chronologiques sont ajustés dans un cadre bayésien hiérarchique, à l'aide de l'échantillonneur de Gibbs. Les modèles comportant différentes combinaisons d'effets fixes et aléatoires sont comparés en fonction du critère d'information de déviance (DIC). Les estimations fondées sur les modèles multiniveaux et espace-état et leurs erreurs-types sont comparées graphiquement et opposées aux estimations initiales par la régression. La modélisation de la corrélation transversale dans les modèles multivariés de séries chronologiques accroît rapidement le nombre d'hyperparamètres qu'il faut estimer. Pour obtenir des modèles plus parcimonieux, on peut notamment utiliser des modèles de facteur commun. Dans le présent article, nous proposons une autre méthode pour modéliser des corrélations entre les composantes des séries chronologiques indirectement, qui se fonde sur une tendance globale commune et des tendances locales pour les écarts propres au domaine.

L'article est structuré selon le plan suivant. La section 2 décrit les données de l'EPA utilisées dans l'étude. La section 3 décrit la façon dont l'estimateur par la régression (Battese, Harter et Fuller, 1988) est utilisé pour le calcul des estimations initiales. Ces estimations initiales sont les données d'entrée des MSCS; elles sont présentées à la section 4. À la section 5, les résultats fondés sur plusieurs modèles multiniveaux et espace-état sont comparés, y compris les estimations de la variation d'une période à l'autre pour les données mensuelles. La section 6 analyse les résultats et présente des pistes de recherche futures. Tout au long de l'article, nous renvoyons au rapport technique de Boonstra et van den Brakel (2016), qui contient plus de détails et de résultats sur ces travaux.

2 L'Enquête sur la population active des Pays-Bas

L'EPA néerlandaise est une enquête auprès des ménages menée selon un plan d'échantillonnage à panel rotatif dans le cadre de laquelle les répondants sont interviewés cinq fois par trimestre. Chaque mois, un échantillon stratifié à deux degrés d'adresses est sélectionné. Tous les ménages résidant à une adresse sont compris dans l'échantillon. Notre étude s'appuie sur 72 mois de données de l'EPA, soit de 2003 à 2008. Au cours de cette période, le plan d'échantillonnage était autopondéré. La première vague du panel se compose de données collectées au moyen d'interviews sur place assistées par ordinateur (IPAO), alors que les quatre vagues de suivi contiennent des données recueillies au moyen d'interviews téléphoniques assistées par ordinateur (ITAO).

Les Pays-Bas sont divisés en 12 provinces, qui servent de domaines aux fins de l'estimation des chiffres mensuels du chômage. La taille des échantillons nationaux mensuels varie de 5 000 à 7 000 personnes pour la première vague et de 3 000 à 5 000 personnes pendant la cinquième vague. Les tailles des échantillons provinciaux varient de 31 à 1 949 personnes pour les échantillons mensuels d'une seule vague.

Nous disposons aussi des données de l'EPA au niveau des unités, c'est-à-dire des personnes. Nous disposons également d'une mine de données auxiliaires provenant de plusieurs enregistrements au niveau des unités. Ces variables auxiliaires comprennent le chômage enregistré, qui est un bon prédicteur de la

variable de chômage d'intérêt Ces prédictors sont utilisés dans le calcul des estimations initiales, qui sont des données d'entrée des modèles de séries chronologiques.

La variable cible considérée dans l'étude est la fraction de personnes au chômage dans un domaine et elle est définie par $\bar{Y}_{it} = \sum_{j \in i} y_{ijt} / N_{it}$, avec y_{ijt} égal à un si une personne j d'une province i dans la période t est au chômage et égal à zéro autrement, et N_{it} étant la taille de la population dans la province i et la période t .

3 Estimations initiales

Soit $\hat{Y}_{itp}$ l'estimation initiale de Y_{it} fondée sur les données de la vague p . Les estimations initiales utilisées comme données d'entrée des modèles sur petits domaines des séries chronologiques sont des estimations par la régression des données d'enquête (Woodruff, 1966; Battese et coll., 1988; Särndal et coll., 1992).

$$\hat{Y}_{itp} = \bar{y}_{itp} + \hat{\beta}'_{tp} (\bar{X}_{it} - \bar{x}_{itp}), \quad (3.1)$$

où $\bar{y}_{itp}$, $\bar{x}_{itp}$ désigne les moyennes d'échantillon, $\bar{X}_{it}$ est le vecteur des moyennes de population des covariables x , et $\hat{\beta}_{tp}$ sont les coefficients de régression estimés. Les coefficients sont estimés séparément pour chaque période et chaque vague, mais ils se basent sur les échantillons nationaux combinant des données de toutes les régions. L'estimateur par la régression des données d'enquête est un estimateur approximativement sans biais sous le plan pour les paramètres de population qui, comme l'estimateur GREG, utilise des données auxiliaires pour réduire le biais de non-réponse. Voir Boonstra et van den Brakel (2016) pour de plus amples renseignements sur le modèle choisi de calcul des estimations par la régression des données d'enquête. Bien que les estimations de coefficient par la régression dans (3.1) ne soient pas selon le domaine, l'estimateur par la régression des données d'enquête est un estimateur de domaine direct dans le sens où il est principalement basé sur les données obtenues dans un domaine et un mois en particulier, et qu'il a par conséquent des erreurs-types inacceptablement élevées en raison de la petite taille des échantillons mensuels de domaine.

Les estimations initiales pour les différentes vagues donnent systématiquement lieu à des différences dans les estimations du chômage, généralement appelées biais de renouvellement (BR) (Bailar, 1975). Les estimations initiales du chômage pour les vagues 2 à 5 sont systématiquement plus petites que celles de la première vague. Ce biais de renouvellement peut s'expliquer par plusieurs causes, dont la sélection, le mode et les effets de panel (van den Brakel et Krieg, 2009). Voir Boonstra et van den Brakel (2016) pour plus de détails et des illustrations graphiques.

Les modèles de séries chronologiques nécessitent également des estimations de la variance correspondant aux estimations initiales. Nous utilisons les estimations lissées transversales suivantes des variances par rapport au plan des estimations par la régression des données d'enquête,

$$v(\hat{Y}_{itp}) = \frac{1}{n_{itp}} \frac{1}{(n_{itp} - m_A)} \sum_{i=1}^{m_A} (n_{itp} - 1) \hat{\sigma}_{itp}^2 \equiv \hat{\sigma}_{itp}^2 / n_{itp}, \text{ avec } \hat{\sigma}_{itp}^2 = \frac{1}{(n_{itp} - 1)} \sum_{j=1}^{n_{itp}} \hat{e}_{ijt_p}^2. \quad (3.2)$$

Ici m_A désigne le nombre de domaines, n_{itp} est le nombre de répondants dans le domaine i , la période t et la vague p , $n_{tp} = \sum_{i=1}^{m_A} n_{itp}$, et $\hat{e}_{ijt_p}$ sont des résidus de l'estimateur par la régression des données d'enquête. Les variances à l'intérieur du domaine $\hat{\sigma}_{itp}^2$ sont regroupées dans les domaines pour obtenir des approximations de variance plus stables. On peut aussi motiver l'utilisation de (3.2) de la façon suivante. Il faut garder à l'esprit que le plan d'échantillonnage est autopondéré. Le calcul des variances à l'intérieur du domaine $\hat{\sigma}_{itp}^2$ tient donc approximativement compte de la stratification, qui est une variable régionale légèrement plus détaillée que la province. L'approximation de la variance tient également compte du calage et de la correction de la non-réponse, puisque les variances à l'intérieur du domaine sont calculées par rapport aux résidus de l'estimateur par la régression des données d'enquête. L'approximation de la variance ne tient pas explicitement compte de la corrélation intra-grappe de personnes au sein des ménages. Toutefois, la corrélation intra-grappe du chômage est faible. De plus, le chômage enregistré est utilisé comme covariable dans l'estimateur par la régression des données d'enquête. Comme cette covariable explique une grande part de la variation du chômage, la corrélation intra-grappe entre résidus est encore plus réduite.

Le plan de sondage entraîne plusieurs corrélations non nulles dans les estimations initiales pour une même province et différentes périodes et vagues. Ces corrélations sont attribuables au chevauchement partiel des ensembles d'unités d'échantillonnage sur lesquels les estimations se fondent. De telles corrélations existent entre les estimations pour une même province dans les mois t_1, t_2 et fondées sur les vagues p_1, p_2 chaque fois que $t_2 - t_1 = 3(p_2 - p_1) \leq 12$. Les covariances entre $\hat{Y}_{it_1p_1}$ et $\hat{Y}_{it_2p_2}$ sont estimées comme étant (voir par exemple, Kish (1965))

$$v(\hat{Y}_{it_1p_1}, \hat{Y}_{it_2p_2}) = \frac{n_{it_1p_1t_2p_2}}{\sqrt{n_{it_1p_1} n_{it_2p_2}}} \hat{\rho}_{t_1p_1t_2p_2} \sqrt{v(\hat{Y}_{it_1p_1}) v(\hat{Y}_{it_2p_2})}, \quad (3.3)$$

avec

$$\hat{\rho}_{t_1p_1t_2p_2} = \frac{1}{(n_{t_1p_1t_2p_2} - m_A)} \sum_{i=1}^{m_A} \sum_{j=1}^{n_{it_1p_1t_2p_2}} \hat{e}_{ijt_1p_1} \hat{e}_{ijt_2p_2},$$

où $n_{it_1p_1t_2p_2}$ est le nombre d'unités dans le chevauchement, c'est-à-dire le nombre d'observations sur les mêmes unités dans le domaine i entre les combinaisons de période et de vague (t_1, p_1) et (t_2, p_2) , et $n_{t_1p_1t_2p_2} = \sum_{i=1}^{m_A} n_{it_1p_1t_2p_2}$. Le coefficient d'(auto)corrélation estimé $\hat{\rho}_{t_1p_1t_2p_2}$ est calculé comme étant la corrélation entre les résidus des modèles de régression linéaire qui sous-tendent les estimateurs par la régression des données d'enquête à (t_1, p_1) et (t_2, p_2) , d'après le chevauchement des deux échantillons pour tous les domaines. De cette façon, ils sont regroupés sur les domaines de la même manière que les variances $\hat{\sigma}_{itp}^2$. Ensemble, (3.2) et (3.3) estiment (une approximation de) la matrice de covariance fondée

sur le plan des estimations initiales par la régression des données d'enquête. Voir Boonstra et van den Brakel (2016) pour de plus de renseignements.

Les estimations de modèle de séries chronologiques pour les chiffres mensuels du chômage provincial seront comparées aux estimations directes. La procédure de calcul des estimations directes mensuelles se fonde sur la méthode utilisée avant 2010 pour calculer les chiffres trimestriels mobiles officiels de la population active. Soit $\hat{Y}_{it}$, l'estimation directe mensuelle pour les provinces, calculée comme étant la moyenne pondérée sur les estimations par la régression des cinq enquêtes par panel, où les poids sont fondés sur les estimations de la variance. Pour corriger le biais de renouvellement, ces estimations directes sont multipliées par un ratio, disons f_{it} , où le numérateur est la moyenne des estimations par la régression des données d'enquête (3.1) pour la première vague sur les trois dernières années, et le dénominateur est la moyenne des estimations directes mensuelles $\hat{Y}_{it}$, également sur les trois dernières années, soit $\tilde{Y}_{it} = f_{it} \hat{Y}_{it}$. Voir Boonstra et van den Brakel (2016) pour une information plus détaillée sur le calcul de $\tilde{Y}_{it}$ et $\hat{Y}_{it}$, y compris une approximation de la variance.

4 Estimation sur petits domaines de séries chronologiques

Les estimations mensuelles initiales du domaine pour les vagues distinctes, accompagnées d'estimations de la variance et de la covariance, sont les données d'entrée pour les modèles de séries chronologiques. À l'étape suivante, on appliquera les MSCS pour lisser les estimations initiales et corriger le biais de renouvellement. On utilise les modèles estimés pour prévoir des fractions du chômage provincial, des tendances du chômage provincial et des variations de ces tendances d'un mois à l'autre. Dans la sous-section 4.1, les MSCS sont définis, puis exprimés en tant que modèles espace-état ajustés dans un cadre fréquentiste. La sous-section 4.2 explique comment ces MSCS peuvent être exprimés sous la forme de modèles multiniveaux de séries chronologiques ajustés dans un cadre bayésien hiérarchique.

4.1 Modèle espace-état

Dans la présente section, on élabore un modèle de séries chronologiques structurel pour les données mensuelles à l'échelle provinciale de 12 provinces simultanément afin de tirer parti des données d'échantillonnage temporelles et transversales. Soit $\hat{Y}_{it} = (\hat{Y}_{it1}, \dots, \hat{Y}_{it5})^t$ le vecteur à cinq dimensions contenant les estimations par la régression des données d'enquête $\hat{Y}_{itp}$ définies par (3.1) dans la période t et le domaine i . On peut modéliser ce vecteur à l'aide du modèle de séries chronologiques structurel suivant (Pfeffermann, 1991; van den Brakel et Krieg, 2009, 2015) :

$$\hat{Y}_{it} = \iota_5 \theta_{it} + \lambda_{it} + e_{it}, \quad (4.1)$$

où $\iota_5 = (1, 1, 1, 1, 1)^t$, θ_{it} est un scalaire désignant le paramètre de population vrai pour la période t dans un domaine i , λ_{it} est un vecteur à cinq dimensions qui modélise le biais de renouvellement et e_{it} un vecteur à cinq dimensions avec erreurs d'échantillonnage. Le paramètre de population θ_{it} dans (4.1) est modélisé comme suit :

$$\theta_{it} = L_{it} + S_{it} + \epsilon_{it}, \quad (4.2)$$

où L_{it} désigne un modèle de tendance stochastique pour saisir la variation de faible fréquence (tendance plus cycle économique), S_{it} une composante saisonnière stochastique pour modéliser les fluctuations mensuelles et ϵ_{it} un bruit blanc pour la variation inexpliquée dans θ_{it} . Pour la composante de tendance stochastique, on utilise le modèle dit de tendance lisse, défini par l'ensemble suivant d'équations :

$$L_{it} = L_{it-1} + R_{it-1}, \quad R_{it} = R_{it-1} + \eta_{R,it}, \quad \eta_{R,it} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{Ri}^2). \quad (4.3)$$

Pour la composante saisonnière stochastique, on utilise la formule trigonométrique, voir Boonstra et van den Brakel (2016) pour plus de détails. Le bruit blanc dans (4.2) est défini comme étant $\epsilon_{it} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{\epsilon_i}^2)$.

Le biais de renouvellement entre les séries d'estimations par la régression des données d'enquête est modélisé dans (4.1) avec $\lambda_{it} = (\lambda_{it1}, \lambda_{it2}, \lambda_{it3}, \lambda_{it4}, \lambda_{it5})^t$. On identifie le modèle en prenant $\lambda_{it1} = 0$. Cela implique que le biais relatif dans les vagues de suivi par rapport à la première vague est estimé et qu'on suppose que les estimations par la régression des données d'enquête de la première vague sont les approximations les plus fiables de θ_{it} , voir la justification dans van den Brakel et Krieg (2009). Les autres composantes modélisent la différence systématique entre la vague p par rapport à la première vague et elles sont modélisées comme des marches aléatoires pour permettre des profils dépendant du temps dans le biais de renouvellement.

$$\lambda_{itp} = \lambda_{it-1;p} + \eta_{\lambda,itp}, \quad \eta_{\lambda,itp} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{\lambda_i}^2), \quad p = 2, 3, 4, 5. \quad (4.4)$$

Enfin, on a élaboré un modèle de séries chronologiques pour les erreurs d'enquête. Soit $e_{it} = (e_{it1}, e_{it2}, e_{it3}, e_{it4}, e_{it5})^t$ le vecteur à cinq dimensions contenant les erreurs d'enquête des cinq vagues. On utilise les estimations de la variance des estimations par la régression des données d'enquête comme information a priori dans le modèle de séries chronologiques pour tenir compte de l'hétéroscédasticité causée par les tailles d'échantillon variables dans le temps au moyen du modèle d'erreur d'enquête suivant :

$$e_{itp} = \sqrt{v(\hat{Y}_{itp})} \tilde{e}_{itp}, \quad (4.5)$$

et $v(\hat{Y}_{itp})$ défini par (3.2). Comme la première vague est observée pour la première fois, il n'y a pas d'autocorrélation avec des échantillons observés auparavant. Pour modéliser l'autocorrélation entre les erreurs d'enquête des vagues de suivi, on dérive des modèles AR appropriés pour $\tilde{e}_{itp}$, en appliquant les équations de Yule-Walker aux coefficients de corrélation

$$\frac{n_{it_1 p_1 t_2 p_2}}{\sqrt{n_{it_1 p_1} n_{it_2 p_2}}} \hat{\rho}_{t_1 p_1 t_2 p_2}, \quad (4.6)$$

qui sont dérivés des microdonnées décrites à la section 3. À partir de cette analyse, on suppose un modèle AR(1) pour les vagues 2 à 5 où les coefficients d'autocorrélation dépendent de la vague et du mois. Ces considérations donnent le modèle suivant pour les erreurs d'enquête :

$$\begin{aligned}\tilde{e}_{it1} &= \nu_{it1}, \quad \nu_{it1} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{\nu_{i1}}^2), \\ \tilde{e}_{itp} &= \varrho_{it(p-1)p} \tilde{e}_{i(t-3)(p-1)} + \nu_{itp}, \quad \nu_{itp} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{\nu_{ip}}^2), \quad p = 2, \dots, 5,\end{aligned}\quad (4.7)$$

avec $\varrho_{it(p-1)p}$ les coefficients d'autocorrélation partielle dépendant du temps entre les vagues p et $p-1$ dérivés de (4.6). Par conséquent, $\text{Var}(e_{it1}) = v(\hat{Y}_{it1}) \sigma_{\nu_{i1}}^2$ et $\text{Var}(e_{itp}) = v(\hat{Y}_{itp}) \sigma_{\nu_{ip}}^2 / (1 - \varrho_{it(p-1)p}^2)$ pour $p = 2, \dots, 5$. Les variances $\sigma_{\nu_{ip}}^2$ sont des paramètres d'échelle avec des valeurs proches de 1 pour la première vague et proches de $\frac{1}{T} \sum_{t=1}^T (1 - \varrho_{it(p-1)p}^2)$ pour les autres vagues, où T indique la longueur de la série observée.

Le modèle (4.1) utilise les données d'échantillonnage observées pendant les périodes précédentes dans chaque domaine pour améliorer la précision de l'estimateur par la régression des données d'enquête et il tient compte du biais de renouvellement et de la corrélation sérielle induite par le plan d'échantillonnage à panel rotatif. Pour tirer parti des données d'échantillonnage entre domaines, on peut combiner le modèle (4.1) conçu pour les domaines séparés dans un modèle multivarié unique :

$$\begin{pmatrix} \hat{Y}_{1t} \\ \vdots \\ \hat{Y}_{m_A t} \end{pmatrix} = \begin{pmatrix} \iota_5 \theta_{1t} \\ \vdots \\ \iota_5 \theta_{m_A t} \end{pmatrix} + \begin{pmatrix} \lambda_{1t} \\ \vdots \\ \lambda_{m_A t} \end{pmatrix} + \begin{pmatrix} e_{1t} \\ \vdots \\ e_{m_A t} \end{pmatrix}, \quad (4.8)$$

où m_A désigne le nombre de domaines, qui est égal à 12 dans la présente application. Ce cadre multivarié permet d'utiliser l'information de l'échantillon entre les domaines en modélisant la corrélation entre les termes de perturbation des différentes composantes structurelles des séries chronologiques (tendance, saisonnière, biais de renouvellement) ou en définissant les hyperparamètres ou les variables d'état de ces composantes comme étant égaux dans tous les domaines. Dans l'article, les modèles présentant une corrélation transversale entre les termes de perturbation de la pente de la tendance (4.3) sont pris en compte, c'est-à-dire que

$$\text{Cov}(\eta_{R,it}, \eta_{R,i't'}) = \begin{cases} \sigma_{Ri}^2 & \text{si } i = i' \text{ et } t = t' \\ \varsigma_{Rii'} & \text{si } i \neq i' \text{ et } t = t' \\ 0 & \text{si } t \neq t' \end{cases} \quad (4.9)$$

La structure de covariance la plus parcimonieuse est une matrice diagonale où tous les domaines partagent une même composante de variance, soit $\sigma_{Ri}^2 = \sigma_R^2$ pour tous les i et $\varsigma_{Rii'} = 0$ tous les i et i' . Il s'agit de modèles de séries chronologiques structurels dits apparemment non liés, qui constituent une méthode synthétique d'utilisation des données d'échantillonnage entre domaines. Une structure de covariance légèrement plus complexe et réaliste serait une matrice diagonale où chaque domaine a une composante de

variance distincte, soit $\varsigma_{Rii'} = 0$ pour tous les i et i' . Dans ce cas, le modèle emprunte de l'information seulement dans le temps et ne tire pas parti de l'information transversale. La structure de covariance la plus complexe permet une matrice de covariance complète. Une forte corrélation entre les perturbations de pente dans tous les domaines peut produire des tendances cointégrées. Cela implique que les tendances communes $q < m_A$ sont nécessaires à la modélisation de dynamique des tendances pour les domaines m_A et permettent la spécification des modèles dits à tendance commune (Koopman, Harvey, Doornik et Shephard, 1999; Krieg et van den Brakel, 2012). Les premières analyses de MSCS ont montré que la composante saisonnière et la composante de biais de renouvellement sont indépendantes du temps. Par conséquent, il ne serait pas judicieux de modéliser des corrélations entre les termes de perturbation saisonnière et de biais de renouvellement. Étant donné que les hyperparamètres des paramètres du domaine de population de bruit blanc tendent à être nuls, il semble préférable de supprimer complètement cette composante du modèle, ce qui implique que les corrélations de modélisation entre le bruit de la population ne sont pas prises en compte. Les corrélations entre les erreurs d'enquête pour différents domaines ne sont pas non plus prises en compte, car les domaines sont des régions géographiques dont les échantillons sont tirés indépendamment.

Comme solution de substitution à un modèle avec matrice de covariance complète pour les perturbations de pente, on considère un modèle de tendance qui a un modèle de tendance lisse commun à toutes les provinces plus les composantes de tendance $m_A - 1$ qui décrivent l'écart de chaque domaine par rapport à cette tendance globale. Dans ce cas, (4.2) est donné par

$$\begin{aligned}\theta_{1t} &= L_t + S_{1t} + \epsilon_{1t}, \\ \theta_{it} &= L_t + L_{it}^* + S_{it} + \epsilon_{it}, \quad i = 2, \dots, m_A.\end{aligned}\quad (4.10)$$

Ici, L_t est la composante de tendance lisse globale, définie par (4.3), et L_{it}^* est l'écart par rapport à la tendance globale pour les domaines distincts, définis comme échelles locales

$$L_{it}^* = L_{it-1}^* + \eta_{L,it}, \quad \eta_{L,it} \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma_{Li}^2), \quad (4.11)$$

ou comme des tendances lisses comme dans (4.3). Ces modèles de tendance permettent implicitement des corrélations (positives) entre les tendances des différents domaines.

Les paramètres qu'il faut estimer dans la méthode de modélisation de séries chronologiques sont la tendance et le signal. Ce dernier se définit comme la tendance plus la composante saisonnière. La méthode des séries chronologiques convient particulièrement à l'estimation des variations d'un mois à l'autre. Les profils saisonniers compliquent l'interprétation des variations d'un mois à l'autre des estimations directes et des signaux lissés. C'est pourquoi les variations d'un mois à l'autre sont calculées seulement pour les tendances. En raison de la forte corrélation positive entre les niveaux des périodes consécutives, les erreurs-types des variations d'un mois à l'autre du niveau des tendances sont considérablement plus faibles que celles des variations d'un mois à l'autre des estimations directes, par exemple. La variation d'un mois à l'autre de la tendance est définie comme étant $\Delta_{it}(1) = L_{it} - L_{it-1}$ pour les modèles avec tendances séparées pour les domaines ou $\Delta_{it}(1) = L_t - L_{t-1} + L_{it}^* - L_{it-1}^*$ pour les modèles avec une tendance globale et des

tendances $m_A - 1$ pour l'écart des domaines distincts par rapport à la tendance globale. Cette méthode de modélisation est également utile pour l'estimation des évolutions d'une année à l'autre pour les tendances définies comme étant $\Delta_{it}(12) = L_{it} - L_{it-12}$ ou $\Delta_{it}(12) = L_t - L_{t-12} + L_{it}^* - L_{it-12}^*$. Il est également judicieux de calculer les variations d'une année à l'autre pour les signaux, puisque la plus grande part de la composante saisonnière s'annule. Ces évolutions sont définies d'une manière équivalente aux évolutions de la tendance d'une année à l'autre.

On analyse les modèles de séries chronologiques structurels ci-dessus en les mettant sous une forme dite d'espace-état. On applique ensuite le filtre de Kalman pour ajuster les modèles, où les hyperparamètres inconnus sont remplacés par leurs estimations par le MV. L'analyse est réalisée au moyen d'un logiciel développé dans OxMetrics en combinaison avec les sous-routines de SsfPack 3.0 (Doornik, 2009; Koopman, Shephard et Doornik, 1999, 2008). Les estimations par le MV des hyperparamètres sont obtenues à l'aide de la procédure d'optimisation numérique maxBFGS dans OxMetrics. Voir Boonstra et van den Brakel (2016) pour plus de détails sur la représentation espace-état, l'initialisation du filtre de Kalman et le logiciel utilisé aux fins d'ajustement de ces modèles.

4.2 Modèle multiniveau de séries chronologiques

Pour la description de la représentation des séries chronologiques multiniveau des MSCS, les estimations initiales $\hat{Y}_{itp}$ sont combinées dans un vecteur $\hat{Y} = (\hat{Y}_{111}, \hat{Y}_{112}, \dots, \hat{Y}_{115}, \hat{Y}_{121}, \dots)'$, c'est-à-dire que l'indice de vague s'exécute plus rapidement que l'indice temporel qui est plus rapide que l'indice de domaine. Le nombre de domaines, de périodes et de vagues est représenté par m_A , m_T et m_p , respectivement. La longueur totale de $\hat{Y}$ est par conséquent $m = m_A m_T m_p = 12(\text{domaines}) * 72(\text{mois}) * 5(\text{vagues}) = 4\,320$. De même, les estimations de la variance $v(\hat{Y}_{itp})$ sont mises dans le même ordre sur la diagonale d'une matrice de covariance $m \times m$ Φ .

La matrice de covariance Φ n'est pas diagonale en raison des corrélations induites par le plan d'échantillonnage par panel. Il s'agit d'une matrice bande parcimonieuse, et l'ordre du vecteur $\hat{Y}$ est tel qu'il produit une largeur de bande minimale, ce qui est avantageux pour les calculs.

Les modèles multiniveaux envisagés pour la modélisation du vecteur des estimations directes $\hat{Y}$ prennent une forme additive linéaire générale

$$\hat{Y} = X\beta + \sum_{\alpha} Z^{(\alpha)}v^{(\alpha)} + e, \quad (4.12)$$

où X est une matrice du plan d'expérience $m \times p$ pour les effets fixes β , et les $Z^{(\alpha)}$ sont les matrices du plan d'expérience $m \times q^{(\alpha)}$ pour les vecteurs d'effets aléatoires $v^{(\alpha)}$. Ici, la somme sur α s'exécute sur plusieurs termes possibles d'effet aléatoire à différents niveaux, comme une tendance lisse à l'échelle nationale, les tendances d'échelle locale provinciales, le bruit blanc, etc. Cela est expliqué plus en détail ci-dessous. Les erreurs d'échantillonnage $e = (e_{111}, e_{112}, \dots, e_{115}, e_{121}, \dots)'$ sont considérées comme distribuées normalement comme suit :

$$e \sim \mathcal{N}(0, \Sigma) \quad (4.13)$$

où $\Sigma = \bigoplus_{i=1}^m \lambda_i \Phi_i$ avec Φ_i comme matrice de covariance pour les estimations initiales de la province i , et λ_i un paramètre d'échelle de la variance propre à la province qu'il faut estimer. Comme on l'a décrit à la section 3, les variances par rapport au plan dans $\Phi = \bigoplus_i \Phi_i$ sont regroupées sur toutes les provinces et, en raison de la nature discrète des données sur le chômage, elles perdent ainsi une partie de leur dépendance à l'égard du niveau de chômage. On a constaté que l'intégration des facteurs d'échelle de la variance λ_i permet au modèle de rééchelonner les variances par rapport au plan estimées à un niveau améliorant l'ajustement des données.

Pour décrire le modèle général de chaque vecteur $v^{(\alpha)}$ d'effets aléatoires, nous supprimons l'exposant α . Chaque vecteur v a $q = dl$ composantes correspondant à d effets autorisés à varier sur l niveaux d'une variable de facteur. En particulier,

$$v \sim \mathcal{N}(0, A \otimes V), \quad (4.14)$$

où V et A sont des matrices de covariance $d \times d$ et $l \times l$, respectivement. Comme à la section 4.1, la matrice de covariance V peut être paramétrée de trois manières différentes. La plupart du temps, il s'agit d'une matrice de covariance non structurée, c'est-à-dire entièrement paramétrée. Les formes plus parcimonieuses sont $V = \text{diag}(\sigma_{v;1}^2, \dots, \sigma_{v;d}^2)$ ou $V = \sigma_v^2 I_d$. Si $d = 1$, les trois paramétrisations sont équivalentes. La matrice de covariance A décrit la structure de covariance entre les niveaux de la variable de facteur et elle est supposée connue. Il est généralement plus pratique d'utiliser la matrice de précision $Q_A = A^{-1}$, car elle est parcimonieuse pour de nombreuses structures courantes de corrélation temporelle et spatiale (Rue et Held, 2005).

4.2.1 Relations entre les représentations espace-état et les représentations multiniveaux des séries chronologiques

On peut représenter une tendance lisse unique comme une ordonnée à l'origine aléatoire ($d = 1$) variant dans le temps ($l = m_T$), avec une corrélation temporelle déterminée par une matrice de précision parcimonieuse Q_A de bande $m_T \times m_T$ associée à une marche aléatoire du second ordre (Rue et Held, 2005). Dans ce cas $V = \sigma_v^2$ et la matrice du plan d'expérience Z est la matrice des indicateurs $m \times m_T$ pour le mois, c'est-à-dire la matrice avec un seul 1 dans chaque rang pour le mois correspondant et des 0 ailleurs. La parcimonie de Q_A et Z peut être exploitée dans les calculs. La matrice de précision pour la composante de tendance lisse a deux vecteurs singuliers, $\iota_{m_T} = (1, 1, \dots, 1)$ et $(1, 2, \dots, m_T)'$. Cela signifie que la spécification correspondante (4.14) est complètement non informative sur le niveau global et la tendance linéaire. Afin d'éviter la non-identifiabilité entre les différents termes du modèle, on peut supprimer le niveau global et la tendance de v en imposant les contraintes $Rv = 0$, où R est la matrice $2 \times m_T$ avec les deux vecteurs singuliers comme rangs. Le niveau global et la tendance sont ensuite inclus dans le vecteur β des effets fixes. Dans la représentation espace-état, on obtient ce modèle en définissant un modèle de tendance (4.3) pour tous les domaines, c'est-à-dire $L_{it} = L_t$ et $R_{it} = R_t$ pour tous les i . Définir les variables

d'état pour la tendance comme étant égales dans les domaines est une méthode très synthétique d'utilisation des données d'échantillonnage provenant d'autres domaines, qui se fonde sur des hypothèses qui, dans la plupart des cas, ne se vérifient pas.

On obtient une tendance lisse pour chaque province avec $d = m_A$, $l = m_T$, et V une matrice de covariance $m_A \times m_A$, soit diagonale avec un seul paramètre de variance, diagonale avec des paramètres de variance m_A , ou non structurée, c'est-à-dire entièrement paramétrée pour ce qui est des paramètres de variance m_A et des paramètres de corrélation $m_A(m_A - 1)/2$. Dans ce cas, la matrice du plan d'expérience est $I_{m_A} \otimes I_{m_T} \otimes I_{m_P}$. Dans la représentation espace-état, on obtient ces modèles avec le modèle de tendance (4.3) et la structure de covariance (4.9).

Un autre modèle de tendance consisterait en une seule tendance globale lisse (marche aléatoire de second ordre) complétée par une tendance d'échelle locale, c'est-à-dire une marche aléatoire ordinaire (de premier ordre) pour chaque province. Ce dernier peut être modélisé de la manière décrite dans le paragraphe précédent, mais avec une matrice de précision associée à une marche aléatoire de premier ordre. Ce modèle de tendance correspond aux modèles (4.10) et (4.11) dans le contexte espace-état. À la différence de la méthode espace-état, il n'est pas nécessaire de supprimer du modèle une des tendances de la marche aléatoire provinciale aux fins d'identifiabilité. Cela s'explique par le fait que les contraintes de la méthode multiniveau sont imposées de façon à ce que la tendance lisse globale ainsi que toutes les tendances de marche aléatoire provinciale soient à somme nulle dans le temps. Les composantes contraintes correspondent aux ordonnées à l'origine globales et provinciales, qui sont incluses séparément dans le modèle en tant qu'effets fixes, avec l'exclusion d'un effet fixe provincial.

On peut également exprimer les effets saisonniers en termes d'effets aléatoires corrélés (4.14). La composante saisonnière trigonométrique est équivalente au modèle saisonnier de la variable nominale équilibrée (Proietti, 2000; Harvey, 2006), qui correspond à des marches aléatoires de premier ordre dans le temps pour chaque mois, sujettes à une contrainte à somme nulle pour tous les mois. Dans ce cas $d = 12$ (saisons), $V = \sigma_v^2 I_{12}$ et $l = m_T$, Q_A étant la matrice de précision d'une marche aléatoire de premier ordre. Les contraintes à somme nulle sur les saisons pour chaque temps, ainsi que les contraintes à somme nulle dans le temps pour chaque marche aléatoire peuvent être imposées comme $Rv = 0$ avec R étant la matrice $(m_T + 12) \times 12m_T$.

$$R = \begin{pmatrix} \iota'_{12} \otimes I_{m_T} \\ I_{12} \otimes \iota'_{m_T} \end{pmatrix}. \quad (4.15)$$

Accompagné d'effets fixes pour chaque saison (encore une fois avec une contrainte à somme nulle imposée), ce terme d'effet aléatoire est équivalent à la saisonnalité trigonométrique. Il peut être étendu à une composante saisonnière pour chaque province, avec un paramètre de variance distinct pour chaque province.

Pour tenir compte du biais de renouvellement, le modèle multiniveau comprend des effets fixes pour les vagues 2 à 5. Facultativement, on peut modéliser ces effets dynamiquement en ajoutant des marches

aléatoires dans le temps pour chaque vague. Il faut aussi faire un autre choix, à savoir si les effets fixes et aléatoires sont croisés avec la province.

On peut inclure d'autres effets fixes dans le modèle, par exemple ceux associés aux variables auxiliaires utilisées dans les estimations par la régression des données d'enquête. On pourrait aussi modéliser certaines interactions à effet fixe, par exemple saison $\times$ province ou vague $\times$ province, en tant qu'effets aléatoires pour réduire le risque de surajustement.

Enfin, on peut ajouter un terme de bruit blanc au modèle pour tenir compte des variations inexpliquées par domaine et par temps dans le signal.

Le modèle (4.12) peut être considéré comme une généralisation du modèle au niveau du domaine de Fay-Herriot. Le modèle de Fay-Herriot ne comprend qu'un seul vecteur d'effets aléatoires non corrélés sur les niveaux d'une variable à un seul facteur (généralement des régions). Les modèles utilisés dans le présent article comprennent diverses combinaisons d'effets aléatoires non corrélés et corrélés sur des régions et des mois. Des modèles de séries chronologiques multiniveaux étendant le modèle de Fay-Herriot ont été présentés dans Rao et Yu (1994), Datta et coll. (1999), et You (2008). Dans Datta et coll. (1999) et You (2008) des modèles de séries chronologiques sont utilisés avec des effets de domaine indépendants et des marches aléatoires de premier ordre dans le temps pour chaque domaine. Rao et Yu (1994) utilisent un modèle avec effets aléatoires indépendants et un processus autorégressif AR(1) stationnaire au lieu d'un modèle de marche aléatoire dans le temps. You et coll. (2003) ont constaté que le modèle de marche aléatoire correspondait légèrement mieux aux données canadiennes sur le chômage que les modèles AR(1) avec un paramètre d'autocorrélation fixé à 0,5 ou 0,75. Nous ne nous intéresserons pas aux modèles AR(1) dans cet article et nous renvoyons à Diallo (2014) pour une approche permettant à la fois des tendances stationnaires et non stationnaires. Comparé aux références ci-dessus, notre modèle présente une nouvelle caractéristique, à savoir que l'on tient compte des tendances lisses plutôt que des marches aléatoires de premier ordre ou des composantes autorégressives, ou en plus de celles-ci. Nous incluons également les effets aléatoires indépendants de zone dans le temps comme terme de bruit blanc tenant compte de la variation inexpliquée au niveau d'agrégation d'intérêt.

4.2.2 Estimation de modèles multiniveaux de séries chronologiques

On utilise une approche bayésienne pour adapter le modèle (4.12)-(4.14). Cela signifie que nous avons besoin de distributions a priori pour tous les (hyper)paramètres du modèle. Les lois a priori suivantes sont utilisées :

- On attribue aux paramètres de variance au niveau des données λ_i pour $i = 1, \dots, m_A$ des lois du chi carré inverses avec des degrés de liberté et des paramètres d'échelle égaux à 1.
- Les effets fixes se voient attribuer une loi a priori normale avec une matrice de variance diagonale fixe et une moyenne nulle, présentant de très grandes valeurs (1e10).

- Pour une matrice de covariance entièrement paramétrée V dans (4.14), nous utilisons la loi de de Wishart inverse a priori mise à l'échelle proposée dans O'Malley et Zaslavsky (2008) et recommandée par Gelman et Hill (2007). Conditionnellement à un paramètre vectoriel dimensionnel d ξ ,

$$V | \xi \sim \text{Inv - Wishart}(V | \nu, \text{diag}(\xi) \Psi \text{diag}(\xi)) \quad (4.16)$$

où $\nu = d + 1$ est choisi et $\Psi = I_d$. On attribue au vecteur ξ une distribution normale $\mathcal{N}(0, I_d)$.

- On attribue à tous les autres paramètres de variance apparaissant dans une matrice diagonale V dans (4.14), conditionnellement à un paramètre auxiliaire ξ , des lois a priori du chi carré inverses avec un degré de liberté de 1 et un paramètre d'échelle ξ^2 . On attribue à chaque paramètre ξ une loi a priori $\mathcal{N}(0, 1)$. Marginalement, les paramètres de l'écart-type ont des demi-distributions de Cauchy a priori. Gelman (2006) montre que ces lois a priori sont de meilleures lois a priori par défaut que les lois a priori du chi carré inverses, plus courantes.

Le modèle est ajusté au moyen d'un échantillonnage par méthode de Monte-Carlo par chaînes de Markov (MCMC), en particulier l'échantillonneur de Gibbs (Geman et Geman, 1984; Gelfand et Smith, 1990). Les modèles multiniveaux considérés appartiennent à la classe des modèles gaussiens latents additifs dont les termes d'effet aléatoire sont des champs aléatoires gaussiens de Markov (GMRF), et nous utilisons la matrice parcimonieuse et les techniques d'échantillonnage par blocs décrites dans Rue et Held (2005) pour ajuster efficacement ces modèles aux données. De plus, la paramétrisation selon les paramètres auxiliaires susmentionnés ξ (Gelman, Van Dyk, Huang et Boscardin, 2008) améliore considérablement la convergence de l'échantillonneur de Gibbs utilisé. Voir Boonstra et van den Brakel (2016) pour plus de détails sur l'échantillonneur de Gibbs utilisé, y compris les spécifications des distributions conditionnelles complètes. Les méthodes sont mises en œuvre en R au moyen du paquet en R *mcmcsm* (Boonstra, 2016).

Pour chaque modèle considéré, l'échantillonneur de Gibbs est exécuté dans trois chaînes indépendantes avec des valeurs de départ aléatoires. Chaque chaîne s'exécute pendant 2 500 itérations. Les 500 premiers tirages sont supprimés en tant que « qu'échantillon de rodage ». Sur les 2 000 tirages restants de chaque chaîne, nous conservons un tirage sur cinq pour économiser de la mémoire tout en réduisant l'effet d'autocorrélation entre les tirages successifs. Cela nous laisse $3 * 400 = 1\,200$ tirages pour calculer les estimations et les erreurs-types. On a constaté que le nombre réel de tirages indépendants était de près de 1 200 pour la plupart des paramètres du modèle, ce qui signifie que la plupart des autocorrélations sont éliminées par la réduction. On évalue la convergence de la simulation par MCMC au moyen de tracés et de graphiques d'autocorrélation ainsi que du facteur de réduction d'échelle potentiel de Gelman-Rubin (Gelman et Rubin, 1992), qui diagnostique le mélange de chaînes. Les diagnostics suggèrent que toutes les chaînes convergent correctement à l'étape de rodage, et que les chaînes se mélangent bien, puisque tous les facteurs de Gelman-Rubin sont proches de 1. De plus, les erreurs de simulation de Monte-Carlo estimées (qui tiennent compte

de toute autocorrélation restant dans les chaînes) sont faibles comparativement aux erreurs-types a posteriori pour tous les paramètres, si bien que le nombre de tirages conservés suffit à nos besoins.

On peut exprimer les paramètres à estimer d'intérêt sous forme de fonctions des paramètres, et l'application de ces fonctions aux données de sortie de la méthode MCMC des paramètres donne des résultats tirés des lois a posteriori pour ces paramètres à estimer. Dans l'article, nous synthétisons ces tirages en fonction de leur écart moyen et de leur écart-type, qui servent respectivement d'estimations et d'erreurs-types. Toutes les estimations prises en compte peuvent être exprimées sous forme de prédicteurs linéaires, c'est-à-dire sous forme de combinaisons linéaires des paramètres du modèle. On calcule les estimations et les erreurs-types pour les paramètres à estimer suivants :

- **Signal** : le vecteur θ_{it} , y compris tous les effets fixes et aléatoires, sauf ceux associés aux vagues 2 à 5. Ces valeurs correspondent aux valeurs ajustées $X\beta + \sum_{\alpha} Z^{(\alpha)}v^{(\alpha)}$ associées à une rangée sur cinq 1, 6, 11, ... de $\hat{Y}$ et des matrices du plan d'expérience.
- **Tendance** : prédiction de la tendance à long terme. On réalise ce calcul en incorporant seulement les composantes des tendances de chaque modèle dans le prédicteur linéaire. Pour la plupart des modèles considérés, la tendance correspond aux chiffres désaisonnalisés, c'est-à-dire à des prédictions du signal dont tous les effets saisonniers sont éliminés.
- **Croissance de la tendance** : les différences entre les tendances au cours de deux mois consécutifs.

5 Résultats

Les résultats obtenus avec les représentations de séries chronologiques espace-état et multiniveaux des MSCS sont décrits aux sous-sections 5.1 et 5.2, respectivement. Premièrement, on définit deux mesures des divergences pour évaluer et comparer les différents modèles. La première mesure est le biais relatif moyen (BRM), qui résume les différences entre les estimations du modèle et les estimations directes moyennes sur la période, en pourcentage de ces dernières. Pour un modèle donné M , le BRM_i est défini comme étant

$$BRM_i = \frac{\sum_t (\hat{\theta}_{it}^M - \tilde{Y}_{it.})}{\sum_t \tilde{Y}_{it.}} \times 100 \%, \quad (5.1)$$

où $\tilde{Y}_{it.}$ sont les estimations directes par province et par mois incorporant l'ajustement du biais de renouvellement par le quotient mentionné à la fin de la section 3. Cette mesure de référence montre pour chaque province dans quelle mesure les estimations fondées sur un modèle s'écartent des estimations directes. Les divergences ne devraient pas être trop grandes, car on peut s'attendre à ce que les estimations directes moyennes sur la période soient proches du véritable niveau moyen du chômage. La deuxième mesure de la divergence est la réduction relative des erreurs-types (RRET); elle mesure les pourcentages de réduction des erreurs-types estimées entre les estimations fondées sur un modèle et les estimations directes, à savoir

$$\text{RRET}_i = 100 \% \times \frac{1}{m_T} \sum_t \left(\text{et}(\tilde{Y}_{it.}) - \text{et}(\hat{\theta}_{it}^M) \right) / \text{et}(\tilde{Y}_{it.}), \quad (5.2)$$

pour un modèle donné M . Ici, les erreurs-types estimées pour les estimations directes découlent d'une approximation de la variance pour $\tilde{Y}_{it.}$, alors que les erreurs-types fondées sur un modèle sont des écarts-types a posteriori ou découlent du filtre/lisseur de Kalman. Les écarts-types a posteriori, les erreurs-types obtenues au moyen du filtre de Kalman et les erreurs-types des estimateurs directs proviennent de différents cadres et sont formellement exprimées comme non comparables. On les utilise dans (5.2) pour quantifier la réduction par rapport à l'estimateur direct seulement, mais non comme critères de sélection du modèle.

5.1 Résultats des modèles espace-état

Dix modèles espace-état différents sont comparés. Quatre modèles de tendance différents sont distingués. La première composante de tendance est un modèle de tendance lisse sans corrélation entre les domaines (4.3), abrégée T1. Le deuxième modèle de tendance, T2, est un modèle de tendance lisse (4.3) avec matrice de corrélation complète pour les perturbations de la pente (4.9). La troisième composante de tendance, le T3, est un modèle de tendance lisse commun à toutes les provinces, avec 11 modèles de tendance d'échelle locale pour l'écart des domaines par rapport à cette tendance globale ((4.10) en combinaison avec (4.11)). Le quatrième modèle de tendance, le T4, est un modèle de tendance lisse commun à toutes les provinces, avec 11 modèles de tendance lisse pour l'écart des domaines par rapport à cette tendance globale ((4.10) en combinaison avec (4.3)). Dans T3 et T4, la province de Groningen est considérée comme égale à la tendance globale. La composante du biais de renouvellement (4.4) peut être propre au domaine (ce qui est indiqué par la lettre « R » dans le nom du modèle) ou choisie comme étant égale pour tous les domaines (pas de « R » dans le nom du modèle). Une autre façon de simplifier consiste à supposer que le biais de renouvellement pour les vagues 2, 3, 4 et 5 est égal, mais propre au domaine (ce qui est indiqué par « R2 »). De même, la composante saisonnière peut être choisie comme étant propre au domaine (ce qui est indiqué par « S ») ou égale pour tous les domaines. Tous les modèles partagent la même composante pour l'erreur d'enquête, c'est-à-dire un modèle AR(1) avec des coefficients d'autocorrélation variables dans le temps pour les vagues 2 à 5 afin que soit modélisée l'autocorrélation dans les erreurs d'enquête. Les modèles espace-état suivants sont comparés :

- T1SR : Modèle de tendance lisse sans corrélation entre les perturbations de la pente; saisonnière et propre au domaine du biais de renouvellement.
- T2SR : Modèle de tendance lisse avec matrice de corrélation complète pour les perturbations de la pente; saisonnière et propre au domaine du biais de renouvellement.
- T2S : Modèle de tendance lisse avec matrice de corrélation complète pour les perturbations de la pente; saisonnière et propre au domaine, biais de renouvellement égal pour tous les domaines.
- T2R : Modèle de tendance lisse avec matrice de corrélation complète pour les perturbations de la pente; saisonnière égale pour tous les domaines, propre aux domaines du biais de renouvellement.

- T3SR : Un modèle de tendance lisse commun pour tous les domaines plus 11 modèles de tendance d'échelle locale pour les écarts par rapport à la tendance globale; saisonnière et propre au domaine du biais de renouvellement.
- T3R : Un modèle de tendance lisse commun pour tous les domaines plus des modèles de tendance d'échelle locale pour les écarts par rapport à la tendance globale; saisonnière égale pour tous les domaines, propre aux domaines du biais de renouvellement.
- T3R2 : Un modèle de tendance lisse commun pour tous les domaines plus 11 modèles de tendance d'échelle locale pour les écarts par rapport à la tendance globale; saisonnière égale pour tous les domaines, le biais de renouvellement est propre à un domaine, mais il est supposé égal pour les quatre vagues de suivi.
- T3 : Un modèle de tendance lisse commun pour tous les domaines plus des modèles de tendance d'échelle locale pour les écarts par rapport à la tendance globale; saisonnière et biais de renouvellement égal pour tous les domaines.
- T4SR : Un modèle de tendance lisse commun pour tous les domaines plus 11 modèles de tendance lisse pour les écarts par rapport à la tendance globale; saisonnière et propre au domaine du biais de renouvellement.
- T4R : Un modèle de tendance lisse commun pour tous les domaines et 11 modèles de tendance lisse pour les écarts par rapport à la tendance globale; saisonnière égale pour tous les domaines, propre au domaine du biais de renouvellement.

Pour tous les modèles, les estimations par le MV pour les hyperparamètres du biais de renouvellement et les effets saisonniers tendent à zéro, ce qui implique que ces composantes sont invariantes dans le temps. De plus, les estimations par le MV des composantes de la variance du bruit blanc des paramètres du domaine de la population tendent à zéro. C'est pourquoi cette composante est supprimée du modèle (4.2). Les estimations par le MV des composantes de la variance des erreurs d'enquête de la première vague varient entre 0,93 et 1,90. Pour les vagues de suivi, les estimations par le MV varient entre 0,86 et 1,80. Les variances des estimations directes sont regroupées dans les domaines (3.2), ce qui pourrait introduire un biais, par exemple la sous-estimation de la variance dans les domaines où les taux de chômage sont élevés. Il est nécessaire d'échelonner les variances des erreurs d'enquête avec les estimations par le MV $\sigma_{v_{ip}}^2$ pour corriger ce biais. Les estimations par le MV des hyperparamètres pour les composantes de tendance se trouvent dans Boonstra et van den Brakel (2016).

Les modèles sont comparés par des fonctions de log-vraisemblance. Pour tenir compte des différences de complexité des modèles, on utilise des critères d'information d'Akaike (AIC) et des critères d'information bayésiens (BIC), voir Durbin et Koopman (2012), section 7.4. Les résultats sont synthétisés dans le tableau 5.1. Des modèles parcimonieux dans lesquels les effets saisonniers ou le biais de renouvellement sont égaux dans tous les domaines sont privilégiés par les critères AIC et BIC. Notons toutefois que les log-vraisemblances ne sont pas entièrement comparables entre modèles. Pour obtenir des

vraisemblances comparables, on ignore les 24 premiers mois de la série dans le calcul de la vraisemblance pour tous les modèles. Néanmoins, certaines vraisemblances sont étranges. Par exemple, la vraisemblance de T2SR est plus petite que la vraisemblance de T2S, bien que T2SR comprenne plus de paramètres de modèle. Cela résulte probablement de la combinaison de modèles de séries chronologiques grands et complexes avec des séries chronologiques relativement courtes, ce qui donne lieu à des fonctions de vraisemblance plates. De ce point de vue également, les modèles parcimonieux évitant les surajustements sont toujours favorables, ce qui est conforme aux résultats des valeurs AIC et BIC du tableau 5.1.

Tableau 5.1
Critères AIC et BIC pour les modèles espace-état

Modèle	Log- vraisemblance	États	Hyperparamètres	AIC	BIC
T1SR	9 813,82	204	24	-399,41	-390,52
T2SR	9 862,86	204	35	-400,99	-391,68
T2S	9 879,03	160	35	-403,50	-395,90
T2R	9 859,97	83	35	-405,92	-401,32
T3SR	9 855,35	193	24	-401,60	-393,14
T3R	9 851,62	72	24	-406,48	-402,74
T3R2	9 871,65	36	24	-408,82	-406,48
T3	9 881,16	28	24	-409,55	-407,52
T4SR	9 857,47	204	24	-401,23	-392,34
T4R	9 853,65	83	24	-406,11	-401,94

La modélisation des corrélations entre les perturbations de pente de la tendance améliore considérablement le modèle. Le modèle T1SR, par exemple, est emboîté dans le modèle T2SR et un test du rapport de vraisemblance favorise clairement ce dernier. Pour le modèle T2SR, il s'ensuit qu'on peut modéliser la dynamique des tendances de ces 12 domaines au moyen de seulement 2 tendances communes sous-jacentes, puisque le rang de la matrice de covariance 12×12 est égal à 2. Par conséquent, la matrice de covariance complète pour les perturbations de la pente des 12 domaines est en fait modélisée avec 23 hyperparamètres au lieu de 78. Cela montre que les corrélations entre les perturbations de pente sont très fortes. En effet, les corrélations varient entre 1,00 et 0,98. Voir Boonstra et van den Brakel (2016) pour les estimations par le MV de la matrice de covariance complète.

Le tableau 5.2 présente le BRM, défini par (5.1). Les modèles qui supposent que le biais de renouvellement est égal dans tous les domaines, c'est-à-dire T2S et T3, ont des biais relatifs importants pour certains des domaines. On constate des biais élevés dans les domaines où le chômage est élevé (par exemple, Groningue) ou faible (par exemple, Utrecht) comparativement à la moyenne nationale. Un compromis possible entre la parcimonie et le biais consiste à supposer que le biais de renouvellement est égal pour les quatre vagues de suivi, mais toujours propre au domaine (T3R2). Pour ce modèle, le biais est faible, sauf dans la province de Gelderland.

Tableau 5.2
Biais relatif moyen (5.1) dans le temps (%), par province pour les modèles espace-état

	Grn	Frs	Drn	Ovr	Flv	Gld	Utr	N-H	Z-H	Zln	N-B	Lmb
T1SR	1,1	0,5	2,0	-0,2	0,1	3,4	0,1	0,6	1,7	-2,1	0,5	2,1
T2SR	1,2	0,7	2,2	-0,1	0,2	3,5	0,2	0,6	1,7	-2,1	0,5	2,1
T2S	-3,1	3,1	0,7	0,9	-4,4	2,8	2,4	0,8	0,5	1,7	1,8	1,5
T2R	0,9	0,8	1,8	-0,2	-0,4	3,4	0,1	0,6	1,7	-1,6	0,6	2,2
T3SR	0,8	0,6	2,0	-0,2	-0,3	3,5	0,3	0,5	1,7	-2,0	0,6	2,0
T3R2	-0,1	1,3	2,1	-0,6	-0,8	3,6	0,9	0,6	1,5	-1,1	1,0	1,2
T3R	0,5	0,7	1,8	-0,2	-0,8	3,5	0,3	0,5	1,6	-1,5	0,7	2,1
T3	-4,0	2,5	0,1	0,9	-5,0	2,8	2,3	0,7	0,6	2,5	2,0	1,3
T4SR	0,8	0,7	2,1	-0,2	-0,0	3,5	0,2	0,6	1,7	-1,9	0,5	2,1
T4R	0,6	0,7	1,8	-0,2	-0,6	3,4	0,1	0,6	1,7	-1,3	0,7	2,1

La figure 5.1 compare les tendances lisses et les erreurs-types des modèles T1SR, T2SR et T2S. La figure 5.2 compare l'évolution d'un mois à l'autre de la tendance et les erreurs-types pour ces trois modèles. Les tendances lisses obtenues par le modèle de tendance commun sont légèrement plus souples comparativement à un modèle sans corrélation entre les perturbations de la pente. On le constate manifestement dans la variation d'un mois à l'autre des tendances. La modélisation de la corrélation entre les perturbations de pente réduit clairement l'erreur-type de la tendance et la variation de la tendance d'un mois à l'autre. Le fait de supposer que le biais de renouvellement est égal pour tous les domaines (modèle T2S) a une incidence sur le niveau de la tendance et réduit davantage l'erreur-type, principalement parce que le nombre de variables d'état est réduit. La différence entre la tendance dans T2SR et T2S est une variation de niveau. Cela découle des variations de la tendance d'un mois à l'autre dans les modèles T2SR et T2S, qui sont exactement égales. Selon les critères AIC et BIC, la réduction du nombre de variables d'état quand on suppose un biais de renouvellement égal pour tous les domaines constitue une amélioration du modèle. Cependant, dans cette application, l'intérêt est axé sur le modèle ajusté aux domaines distincts. Le fait de supposer que le biais de renouvellement est égal dans tous les domaines est efficace en moyenne pour les mesures globales d'adéquation de l'ajustement, comme AIC et BIC, mais pas nécessairement pour tous les domaines distincts. Le biais introduit dans les tendances de certains domaines quand on suppose que le biais de renouvellement est égal sur les domaines n'est pas souhaitable.

La figure 5.3 compare les tendances lisses et les erreurs-types des modèles T2SR, T3SR et T4SR. L'évolution de la tendance d'un mois à l'autre et ses erreurs-types se trouvent dans Boonstra et van den Brakel (2016). Les tendances obtenues avec une tendance lisse globale plus 11 tendances pour les écarts de domaine par rapport à la tendance globale ressemblent aux tendances obtenues avec le modèle de tendance commun. Dans cette application, la dynamique fondée sur les deux tendances communes du modèle T2SR est raisonnablement bien établie approximativement par les autres tendances des modèles T3SR et T4SR. Il s'agit d'une constatation empirique qui n'est pas nécessairement généralisable à d'autres situations, surtout lorsque plus de facteurs communs sont requis. Toutefois, le modèle de tendance commun comporte les plus petites erreurs-types pour la tendance. De plus, les tendances du modèle avec une échelle

locale pour les écarts de domaine par rapport à la tendance globale sont dans certains domaines plus volatiles que les deux autres modèles. Cela est particulièrement évident dans les variations de la tendance d'un mois à l'autre. En général, les modèles de tendance avec niveaux aléatoires se caractérisent par des tendances plus volatiles, voir Durbin et Koopman (2012), chapitre 3. Le modèle de tendance plus souple de T3 produit également une erreur-type plus élevée des variations d'un mois à l'autre.

Le fait de supposer que les effets saisonniers sont égaux pour tous les domaines est une autre façon de réduire le nombre de variables d'état et d'éviter de surajuster les données. Cette hypothèse n'influe pas sur le niveau de la tendance puisque le BRM est petit (voir le tableau 5.2) et produit une amélioration importante du modèle selon les critères AIC et BIC. Particulièrement quand l'intérêt porte sur les estimations de tendance, un biais dans les profils saisonniers est acceptable, et un modèle ayant une tendance fondée sur T2 ou T4, avec une composante saisonnière supposée égale dans tous les domaines pourrait être un bon compromis entre un modèle qui prend suffisamment en compte les différences entre domaines et la parcimonie du modèle pour éviter le surajustement des données.

Le modèle T3 est le modèle le plus parcimonieux qui soit, selon les critères AIC et BIC. En particulier, l'hypothèse d'un biais de renouvellement égal donne des estimations de tendance biaisées dans certains domaines (voir le tableau 5.2). Voir Boonstra et van den Brakel (2016) pour une comparaison de la tendance et de l'évolution d'un mois à l'autre de la tendance des modèles T2R, T3 et T4R. En supposant que les effets saisonniers sont égaux dans tous les domaines, on obtient un profil saisonnier moins prononcé. Voir Boonstra et van den Brakel (2016) pour une comparaison des signaux des modèles T2SR et T2R.

Boonstra et van den Brakel (2016) donne les résultats de la variation d'une année à l'autre des tendances dans les modèles T2R et T3R2. Les estimations de séries chronologiques pour la variation d'une année à l'autre sont très stables et précises et elles améliorent grandement les estimations directes de la variation d'une année à l'autre.

Le tableau 5.3 montre la RRET, définie par (5.2), pour les 10 modèles espace-état. Gardons à l'esprit que la RRET quantifie la réduction par rapport à l'estimateur direct et qu'elle n'est pas un critère de sélection du modèle. Le tableau 5.4 présente les moyennes des erreurs-types du signal, de la tendance et de la croissance (différences de tendance d'un mois à l'autre). La moyenne est calculée pour tous les mois et toutes les provinces. La modélisation de la corrélation entre les tendances explicitement (T2) ou implicitement (T3 ou T4) réduit significativement les erreurs-types pour la tendance et le signal. L'approche de la modélisation de séries chronologiques est particulièrement adaptée à l'estimation des variations d'un mois à l'autre par la composante de la tendance. La précision des variations d'un mois à l'autre dépend toutefois fortement du choix du modèle de tendance. Un modèle de tendance à l'échelle locale (T3) donne des tendances plus volatiles et présente une erreur-type nettement plus grande pour la variation d'un mois à l'autre. Les modèles parcimonieux où l'on suppose que les composantes saisonnières ou le biais de renouvellement sont égaux dans les domaines ont pour résultat d'autres fortes réductions des erreurs-types au prix de l'introduction d'un biais dans la tendance ou les profils saisonniers.

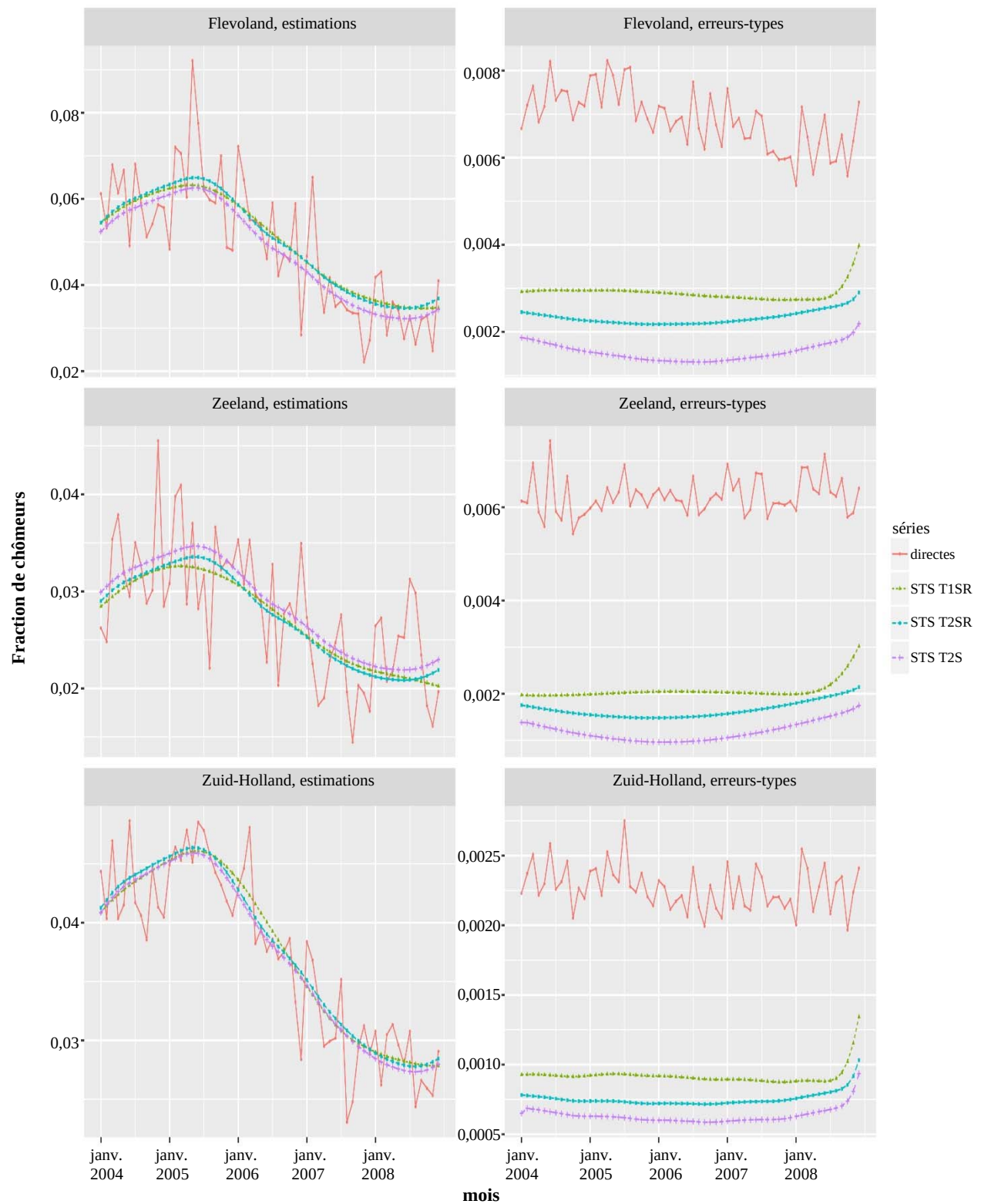


Figure 5.1 Comparaison des estimations directes et des estimations de tendance lisse pour trois modèles (à gauche) et leurs erreurs-types estimées (à droite).

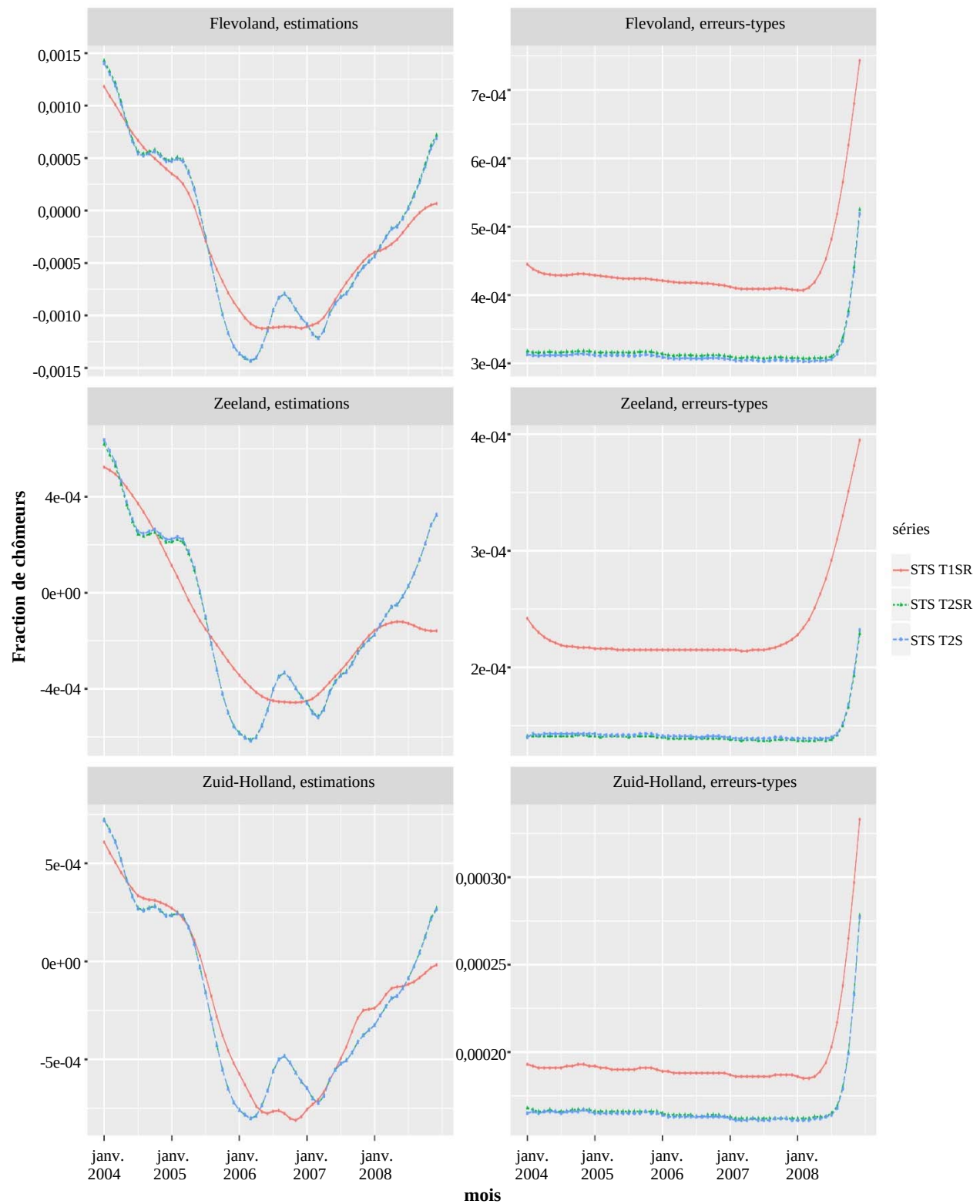


Figure 5.2 Comparaison des évolutions lisses d'un mois à l'autre (à gauche) et de leurs erreurs-types (à droite).

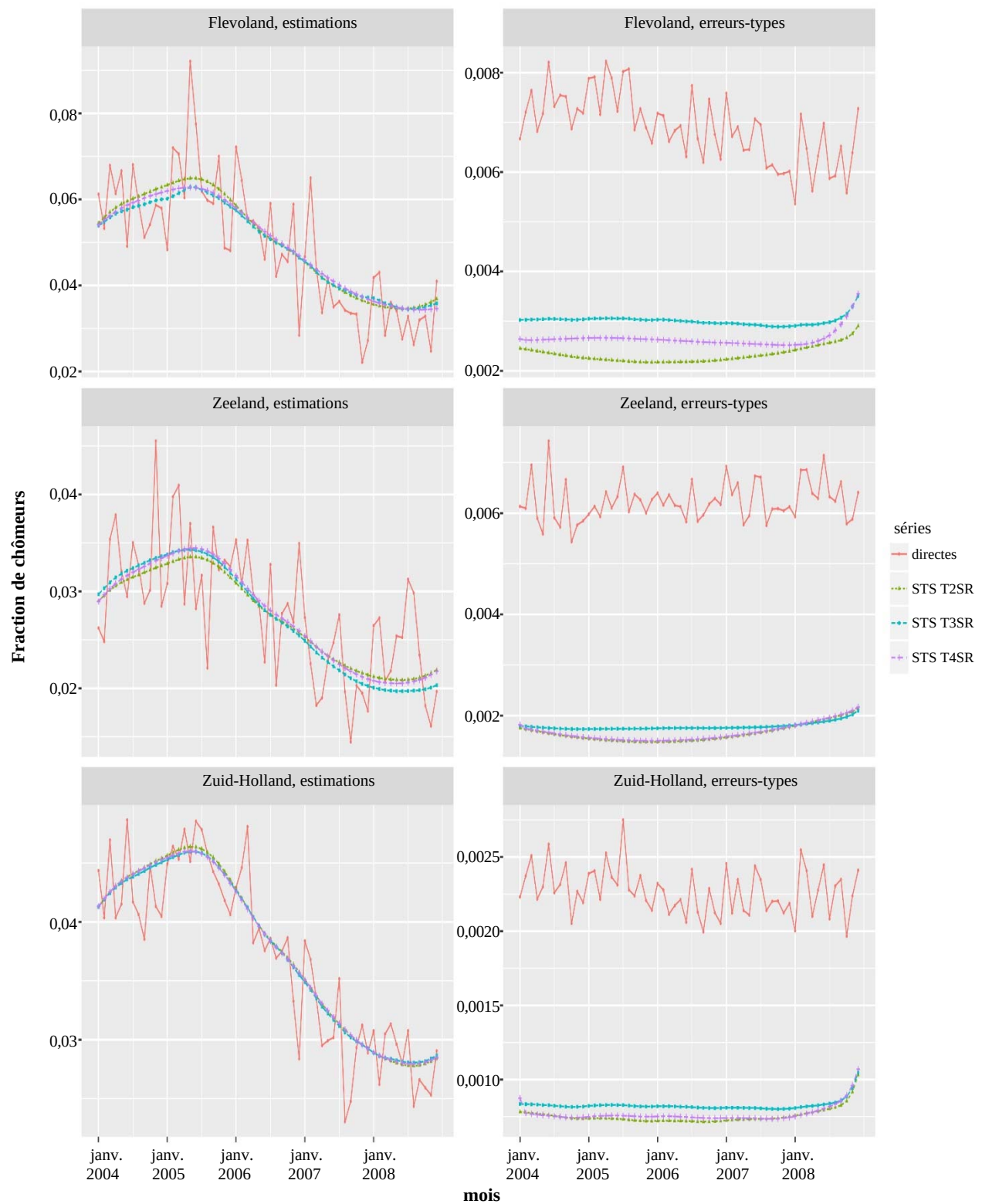


Figure 5.3 Comparaison des estimations directes et des estimations de tendance lisse pour trois modèles (à gauche) et leurs erreurs-types estimées (à droite).

Tableau 5.3

Réductions relatives des erreurs-types (5.2) des estimations des signaux fondées sur des modèles espace-état comparées à celles des estimations directes (%), par province

	Grn	Frs	Drn	Ovr	Flv	Gld	Utr	N-H	Z-H	Zln	N-B	Lmb
T1SR	36	36	38	42	43	44	47	47	45	50	47	43
T2SR	43	42	43	48	49	49	53	53	50	54	53	48
T2S	49	48	51	53	55	54	58	56	54	58	56	54
T2R	64	63	62	65	66	63	68	68	63	73	67	64
T3SR	45	41	45	48	42	51	49	50	48	53	49	50
T3R	67	62	63	66	56	61	62	64	65	70	60	66
T3R2	68	63	64	67	57	62	62	65	65	70	60	67
T3	79	74	76	75	65	69	69	69	69	76	63	76
T4SR	43	41	45	48	45	50	49	53	50	54	51	49
T4R	65	63	64	65	62	62	63	68	63	73	63	65

Tableau 5.4

Moyennes des erreurs-types pour tous les mois et toutes les provinces par rapport à la moyenne des erreurs-types de l'estimateur direct (%) pour les modèles espace-état

	e-t(signal)	e-t(tendance)	e-t(croissance)
directes	100		
T1SR	57	41	6
T2SR	51	33	4
T2S	46	23	4
T2R	34	33	4
T3SR	53	35	9
T3R	36	35	9
T3R2	36	34	9
T3	28	26	9
T4SR	52	34	4
T4R	35	34	4

5.2 Résultats des modèles multiniveaux

Les 10 modèles T1SR à T4R, des pages 436-437, ajustés comme modèle espace-état au moyen du filtre de Kalman, ont également été ajustés au moyen d'une approche bayésienne multiniveau à l'aide d'un échantillonneur Gibbs. Voir Boonstra et van den Brakel (2016) pour une description détaillée des matrices du plan d'expérience à effet fixe, des matrices du plan d'expérience à effet aléatoire et des matrices de précision correspondant à ces modèles. L'approche bayésienne tient compte de l'incertitude des hyperparamètres en considérant leurs distributions a posteriori, ce qui implique que les paramètres de variance ne deviennent en fait pas nuls, comme cela se produit fréquemment pour les estimations par le MV dans la méthode espace-état. Cependant, à des fins de comparaison, les effets absents du modèle espace-état en raison d'estimations par le MV nulles ont également été supprimés dans les modèles multiniveaux correspondants. En plus de ces 10 modèles, nous considérons un autre modèle ayant des termes supplémentaires, notamment une composante dynamique de biais de renouvellement et un terme de bruit blanc.

Les différences entre les estimations des modèles espace-état et multiniveaux fondées sur les 10 modèles considérés peuvent s'expliquer par :

- les différentes méthodes d'estimation, MV par rapport à MCMC;
- la modélisation différente des erreurs d'enquête. Dans les modèles multiniveaux, on suppose que la matrice de covariance des erreurs d'enquête est $\Sigma = \bigoplus_{i=1}^{m_A} \lambda_i \Phi_i$ avec Φ_i la matrice de

covariance des variances par rapport au plan estimées pour les estimations initiales de la province i , et les facteurs d'échelle λ_i , un par province. Dans les modèles espace-état, les erreurs d'enquête sont autorisées à dépendre d'un plus grand nombre de paramètres, bien que, finalement, un modèle AR(1) soit utilisé pour l'estimation de ces dépendances;

- les paramétrisations légèrement différentes des composantes de la tendance. Par exemple, en ce qui concerne la tendance du modèle T3, la province de Groningen se distingue par le modèle espace-état utilisé, car aucune composante d'échelle locale n'est ajoutée pour cette province.

Les estimations et, dans une moindre mesure, les erreurs-types fondées sur les modèles multiniveaux sont très semblables aux résultats obtenus au moyen des modèles espace-état. Nous le montrons seulement pour les signaux lissés du modèle T2R à la figure 5.4, car les différences qualitatives entre les résultats des modèles espace-état et multiniveaux sont assez convergents dans tous les modèles. Boonstra et van den Brakel (2016) présentent d'autres comparaisons des signaux, des tendances et des développements d'un mois à l'autre pour les modèles T2R et T3R2.

Les petites différences entre les estimations de signaux des modèles espace-état et multiniveaux sont attribuables à des tendances légèrement plus souples dans les modèles multiniveaux estimés. On observe des différences plus grandes dans les erreurs-types du signal : les modèles multiniveaux produisent presque toujours des erreurs-types plus grandes pour les provinces au taux de chômage élevé (Flevoland et Zuid-Holland dans la figure), alors que pour les provinces au taux de chômage plus faible (par exemple, Zeeland), les différences sont un peu moins prononcées.

La plus grande souplesse des tendances du modèle multiniveau est fort probablement attribuable à l'incertitude relativement grande sur les paramètres de variance pour la tendance, qui est prise en compte dans l'approche bayésienne multiniveau, mais ignorée dans l'approche du MV des modèles espace-état. Les distributions a posteriori pour les paramètres de la variance de la tendance présentent également un léger étalement à droite. Les moyennes a posteriori des écarts-types sont toujours plus grandes que les estimations par le MV des hyperparamètres correspondants des modèles espace-état (comparer le tableau 2 et le tableau 8 dans Boonstra et van den Brakel (2016)). Pour les modèles ayant la tendance T2, c'est-à-dire avec matrice de covariance entièrement paramétrée sur toutes les provinces, les modèles multiniveaux montrent des corrélations positives entre les provinces, tout comme les estimations par le MV du modèle espace-état, mais ces dernières sont nettement plus concentrées près de 1, alors que les moyennes a posteriori des corrélations dans le modèle multiniveau correspondant T2SR se situent toutes entre 0,45 et 0,8.

Le tableau 5.5 contient les valeurs du critère de sélection du modèle à critère d'information de déviance (DIC) (Spiegelhalter, Best, Carlin et van der Linde, 2002), le nombre réel associé de paramètres du modèle p_{eff} et la moyenne a posteriori de la log-vraisemblance. Le modèle parcimonieux T3 est le modèle le plus favorable selon le critère DIC. Donc, dans ce cas, le critère DIC sélectionne le même modèle que celui sélectionné par les critères AIC et BIC dans les modèles espace-état. Le critère DIC présente l'avantage d'utiliser un nombre réel de paramètres de modèle en fonction de la taille des effets aléatoires, plutôt que seulement le nombre de paramètres de modèle utilisés dans AIC et BIC. Cela dit, les nombres p_{eff} correspondent aux totaux des nombres d'états et d'hyperparamètres du tableau 5.1 pour les modèles espace-état.

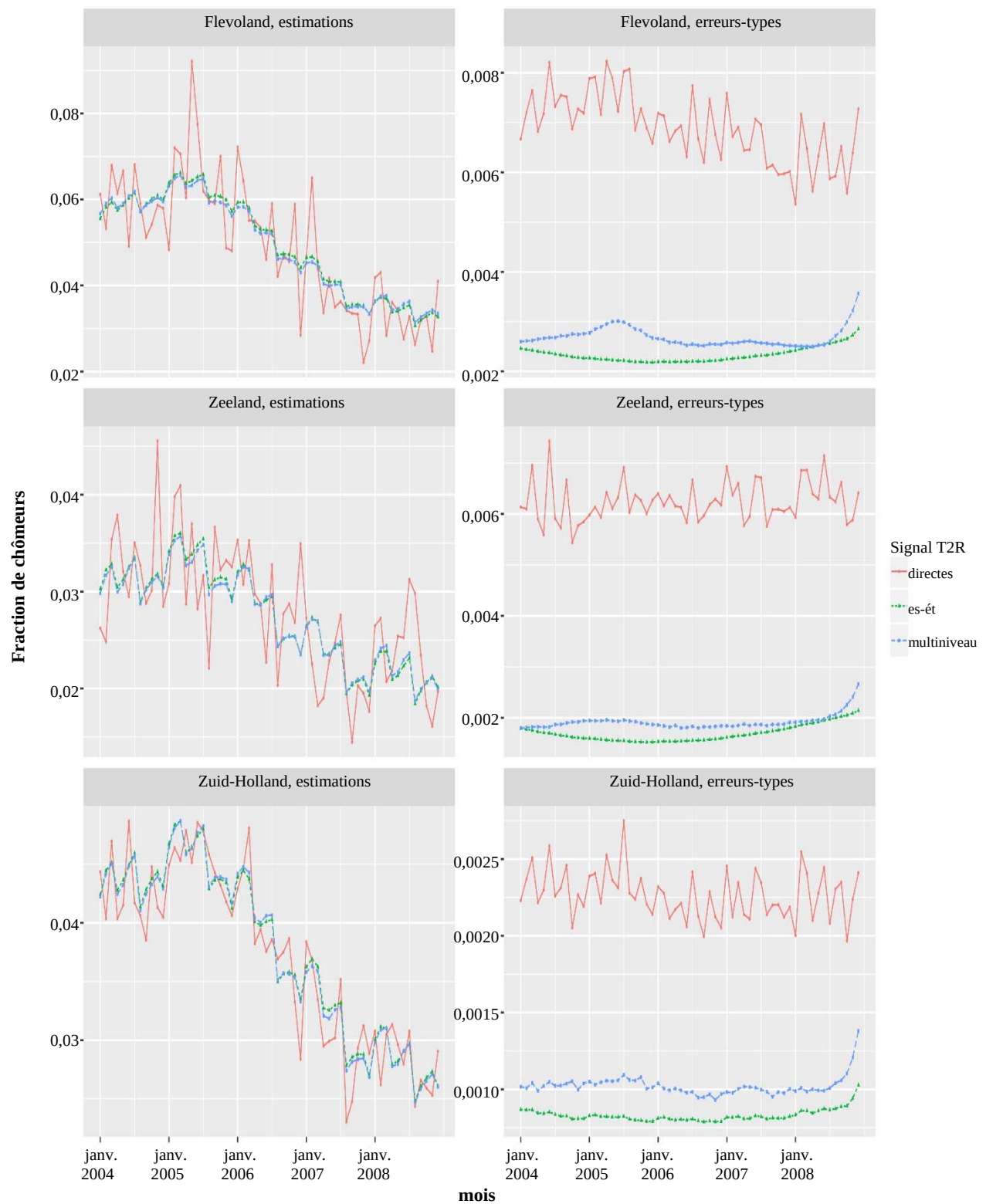


Figure 5.4 Comparaison entre les signaux lissés (à gauche) et leurs erreurs-types (à droite) obtenues au moyen du modèle espace-état (es-ét) T2R et du modèle multiniveau correspondant.

Tableau 5.5
DIC, nombre réel de paramètres du modèle et moyenne a posteriori de la log-vraisemblance

	DIC	p_{eff}	moyenne l-vrais.
T1SR	-29 054	255	14 655
T2SR	-29 076	235	14 656
T2S	-29 129	196	14 662
T2R	-29 164	118	14 641
T3SR	-29 081	242	14 662
T3R	-29 174	126	14 650
T3R2	-29 217	94	14 655
T3	-29 230	82	14 656
T4SR	-29 084	228	14 656
T4R	-29 170	109	14 640

Comme dans les modèles espace-état, le modèle parcimonieux T3 s'accompagne d'un biais moyen plus grand dans le temps pour les provinces de Groningen et Flevoland, qui présentent les taux de chômage les plus élevés. Le modèle T3R2 a des biais moyens nettement plus faibles pour les provinces de Groningen et Flevoland et, comme la valeur de son critère DIC n'est pas beaucoup plus élevée que dans le modèle T3, le modèle T3R2 semble être un bon compromis entre les modèles T3 et T3R, car il est plus parcimonieux que T3R et qu'il respecte mieux les différences provinciales que T3.

Le tableau 5.6 présente les erreurs-types moyennes pour le signal, la tendance et les différences d'un mois à l'autre de la tendance, comparativement à la moyenne des estimations directes. La moyenne est calculée pour tous les mois et toutes les provinces. Encore une fois, les résultats ressemblent à ceux obtenus au moyen des modèles espace-état (voir le tableau 5.4), bien que les erreurs-types des variations d'un mois à l'autre soient plus grandes dans les modèles multiniveaux.

Tableau 5.6
Moyennes des erreurs-types pour tous les mois et toutes les provinces par rapport à la moyenne des erreurs-types de l'estimateur direct (%) dans les modèles de séries chronologiques multiniveaux

	e-t(signal)	e-t(tendance)	e-t(croissance)
directes	100		
T1SR	55	41	8
T2SR	52	37	6
T2S	49	33	7
T2R	39	38	6
T3SR	53	38	15
T3R	39	38	15
T3R2	39	38	15
T3	34	32	15
T4SR	51	36	6
T4R	37	36	6

Enfin, un modèle multiniveau fondé sur le modèle T3R2, mais avec des effets aléatoires supplémentaires, a été ajusté aux données. Ce modèle étendu comprend un terme de bruit blanc, un effet saisonnier nominal

équilibré (équivalent à la composante saisonnière trigonométrique) et une composante dynamique de biais de renouvellement. On a constaté que ces composantes étaient absentes ou indépendantes du temps dans l'approche espace-état en raison d'estimations d'hyperparamètres par le MV nulles. C'est pourquoi elles n'ont pas été incluses dans les modèles multiniveaux examinés jusqu'à maintenant. De plus, le modèle multiniveau étendu comprend des effets aléatoires saisonniers par province, comme compromis entre des effets saisonniers provinciaux fixes et l'absence totale d'effet d'interaction de ce type. Boonstra et van den Brakel (2016) présentent plus de détails et de chiffres comparant les résultats d'estimation donnés par ce modèle étendu à ceux des modèles multiniveaux T3R2 et T3SR. On a constaté que la plupart des effets aléatoires supplémentaires étaient petits, de sorte que les estimations fondées sur le modèle étendu sont assez proches des estimations fondées sur le modèle T3R2, et que les erreurs-types estimées ne sont que légèrement plus grandes que celles du modèle T3R2. On a trouvé une valeur de DIC de -29 260, soit une valeur bien inférieure à celle du DIC du modèle T3R2. On a constaté que cette amélioration du DIC était presque entièrement causée par la composante dynamique du biais de renouvellement. Apparemment, la modélisation du biais de renouvellement comme dépendant du temps donne un meilleur ajustement. Ce résultat semble concorder avec les variations temporelles des différences entre les estimations par la régression d'enquête de la première vague et des vagues de suivi, présentées dans figure 3 de Boonstra et van den Brakel (2016).

6 Discussion

On a appliqué un modèle d'estimation sur petits domaines de séries chronologiques à un grand nombre de données d'enquête, comprenant six années de données de l'EPA néerlandaise, afin d'estimer les fractions mensuelles de chômage dans 12 provinces au cours de cette période. Deux méthodes d'estimation différentes pour modèles de séries chronologiques structurels (MSCS) ont été appliquées et comparées. La première est une approche espace-état utilisant un filtre de Kalman, où les hyperparamètres inconnus sont remplacés par leurs estimations par le MV. La deuxième est une approche bayésienne multiniveau de séries chronologiques, utilisant un échantillonneur de Gibbs.

Les modèles de séries chronologiques qui ne tiennent pas compte des corrélations transversales et qui empruntent de l'information uniquement dans le temps montrent déjà une réduction importante des erreurs-types par rapport aux estimations directes. On obtient une petite diminution supplémentaire des erreurs-types en empruntant de l'information dans l'espace par des corrélations transversales dans les modèles de séries chronologiques. La méthode de modèle de séries chronologiques présente un autre grand avantage en matière d'estimation des variations. Le modèle multiniveau permet de calculer facilement les estimations des variations et de leurs erreurs-types, surtout lorsque l'ajustement du modèle prend la forme d'une simulation par MCMC. Dans l'approche espace-état, les estimations des variations découlent directement de la récursion du filtre de Kalman par conservation des variables d'état requises du passé dans le vecteur d'état. L'estimation de la variation souhaitée, y compris son erreur-type, résulte du contraste des variables

d'état particulières. La variation des données mensuelles d'un mois à l'autre et d'une année à l'autre est très stable et précise, ce qui est une conséquence de la forte corrélation positive entre les estimations de niveau. La stabilité des estimations des variations dépend toutefois fortement du choix du modèle de tendance. Les modèles à échelle locale donnent des estimations de tendance plus volatiles et, par conséquent, des estimations de la variation plus volatiles et ils présentent naturellement une erreur-type plus élevée que les modèles à tendance lisse.

Dans le présent article, on considère différents modèles de tendances comme étant la corrélation du modèle entre les domaines dans le but d'emprunter de l'information dans le temps et l'espace. La méthode la plus complexe consiste à prescrire une matrice de covariance complète pour les termes de perturbation de la composante de tendance. Une des manières de construire des modèles parcimonieux consiste à tirer profit de la cointégration. Dans le cas d'une forte corrélation entre les domaines, la matrice de covariance sera de rang réduit, ce qui signifie que les tendances des domaines m_A sont fondées sur moins que m_A tendances communes. Dans cette application, il suffit de deux tendances communes pour modéliser la dynamique des 12 provinces, ce qui entraîne une forte réduction du nombre d'hyperparamètres requis pour la modélisation des corrélations transversales entre domaines. Afin de réduire davantage encore le nombre d'états et d'hyperparamètres, on considère d'autres modèles de tendance qui tiennent implicitement compte des corrélations transversales. Dans cette approche, tous les domaines partagent une tendance globale. Chaque domaine a une tendance propre au domaine qui permet de tenir compte de l'écart par rapport à la tendance globale. Cela peut être perçu comme une forme simplifiée de modèle de tendance commun. Dans cette application, l'autre modèle de tendance a pour résultat des estimations comparables pour les tendances et les erreurs-types. Cette méthode pourrait donc être une autre solution pratique et attrayante pour les modèles de tendance communs. Par exemple, si le nombre de domaines est élevé ou que le nombre de facteurs communs est plus élevé, les modèles de tendance proposés sont moins complexes que les modèles de tendance commune globale. Il faudrait entreprendre des recherches plus poussées sur les propriétés statistiques de ces modèles de tendance de substitution pour mieux comprendre les structures de covariance implicites.

On peut observer plusieurs différences entre les modèles multiniveaux de séries chronologiques ajustés dans un cadre bayésien hiérarchique et les modèles espace-état ajustés au moyen du filtre de Kalman dans une approche fréquentiste. Dans le cadre bayésien multiniveau, différents MSCS sont comparés au moyen du critère DIC comme critère de sélection formel du modèle. Étant donné que les modèles espace-état sont ajustés dans un cadre fréquentiste, les MSCS sont comparés au moyen du critère AIC ou BIC. L'un des avantages du critère DIC utilisé dans l'approche bayésienne multiniveau est qu'il utilise le nombre réel de degrés de liberté comme pénalité pour la complexité du modèle. Cela signifie que la pénalité pour un effet aléatoire augmente avec la taille des composantes de la variance de ce facteur aléatoire et varie entre zéro – si la composante de la variance est égale à zéro – et le nombre de niveaux de ce facteur si la composante de la variance tend à l'infini. En présence d'un critère AIC ou BIC, la pénalité pour une composante aléatoire est toujours égale à 1, quelle que soit la taille de sa composante de variance, et elle ne tient donc pas

correctement compte de la complexité du modèle. Il convient de noter que pour les modèles multiniveaux ajustés dans un cadre fréquentiste, le critère AIC dit conditionnel est proposé (Vaida et Blanchard, 2005), dans lequel la pénalité pour la complexité du modèle est également fondée sur les degrés de liberté réels. Dans ce cas, la pénalité pour un effet aléatoire augmente à mesure que la taille de sa composante de variance augmente, de la même façon qu'avec le critère DIC. Pour les modèles espace-état ajustés dans un cadre fréquentiste, ces critères de sélection du modèle semblent moins facilement disponibles.

Les modèles multiniveaux et les modèles espace-état diffèrent notamment parce que dans le premier type de modèle, on constate plus souvent que les composantes varient en fonction du temps, alors que, dans l'approche espace-état, la plupart des composantes, sauf la tendance, sont estimées comme invariantes dans le temps. Cela résulte de la méthode d'ajustement du modèle. Dans l'approche fréquentiste appliquée aux modèles espace-état, les estimations par le MV de nombreux hyperparamètres se trouvent à la limite de l'espace de paramètre, c'est-à-dire 0 pour les composantes de la variance et 1 pour les corrélations entre les termes de perturbation de la pente. Dans l'approche bayésienne hiérarchique, toute la distribution des paramètres de (co)variance est simulée, ce qui donne des valeurs moyennes pour ces hyperparamètres qui ne sont jamais exactement à la limite de l'espace de paramètre, par exemple, toujours positives dans le cas des composantes de la variance. Cette caractéristique a notamment pour conséquence que les variances des hyperparamètres de tendance sont plus élevées et que les covariances entre les perturbations de tendance sont plus petites que dans l'approche bayésienne hiérarchique. Notons en outre que le critère DIC préfère les modèles avec biais de renouvellement variant dans le temps et des composantes saisonnières variant dans le temps ainsi qu'un terme de bruit blanc pour le paramètre de population. Dans cette application, il en résulte des modèles présentant un degré de complexité plus élevé dans les modèles hiérarchiques bayésiens multiniveaux que dans les modèles espace-état ajustés dans une approche fréquentiste. Toutefois, les différences dans les estimations de la tendance et des signaux sont minimales.

L'un des avantages de l'approche bayésienne hiérarchique est que les erreurs-types des prédictions de domaine tiennent compte de l'incertitude sur les hyperparamètres. Par conséquent, les erreurs-types obtenues selon l'approche bayésienne hiérarchique des modèles comparables sont légèrement plus élevées et moins biaisées que celles obtenues dans l'approche espace-état. Pour la méthode espace-état, on dispose de plusieurs méthodes bootstrap pour tenir compte de l'incertitude des hyperparamètres (Pfeffermann et tiller, 2005), mais ces méthodes augmentent de façon notable le coût de calcul.

Les méthodes sont aussi différentes du point de vue du calcul. On peut utiliser en ligne l'approche du filtre de Kalman appliquée aux modèles espace-état, ce qui permet de produire de nouvelles estimations filtrées par la mise à jour des prédictions antérieures à l'arrivée des données d'un nouveau mois et ce qui rend, de ce point de vue, le calcul très efficace. En revanche, la procédure d'optimisation numérique de l'estimation par le MV des hyperparamètres peut être lourde pour les grands modèles multivariés si le nombre d'hyperparamètres est élevé. L'approche multiniveau avec échantillonneur de Gibbs utilisée ici produit simultanément des estimations pour toute la série chronologique. Elle doit être estimée de nouveau entièrement quand les données d'un nouveau mois arrivent. Il reste qu'en raison de l'utilisation de matrices

parcimonieuses et d'une paramétrisation redondante, l'approche multiniveau est assez compétitive du point de vue des calculs, voir aussi Knorr-Held et Rue (2002). L'un des avantages de l'estimation multiniveau simultanée est qu'on peut facilement imposer des contraintes dans le temps. Par exemple, l'imposition de contraintes à somme nulle dans le temps permet d'inclure les tendances provinciales à l'échelle locale pour toutes les provinces qui s'ajoutent à une tendance lisse globale sans que cela pose de problème d'identification.

Dans cette application, une préférence s'est dégagée pour les modèles multiniveaux de séries chronologiques dans le cadre bayésien hiérarchique. Cette préférence s'explique notamment par le calcul relativement simple du critère DIC, qui tient mieux compte de la complexité du modèle que les critères AIC ou BIC. De plus, l'échantillonneur de Gibbs dans l'approche bayésienne convient mieux à des MSCS multivariés complexes ayant un grand nombre d'hyperparamètres. En outre, les erreurs-types des prédictions de domaine obtenues dans les modèles multiniveaux tiennent compte de l'incertitude sur les hyperparamètres, et ce de façon simple.

Les estimations des séries chronologiques sont assez lisses, et il faudrait évaluer le modèle plus minutieusement pour déterminer si cela est adéquat ou si le modèle de séries chronologiques ajuste insuffisamment les données sur le chômage et qu'il est susceptible d'être amélioré par d'autres moyens. Il existe de nombreuses façons d'étendre le modèle d'EPD de séries chronologiques pour améliorer les estimations et les erreurs-types. Ainsi, il pourrait être préférable d'utiliser une fonction de liaison logarithmique dans la formulation du modèle comme dans You (2008). Les effets seraient alors multiplicatifs et non additifs. Une autre amélioration possible serait une modélisation plus poussée des variances d'échantillonnage (You et Chapman, 2006; You, 2008; Gómez-Rubio, Best, Richardson, Li et Clarke, 2010). On peut également perfectionner les modèles en incluant de l'information auxiliaire supplémentaire de la province par mois, par exemple, le chômage enregistré. Dans Datta et coll. (1999), les effets similaires associés à l'assurance-chômage sont modélisés comme variant selon les régions, mais pas dans le temps.

Remerciements

Les opinions exprimées dans l'article sont celles des auteurs et ne correspondent pas nécessairement à la politique du bureau central de la statistique des Pays-Bas. Ces travaux ont bénéficié d'un financement de l'Union européenne (subvention n° 07131.2015.001-2015.257). Nous remercions les relecteurs anonymes et le rédacteur associé pour leur lecture attentive d'une version antérieure du manuscrit et pour leurs précieux commentaires.

Bibliographie

Andrieu, C., Poucet, R. et Holenstein, R. (2010). Particle markov chain monte carlo methods. *Journal of the Royal Statistical Society, Series B*, 72, 269-342.

- Bailar, B.A. (1975). The effects of rotation group bias on estimates from panel surveys. *Journal of the American Statistical Association*, 70, 349, 23-30.
- Battese, G.E., Harter, R.M. et Fuller, W.A. (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.
- Bollineni-Balabay, O., van den Brakel, J., Palm, F. et Boonstra, H.J. (2016). Multilevel hierarchical bayesian vs. state-space approach in time series small area estimation; the dutch travel survey. Rapport technique 2016-03, <https://www.cbs.nl/en-gb/our-services/methods/papers>, Statistics Netherlands.
- Boonstra, H.J. (2016). *mcmcscsae: MCMC Small Area Estimation*. R package version 0.7, disponible sur demande auprès de l'auteur.
- Boonstra, H.J., et van den Brakel, J. (2016). Estimation of level and change for unemployment using multilevel and structural time-series models. Rapport technique 2016-10, <https://www.cbs.nl/en-gb/our-services/methods/papers>, Statistics Netherlands.
- Chan, J.C.C., et Jeliaskov, I. (2009). Efficient simulation and integrated likelihood estimation in state space models. *International Journal of Mathematical Modelling and Numerical Optimisation*, 1, 101-120.
- Datta, G.S., Lahiri, P., Maiti, T. et Lu, K.L. (1999). Hierarchical bayes estimation of unemployment rates for the states of the U.S. *Journal of the American Statistical Association*, 94, 448, 1074-1082.
- Diallo, M.S. (2014). *Small Area Estimation Under Skew-Normal Nested Error Models*. Thèse de doctorat, School of Mathematics and Statistics, Carleton University, Ottawa, Canada.
- Doornik, J.A. (2009). *An Object-oriented Matrix Programming Language Ox 6*. Timberlake Consultants Press.
- Durbin, J., et Koopman, S.J. (2012). *Time Series Analysis by State Space Methods*. Oxford University Press.
- 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.
- Gelfand, A.E., et Smith, A.F.M. (1990). Sampling based approaches to calculating marginal densities. *Journal of the American Statistical Association*, 85, 398-409.
- Gelman, A. (2006). Prior distributions for variance parameters in hierarchical models. *Bayesian Analysis*, 1, 3, 515-533.
- Gelman, A., et Hill, J. (2007). *Data Analysis Using Regression and Multilevel/Hierarchical Models*. Cambridge University Press.
- Gelman, A., et Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7, 4, 457-472.
- Gelman, A., Van Dyk, D.A., Huang, Z. et Boscardin, W.J. (2008). Using redundant parameterizations to fit hierarchical models. *Journal of Computational and Graphical Statistics*, 17, 1, 95-122.

- Geman, S., et Geman, D. (1984). Stochastic relaxation, gibbs distributions and the bayesian restoration of images. *IEEE Trans. Pattn Anal. Mach. Intell.*, 6, 721-741.
- Gómez-Rubio, V., Best, N., Richardson, S., Li, G. et Clarke, P. (2010). Bayesian statistics for small area estimation. Rapport technique <http://www.bias-project.org.uk/research.htm>.
- Harvey, A.C. (2006). Seasonality and unobserved components models: an overview. Rapport technique ISSN 1725-4825, Conference on seasonality, seasonal adjustment and their implications for shortterm analysis and forecasting, 10 au 12 mai 2006, Luxembourg.
- Kish, L. (1965). *Survey Sampling*. Wiley.
- Knorr-Held, L., et Rue, H. (2002). On block updating in markov random field models for disease mapping. *Scandinavian Journal of Statistics*, 29, 4, 597-614.
- Koopman, S.J., Harvey, A.C., Doornik, J.A. et Shephard, A. (1999). *STAMP: Structural Time Series Analyser, Modeller and Predictor*. Timberlake Consultants, Press London.
- Koopman, S.J., Shephard, A. et Doornik, J.A. (1999). Statistical algorithms for models in state-space form using ssfpack 2.2. *Econometrics Journal*, 2, 113-166.
- Koopman, S.J., Shephard, A. et Doornik, J.A. (2008). *Ssfpack 3.0: Statistical algorithms for models in state-space form*. Timberlake Consultants, Press London.
- Krieg, S., et van den Brakel, J.A. (2012). Estimation of the monthly unemployment rate for six domains through structural time series modelling with cointegrated trends. *Computational Statistics and Data Analysis*, 56, 2918-2933.
- McCausland, W.J., Miller, S. et Pelletier, D. (2011). Simulation smoothing for state-space models: A computational efficiency analysis. *Computational Statistics and Data Analysis*, 55, 199-212.
- O'Malley, A.J., et Zaslavsky, A.M. (2008). Domain-level covariance analysis for multilevel survey data with structured nonresponse. *Journal of the American Statistical Association*, 103, 484, 1405-1418.
- Pfeffermann, D. (1991). Estimation and seasonal adjustment of population means using data from repeated surveys. *Journal of Business & Economic Statistics*, 163-175.
- Pfeffermann, D., et Burck, L. (1990). Estimation robuste pour petits domaines par la combinaison de données chronologiques et transversales. *Techniques d'enquête*, 16, 2, 229-249. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/1990002/article/14534-fra.pdf>.
- Pfeffermann, D., et Tiller, R. (2005). Bootstrap approximation to prediction mse for state-space models with estimated parameters. *Journal of Time Series Analysis*, 26, 893-916.
- Pfeffermann, D., et Tiller, R. (2006). Small area estimation with state-space models subject to benchmark constraints. *Journal of the American Statistical Association*, 101, 1387-1397.
- Piepho, H.-P., et Ogutu, J.O. (2014). Simple state-space models in a mixed model framework. *The American Statistician*, 61, 3, 224-232.
- Proietti, T. (2000). Comparing seasonal components for structural time series models. *International Journal of Forecasting*, 16, 247-260.

- Rao, J.N.K., et Molina, I. (2015). *Small Area Estimation*. Wiley-Interscience.
- Rao, J.N.K., et Yu, M. (1994). Small area estimation by combining time series and cross-sectional data. *The Canadian Journal of Statistics*, 22, 511-528.
- Rue, H., et Held, L. (2005). *Gaussian Markov Random Fields: Theory and Applications*. Chapman and Hall/CRC.
- Ruiz-Cárdenas, R., Krainski, E.T. et Rue, H. (2012). Direct fitting of dynamic models using integrated nested laplace approximations - inla. *Computational Statistics and Data Analysis*, 56, 1808-1828.
- Särndal, C.-E., Swensson, B. et Wretman, J. (1992). *Model Assisted Survey Sampling*. Springer.
- Spiegelhalter, D.J., Best, N.G. Carlin, B.P. et van der Linde, A. (2002). Bayesian measures of model complexity and fit. *Journal of the Royal Statistical Society B*, 64, 4, 583-639.
- Vaida, F., et Blanchard, S. (2005). Conditional akaike information for mixed-effects models. *Biometrika*, 95, 2, 351-370.
- van den Brakel, J.A., et Krieg, S. (2009). Estimation du taux de chômage mensuel par modélisation structurelle de séries chronologiques dans un plan de sondage avec renouvellement de panel. *Techniques d'enquête*, 35, 2, 193-207. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2009002/article/11040-fra.pdf>.
- van den Brakel, J.A., et Krieg, S. (2015). Remédier aux petites tailles d'échantillon, au biais de groupe de renouvellement et aux discontinuités dans les plans de sondage avec renouvellement de panel. *Techniques d'enquête*, 41, 2, 281-312. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2015002/article/14231-fra.pdf>.
- van den Brakel, J.A., et Krieg, S. (2016). Small area estimation with state-space common factor models for rotating panels. *Journal of the Royal Statistical Society*, 179, 763-791.
- Woodruff, R. (1966). Use of a regression technique to produce area breakdowns of the monthly national estimates of retail trade. *Journal of the American Statistical Association*, 61, 314, 496-504.
- You, Y. (2008). Une approche intégrée de modélisation de l'estimation du taux de chômage pour les régions infraprovinciales au Canada. *Techniques d'enquête*, 34, 1, 21-31. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2008001/article/10614-fra.pdf>.
- You, Y., et Chapman, B. (2006). Estimation pour petits domaines au moyen de modèles régionaux et d'estimations des variances d'échantillonnage. *Techniques d'enquête*, 32, 1, 107-114. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2006001/article/9263-fra.pdf>.
- You, Y., Rao, J.N.K. et Gambino, J. (2003). Estimation du taux de chômage fondée sur un modèle pour l'Enquête sur la population active du Canada : une approche bayésienne hiérarchique. *Techniques d'enquête*, 29, 1, 27-36. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2003001/article/6602-fra.pdf>.