

Techniques d'enquête

Modélisation de séries chronologiques multiniveaux de la couverture des soins prénataux au Bangladesh à des niveaux administratifs désagrégés

par Sumonkanti Das, Jan van den Brakel, Harm Jan Boonstra
et Stephen Haslett

Date de diffusion : le 15 décembre 2022



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 infostats@statcan.gc.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 |

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 le Roi du chef du Canada, représenté par le ministre de l'Industrie 2022

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.

Modélisation de séries chronologiques multiniveaux de la couverture des soins prénataux au Bangladesh à des niveaux administratifs désagrégés

Sumonkanti Das, Jan van den Brakel, Harm Jan Boonstra et Stephen Haslett¹

Résumé

Des modèles de séries chronologiques multiniveaux sont appliqués pour estimer les tendances de séries chronologiques de la couverture des soins prénataux à plusieurs niveaux administratifs du Bangladesh, d'après les cycles répétés de la *Bangladesh Demographic and Health Survey* (BDHS, Enquête démographique et sur la santé du Bangladesh) pendant la période allant de 1994 à 2014. Les modèles de séries chronologiques multiniveaux sont exprimés dans un cadre bayésien hiérarchique et ajustés au moyen de simulations Monte Carlo par chaînes de Markov. Les modèles tiennent compte des intervalles variables de trois ou quatre ans entre les cycles de la BDHS et fournissent aussi des prédictions pour les années intermédiaires. Il est proposé d'appliquer les modèles transversaux de Fay-Herriot aux années d'enquête séparément au niveau des districts, soit l'échelle régionale la plus détaillée. Les séries chronologiques de ces prédictions pour petits domaines au niveau des districts et leurs matrices de variance-covariance sont utilisées comme séries de données d'entrée pour les modèles de séries chronologiques multiniveaux. Dans ces modèles, on examine les corrélations spatiales entre les districts, la pente et l'ordonnée à l'origine aléatoires au niveau des districts, ainsi que les différents modèles de tendance au niveau des districts et aux niveaux régionaux plus élevés pour l'emprunt d'information dans le temps et l'espace. Les estimations des tendances au niveau des districts sont obtenues directement à partir des résultats des modèles, tandis que les estimations des tendances à des échelons régionaux et nationaux plus élevés sont obtenues par agrégation des prédictions au niveau des districts, ce qui donne un ensemble cohérent d'estimations des tendances sur le plan numérique.

Mots-clés : Approche bayésienne hiérarchique; enquêtes démographiques et sur la santé; estimation sur petits domaines; modèle transversal de Fay-Herriot; simulation par la méthode de Monte Carlo par chaînes de Markov (MCMC).

1. Introduction

Dans plus de 90 pays, les enquêtes démographiques et sur la santé ont été largement utilisées pour estimer des indicateurs nationaux et infranationaux concernant la fécondité, la planification familiale, la mortalité infantile, la santé infantile, la santé maternelle et la nutrition des enfants et des femmes (DHS, 2021). Le plan de sondage de l'Enquête démographique et sur la santé du Bangladesh (BDHS) ne prend pas en compte les unités administratives inférieures à l'échelle infranationale (7 divisions), notamment 64 districts et plus de 450 sous-districts (deuxième et troisième niveaux de hiérarchie administrative respectivement). Par conséquent, la taille des échantillons est trop petite pour permettre d'estimer les indicateurs sous le niveau de la division au moyen d'estimateurs classiques fondés sur le plan de sondage. Au total, sept enquêtes ont été réalisées pour la période allant de 1994 à 2014. Elles ont fourni des séries chronologiques d'estimations directes à l'échelle nationale et au niveau de la division concernant les indicateurs susmentionnés, qui permettent de suivre les progrès de la diminution de la mortalité maternelle et néonatale au Bangladesh. Toutefois, pour optimiser l'affectation des ressources et l'élaboration des politiques, des renseignements statistiques plus détaillés et fiables à l'échelle régionale des districts sont

1. Sumonkanti Das, École de démographie, Australian National University, Département d'économie quantitative, Maastricht University; Jan van den Brakel, Département d'économie quantitative, Maastricht University, Département de méthodologie, Bureau central de la statistique des Pays-Bas (CBS). Courriel : ja.vandenbrakel@cbs.nl; Harm Jan Boonstra, Département de méthodologie, Bureau central de la statistique des Pays-Bas (CBS); Stephen Haslett, École des sciences fondamentales et Centre de recherche en santé publique, Massey University, École de recherche en finances, Études d'actuariat et statistique, Australian National University.

nécessaires. C'est à cette fin que des modèles d'estimation sur petits domaines sont élaborés dans le présent article. L'estimation sur petits domaines désigne une classe de procédures d'estimation fondées sur un modèle qui améliorent l'exactitude des estimations directes pour un domaine en augmentant la taille réelle de l'échantillon dans chaque domaine distinct au moyen de l'information de l'échantillon observée dans d'autres domaines ou dans la période de référence précédente. Cette méthode est souvent décrite comme un emprunt d'information dans l'espace ou dans le temps, respectivement (Rao et Molina, 2015).

La BDHS est réalisée de façon répétée avec des intervalles de trois ou quatre ans entre deux enquêtes consécutives. La présente étude porte sur sept éditions couvrant la période de 1994 à 2014. Dans le présent article, des modèles multivariés de séries chronologiques multiniveaux sont élaborés pour produire des estimations fiables des tendances de la couverture des soins prénataux au niveau des districts, des divisions et du pays. Ces modèles sont élaborés dans un cadre bayésien hiérarchique et ajustés au moyen de simulations Monte Carlo par chaînes de Markov. L'avantage d'une approche multivariée des séries chronologiques est qu'elle tire parti de toute l'information disponible en modélisant les corrélations transversales et temporelles entre districts et périodes de référence. Étant donné que les modèles sont définis à une fréquence annuelle, ils prennent en compte les intervalles variables de trois ou quatre ans entre cycles d'enquête. De plus, les modèles de séries chronologiques multiniveaux fournissent des prédictions pour les années sans données d'enquête.

Deux variables de réponses connexes sont examinées dans le présent article : si aucune consultation prénatale n'a été reçue ou si au moins quatre consultations prénatales ont été reçues, variables désignées respectivement par l'abréviation ANC0 (*antenatal care*) et ANC4. Les estimations directes et les estimations de la variance sont calculées à partir des données transversales des sept enquêtes BDHS à l'échelle régionale la plus détaillée des districts et sont utilisées comme données d'entrée dans le modèle de séries chronologiques multiniveaux. Cette méthode présente toutefois un inconvénient : d'autres renseignements auxiliaires, tirés de deux recensements, ne peuvent pas être inclus dans les modèles de séries chronologiques multiniveaux. Le recensement est mené tous les 10 ans. Cela signifie que les mêmes valeurs d'information auxiliaire, disponibles dans un recensement donné, sont utilisées dans deux ou même trois éditions ultérieures de la BDHS menées après ce recensement, mais avant le prochain. Cela produit des chocs dans les prédictions des séries chronologiques multiniveaux pendant les périodes pour lesquelles les renseignements d'un nouveau recensement sont publiés. On contourne ce problème en élaborant la méthode suivante en deux étapes. Dans un premier temps, on applique des modèles transversaux de Fay-Herriot (F-H) (Fay et Herriot, 1979) à chaque cycle d'enquête en utilisant les estimations directes au niveau du district et leurs erreurs-types lissées. L'information auxiliaire de recensement sert à améliorer ces modèles transversaux de F-H. Dans une deuxième étape, ces estimations transversales de F-H sont utilisées comme données d'entrée dans le modèle de séries chronologiques multiniveau. Notons que les prédictions transversales de F-H pour une année d'enquête donnée sont corrélées. Les modèles de séries chronologiques multiniveaux tiennent compte de cette corrélation en utilisant les matrices de variance-covariance complètes des prédictions de F-H transversales comme données d'entrée dans le modèle de séries chronologiques multiniveau. Cette approche en deux étapes présente l'avantage d'éliminer les grandes erreurs d'échantillonnage des estimations directes et de

stabiliser la série de données d'entrée pour les modèles de séries chronologiques multiniveaux. Cela repose toutefois sur l'hypothèse que les séries d'entrées des modèles de séries chronologiques ne présentent pas de biais en raison d'une erreur de spécification des modèles de F-H transversaux. Pour éviter cela, il faut un processus de sélection et d'évaluation minutieux des modèles transversaux de F-H à la première étape.

Les modèles de séries chronologiques multiniveaux empruntent de l'information dans le temps et l'espace de plusieurs manières. Les relations transversales sont modélisées au moyen d'effets fixes ainsi que de pentes et d'ordonnées à l'origine aléatoires au niveau du district, indépendantes ou corrélées. Les corrélations spatiales entre districts sont aussi prises en compte. Les tendances lisses et les tendances locales au niveau des districts, des divisions et du pays servent à modéliser les corrélations temporelles et transversales. Au lieu que soit définie une matrice de corrélation complète entre les termes de perturbation des tendances au niveau du district, les tendances sont définies au niveau de la division, et sont partagées par tous les districts sous-jacents. Les écarts par rapport à cette tendance générale sont modélisés au moyen des tendances au niveau du district. C'est une façon parcimonieuse de modéliser les relations transversales entre districts (Boonstra et van den Brakel, 2019). Les estimations des tendances au niveau des districts sont obtenues directement à partir des résultats des modèles, tandis que les tendances au niveau des divisions et du pays sont obtenues par agrégation des prédictions au niveau des districts. La production d'estimations pour des niveaux d'agrégation plus élevés par agrégation des prédictions à partir de l'échelle régionale la plus détaillée permet de disposer de tableaux de publication numériquement cohérents par définition. Les estimations relatives aux districts pour les années non couvertes par l'enquête et aux districts non couverts l'année de l'enquête sont aussi prédites à partir des modèles de séries chronologiques estimés.

Les modèles de séries chronologiques multiniveaux élaborés dans le présent article sont des extensions du modèle de F-H. Rao et Yu (1994) ont étendu le modèle de F-H pour emprunter de l'information dans le temps en supposant que les effets aléatoires propres à une zone suivent un modèle autorégressif de premier ordre $Ar(1)$ dans le temps, de manière indépendante entre les domaines. Datta, Lahiri, Maiti et Lu (1999), et You et Rao (2000) ont généralisé une extension de la série chronologique du modèle de F-H à partir de Rao et Yu (1994) dans le cadre hiérarchique de Bayes en considérant que les termes d'erreur propres à un domaine suivent un modèle de marche aléatoire de premier ordre dans le temps plutôt qu'un processus $Ar(1)$. Datta, Lahiri et Maiti (2002) ont également utilisé un modèle de marche aléatoire pour la composante temps afin d'estimer le revenu médian des familles de quatre personnes par État en s'appuyant sur des séries chronologiques et des données transversales en fonction de la méthode d'estimation empirique de Bayes. Marhuenda, Molina et Morales (2013) ont étendu le modèle de Rao et Yu (1994) à une version spatio-temporelle du modèle de F-H en intégrant une hypothèse supplémentaire selon laquelle les effets aléatoires au niveau du domaine suivent un processus autorégressif simultané de premier ordre (Pratesi et Salvati, 2008) pour inclure la corrélation spatiale entre les données des domaines voisins. Ces extensions sont très spécifiques à la seule composante des effets aléatoires au niveau du domaine. Dans le modèle spatio-temporel de F-H examiné dans la présente étude, il est possible de spécifier les effets aléatoires à divers niveaux de désagrégation à côté des domaines des niveaux détaillés cibles afin d'intégrer des corrélations spatiales, temporelles et spatio-temporelles entre données. À cet

égard, le modèle bayésien hiérarchique considéré est plus souple que l'autre extension du modèle de F-H. D'autres études pertinentes de modèles de séries chronologiques multiniveaux et de modèles d'espaces d'état qui étendent le modèle de F-H pour emprunter de l'information dans le temps et l'espace se trouvent notamment dans You, Rao et Gambino (2003); You (2008); Pfeffermann et Burck (1990); Pfeffermann et Tiller (2006); Bollineni-Balabay, van den Brakel, Palm et Boonstra (2017); Boonstra et van den Brakel (2022, 2019) et Boonstra, van den Brakel et Das (2021).

Le reste du présent l'article s'organise comme suit. La section 2 décrit le besoin de données statistiques fiables à un échelon régional peu élevé pour évaluer les objectifs de développement durable liés à la mortalité maternelle et néonatale au Bangladesh. La section 3 décrit brièvement les sources de données et le calcul des estimations directes et des estimations de la variance à partir des données de l'enquête BDHS, ainsi que les transformations des estimations directes et l'approche de la fonction généralisée de variance (FGV) aux fins de lissage des estimations de la variance, deux méthodes qui améliorent l'ajustement du modèle. La section 4 décrit le cadre bayésien hiérarchique de la modélisation multiniveau des séries chronologiques. Les modèles retenus pour ANC0 et ANC4 sont présentés à la section 5, accompagnés d'une brève discussion sur le processus de construction des modèles. La section 6 présente une analyse des estimations des tendances fondées sur les modèles élaborés et certains résultats d'évaluation de modèle sont illustrés à la section 7. Le document se termine par une discussion à la section 8.

2. La nécessité de statistiques régionales fiables sur la mortalité maternelle et néonatale au Bangladesh

Le Bangladesh a fait des progrès remarquables dans la réduction du taux de mortalité maternelle et du taux de mortalité néonatale conformément aux objectifs 4 et 5 du rapport intitulé *Objectifs du Millénaire pour le développement*. Cependant, les indicateurs du taux de mortalité maternelle (170 pour 100 000 naissances vivantes [OMS, UNICEF et autres, 2014]) et du taux de mortalité néonatale (28 pour 1 000 naissances vivantes [NIPORT, 2015a]) restent relativement hauts comparativement à l'objectif de développement durable (ODD) de réduction du taux de mortalité maternelle à 70 pour 100 000 naissances vivantes et du taux de mortalité néonatale à 12 décès pour 1 000 naissances vivantes au Bangladesh (BBS, 2020). Ces taux de mortalité élevés s'expliquent principalement par le recours insuffisant aux services de santé maternelle, comme les soins prénataux, les naissances assistées par du personnel soignant qualifié et les soins postnataux (NIPORT, 2016). En effet, il est important de recevoir suffisamment de soins prénataux pendant la grossesse, car cela augmente également le recours à du personnel soignant qualifié pour l'accouchement et à des soins postnataux (Mrisho, Obrist, Schellenberg, Haws, Mushi, Mshinda, Tanner et Schellenberg, 2009).

L'enquête auprès des ménages la plus récente indique que la majorité des femmes enceintes (75 %) au Bangladesh reçoivent des soins prénataux donnés par des fournisseurs ayant reçu une formation médicale. Cependant, la proportion de femmes qui reçoivent au moins 4 consultations prénatales comme recommandés par l'OMS est nettement moins élevée, soit 37 % (BBS et UNICEF, 2019). Ces données donnent à penser que le Bangladesh accuse un retard dans l'atteinte de l'objectif national de 50 % de

femmes qui reçoivent au moins 4 consultations prénatales d'ici 2016. Pour combler cet écart et atteindre l'objectif 3 de l'ODD, qui consiste à porter à plus de 98 % d'ici 2030 (NIPORT, 2015b) le nombre de femmes qui reçoivent au moins 4 consultations prénatales, le pays a besoin d'une stratégie globale et de jalons précis. Les tendances nationales de la couverture des soins prénataux indiquent que la proportion de femmes n'ayant pas reçu de soins prénataux (ANC0) s'est améliorée pour atteindre seulement 17,2 % en 2019, comparativement à 85 % en 1994, tandis que la proportion de femmes ayant reçu au moins quatre visites de soins prénataux (ANC4) a augmenté à 37 % en 2019, alors qu'elle était de 6 % en 1994. L'amélioration des indicateurs pendant cette période varie selon les divisions. L'amélioration la plus marquée est observée dans la division de *Khulna*, où ANC0 et ANC4 sont passés d'environ 70 % et 5 % à environ 12 % et 40 %, respectivement. L'évolution la plus faible a été observée dans la division de *Sylhet*.

Les installations proposant des services de soins prénataux présentent des variations considérables au Bangladesh. On trouve des cliniques communautaires et des centres de protection de la famille au niveau des unions (ainsi que des cliniques d'organisations non gouvernementales), des complexes de santé au niveau des upazila, qui sont des sous-districts, et des hôpitaux de district et tertiaires au niveau des districts. De plus, l'accès aux médecins privés varie selon le degré d'urbanisation ainsi que la distance entre le district ou le sous-district et les villes métropolitaines les plus proches, en particulier la capitale *Dhaka*. L'inégalité d'accès aux soins prénataux est aussi explicitement visible au niveau de la division. On peut s'attendre à ce que les inégalités soient encore plus grandes aux niveaux administratifs désagrégés, comme les districts et les sous-districts. Aucune étude n'a toutefois confirmé cette hypothèse, principalement en raison de l'absence de données d'enquête suffisamment détaillées à ces niveaux. Des données récentes tirées d'études à des niveaux désagrégés sur la pauvreté, la nutrition des enfants et la morbidité indiquent des niveaux élevés d'inégalité au niveau des districts et des sous-districts (Haslett et Jones, 2004; Haslett, Jones et Isidro, 2014; Das, Kumar et Kawsar, 2020; Hossain, Das et Chandra, 2020).

3. Sources des données et estimations des entrées

3.1 Sources des données

Depuis 1993-1994, la BDHS est réalisée sous l'autorité de l'Institut national de recherche et de formation en démographie (NIPORT) du ministère de la Santé et de la Protection de la famille (MOHFW). L'enquête évalue les programmes sociaux et sanitaires en place afin d'élaborer de nouvelles stratégies d'amélioration de l'état de santé des femmes et des enfants du pays. En 2018, huit enquêtes BDHS avaient été menées : en 1993-1994, 1996-1997, 2000, 2004, 2007, 2011, 2014 et 2017-2018. La présente étude s'appuie sur les données d'enquête relatives à la période de 1994 à 2014 puisque l'emplacement au niveau du district des grappes sondées n'est pas divulgué dans la BDHS la plus récente de 2017-2018. De 1994 à 2014, trois recensements de la population et du logement ont été réalisés, en 1991, 2001 et 2011. Les données complètes du recensement ne sont pas disponibles : seulement 10 % des données du Recensement de 1991, 10 % des données du Recensement de 2001 et 5 % des données du Recensement de 2011 sont accessibles au public par l'intermédiaire d'IPUMS-International (<https://international.ipums.org>). Un certain nombre de variables contextuelles au niveau des districts ont été générées et utilisées dans

l'élaboration de modèles transversaux de F-H afin de produire des estimations d'entrées pour les modèles de séries chronologiques multiniveaux.

3.2 Estimations directes

Les variables analysées dans le présent article sont ANC0 et ANC4. Le Bangladesh est divisé en sept régions infranationales appelées divisions. Ces divisions sont ensuite divisées en 64 districts, soit l'échelle régionale la plus détaillée examinée dans la présente étude. Dans un premier temps, les estimations et les estimations de la variance des deux variables cibles au niveau des districts sont obtenues à partir des données au niveau de l'unité de chaque année d'enquête au moyen d'un estimateur d'enquête direct classique fondé sur le plan (désigné par DIR ci-après), dans lequel les poids d'enquête servent à prendre en compte le plan de sondage et la non-réponse.

Dans la présente étude, les femmes en âge de procréer qui ont été mariées et qui ont accouché au cours des trois années précédant une année d'enquête sont considérées comme la population cible. Étant donné qu'on ne dispose pas de ces renseignements sur la grossesse dans la population du recensement, la taille de la population par domaine est estimée par le nombre de femmes en âge de procréer ayant déjà été mariées qui se trouve dans les trois recensements. Cela signifie que même si les tailles d'échantillon d'un domaine reposent sur un recensement, il y a une certaine incertitude à leur sujet, qui est ignorée dans les modèles d'estimation sur petits domaines (EPD). Pour une analyse plus détaillée sur la taille des populations propres aux divisions et aux districts, voir Das, van den Brakel, Boonstra et Haslett (2021).

La BDHS utilise un échantillon stratifié des ménages à deux degrés. Les strates sont formées à partir des divisions et sous-divisions selon leur caractérisation urbaine-rurale. Les unités primaires d'échantillonnage (UPE) sont les secteurs de dénombrement du recensement de la population et du logement créés pour comporter en moyenne environ 120 ménages (légère variation par rapport au recensement). À la première étape, les UPE sont sélectionnées avec des probabilités proportionnelles à la taille de l'UPE, c'est-à-dire au nombre de ménages. À la deuxième étape, le listage complet des ménages est effectué dans toutes les UPE sélectionnées, puis une trentaine de ménages sont sélectionnés dans chaque UPE au moyen d'un échantillonnage systématique. Le taux de réponse chez les femmes admissibles a dépassé 95 % toutes les années de la BDHS. Bien que la taille de l'échantillon des femmes qui ont été mariées soit supérieure à 10 000 dans toutes les enquêtes, la taille de l'échantillon de la présente étude est plus petite, car seules les femmes déjà mariées qui ont eu un enfant au cours des trois années précédant l'année d'enquête sont prises en compte. Au niveau des districts, la taille moyenne des échantillons varie entre 60 et 114, et certains districts comptent moins de 10 femmes observées, voire aucune.

Les poids de sondage sont calculés en fonction des probabilités de sélection. On a ensuite ajusté ces poids pour tenir compte de la non-réponse des ménages et des personnes. L'estimation directe de la proportion de la population dans un domaine i pour l'année d'enquête t est calculée comme étant la moyenne de l'échantillon.

$$\hat{Y}_{it} = \frac{\sum_{j \in s_{it}} w_{ijt} y_{ijt}}{\sum_{j \in s_{it}} w_{ijt}}, \quad (3.1)$$

où y est la variable de réponse d'intérêt, s_{it} est l'ensemble des femmes déjà mariées dans le domaine i pour lequel y est observé l'année t , et w_{ijt} est le poids de sondage pour la personne j vivant dans le domaine i l'année t . Notons que les poids w_{ijt} sont mis à l'échelle de façon à ce que la somme sur les poids dans l'échantillon soit égale à la taille nette de l'échantillon. Les estimations de la variance correspondantes sont obtenues par approximation comme suit :

$$\text{var}(\hat{Y}_{it}) = \frac{1}{n_{it}(n_{it} - 1)} \sum_{j \in s_{it}} w_{ijt} (y_{ijt} - \hat{Y}_{it})^2, \quad (3.2)$$

où n_{it} est le nombre de femmes qui ont été mariées observé dans le domaine i au cours de l'année de l'enquête t . Au départ, on estimait la variance en calculant la variance parmi les totaux estimés d'UPE comme si la sélection avait été réalisée au moyen d'un échantillonnage stratifié avec remise, appelé approximation de la variance de l'unité d'échantillonnage finale. Cela a donné des estimations de variance nulle pour quelques domaines. L'approximation de la variance (3.2) évite ces estimations de variance nulle, et donne des estimations de la variance comparables à l'approximation initiale où l'on supposait que les UPE étaient sélectionnées avec remise. Dans le premier modèle de séries chronologiques multiniveaux, représenté par MTS-I, ces estimations directes servent de séries de données d'entrée.

3.3 Estimations transversales de Fay-Herriot

Un des problèmes que pose le modèle MTS-I est qu'il utilise des données du recensement comme variables auxiliaires dans le modèle de séries chronologiques multiniveau. Étant donné que l'intervalle entre deux recensements est de 10 ans, alors que la BDHS est réalisée tous les trois ou quatre ans, les covariables du recensement restent identiques jusqu'à ce que de nouvelles données du recensement soient disponibles. L'inclusion de ces données de recensement comme covariables dans les modèles MTS-I faussera les estimations des tendances et des changements d'une période à l'autre. L'une des façons de tirer parti des données du recensement consiste à modéliser les estimations directes au niveau des districts dans des modèles transversaux de F-H distincts en utilisant des variables contextuelles pertinentes extraites des données de recensement. On s'attend aussi à ce que l'utilisation de variables auxiliaires de recensement disponibles à temps dans des modèles transversaux de F-H répétitifs puisse avoir une incidence sur les coefficients de régression et l'exactitude des prédictions de modèle de la variable dépendante, mais non sur les prédictions de la variable dépendante elle-même. Comparativement aux estimations directes utilisées dans MTS-I, ces modèles transversaux de F-H fournissent aussi de meilleures estimations parce qu'ils empruntent certaines informations dans les districts.

Les estimations transversales de F-H et leurs erreurs-types sont utilisées comme données d'entrée pour un deuxième modèle, désigné MTS-II. Les estimations transversales de F-H sont corrélées en raison de leurs composantes communes à effet fixe, qui sont ignorées dans MTS-II. Par conséquent, un troisième modèle de séries chronologiques multiniveaux, désigné MTS-III, est élaboré au moyen des estimations transversales de F-H et de leur matrice de covariance complète comme données d'entrée.

Les composantes à effet fixe et aléatoire pour les modèles transversaux de F-H spécifiques à l'enquête sont présentées dans les tableaux A.2 et A.3 de l'annexe. Pour tous les modèles, on suppose que les effets aléatoires suivent une loi normale. On a examiné des modèles non normaux pour les effets aléatoires (Laplace et *horseshoe*) et l'erreur d'échantillonnage (distribution *t*) comme solutions de rechange à la loi normale. Cela n'a toutefois pas amélioré l'ajustement du modèle.

3.4 Fonctions généralisées de variance

Dans les modèles de F-H et de séries chronologiques multiniveaux, les estimations de la variance des estimations directes sont en grande partie traitées comme des quantités fixes. Comme ces estimations de la variance peuvent présenter un bruit important, elles sont lissées au moyen d'une fonction généralisée de variance (FGV) avant d'être utilisées dans les modèles de F-H et de séries chronologiques multiniveaux. Il est entendu qu'un district sans données d'échantillon est considéré comme manquant et qu'il n'est donc pas pris en compte dans la méthode d'élaboration du modèle. Le modèle de F-H transversal peut produire des estimations et des erreurs-types pour ces domaines hors échantillon. On n'utilise toutefois pas ces estimations synthétiques dans l'élaboration des modèles MTS-II et MTS-III pour obtenir une meilleure comparaison avec le modèle MTS-I.

Les FGV sont des modèles de régression qui relient les estimations de la variance à des prédicteurs comme la taille de l'échantillon, les variables du plan de sondage et les estimations ponctuelles (Wolter (2007), chapitre 7). Pour les variables ANC0 et ANC4, on utilise la FGV suivante :

$$\log se(\hat{Y}_{it}) = \alpha + \beta \log \tilde{Y}_{it} + \gamma \log(m_{it} + 1) + \delta \text{Division} + \epsilon_{it}, \quad (3.3)$$

où $se(\hat{Y}_{it})$ est l'erreur-type de $\hat{Y}_{it}$ dans (3.1), m_{it} le nombre d'unités d'échantillonnage contribuant au district i l'année t et Division est une variable catégorique ayant 7 niveaux. Étant donné que nous ne pouvons pas nous fier aux estimations directes pour les valeurs très petites de m_{it} , les $\tilde{Y}_{it}$ du côté droit de (3.3) sont des estimations lissées simples.

$$\begin{aligned} \tilde{Y}_{it} &= \lambda_{it} \hat{Y}_{it} + (1 - \lambda_{it}) \bar{Y}_{d[i]t}, \\ \lambda_{it} &= \frac{m_{it}}{m_{it} + 1}, \end{aligned} \quad (3.4)$$

où $\bar{Y}_{d[i]t}$ désigne la moyenne pour la division d ($d = 1$ à 7) à laquelle appartient le district i , pour l'année t . Comme l'a mentionné un réviseur, un estimateur par la régression composite peut être utilisé comme solution de rechange pour (3.4).

Les erreurs de régression ϵ_{it} , sont censées être indépendantes et normalement distribuées avec un paramètre de variance commun σ^2 . Les FGV sont ajustées uniquement aux districts pour lesquels les erreurs-types des estimations directes ne sont pas nulles. Les erreurs-types prédites (lissées) qui reposent sur les modèles ajustés sont

$$se_{\text{pred}}(\hat{Y}_{it}) = \exp\left(\hat{\alpha} + \hat{\beta} \log \tilde{Y}_{it} + \hat{\gamma} \log(m_{it} + 1) + \hat{\delta} \text{Division} + \hat{\sigma}^2/2\right), \quad (3.5)$$

où $\hat{\sigma}$ est 0,03 pour ANC0 et 0,003 pour ANC4, respectivement. Les valeurs de R au carré pour les deux modèles sont assez élevées : 0,79 pour ANC0 et 0,99 pour ANC4. Notons que la rétro transformation exponentielle dans (3.5) comprend une correction de biais qui, dans ce cas, n'a qu'un petit effet. Cette méthode sert à obtenir des erreurs-types lissées pour les modèles de F-H transversaux et le modèle MTS-I.

3.5 Transformations des séries de données d'entrée

La transformation par la racine carrée, la transformation logarithmique et la transformation log-ratio sont considérées comme une transformation stabilisatrice de la variance, voir Sakia (1992). La transformation par la racine carrée est appliquée aux données de ANC4 (les modèles de séries chronologiques multiniveaux et les modèles de F-H transversaux), car cette transformation réduit la corrélation entre les estimations ponctuelles et leurs erreurs-types des séries de données d'entrée, diminue l'hétérogénéité, améliore la convergence de la simulation par la méthode de Monte Carlo par chaînes de Markov, et réduit l'asymétrie des données de proportion si elles prennent des valeurs proches de la borne inférieure de zéro. Pour ce qui est de ANC0, la transformation par la racine carrée est utilisée uniquement pour les modèles transversaux de F-H spécifiques aux années 2011 et 2014. Aucune transformation n'est appliquée aux autres années. Dans les trois modèles de séries chronologiques multiniveaux, aucune transformation n'est appliquée pour ANC0, car la transformation par la racine carrée des séries de données d'entrée augmente la dépendance entre les estimations directes et les erreurs-types.

Soit $\hat{Y}_{it} = \sqrt{(\hat{Y}_{it} + \varepsilon)}$ désigne les estimations directes transformées par la racine carrée, où ε est un petit nombre (0,005), ce qui est nécessaire parce que pour certains districts, les estimations directes sont égales à zéro. Au moyen d'une approximation de Taylor du premier degré, on peut montrer que $se(\hat{Y}_{it}) \approx se(\hat{Y}_{it}) / (2\sqrt{\hat{Y}_{it} + \varepsilon})$.

Si la FGV (3.3) est appliquée aux erreurs-types des estimations directes non transformées, alors les erreurs-types pour les domaines comportant un très petit nombre d'unités d'échantillonnage peuvent devenir excessivement grandes en raison de l'approximation de la linéarisation. On évite ce problème en appliquant la FGV aux erreurs-types des estimations transformées, c'est-à-dire $se(\hat{Y}_{it})$.

4. Modélisation multiniveau de séries chronologiques

Dans la présente étude, les estimations directes et leurs erreurs-types sont disponibles pour les années d'enquête 1994, 1997, 2000, 2004, 2007, 2011 et 2014. Pour tenir compte des intervalles variables de trois ou quatre ans entre les années d'enquête, les modèles de séries chronologiques multiniveaux sont définis à une fréquence annuelle (c'est-à-dire que les valeurs désignent une période de référence d'un an) à l'échelle régionale la plus détaillée des 64 districts. Sur une durée de 21 ans, on a 1 344 combinaisons domaine-année. À partir des 7 années d'enquête disponibles, le modèle est adapté aux 448 observations domaine-année. Les années entre deux enquêtes consécutives sont définies comme étant manquantes dans le modèle. Ainsi, l'évolution de la tendance d'une période à l'autre est précisée correctement et le modèle fournit des prédictions pour les combinaisons domaine-année manquantes.

Pour plus de commodité, nous désignerons par $\hat{Y}_{it}$ la série d'entrées des modèles de séries chronologiques pour ANC0 ou ANC4 pour l'année t et le domaine i . Il peut s'agir des estimations directes non transformées, des estimations directes transformées par la racine carrée ou des prédictions du modèle obtenues au moyen de modèles transversaux de F-H. Dans le cas présent, l'indice de domaine i s'exécute de 1 à $M_d = 64$ et l'indice temporel t de 1 à $T = 21$. Nous combinons ensuite ces estimations dans un vecteur $\hat{Y} = (\hat{Y}_{11}, \dots, \hat{Y}_{M_d 1}, \dots, \hat{Y}_{1T}, \dots, \hat{Y}_{M_d T})'$, un vecteur de dimension $M = M_d T$.

4.1 Structure du modèle

Les modèles multiniveaux examinés prennent une forme additive linéaire générale

$$\hat{Y} = X\beta + \sum_{\alpha} Z^{(\alpha)} v^{(\alpha)} + e, \quad (4.1)$$

où X est une matrice de plan $M \times p$ pour un vecteur d'effets fixes β de dimension p et les $Z^{(\alpha)}$ sont des matrices de plan $M \times q^{(\alpha)}$ pour les vecteurs d'effets aléatoires $v^{(\alpha)}$ de dimension $q^{(\alpha)}$. Dans le cas présent, la somme sur α s'exécute sur plusieurs termes possibles d'effet aléatoire à différents niveaux, comme le niveau local et des tendances lisses au niveau du district et de la division, le bruit blanc au niveau le plus détaillé des domaines M , etc. Des explications plus détaillées sont données ci-dessous. Dans la formule (4.1), $e = (e_{11}, \dots, e_{M_d 1}, \dots, e_{M_d T})'$ désigne, selon la série de données d'entrée, les erreurs d'échantillonnage des estimations directes ou les erreurs de prédiction du modèle transversal de F-H. Les erreurs d'échantillonnage sont considérées comme normalement distribuées, soit $e \sim N(0, \Sigma)$ où $\Sigma = \bigoplus_{t=1}^T \Sigma_t$. Si les séries de données d'entrée sont les estimations directes non transformées, alors Σ_t est la matrice de covariance pour les estimations directes non transformées observées l'année t . Si les séries de données d'entrée sont transformées, alors Σ_t est la matrice de covariance pour les estimations directes transformées, comme cela est décrit dans la sous-section 3.5. Si les séries de données d'entrée sont les prédictions qui reposent sur des modèles transversaux de F-H, alors Σ_t contient les erreurs quadratiques moyennes estimées des prédictions de F-H. Dans MTS-II, Σ_t est diagonale et ignore les corrélations entre les prédictions de domaine. Dans MTS-III, Σ_t est une matrice de covariance complète qui prend aussi en compte les corrélations entre les prédictions de domaine.

Selon la répartition des erreurs d'échantillonnage e dans (4.1), la fonction de vraisemblance conditionnelle à des paramètres d'effets fixes et aléatoires peut être définie comme étant

$$p(\hat{Y} | \eta, \Sigma) = N(\hat{Y} | \eta, \Sigma), \quad (4.2)$$

où $\eta = X\beta + \sum_{\alpha} Z^{(\alpha)} v^{(\alpha)}$ est le prédicteur linéaire. Pour les erreurs e , on peut considérer qu'une loi t de Student au lieu d'une loi normale donne moins de poids à un plus grand nombre d'observations aberrantes, d'après West (1984).

La partie à effets fixes de η peut contenir des composantes comme une ordonnée à l'origine, une tendance linéaire, les effets principaux pour le niveau de la division et du district et peut-être les interactions de second ordre pour les tendances linéaires et la division ou le district. On attribue au vecteur

β d'effets fixes une loi *a priori* normale $p(\beta) = N(0, 100I_p)$, avec I_x la matrice d'identité de dimension $x \times x$. Cela n'est que très peu informatif, car une erreur-type de 10 est très grande par rapport aux échelles des estimations directes (transformées) et des covariables utilisées.

Le deuxième terme du côté droit de (4.1) consiste en une somme des contributions au prédicteur linéaire par des effets aléatoires ou des termes de coefficient variables. On suppose que les vecteurs d'effets aléatoires $v^{(\alpha)}$ pour différents α sont indépendants, mais que les composantes à l'intérieur d'un vecteur $v^{(\alpha)}$ peuvent être corrélées pour tenir compte de la corrélation temporelle ou transversale. Afin de décrire le modèle général de chaque vecteur d'effets aléatoires $v^{(\alpha)}$, nous supprimons l'exposant α dans ce qui suit aux fins de simplification de la notation.

On suppose que chaque vecteur d'effets aléatoires v est distribué ainsi

$$v \sim N(0, A \otimes V), \quad (4.3)$$

où V et A sont les matrices de covariance de dimension $d \times d$ et $l \times l$, respectivement, et $A \otimes V$ désigne le produit de Kronecker de A par V . La longueur totale de v est $q = dl$, et ces coefficients peuvent être considérés comme correspondant aux effets d autorisés à varier sur l niveaux d'une variable de facteur. Si, par exemple, V correspond à la division, alors V définit $d = 7$ différents effets aléatoires qui correspondent aux 7 catégories de division. Si ensuite A correspond au temps, alors $l = 21$ années. Dans ce cas, chacun des 7 effets peut varier sur ses 21 niveaux (des années en l'occurrence). Chaque effet aléatoire généré pour une combinaison division $\times$ année est partagé par tous les districts appartenant à cette division pour cette année particulière.

La matrice de covariance A décrit la structure de covariance entre les niveaux de la variable de facteur et est censée être connue. Au lieu des matrices de covariance, on utilise en fait les matrices de précision $Q_A = A^{-1}$, pour un calcul plus efficace (Rue et Held, 2005). La matrice de covariance V pour les d effets variables peut être paramétrée de l'une des trois façons suivantes : i) une matrice de covariance paramétrée complète; ii) une matrice diagonale avec des éléments diagonaux inégaux; iii) une matrice diagonale avec des éléments diagonaux égaux. Une loi *a priori* de Wishart inversée et mise à l'échelle est utilisée comme ce qui est proposé dans O'Malley et Zaslavsky (2008) et recommandé par Gelman et Hill (2007) quand on suppose une matrice de covariance complète, tandis que les lois *a priori* demi-Cauchy sont utilisées pour les écarts-types quand la matrice de covariance est supposée diagonale avec des éléments égaux ou inégaux. En cas de variances diagonales, les lois *a priori* demi-Cauchy sont de meilleures lois *a priori* par défaut que les lois *a priori* gamma inverses plus courantes (Gelman, 2006).

Les structures d'effet aléatoire suivantes sont examinées dans la procédure de sélection du modèle :

1. Ordonnées à l'origine aléatoires pour les domaines M_d . Dans ce cas, $A = I_{M_d}$ et V est un paramètre de variance scalaire. Cela suppose que $v_{it} = v_i, \forall t$ et $v_i \sim N(0, \sigma_i^2)$.
2. Marches aléatoires de premier ou de second ordre à différents niveaux d'agrégation. Une marche aléatoire de premier ordre ou une tendance locale au niveau du district est définie comme étant $v_{it} = L_{it}$, avec $L_{it} = L_{i,t-1} + \eta_{it}$ et $\eta_{it} \sim N(0, \sigma_{R1,i}^2)$. Une marche aléatoire de second ordre ou un modèle de tendance lissée au niveau du district est définie comme étant $v_{it} = L_{it}$, avec $L_{it} = L_{i,t-1} + R_{i,t-1}$ et $\eta_{it} \sim N(0, \sigma_{R2,i}^2)$. Les deux types de tendances peuvent être définis

similairement au niveau de la division ou à l'échelle nationale. Voir dans Rue et Held (2005) la spécification de la matrice de précision Q_A pour les marches aléatoires de premier et de second ordre. On peut considérer une matrice de covariance complète pour les innovations de tendance afin de permettre des corrélations transversales en plus des corrélations temporelles, ou une matrice diagonale avec des paramètres de variance différents ou égaux pour permettre des corrélations temporelles seulement. Dans le cas de variances égales, $\sigma_{R1,i}^2 = \sigma_{R1}^2$ et $\sigma_{R2,i}^2 = \sigma_{R2}^2, \forall i$. Les composantes des marches aléatoires de premier et de second ordre au niveau du district sont désignées ci-dessous respectivement par RW1_District et RW2_District. Au niveau de la division, elles sont désignées par RW1_Division et RW2_Division.

3. Les marches aléatoires de premier ordre utilisées dans nos modèles ne peuvent pas saisir de niveau global, car les effets aléatoires correspondants sont soumis à une contrainte de somme nulle dans le temps. De même, les marches aléatoires de second ordre ne peuvent pas saisir à la fois le niveau et la tendance linéaire. Cela signifie que le niveau et la tendance linéaire doivent être pris en compte par d'autres termes du modèle, comme les effets fixes ou aléatoires. Les ordonnées à l'origine au niveau du district ont été traitées au point 1. Pour inclure aussi les tendances linéaires par district, cette composante peut être étendue aux ordonnées à l'origine et aux pentes aléatoires linéaires dans le temps. Dans ce cas, V peut être soit une matrice de covariance générale 2×2

$$V = \begin{pmatrix} \sigma_I^2 & \rho_{IS} \sigma_I \sigma_S \\ \rho_{IS} \sigma_I \sigma_S & \sigma_S^2 \end{pmatrix},$$

qui prend en compte les corrélations entre les ordonnées à l'origine et les pentes, ou une matrice diagonale avec des éléments diagonaux σ_I^2 et σ_S^2 les variances de l'ordonnée à l'origine et des pentes aléatoires, respectivement. Cette composante du modèle est désignée par RIS_District ci-dessous.

4. Effets aléatoires spatiaux : ordonnées à l'origine aléatoires qui varient selon l'emplacement spatial des districts selon un modèle autorégressif conditionnel intrinsèque (ICAR) (Besag et Kooperberg, 1995), défini comme étant $v_i | v_{-i} \sim N\left(\left(\sum_{i' \in nb(i)} v_{i'}\right) / a_i, \sigma_{sp}^2 / a_i\right)$ pour chaque effet spatial conditionnel aux autres. Dans le cas présent, $nb(i)$ est l'ensemble des domaines voisins du domaine i et a_i le nombre de domaines voisins du domaine i . Voir dans Rue et Held (2005) la spécification de la matrice de précision Q_A . Cette composante spatiale est désignée ultérieurement par Spatial_District.
5. Bruit blanc : pour permettre une variation aléatoire inexpliquée, on peut inclure le bruit blanc au niveau de domaine le plus détaillé par année. Dans ce cas $A = I_M$ et V est un paramètre de variance scalaire. Cela suppose que $v_{it} \sim N(0, \sigma_w^2)$.

Nous avons aussi étudié des généralisations de (4.3) en répartitions anormales d'effets aléatoires en mettant en œuvre une loi de Student, une répartition *a priori horseshoe* (Carvalho, Polson et Scott, 2010)

et Laplace (Tibshirani, 1996; Park et Casella, 2008). Les queues plus lourdes de ces répartitions de rechange permettent de grands effets occasionnels. Cependant, ces répartitions n'ont pas amélioré les résultats pour les variables cibles considérées en ce qui concerne les critères d'information du modèle ainsi que les prédictions des tendances sous-jacentes. C'est pourquoi la loi normale est utilisée pour toutes les composantes à effet aléatoire. La présentation exacte des modèles de séries chronologiques multiniveaux définitifs pour ANC0 et ANC4 est précisée respectivement dans les sous-sections 5.1 et 5.2.

4.2 Estimation des modèles

Les modèles sont ajustés au moyen d'un échantillonnage par la méthode de Monte Carlo par chaînes de Markov (MCMC), en particulier l'échantillonneur de Gibbs (Geman et Geman, 1984; Gelfand et Smith, 1990). Voir dans Boonstra et van den Brakel (2022) la spécification des lois conditionnelles complètes. Les modèles précisés dans la sous-section 4.1 sont exécutés dans R (R Core Team, 2015) au moyen du progiciel `mcmc_sae` (Boonstra, 2021). L'échantillonneur de Gibbs est exécuté en parallèle pour trois chaînes indépendantes avec des valeurs de départ générées aléatoirement. Dans l'étape de construction du modèle, 1 000 itérations sont utilisées, en plus d'une période de « rodage » de 100 itérations. Cela suffisait pour les estimations raisonnablement stables de Monte Carlo des paramètres du modèle et les prédictions de tendances. Pour le modèle sélectionné, nous utilisons un cycle plus long de 1 000 exécutions de rodage plus 5 000 itérations sur lesquelles les tirages d'une itération sur cinq sont stockés. Cela donne $3 \times 1\,000 = 3\,000$ tirages pour calculer les estimations et les erreurs-types. 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 décèle le mélange de chaînes. Pour la simulation plus longue du modèle sélectionné, tous les paramètres du modèle et les prédictions du modèle ont des facteurs de réduction d'échelle potentiels inférieurs à 1,01 et un nombre suffisamment efficace de tirages indépendants.

De nombreux modèles de la forme (4.1) ont été ajustés aux données. Afin de comparer les modèles au moyen des mêmes données d'entrée, nous utilisons le critère d'information largement applicable aussi appelé critère d'information Watanabe-Akaike (WAIC) (Watanabe, 2010, 2013) et le critère d'information de déviance (DIC) (Spiegelhalter, Best, Carlin et van der Linde, 2002). Nous comparons aussi les modèles graphiquement selon leurs ajustements et prédictions de tendances à trois niveaux d'agrégation.

5. Modèles sélectionnés et prédiction des modèles

5.1 Modèle de séries chronologiques multiniveaux pour ANC0

Aucune transformation n'est examinée pour les séries de données d'entrée des estimations directes ou des estimations de F-H. Les composantes à effet fixe suivantes sont incluses dans les modèles sélectionnés pour MTS-I, MTS-II et MTS-III :

$$1 + \text{Division} + \text{yr.c} + \text{Division} * \text{yr.c}, \quad (5.1)$$

où $yr.c$ désigne la variable d'année quantitative standardisée, qui définit une tendance linéaire à effet fixe. De même, $Division * yr.c$ définit une tendance linéaire à effet fixe pour chaque division distincte. Le tableau 5.1 présente la partie des effets aléatoires des trois modèles. Si de multiples effets variables sont modélisés, il y a donc un choix entre une matrice V scalaire, diagonale ou de covariance complète dans (4.3). Pour les variations au fil du temps, les marches aléatoires de second ordre $RW2_Division$ et $RW2_District$ ont finalement été sélectionnées. Les composantes de bruit blanc sont considérées, mais ne sont pas incluses dans le modèle final, car elles n'amélioreraient pas l'ajustement du modèle.

Tableau 5.1

Résumé des composantes à effet aléatoire pour le modèle de séries chronologiques multiniveau sélectionné pour ANC0. Les deuxième et troisième colonnes font référence aux effets variables avec la matrice de covariance V dans (4.3), tandis que la quatrième colonne fait référence à la variable de facteur associée avec A dans (4.3). La dernière colonne contient le nombre total d'effets aléatoires pour chaque composante

Composante du modèle	Formule V	Structure de la variance	Facteur A	Nbre d'effets
RIS_District	$1 + yr.c$	complète	District	128
RW2_Division	Division	scalaire	RW2(année)	147
RW2_District	District	scalaire	RW2(année)	1 344
District spatial	1	scalaire	Spatial(District)	64

Le prédicteur linéaire du modèle sélectionné peut être écrit, selon les éléments pour le district i et l'année t , comme suit :

$$\eta_{it} = \beta' x_{it} + \nu_i + z_i \nu_i^{(yr)} + u_{it} + u_{j[i]t}^{(div)} + s_i, \quad (5.2)$$

où β est le vecteur des effets fixes correspondant aux covariables x_{it} conformément à la spécification dans (5.1), ν_i désigne des ordonnées à l'origine aléatoires variant selon le district, z_i désigne la $yr.c$ variable pour l'année t , et $\nu_i^{(yr)}$ les pentes aléatoires correspondantes variant selon le district. Ces ordonnées à l'origine et pentes aléatoires sont distribuées conjointement comme suit :

$$\begin{pmatrix} \nu_i \\ \nu_i^{(yr)} \end{pmatrix} \stackrel{iid}{\sim} N \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_I^2 & \rho \sigma_I \sigma_S \\ \rho \sigma_I \sigma_S & \sigma_S^2 \end{pmatrix} \right). \quad (5.3)$$

Les effets de marche aléatoire de second ordre au niveau du district et de la division sont distribués ainsi :

$$\begin{aligned} u_{it} - 2u_{i(t-1)} + u_{i(t-2)} &\stackrel{iid}{\sim} N(0, \sigma_{R2}^2) \\ u_{j[i]t}^{(div)} - 2u_{j[i](t-1)}^{(div)} + u_{j[i](t-2)}^{(div)} &\stackrel{iid}{\sim} N(0, (\sigma_{R2}^{(div)})^2), \end{aligned} \quad (5.4)$$

où $j[i]$ doit être lu comme la division j à laquelle le district i appartient. Enfin, les effets spatiaux s_i sont distribués comme suit :

$$s_i | s_{i' \neq i} \stackrel{\text{ind}}{\sim} N\left(\frac{1}{a_i} \sum_{i' \in nb(i)} s_{i'}, \frac{1}{a_i} \sigma_{sp}^2\right), \quad (5.5)$$

où a_i est la taille de l'ensemble $nb(i)$ des districts voisins du district i . Les distributions *a priori* pour la matrice de covariance de (5.3) et les autres paramètres de variance sont choisis de la façon décrite à la section 4.1. Pour que les composantes du modèle soient identifiables, les contraintes suivantes sont imposées :

$$\begin{aligned} \sum_{t=1}^T u_{it} &= 0 \quad \text{et} \quad \sum_{t=1}^T t u_{it} = 0 \quad \text{pour tous les districts } i, \\ \sum_{t=1}^T u_{j[i]t}^{(div)} &= 0 \quad \text{et} \quad \sum_{t=1}^T t u_{j[i]t}^{(div)} = 0 \quad \text{pour toutes les divisions } j, \\ \sum_{i=1}^{M_d} s_i &= 0. \end{aligned} \quad (5.6)$$

Notons que les tendances de RW2 sont précisées au niveau des divisions et des districts, chaque fois avec une structure de variance scalaire. Toute tendance au niveau de la division est partagée par tous les districts sous-jacents. Les écarts de chaque district par rapport à cette tendance au niveau de la division sont modélisés avec les tendances de la RW2 au niveau du district. Il s'agit d'une solution de rechange parcimonieuse permettant d'emprunter de l'information dans le temps et l'espace, comparativement à ce qu'apporterait la modélisation des tendances de RW2 au niveau du district seulement avec une matrice de covariance complète (Boonstra et van den Brakel, 2019).

5.2 Modèle de séries chronologiques pour ANC4

La transformation par la racine carrée est appliquée aux séries de données d'entrée des estimations directes et de F-H de ANC4 pour les modèles MTS-I, MTS-II et MTS-III. Pour MTS-I, on applique la FGV (3.3) aux erreurs-types transformées afin d'obtenir la matrice de la variance Σ , comme cela est expliqué à la fin de la sous-section 3.5. Pour la composante à effet fixe, une variable de facteur appelée « Région » a été créée en fonction du degré d'urbanisation d'après Rahman, Mohiuddin, Kafy, Sheel et Di (2019). La variable comporte quatre niveaux : 1) pour trois grandes villes, *Dhaka*, *Chittagong* et *Gazipur*; 2) pour neuf autres grandes villes régionales (*Barisal*, *Bogra*, *Comilla*, *Khulna*, *Mymensing*, *Narayanganj*, *Rajshahi*, *Rangpur*, *Sylhet*); 3) pour trois districts vallonnés (*Bandarban*, *Khagrachhari* et *Rangamati*); 4) pour les districts restants. Cette variable a surtout aidé dans l'ajustement des estimations pour les trois districts vallonnés sur lesquels les sept enquêtes examinées comportent très peu (voire pas) d'information. Le modèle final comprend les composantes à effets fixes suivantes :

$$1 + \text{Division} + \text{yr.c} + \text{Région}. \quad (5.7)$$

On constate que l'interaction entre « Division » et « yr.c » (comme dans le modèle ANC0) est négligeable dans le modèle ANC4. Les composantes à effet aléatoire du modèle ANC4 présentées au tableau 5.2 sont

très semblables à celles utilisées pour le modèle ANC0 (voir le tableau 5.1). On a considéré une tendance au niveau local plutôt qu'une tendance lisse au niveau de la division (RW1_Division dans le tableau 5.2) puisque la composante de tendance lisse (RW2_Division, comme dans le tableau 5.1) produisait un certain biais dans les tendances à l'échelle du pays et de la division. De plus, le modèle avec la composante RW1_Division donne de meilleurs résultats pour les critères d'information que le modèle avec la composante RW2_Division. Les composantes de bruit blanc sont considérées, mais ne sont pas incluses dans le modèle final, car elles n'amélioreraient pas l'ajustement du modèle.

Tableau 5.2

Résumé des composantes à effet aléatoire pour le modèle de séries chronologiques multiniveau sélectionné pour ANC4. Les deuxième et troisième colonnes font référence aux effets variables avec la matrice de covariance V dans (4.3), tandis que la quatrième colonne fait référence à la variable de facteur associée avec A dans (4.3). La dernière colonne contient le nombre total d'effets aléatoires pour chaque terme

Composante du modèle	Formule V	Structure de la variance	Facteur A	Nbre d'effets
RIS_District	$1 + yr.c$	complète	District	128
RW1_Division	Division	scalaire	RW1(année)	147
RW2_District	District	scalaire	RW2(année)	1 344
District spatial	1	scalaire	Spatial(District)	64

Par ailleurs, le modèle peut être exprimé comme dans (5.2), où désormais β et x_{it} correspondent à la spécification des effets fixes (5.7). La seule autre différence est que les tendances au niveau de la division sont maintenant modélisées comme une marche aléatoire de premier ordre :

$$u_{j[i]t} - u_{j[i](t-1)} \stackrel{iid}{\sim} N(0, (\sigma_{R1}^{(div)})^2), \quad (5.8)$$

où pour des raisons d'identifiabilité, la contrainte $\sum_{t=1}^T u_{jt}^{(div)} = 0$ est imposée pour toute division j . Comme dans le cas de ANC0, les tendances de RW1 sont précisées au niveau des divisions et les tendances de RW2 au niveau des districts, les deux avec une structure de variance scalaire comme moyen parcimonieux d'emprunter de l'information dans le temps et l'espace.

5.3 Estimation des tendances

Les estimations des tendances sont calculées à partir des résultats de la simulation MCMC. Dans une première étape, pour chaque réplique de MCMC, un vecteur de dimension M contenant des prédictions au niveau le plus détaillé de toutes les combinaisons année-district est calculé comme suit :

$$\eta^{(r)} = X\beta^{(r)} + \sum_{\alpha} Z^{(\alpha)} v^{(\alpha,r)}, \quad (5.9)$$

où l'exposant (r) indexe les tirages MCMC retenus. Notons que $\eta^{(r)}$ comprend aussi les prédictions pour les années sans observations d'enquête. Étant donné qu'une transformation par la racine carrée a été

appliquée à aux séries de ANC4, la rétrotransformation suivante pour les vecteurs $\eta^{(r)}$ a été examinée, d'après Boonstra et coll. (2021) :

$$\theta^{(r)} = (\eta^{(r)})^2 + \left(\text{se}(\hat{\mathcal{Y}}_t)\right)^2. \quad (5.10)$$

Le deuxième terme du côté droit est une correction de biais (relativement petite) utilisant les erreurs-types transformées et lissées. La correction du biais découle du fait que l'espérance par rapport au plan de sondage des estimations directes peut être écrite comme suit :

$$E(\hat{Y}) = E((\hat{\mathcal{Y}})^2) = E((\eta + \mathcal{E})^2) = \eta^2 + \text{var}(\mathcal{E}),$$

où $\mathcal{E}$ est le vecteur des erreurs d'échantillonnage après transformation, censé être normalement distribué avec des erreurs-types $\text{se}(\hat{\mathcal{Y}}_t)$. Le problème posé par les données dont nous disposons est que la correction du biais ne peut être appliquée qu'aux années d'enquête, étant donné que les erreurs-types ne sont disponibles que pour ces années. L'application de la correction du biais uniquement pour les années d'enquête fausse les estimations des tendances, comme le montrent Das, van den Brakel, Boonstra et Haslett (2021). Dans le cas du modèle MTS-I, l'effet de cette correction de biais est plus évident pour les domaines pour lesquels il n'y a pas d'estimations directes, en particulier pour les districts vallonnés de *Chittagong*. L'effet de la correction du biais est moindre dans le cas des modèles MTS-II et MTS-III, puisque les erreurs-types estimées par les estimations de F-H sont déjà suffisamment lissées et convergentes. Toutefois, à l'échelle nationale et au niveau de la division, cette correction de biais entraîne une certaine surestimation dans certaines années d'enquête pour toutes les tendances qui reposent sur des modèles MTS. Par conséquent, la correction du biais pour la transformation par la racine carrée n'est pas appliquée dans les estimations de tendances, mais seulement utilisée dans le calcul des estimations transversales de F-H.

On obtient les estimations des tendances avec leurs erreurs-types au niveau le plus détaillé des districts pour toutes les années en prenant la moyenne et l'écart-type sur les répliques MCMC $\eta^{(r)}$, respectivement. Pour obtenir les tendances au niveau de la division et du pays, on agrège chaque réplique MCMC à partir de l'échelle régionale la plus détaillée des districts, en utilisant le nombre de femmes qui ont été mariées comme variable de pondération. Ensuite, on obtient les estimations de tendance et leurs erreurs-types en prenant la moyenne et l'écart-type sur ces répliques MCMC agrégées.

6. Résultats

Les tendances de ANC0 et de ANC4 présentées dans les figures comprennent cinq types d'estimations dont les intervalles de confiance sont d'environ 95 % : i) les estimations directes pondérées pour l'année d'enquête (ligne de barre d'erreur noire); ii) les estimations transversales de F-H pour l'année d'enquête (ligne de barre d'erreur verte); iii) les estimations qui reposent sur le modèle MTS-I (ligne rouge); iv) les estimations qui reposent sur le modèle MTS-II (ligne verte); v) les estimations qui reposent sur le modèle MTS-III (ligne bleue).

6.1 ANC0

Les tendances nationales de ANC0 sont présentées à la figure 6.1. La figure montre que les estimations directes et de F-H transversales sont très semblables pour les années d'enquête, selon un intervalle de confiance (IC) d'environ 95 %. On peut s'y attendre pour les chiffres à l'échelle nationale, puisque le gain de précision obtenu avec un modèle de prédiction sur petits domaines pour ce qui est d'un estimateur direct diminue à mesure que la taille de l'échantillon augmente. Pendant la première période allant de 1994 à 2000, la tendance nationale qui repose sur le modèle MTS-I suit les estimations directes et de F-H transversales, tandis que les tendances qui reposent sur les modèles MTS-II et MTS-III sont légèrement plus élevées. Pour la période allant de 2004 à 2010, la tendance qui repose sur le modèle MTS-I est légèrement supérieure aux tendances qui reposent sur les modèles MTS-II et MTS-III. Les différences sont cependant très minimes.

Les tendances au niveau de la division, illustrées à la figure 6.2, indiquent que les tendances dans MTS-I sont très semblables à celles qui reposent sur les modèles MTS-II et MTS-III, à quelques petites exceptions près dans les divisions de *Dhaka*, de *Khulna* et de *Rajshahi*. Les différences entre les divisions de *Dhaka* et de *Khulna* peuvent être à l'origine de la plupart des différences observées dans les tendances à l'échelle nationale.

Les tendances qui reposent sur les modèles MTS-II et MTS-III sont presque identiques au niveau du pays et de la division. Cela est appuyé par les composantes de la variance estimées de la composante aléatoire à tendance lisse au niveau de la division dans les deux modèles élaborés ($\hat{\sigma}_{R2}^{(div)}$: environ 0,020) présentés dans le tableau 6.1. En revanche, on constate des différences plus importantes dans les tendances des modèles MTS-II et MTS-III au niveau des districts, comme le montrent les figures 6.3 et 6.4. Voir Das, van den Brakel, Boonstra et Haslett (2021) qui présentent des graphiques pour tous les districts. Les tendances qui reposent sur le modèle MTS-III sont plus lisses que celles qui reposent sur le modèle MTS-II, ce qui est le résultat des plus petites valeurs de la composante de la variance estimée $\hat{\sigma}_{R2}$ dans MTS-III (voir le tableau 6.1).

Tableau 6.1

Moyennes *a posteriori* des paramètres d'écart-type des composantes aléatoires des modèles MTS-I, MTS-II, MTS-III pour ANC0. L'absence d'exposant fait référence au niveau du district, les exposants (*div*) font référence au niveau de la division

Modèle	$\hat{\sigma}_t$ (ET)	$\hat{\sigma}_s$ (ET)	$\hat{\rho}_{IS}$ (ET)	$\hat{\sigma}_{sp}$ (ET)	$\hat{\sigma}_{R2}^{(div)}$ (ET)	$\hat{\sigma}_{R2}$ (ET)
MTS-I	0,083 (0,013)	0,054 (0,007)	0,168 (0,171)	0,068 (0,032)	0,019 (0,003)	0,024 (0,002)
MTS-II	0,069 (0,012)	0,033 (0,004)	0,254 (0,180)	0,071 (0,028)	0,020 (0,003)	0,013 (0,002)
MTS-III	0,062 (0,013)	0,027 (0,013)	0,227 (0,201)	0,067 (0,030)	0,020 (0,003)	0,009 (0,001)

Les tendances au niveau du district ont tendance à suivre la tendance du niveau de leur division respective illustrée à la figure 6.2. C'est particulièrement le cas pour les domaines ayant un nombre relativement faible d'observations, comme les districts de *Bandarban*, de *Khagrachhari* et de *Rangamati* à la figure 6.3, qui appartiennent à la division de *Chittagong* dans la figure 6.2. Pour réduire cette tendance, on a élaboré un modèle de séries chronologiques multiniveau en supprimant la composante de tendance lisse RW2_Division au niveau de la division dans le tableau 5.1. Cela a toutefois entraîné des tendances

irréalistes très lisses au niveau du pays et de la division. De même, pour examiner la nécessité d'une composante spatiale, on a élaboré des modèles de séries chronologiques multiniveaux en tenant compte et en ne tenant pas compte de la composante spatiale (*Spatial_District* dans le tableau 5.1). On constate que la composante spatiale rend les estimations plus plausibles pour les districts dont l'échantillon est petit ou nul. Observons par exemple les tendances des districts de *Bandarban* et de *Rangamati* de la division de *Chittagong*.

Le modèle MTS-I montre des tendances à la hausse pour certains districts pendant la période de 1994 à 2000. Ces évolutions ne sont pas probables d'un point de vue spécialisé et sont bien corrigées par les modèles MTS-II et MTS-III qui utilisent des estimations de F-H comme séries de données d'entrée. Voir par exemple les districts de *Noakhali*, de *Bandarban*, de *Rangamati*, de *Narayanganj*, de *Rajbari* et de *Narail* dans la figure 6.3. Certains districts présentent des tendances volatiles selon les estimations directes et le modèle MTS-I pendant toute la période, principalement en raison de la variation de la taille de l'échantillon. Voir par exemple les districts de *Bandarban*, de *Bhola*, de *Khagrachhari*, de *Kishoreganj* et de *Rangamati* à la figure 6.3, et de *Chapai Nababganj*, de *Feni*, de *Jhalokati*, de *Joypurhat* et de *Pabna* à la figure 6.4. D'un point de vue spécialisé, on attend une tendance lisse à la baisse pour la couverture de ANC0. Plus particulièrement, on ne s'attend pas aux points de retournement visibles dans plusieurs districts en 2007 et 2011. Les tendances qui reposent sur les modèles MTS-II et MTS-III ne tiennent pas compte de la plupart de ces volatilités et montrent des tendances lisses raisonnables pour ces districts; elles sont donc plus réalistes que celles de MTS-I. Néanmoins, les ajustements des trois modèles sont compatibles avec les données observées. MTS-II semble être un bon compromis entre les modèles I et III.

Figure 6.1 Tendances nationales d'ANC0 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); et v) MTS-III (ligne bleue).

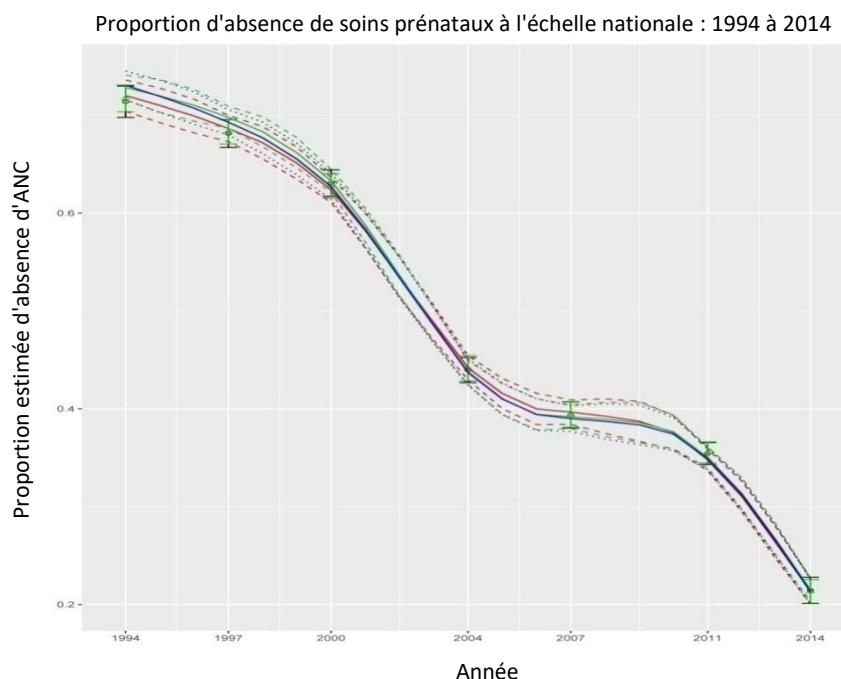


Figure 6.2 Tendances au niveau des divisions d'ANC0 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).

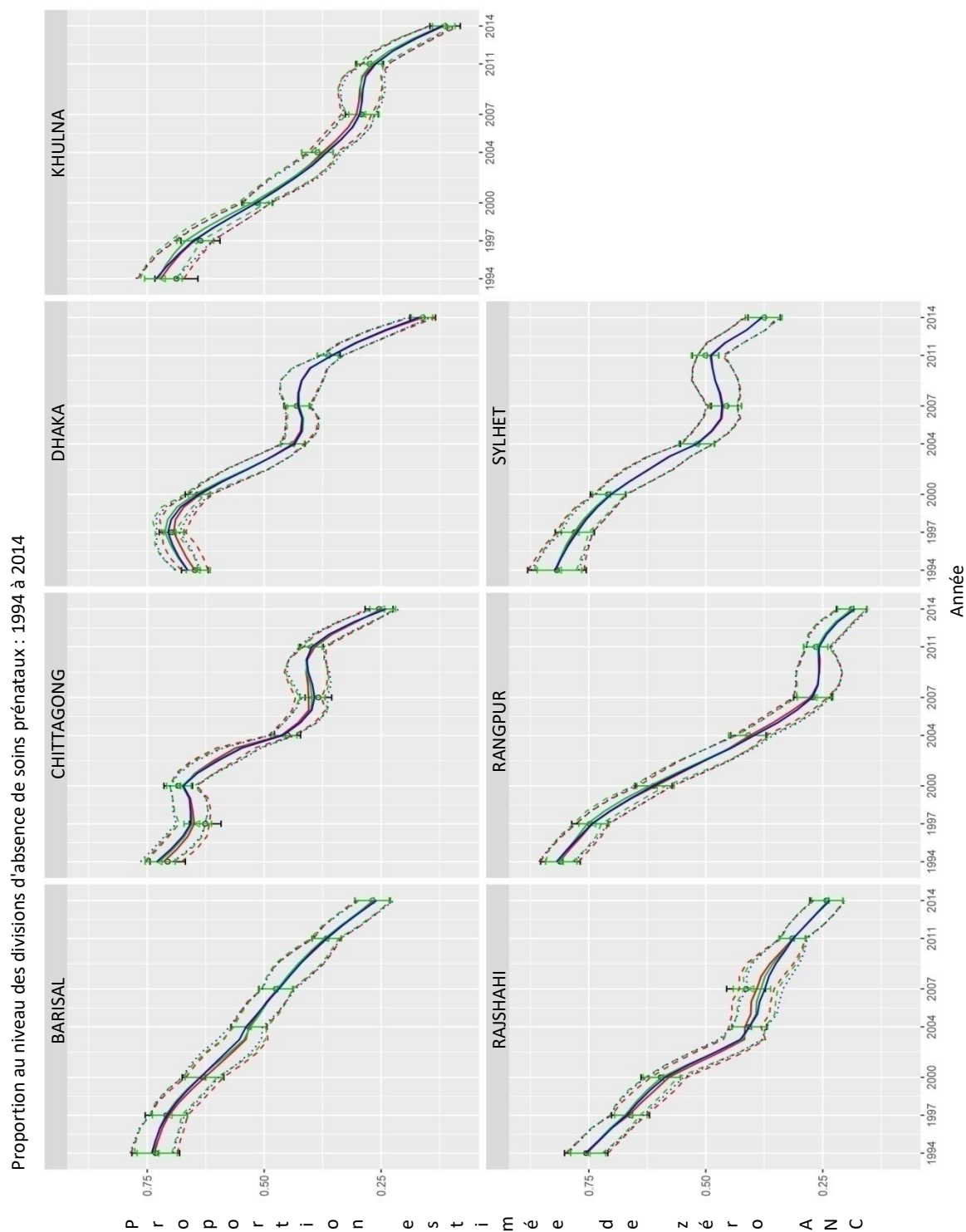


Figure 6.3 Tendances au niveau des districts d'ANC0 au Bangladesh : i) DIR (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).

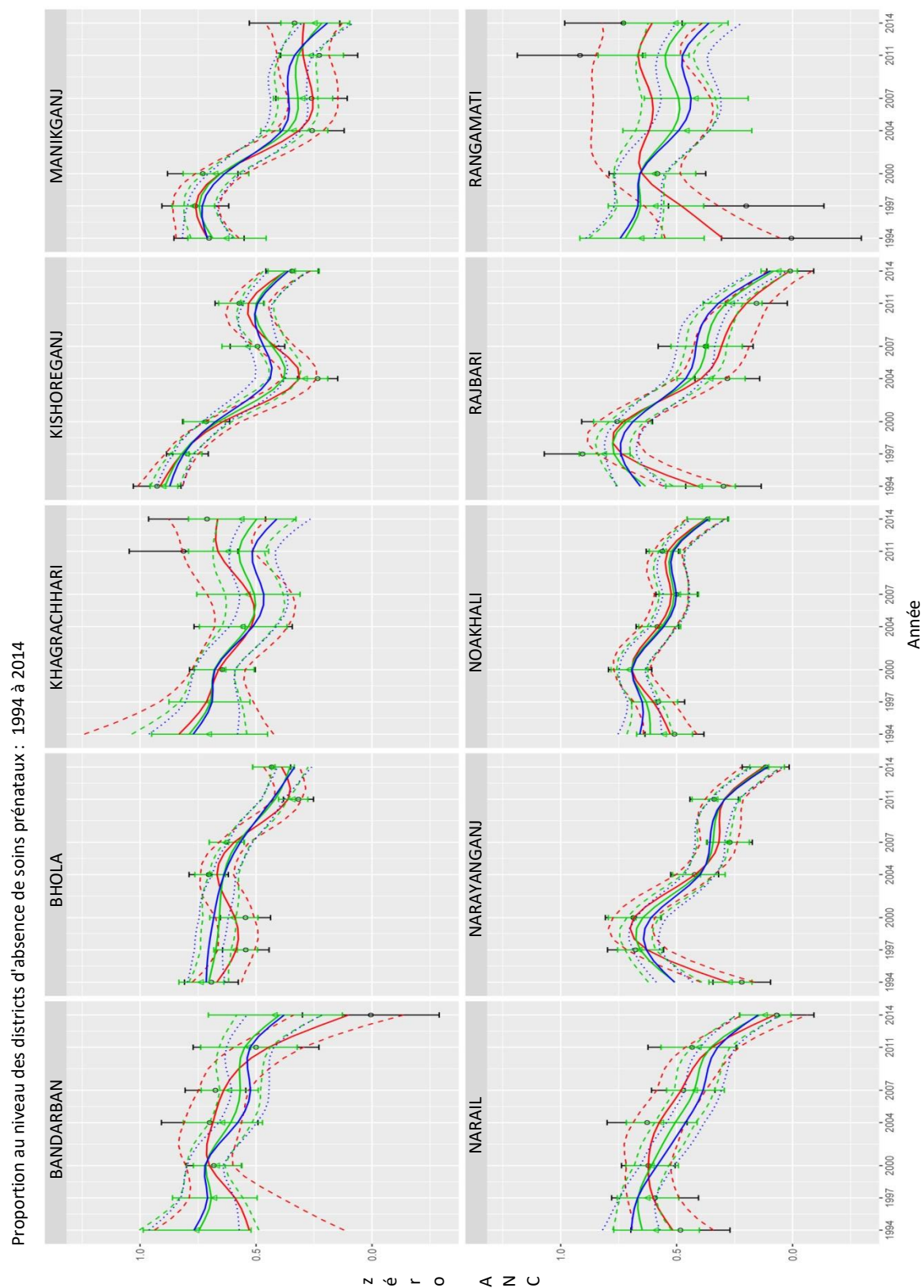
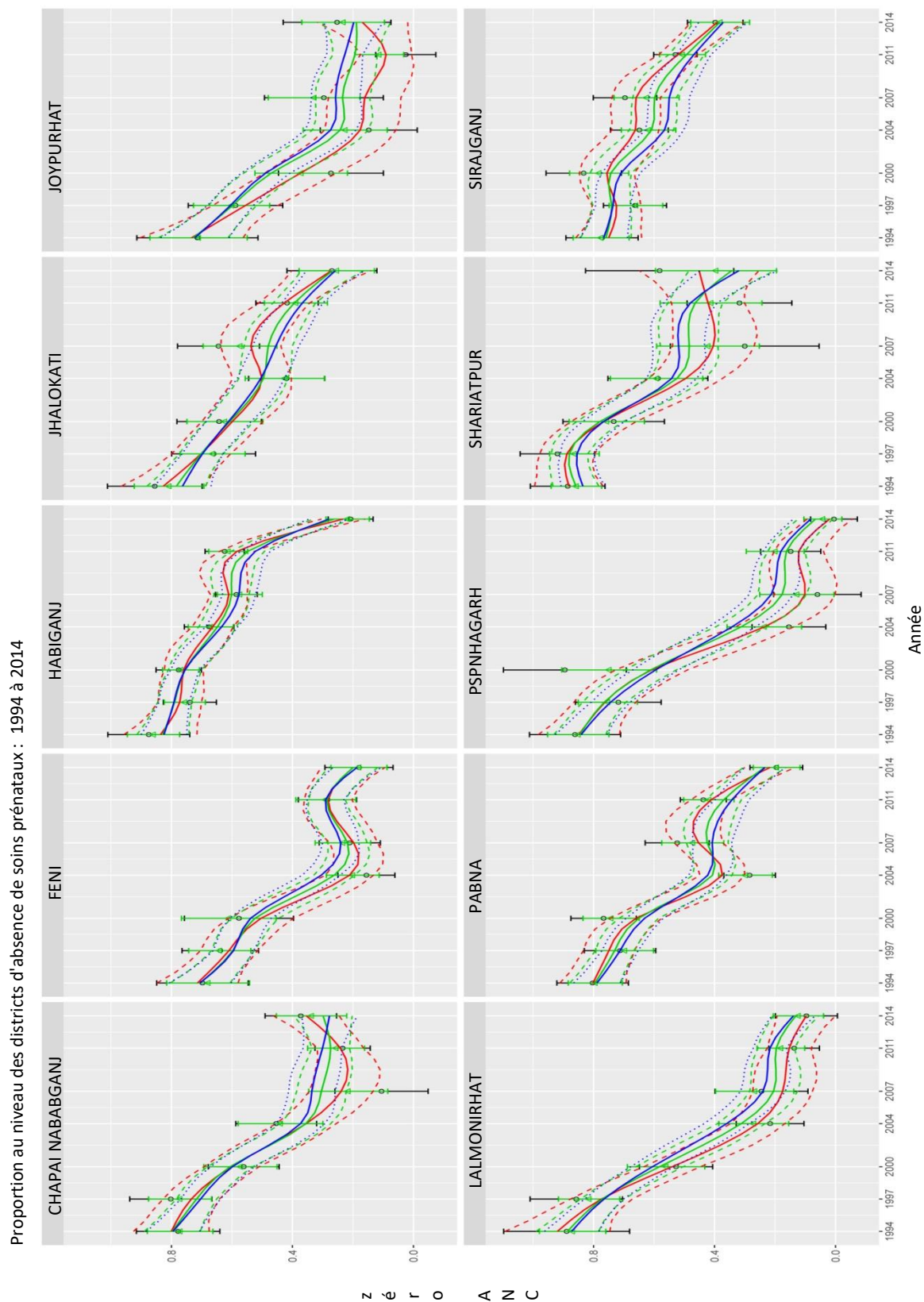


Figure 6.4 Tendances au niveau des districts d'ANC0 au Bangladesh : i) DIR (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).



Dans la plupart des cas, les modèles MTS-II et MTS-III se comportent de la même manière. Toutefois, le modèle MTS-III, qui prend en compte la corrélation dans les estimations de F-H transversales, surestime ANC0 pour certains districts (comme ceux de *Chapai Nababganj*, *Lalmonirhat* et *Shariatpur* dans la figure 6.4) et, de plus, sous-estime légèrement la tendance dans certains districts (comme ceux de *Khagrachari*, *Rangamati* et *Shirajganj* dans la figure 6.3) comparativement aux estimations de F-H transversales. Encore une fois, MTS-II cherche un compromis entre les tendances lisses de MTS-III et les tendances plus volatiles de MTS-I dans la plupart des districts et semble être le modèle privilégié aux fins d'estimation des tendances d'ANC0.

6.2 ANC4

La tendance nationale d'ANC4 illustrée à la figure 6.5 montre une augmentation linéaire de 6 % en 1994, à environ 31 % en 2014. Comme pour ANC0, les estimations directes et de F-H transversales d'ANC4 sont très semblables pour les années d'enquête, avec un IC d'environ 95 %. Les tendances estimées à partir de MTS-I (ligne rouge), MTS-II (ligne verte) et MTS-III (ligne bleue) présentent des tendances très similaires. Comparativement aux estimations directes et de F-H transversales, la tendance de MTS-I est légèrement inférieure en 2007 et 2014. Dans MTS-II et MTS-III, les tendances de l'année d'enquête 2011 sont un peu plus grandes que les estimations directes et de F-H transversales. Les tendances au niveau de la division sont présentées dans la figure 6.6. Les trois modèles de MTS donnent des estimations de tendance très similaires. Certaines différences sont constatées dans les divisions de *Chittagong*, de *Dhaka* et de *Rangpur*. MTS-I donne en effet une tendance légèrement plus élevée que les estimations directes et de F-H pour la division de *Rangpur* dans la période allant de 1994 à 2000. Pour MTS-II et MTS-III, la tendance est un peu plus élevée dans la division de *Rajshahi* dans la période allant de 2011 à 2014 comparativement aux estimations directes et de F-H. Les trois modèles de MTS montrent des bandes d'IC à 95 % légèrement en forme d'arc entre deux années d'enquête consécutives, ce qui indique une incertitude légèrement plus élevée pendant les années sans enquête comparativement aux années d'enquête.

Figure 6.5 Tendances nationales d'ANC4 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).

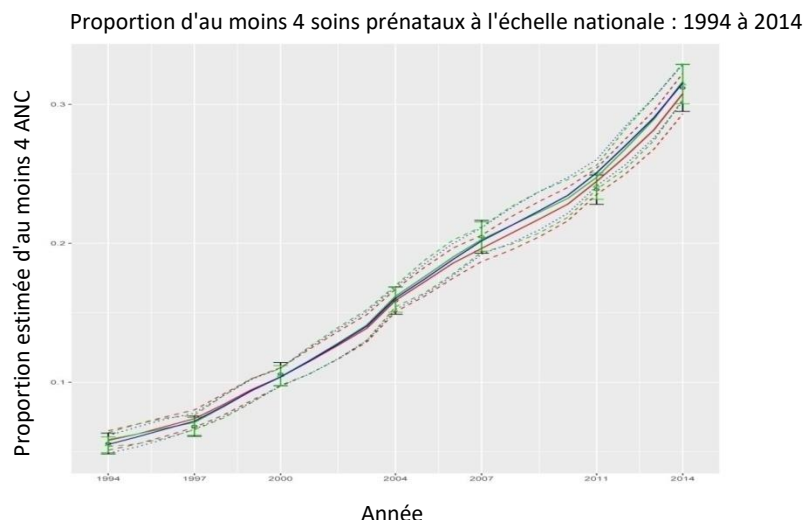
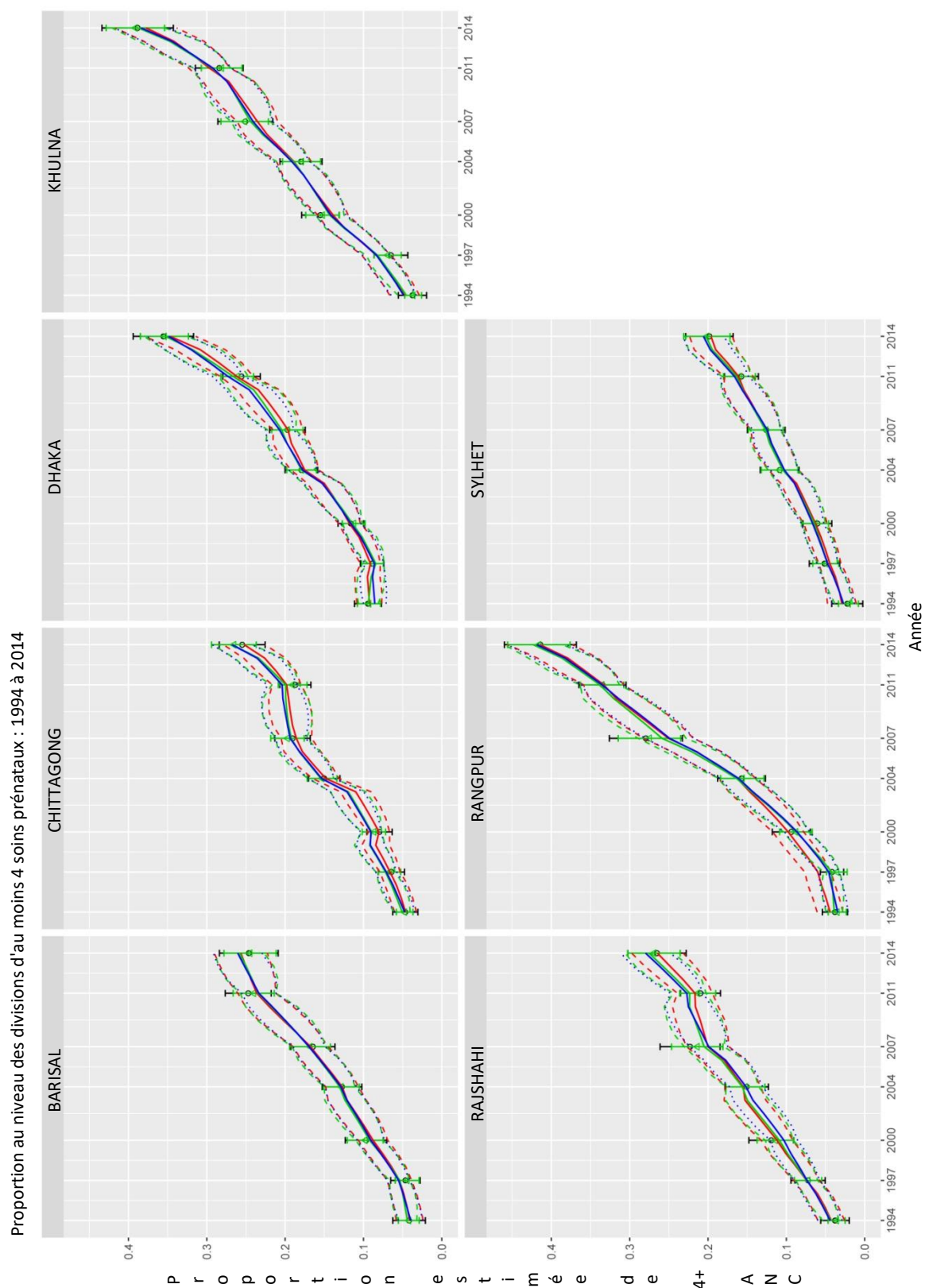


Figure 6.6 Tendances au niveau des divisions d'ANC4 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).



Bien que les tendances qui reposent sur MTS-II et MTS-III soient presque identiques au niveau du pays et des divisions, les composantes de la variance estimée dans les deux modèles diffèrent considérablement comme le montre le tableau 6.2 ci-dessous. Ces différences entraînent des différences importantes dans les estimations des tendances au niveau des districts dans MTS-II et MTS-III. Les figures 6.7 et 6.8 présentent des graphiques pour certains districts. Das, van den Brakel, Boonstra et Haslett (2021) présentent des graphiques pour tous les districts. Comme pour ANC0, les tendances d'ANC4 dans MTS-III sont plus lisses que celles dans MTS-II. Les composantes de variance les plus petites de MTS-III entraînent aussi des bandes de confiance plus étroites que MTS-II.

Tableau 6.2

Moyennes *a posteriori* des paramètres d'écart-type des composantes aléatoires des modèles MTS-I, MTS-II, MTS-III pour ANC4. L'absence d'exposant fait référence au niveau du district, les exposants (*div*) font référence au niveau de la division

Modèle	$\hat{\sigma}_t$ (ET)	$\hat{\sigma}_s$ (ET)	$\hat{\rho}_{ts}$ (ET)	$\hat{\sigma}_{sp}$ (ET)	$\hat{\sigma}_{R1}^{(div)}$ (ET)	$\hat{\sigma}_{R2}$ (ET)
MTS-I	0,060 (0,010)	0,033 (0,005)	0,428 (0,178)	0,047 (0,026)	0,012 (0,004)	0,009 (0,001)
MTS-II	0,046 (0,007)	0,022 (0,003)	0,501 (0,162)	0,035 (0,018)	0,016 (0,003)	0,004 (0,001)
MTS-III	0,038 (0,006)	0,018 (0,006)	0,522 (0,165)	0,027 (0,016)	0,014 (0,003)	0,002 (0,001)

Les estimations de tendance dans MTS-I sont volatiles et montrent des tendances à la baisse inattendues pour certains districts, voir par exemple les districts de *Bhola* et de *Pirojpur* de la division de Barisal, *Gazipur*, *Kishoreganj* et *Manikganj* de la division de *Dhaka*, *Bogra*, *Chapai Nababganj* et *Rajshahi* de la division de *Rajshahi*, et le district de *Habiganj* de la division de *Sylhet* dans la figure 6.7. D'un point de vue spécialisé, on n'attend pas de mouvements brusques ni de points de retournement dans la couverture d'ANC4. Il semblerait donc que MTS-I suive trop rigoureusement les estimations directes. Les tendances de MTS-II ignorent généralement ces volatilités et présentent des tendances raisonnablement stables pour ces districts. Dans MTS-III, les tendances sont encore plus lisses pour certains de ces districts, notamment les districts de *Bogra* et de *Habiganj* dans la figure 6.7, et les districts de *Mymensingh* et de *Sylhet* dans la figure 6.8.

La principale difficulté se pose pour les trois districts vallonnés de la division de *Chittagong*, à savoir *Khagrachhari*, *Rangamati* et *Lakshmipur* (les deux premiers districts sont représentés graphiquement dans la figure 6.8). MTS-I donne des estimations de tendance très médiocres pour ANC4 sur toute la période, principalement en raison d'estimations directes erratiques, qui sont soit nulles soit très incohérentes dans la plupart des enquêtes. Les estimations transversales de F-H sont plus robustes et, par conséquent, MTS-II et MTS-III présentent des tendances à la hausse raisonnables pour ANC4. On s'attend à ce que les femmes habitant dans des régions urbanisées et ayant de meilleures conditions socioéconomiques reçoivent plus de visites d'ANC que celles qui habitent dans des régions rurales aux conditions socioéconomiques défavorables. MTS-I donne, sur toute la période, des estimations inférieures dans certains districts et des estimations de tendance plus élevées que ce qui est attendu dans d'autres districts. Par exemple, la tendance obtenue avec MTS-I pour *Narsingdi*, district fortement urbanisé de la division de *Dhaka*, et présentée dans la figure 6.8 est plus basse que prévu. Ou encore, la tendance obtenue avec MTS-I pour *Munshiganj*, un district de *Dhaka* moins urbanisé, et présentée dans la figure 6.8 est plus élevée que ce qui est attendu. En outre, les estimations de tendance dans MTS-I sont plus élevées que ce

qui est attendu pour toute la période dans le district de *Meherpur* de la division de *Khulna* et les districts de *Lalmonirhat* et de *Panchagarh* de la division de *Rangpur*. Les estimations de tendances dans MTS-II et MTS-III semblent plus plausibles parce que les estimations transversales de F-H semblent plus réalistes que les estimations directes. Dans l'ensemble, comme pour ANC0, MTS-II est un bon compromis entre MTS-I et MTS-III.

Figure 6.7 Tendances au niveau des districts d'ANC4 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).

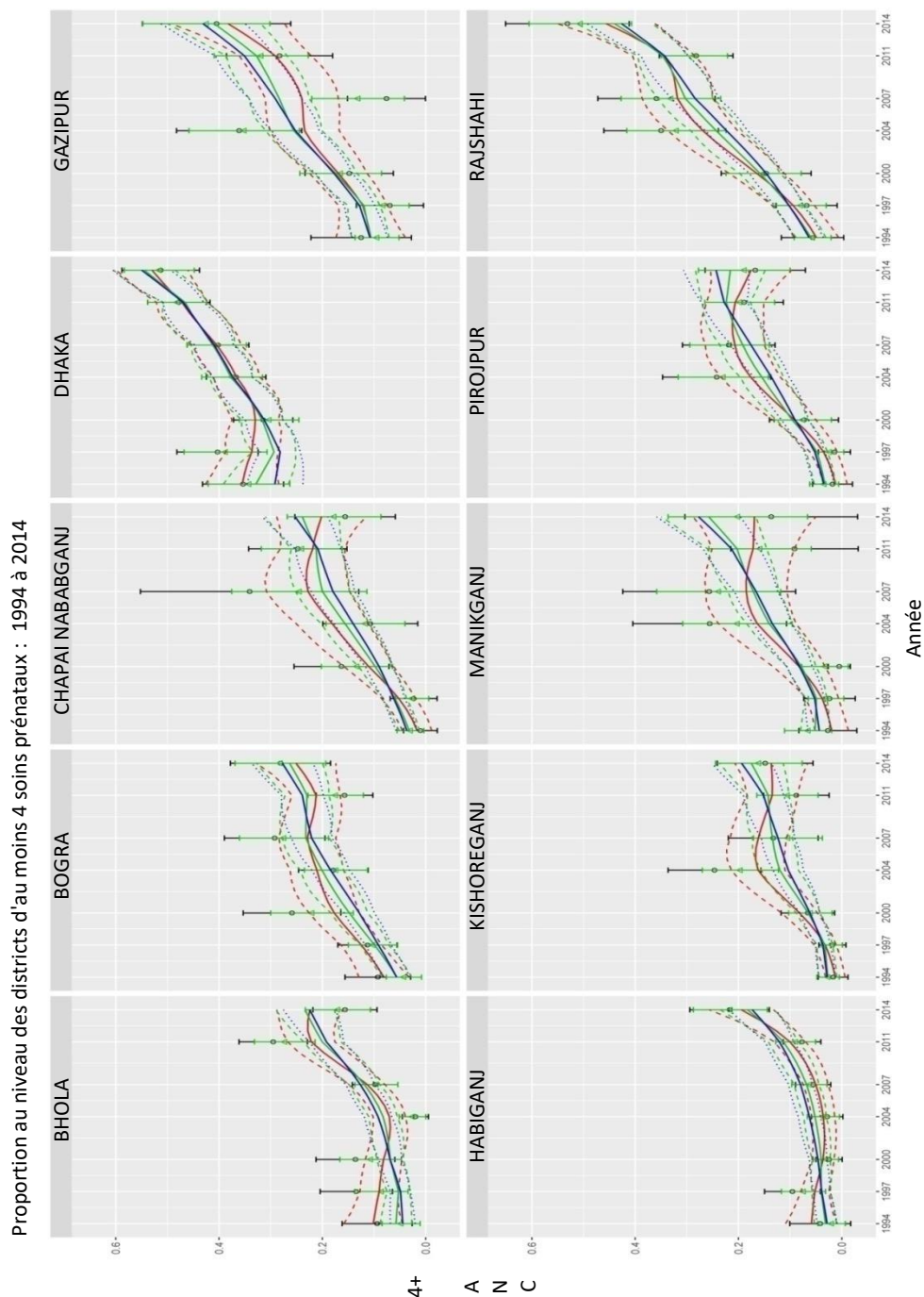
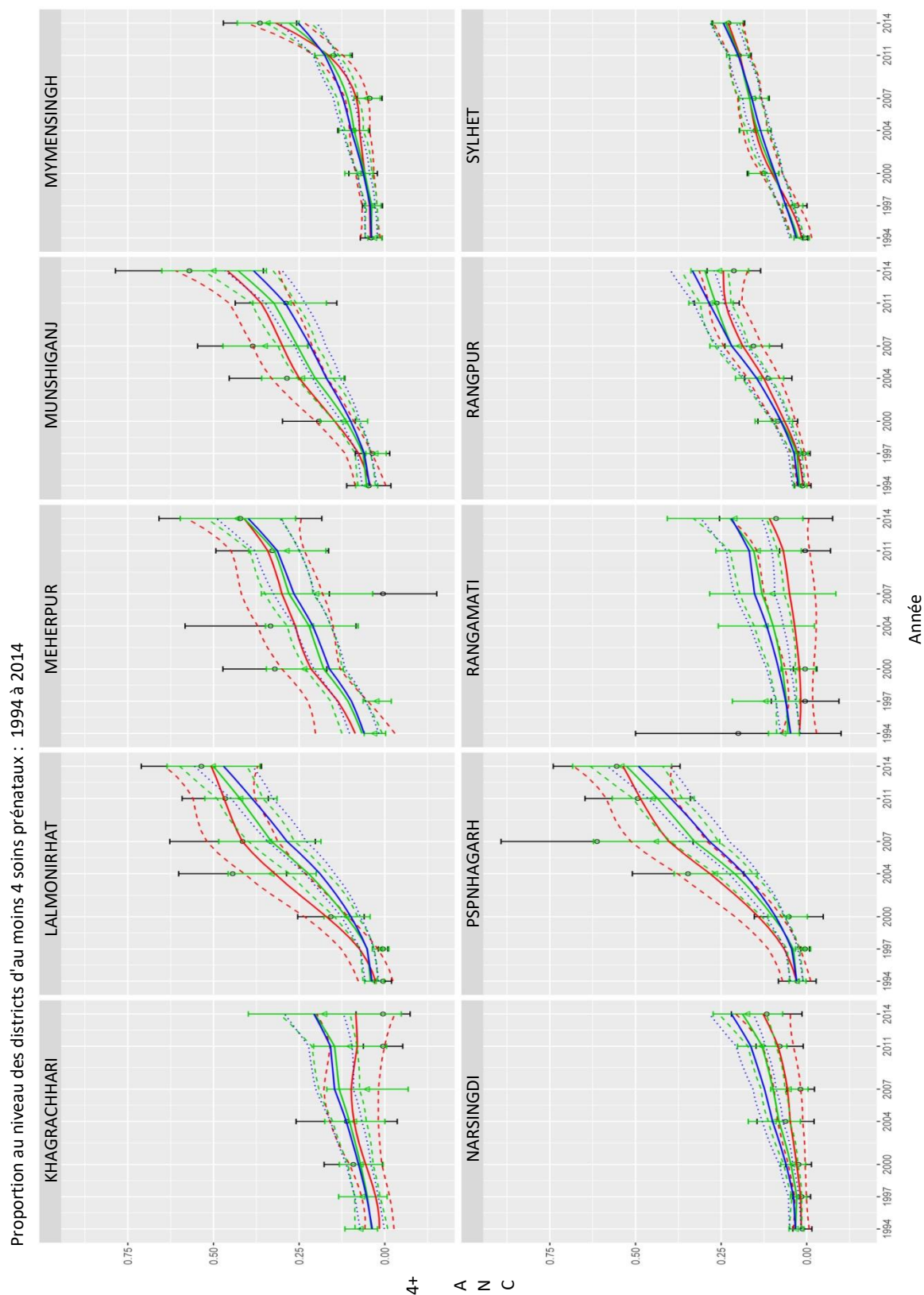


Figure 6.8 Tendances au niveau des districts d'ANC4 au Bangladesh : i) Estimations directes (ligne de barre d'erreur noire); ii) F-H transversale (ligne de barre d'erreur verte); iii) MTS-I (ligne rouge); iv) MTS-II (ligne verte); v) MTS-III (ligne bleue).



7. Évaluation des modèles

Dans la présente étude, les modèles ont été sélectionnés à partir du critère d'information Watanabe-Akaike (WAIC), du critère d'information de défiance (DIC) et de comparaisons graphiques de leurs prédictions de tendances à trois niveaux hiérarchiques. En plus de ces diagnostics de modèles, trois mesures de divergence sont définies pour évaluer et comparer les modèles de séries chronologiques multiniveaux. Les deux premières mesures sont le biais relatif (BR) et le biais relatif absolu (BRA), qui expriment les différences entre les estimations du modèle et les estimations directes, en pourcentage de ces dernières. Pour un modèle donné, BR_{it} et BRA_{it} pour le domaine i et l'année (d'enquête) t sont définis comme suit :

$$BR_{it} = \frac{(\hat{\theta}_{it} - \hat{Y}_{it})}{\hat{Y}_{it}} \times 100 \%, \quad (7.1)$$

$$BRA_{it} = \frac{|\hat{\theta}_{it} - \hat{Y}_{it}|}{\hat{Y}_{it}} \times 100 \% \quad (7.2)$$

$\hat{\theta}_{it}$ étant la prédiction du modèle et $\hat{Y}_{it}$ l'estimation directe. La troisième mesure de divergence est la réduction relative des erreurs-types (RRET), qui mesure le pourcentage de réduction de l'erreur-type des estimations fondées sur un modèle par rapport aux estimations directes, à savoir

$$RRET_{it} = 100 \% \times (se(\hat{Y}_{it}) - se(\hat{\theta}_{it})) / se(\hat{Y}_{it}). \quad (7.3)$$

La mesure de la RRET ne doit pas être interprétée de façon trop stricte, étant donné que les erreurs-types fondées sur le plan et sur le modèle sont des quantités conceptuellement différentes. Il reste que toutes deux sont couramment utilisées comme mesures de l'incertitude, et une fois que des modèles raisonnables qui tiennent suffisamment compte des variations à tous les niveaux d'intérêt ont été sélectionnés, à partir d'autres critères, l'utilisation de la RRET comme une des mesures de comparaison est intéressante.

Ces trois mesures de divergence sont calculées au niveau du pays, de la division et du district (ce dernier étant le plus détaillé). Les distributions de ces mesures sont présentées selon la valeur minimale, le 1^{er} quartile (Q_1), la médiane, la moyenne, le 3^e quartile (Q_3) et la valeur maximale.

De plus, le taux de couverture observé (TC exprimé en %) pour l'intervalle de confiance de 95 % des modèles de MTS et de F-H transversaux considérés est calculé au niveau des divisions et des districts en déterminant si l'IC de 95 % estimé de $\hat{\theta}_{it}$ contient les estimations directes ($\hat{Y}_{it}$). La couverture au niveau du district est le pourcentage de combinaisons de district par année (environ 7×64 domaines) où l'estimation directe est incluse dans l'IC de $\hat{\theta}_{it}$. La couverture au niveau de la division est le pourcentage de combinaisons de division par année (environ 7×7 domaines) où l'estimation directe est incluse dans l'IC de $\hat{\theta}_{it}$. On définit les taux de couverture de la même façon pour chaque année d'enquête en établissant la moyenne sur tous les districts disponibles pour une année d'enquête donnée. Enfin, on calcule la couverture séparément pour chaque division en établissant la moyenne sur les sept années d'enquête.

Les distributions de BR_{it} (7.1), BRA_{it} (7.2) et $RRET_{it}$ (7.3) pour trois niveaux administratifs sont fournies dans les tableaux 7.1, 7.2 et 7.3 pour ce qui est d'ANC0 et ANC4 pour les modèles de F-H transversaux, MTS-I, MTS-II et MTS-III. Le tableau 7.1 montre que les modèles de F-H et de MTS-I fournissent un BR moyen inférieur pour ANC0 et ANC4 aux trois niveaux, tandis que MTS-II donne un BR moyen légèrement inférieur par rapport au modèle de MTS-III au niveau du district. Les distributions du BRA du tableau 7.2 montrent que les performances de MTS-II se situent entre celles de MTS-I et de MTS-III pour tous les niveaux administratifs, sauf à l'échelle nationale pour ANC4. Les valeurs du BRA les plus faibles s'observent dans le modèle de F-H transversal. On constate aussi que le BRA augmente à mesure que la taille de l'échantillon du domaine diminue. Le tableau 7.3 montre que MTS-II a les valeurs de RRET les plus élevées au niveau du pays et de la division, tandis qu'au niveau du district, ce modèle donne une RRET légèrement inférieure au modèle de MTS-III pour ANC0 et ANC4. La réduction de la variance augmente quand la taille des échantillons de domaine diminue. Les erreurs-types pour les tendances au niveau du pays et de la division sont plus petites dans MTS-II que dans MTS-III parce que dans MTS-II, les covariances entre les prédictions transversales de F-H au niveau du district dans les séries de données d'entrée sont ignorées. Ces covariances sont majoritairement positives et, par conséquent, les erreurs-types des tendances aux niveaux agrégés sont plus élevées et plus réalistes dans MTS-III. Les valeurs plus élevées de BR, de BRA et de RRET pour les modèles MTS-II et MTS-III sont une conséquence des tendances plus lisses obtenues au moyen des deux modèles. Les petites variances en présence de tendances lisses signifient un plus grand nombre de biais comparativement aux estimations directes. Comme nous l'avons vu à la section 6, ces tendances sont plus plausibles comparativement au modèle transversal de F-H et au modèle de MTS-I, puisque d'un point de vue spécialisé, on attend une diminution lisse pour ANC0 et une augmentation lisse pour ANC4.

Tableau 7.1

Statistiques sommaires du biais relatif (BR, en %) à différents niveaux d'agrégation pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Niveau d'agrégation	Modèle	Min.	Q_1	Médiane	Moyenne	Q_3	Max.
ANC0	Pays	F-H	-0,48	-0,20	0,23	0,08	0,37	0,47
		MTS-I	-1,84	-0,68	0,54	-0,05	0,73	0,88
		MTS-II	-1,16	-0,55	0,29	0,41	1,19	2,43
		MTS-III	-1,53	-0,80	-0,57	-0,02	0,60	2,35
	Division	F-H	-0,68	-0,48	-0,36	0,05	0,50	1,31
		MTS-I	-0,99	-0,50	-0,31	0,05	0,64	1,41
		MTS-II	-0,77	0,04	0,15	0,59	1,08	2,50
		MTS-III	-1,44	-0,37	0,13	0,15	0,89	1,35
	District	F-H	-8,77	-1,72	0,14	0,31	1,67	12,41
		MTS-I	-10,35	-1,24	-0,49	-0,66	0,30	1,87
		MTS-II	-7,87	-1,15	0,77	1,25	2,89	18,34
		MTS-III	-10,05	-2,63	0,89	1,34	3,91	21,43

Tableau 7.1(suite)

Statistiques sommaires du biais relatif (BR, en %) à différents niveaux d'agrégation pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Niveau d'agrégation	Modèle	Min.	Q ₁	Médiane	Moyenne	Q ₃	Max.
ANC4	Pays	F-H	-1,65	-0,62	0,07	-0,07	0,65	1,04
		MTS-I	-4,09	-1,60	0,05	1,00	3,19	7,88
		MTS-II	-1,85	0,27	1,98	1,91	3,80	5,07
		MTS-III	-2,00	-1,35	1,06	1,11	3,10	5,23
	Division	F-H	-1,33	-0,60	-0,13	0,24	0,43	3,47
		MTS-I	-1,17	-0,25	-0,04	-0,07	0,32	0,59
		MTS-II	-0,50	0,68	1,18	1,55	1,70	5,39
		MTS-III	-2,08	0,31	0,73	1,24	1,92	5,58
	District	F-H	-17,83	-4,85	0,40	2,08	6,78	64,77
		MTS-I	-16,32	-3,80	-0,56	-0,42	2,98	15,57
		MTS-II	-22,00	-5,30	0,57	4,57	12,47	84,31
		MTS-III	-29,92	-8,23	0,57	6,12	14,07	124,63

Cette conclusion est confirmée par les valeurs de TC indiquées dans le tableau 7.4. Les TC pour les modèles de F-H transversaux sont trop élevés, ce qui indique que les prédictions de F-H tendent trop vers les estimations directes. Les niveaux de TC sont raisonnablement bons pour MTS-I, considérablement plus bas pour MTS-II et les plus bas pour MTS-III. Les taux de couverture plus faibles de MTS-II et de MTS-III au niveau du district se traduisent par les valeurs correspondantes plus élevées du BRA et de la RRET. Ces résultats montrent que les prédictions du modèle de MTS-I sont plus volatiles et tendent vers les estimations directes, que les prédictions du modèle MTS-III sont fortement lissées et que celles du modèle MTS-II semblent constituer un compromis raisonnable entre les prédictions du modèle de MTS-I et du modèle de MTS-III, en particulier au niveau des districts.

Tableau 7.2

Statistiques sommaires du biais relatif absolu (BRA, en %) à différents niveaux d'agrégation pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Niveau d'agrégation	Modèle	Min.	Q ₁	Médiane	Moyenne	Q ₃	Max.
ANC0	Pays	F-H	0,04	0,27	0,42	0,34	0,46	0,48
		MTS-I	0,26	0,63	0,75	0,87	0,99	1,84
		MTS-II	0,29	0,44	0,58	1,05	1,59	2,43
		MTS-III	0,49	0,61	0,96	1,18	1,61	2,35
	Division	F-H	0,39	0,50	0,65	0,90	1,20	1,84
		MTS-I	0,48	0,66	0,78	1,39	2,13	2,90
		MTS-II	0,79	0,96	1,56	1,78	2,14	3,88
		MTS-III	1,00	1,14	1,41	1,88	2,43	3,61
	District	F-H	1,08	2,73	4,17	5,12	5,84	15,94
		MTS-I	1,48	3,93	6,58	7,53	9,02	26,67
		MTS-II	3,15	6,46	10,31	11,32	14,50	33,01
		MTS-III	4,15	8,65	12,54	13,49	16,98	38,16

Tableau 7.2(suite)

Statistiques sommaires du biais relatif absolu (BRA, en %) à différents niveaux d'agrégation pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Niveau d'agrégation	Modèle	Min.	Q ₁	Médiane	Moyenne	Q ₃	Max.
ANC4	Pays	F-H	0,07	0,25	0,92	0,76	1,08	1,65
		MTS-I	0,05	1,60	2,47	3,09	4,00	7,88
		MTS-II	0,97	1,68	1,98	2,71	3,80	5,07
		MTS-III	1,06	1,19	1,46	2,46	3,53	5,23
	Division	F-H	0,98	1,40	1,71	1,87	2,06	3,47
		MTS-I	1,96	3,06	4,31	4,07	4,64	6,82
		MTS-II	2,18	3,66	4,68	4,33	5,07	6,00
		MTS-III	3,66	4,60	5,36	5,27	5,63	7,46
	District	F-H	1,93	7,64	12,91	14,29	17,60	64,77
		MTS-I	3,86	14,27	18,72	20,61	28,10	53,45
		MTS-II	7,07	19,47	26,22	28,51	35,88	84,31
		MTS-III	8,62	21,36	29,32	33,13	41,00	124,63

Tableau 7.3

Statistiques sommaires de la réduction relative des erreurs-types (RRET en %) à différents niveaux d'agrégation pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Niveau d'agrégation	Modèle	Min.	Q ₁	Médiane	Moyenne	Q ₃	Max.
ANC0	Pays	F-H	-0,65	4,03	8,10	8,00	12,72	15,01
		MTS-I	-0,03	1,35	4,01	3,67	5,82	7,33
		MTS-II	4,07	7,90	13,71	12,89	17,47	21,68
		MTS-III	-3,52	1,02	3,30	3,69	7,15	9,67
	Division	F-H	2,99	5,82	7,56	7,03	8,64	9,75
		MTS-I	2,66	4,00	5,32	5,16	6,65	6,84
		MTS-II	8,47	12,74	13,70	13,12	14,34	15,53
		MTS-III	3,30	4,71	5,21	5,47	6,12	8,16
	District	F-H	-1,60	7,17	10,20	9,98	12,04	21,61
		MTS-I	7,91	15,16	17,81	18,06	21,15	27,47
		MTS-II	12,60	27,84	34,08	33,80	38,46	48,53
		MTS-III	19,48	32,61	38,40	37,79	41,55	52,71
ANC4	Pays	F-H	8,58	11,22	11,66	13,71	14,49	24,32
		MTS-I	6,64	12,16	14,60	14,75	18,50	20,66
		MTS-II	17,79	22,87	23,56	25,12	27,99	32,75
		MTS-III	10,33	16,58	19,45	18,15	21,04	22,04
	Division	F-H	11,08	11,80	14,07	14,23	16,39	18,08
		MTS-I	11,82	14,31	14,46	15,78	18,18	19,17
		MTS-II	20,32	24,96	27,39	26,34	28,15	30,45
		MTS-III	15,49	20,37	21,75	21,72	24,51	25,05
	District	F-H	0,34	11,62	16,77	17,63	22,60	38,62
		MTS-I	17,79	27,84	30,48	30,93	33,65	43,40
		MTS-II	29,58	43,37	46,86	48,10	54,96	66,75
		MTS-III	35,63	48,88	51,75	52,94	59,31	70,35

Tableau 7.4

Taux de couverture observé (TC en %) des prédictions du modèle pour l'intervalle de confiance de 95 % au niveau du district et de la division, ainsi qu'au niveau du district par année d'enquête pour les estimations sur petits domaines d'ANC0 et d'ANC4

Paramètre	Modèle	TC selon l'année au niveau du district							TC global par niveau	
		1994	1997	2000	2004	2007	2011	2014	District	Division
ANC0	F-H	100,00	98,33	100,00	100,00	100,00	98,36	100,00	99,53	100,00
	MTS-I	100,00	90,00	93,44	88,52	93,22	98,36	100,00	94,81	100,00
	MTS-II	88,33	63,33	70,49	67,21	71,19	75,41	91,53	75,10	95,92
	MTS-III	83,33	53,33	50,82	52,46	61,02	55,74	79,66	62,22	95,92
ANC4	F-H	98,15	98,28	100,00	100,00	100,00	100,00	90,20	98,36	100,00
	MTS-I	87,04	84,48	68,33	76,27	81,97	96,72	100,00	84,58	95,92
	MTS-II	44,44	51,72	50,00	52,54	62,30	65,57	76,47	57,55	97,96
	MTS-III	44,44	41,38	40,00	38,98	50,82	55,74	72,55	48,70	97,96

8. Discussion

Dans la présente étude, nous avons élaboré des modèles de séries chronologiques multiniveaux (MTS) pour obtenir le pourcentage de femmes ne recevant pas de consultation de soins prénataux (ANC0) et le pourcentage de femmes recevant au moins quatre consultations de soins prénataux (ANC4) au Bangladesh, en nous appuyant uniquement sur sept éditions de la *Bangladesh Demographic and Health Survey* (BDHS, Enquête démographique et sur la santé du Bangladesh) couvrant la période de 1994 à 2014. Les modèles de séries chronologiques sont définis à une fréquence annuelle pour laquelle les années sans édition d'enquête sont traitées comme manquantes. Ainsi, le modèle tient compte des intervalles entre éditions de la BDHS et produit des prédictions pour les années sans enquête-échantillon. Les tendances sont produites à trois échelons régionaux, à savoir l'échelle nationale, qui se décompose en 7 divisions, elles-mêmes décomposées en 64 districts.

Dans le premier modèle (MTS-I), des estimations directes selon le domaine et l'année et leurs erreurs-types servent de données d'entrée du modèle de séries chronologiques multiniveau. Les tendances obtenues au moyen de ce modèle sont plutôt volatiles puisque les estimations de tendances tendent à suivre les estimations directes. Un autre inconvénient du modèle MTS-I est qu'il entrave l'utilisation de l'information auxiliaire disponible dans les deux recensements, car les valeurs de l'information auxiliaire d'un recensement donné ne changent pas dans deux ou trois éditions consécutives de l'enquête. Pour utiliser ces données de recensement, on propose d'élaborer des modèles transversaux de Fay-Herriot (F-H) séparément pour chaque année d'enquête. Dans un deuxième modèle, dit MTS-II, ces estimations de F-H et leurs erreurs-types servent de séries de données d'entrée. Dans un troisième modèle, dit MTS-III, les estimations de F-H avec leurs matrices de covariance complètes, sont utilisées comme séries de données d'entrée. Ce modèle de MTS tient correctement compte des corrélations transversales entre les estimations de F-H. Le modèle global pour MTS-II et MTS-III est alors un processus non itératif en deux étapes, dans lequel la première étape consiste à produire les estimations de F-H.

Les modèles sont élaborés à l'échelle régionale la plus détaillée, celui des districts. Les tendances au niveau de la division et à l'échelle nationale sont estimées par agrégation des prédictions de tendances au niveau du district. Ainsi, les chiffres à différents niveaux d'agrégation sont numériquement cohérents par définition.

Comparativement à d'autres modèles d'estimation sur petits domaines de séries chronologiques proposés dans la littérature, nos modèles contiennent plus de structure, puisque les modèles de tendances dynamiques sont précisés à différents niveaux d'agrégation. Cela est nécessaire pour obtenir des prédictions agrégées exactes au niveau des divisions et du pays et constitue un moyen plus parcimonieux de modéliser les corrélations transversales. On a envisagé une régularisation supplémentaire du modèle par la spécification des distributions *a priori* globales et locales. Cela n'a cependant pas permis d'améliorer les ajustements du modèle.

Dans l'estimation sur petits domaines, les estimations pour un domaine sont souvent étalonnées sur les estimations directes à l'échelle nationale pour assurer la cohérence numérique et pour tenter de réduire le biais dans les prédictions par domaine fondées sur le modèle. Dans cette application, les estimations de tendances à l'échelle nationale selon les modèles MTS sont déjà très proches des estimations directes. C'est pourquoi nous n'avons pas examiné d'étape supplémentaire d'étalonnage.

Les trois modèles de séries chronologiques fournissent des estimations plus exactes. Étant donné que MTS-II ignore les corrélations majoritairement positives entre les séries de données d'entrée transversales de F-H, les erreurs-types des tendances aux niveaux agrégés sont de fait trop petites. Puisque MTS-III tient compte de ces corrélations, les erreurs-types pour les tendances au niveau du pays et des divisions sont plus grandes, mais aussi plus réalistes. Le modèle MTS-II semble toutefois fournir les tendances les plus plausibles pour les deux variables réponses, particulièrement au niveau des districts, en trouvant un compromis entre la volatilité des tendances dans le modèle MTS-I et le caractère plat des tendances dans le modèle MTS-III. Ce choix est appuyé par le fait que ces variables sont susceptibles d'être relativement lisses dans le temps. Dans une perspective conceptuelle, l'ajustement de ces modèles aux séries d'ANC0 et d'ANC4 est donc certainement approprié. Cela justifie aussi l'interpolation des tendances pour les années sans enquête-échantillon. De plus, notre méthode peut être utile à de nombreux pays en développement qui réalisent des enquêtes démographiques et sur la santé répétées, car elles sont généralement réalisées à des intervalles variables et dépendent principalement de données de recensement non mises à jour dans les deux ou trois éditions suivantes de l'enquête.

On propose comme solution pratique d'utiliser les prédictions des modèles transversaux de F-H comme séries de données d'entrée dans les modèles de séries chronologiques multiniveaux afin d'exploiter au mieux les données de recensement disponibles. Cette méthode présente un autre avantage : elle stabilise les séries de données d'entrée dans les modèles MTS en éliminant les grandes erreurs d'échantillonnage des estimations directes. Cela exige toutefois un processus de sélection et d'évaluation minutieux des modèles transversaux de F-H, car une erreur de spécification des modèles transversaux de F-H peut donner lieu à des séries de données d'entrée biaisées ayant des erreurs-types estimées qui sous-estiment l'incertitude réelle.

Une des limites de la présente étude est liée à la correction du biais pour la transformation par la racine carrée qui est appliquée à ANC4. La correction du biais ne peut être appliquée qu'aux estimations de tendances pour les années d'enquête. Cela entraîne d'étranges augmentations des estimations ponctuelles si l'erreur d'échantillonnage n'est pas suffisamment lissée, en particulier pour les domaines à petite taille d'échantillon. Cela entrave l'estimation des changements d'une période à l'autre entre les années d'enquête et les autres années. C'est pourquoi la correction de biais n'est appliquée qu'aux modèles de F-H transversaux et non aux modèles MTS.

Étant donné que la prévalence des visites de soins prénataux ANC0 et ANC4 est corrélée négativement, un modèle multivarié peut constituer une solution de rechange intéressante aux modèles univariés utilisés dans la présente étude. Les deux séries pourraient être combinées à la série de la catégorie restante dans un seul modèle multivarié. Toutefois, cela nécessite un modèle multinomial qui a l'avantage d'améliorer la précision des estimations et de garantir que les prédictions prennent des valeurs dans leur fourchette admissible, et que les prédictions dans les catégories s'additionnent pour correspondre à 100 %. La mise en œuvre du modèle multinomial n'est cependant pas facile. Dans la présente étude en particulier, la matrice de variance-covariance peut être difficile à estimer pour les districts ayant un petit nombre d'observations. En outre, Datta et coll. (2002) montrent que les modèles univariés peuvent fournir des résultats aussi bons que les modèles multivariés proposés dans Ghosh, Nangia et Kim (1996), tout en étant plus simples à mettre en œuvre. Nous avons préféré ne pas examiner l'extension de nos modèles univariés à un modèle multinomial, qui pourrait faire l'objet d'une autre recherche.

Pour ANC0, l'échelle nationale montre une tendance à la baisse. Cette baisse s'est arrêtée temporairement pour la période de 2004 à 2011. La tendance d'ANC4 indique une augmentation constante pendant la période examinée dans la présente étude. Les tendances au niveau des divisions pour ANC0 montrent une baisse constante dans toutes les divisions, sauf *Dhaka*, *Chittagong* et *Sylhet*. Les tendances de ces trois divisions sont restées stables pendant la période allant de 2004 à 2011, ce qui est la principale cause de la tendance stable à l'échelle nationale d'ANC0. Par ailleurs, au niveau de la division, ANC4 montre des tendances à la hausse presque linéaires pour la plupart des divisions, sauf *Dhaka* et *Chittagong*. La plus forte amélioration est observée dans les divisions de *Khulna* et de *Rangpur*, où les tendances d'ANC4 atteignent plus de 40 % en 2014. Les tendances au niveau des districts permettent de repérer les districts très vulnérables en ce qui concerne les deux variables réponses étudiées. Bien que la tendance nationale d'ANC0 diminue pour s'établir à environ 21 % en 2014, quelques districts enregistrent un pourcentage inférieur à 10 % (*Dhaka*, *Jhenaidaha* et *Meherpur*), tandis qu'un nombre considérable de districts ont encore une valeur d'ANC0 supérieure à 35 % (*Bhola*, *Cox's Bazar*, *Kishoregonj*, *Noakhali*, *Sunamganj*, *Sirajgonj* et trois districts des *Chittagong Hill Tracts*). Pour ANC4, certains districts ont des estimations supérieures à 50 % (*Dhaka*, *Nilphamari* et *Panchagarh*) et la plupart des districts ayant une ANC0 élevée ont des estimations d'ANC4 inférieures à 20 %. Ces tendances au niveau des districts pourraient aider les décideurs à se concentrer sur les lieux les plus vulnérables où les indicateurs ANC0 et ANC4 restent faibles. De toute évidence, des tendances déterminées à des niveaux détaillés pourraient aider les décideurs à prendre des mesures pour réduire les inégalités de niveaux désagregés dans la course à la réalisation des objectifs de développement durable.

Remerciements

Nous tenons à remercier *Measure Evaluation* et le *National Institute of Population Research and Training* (NIPORT) qui ont rendu les données de la BDHS accessibles au public. Nos remerciements vont également à IPUMS qui a le mérite d'avoir donné accès aux données-échantillons des recensements du Bangladesh de 1991, de 2001 et de 2011. Les opinions exprimées dans le présent article sont celles des auteurs et ne reflètent pas nécessairement la politique du Bureau central de la statistique des Pays-Bas. Les auteurs remercient les deux réviseurs anonymes et le rédacteur adjoint pour leurs précieux commentaires sur une version précédente du présent article.

Annexe

Tableau A.1

Variables contextuelles au niveau du district générées à partir des données du Recensement de 1991, du Recensement de 2001 et du Recensement de 2011 pour ANC0

Variable	Définition
Division	Barishal, Chittagong, Dhaka, Khulna, Rajshahi, Rangpur, Sylhet.
Région	(1) Districts densément peuplés de Dhaka, Chittagong et Gazipur, (2) 9 districts régionaux avec de grandes villes, (3) 3 districts vallonnés (Bandarban, Khagrachhari et Rangamati), (4) 49 autres districts (régions moins urbanisées).
Chittagong	Division de Chittagong ?
Dhaka	Division de Dhaka ?
Khulna	Division de Khulna ?
Rangpur	Division de Rangpur ?
Rajshahi	Division de Rajshahi ?
P_U5	Proportion d'enfants de moins de 5 ans.
P_W	Proportion de femmes âgées de 15 à 49 ans.
P_MW	Proportion de femmes mariées âgées de 15 à 49 ans.
P_MW_Prim_Edu	Proportion de femmes mariées âgées de 15 à 49 ans qui ont suivi un enseignement primaire.
P_MW_Sec_Edu	Proportion de femmes mariées âgées de 15 à 49 ans qui ont suivi au moins un enseignement secondaire.
P_HH_No_Edu_W	Proportion de ménages comptant des femmes analphabètes âgées de 15 à 49 ans.
P_HH_Prim_Edu_W	Proportion de ménages comptant des femmes ayant suivi un enseignement primaire, âgées de 15 à 49 ans.
P_HH_High_Edu_W	Proportion de ménages comptant des femmes ayant suivi un enseignement supérieur, âgées de 15 à 49 ans.
P_HH_Sec_Edu_Head	Proportion de ménages ayant un chef de ménage ayant suivi au moins un enseignement secondaire.
P_Ru_HH_4 ⁺	Proportion de ménages ruraux de 4 personnes et plus.
P_Ru_HH_Elec	Proportion de ménages ruraux ayant l'électricité.

Tableau A.1(suite)

Variables contextuelles au niveau du district générées à partir des données du Recensement de 1991, du Recensement de 2001 et du Recensement de 2011 pour ANC0

Variable	Définition
P_Ru_HH_Sing_Moth	Proportion de ménages ruraux comptant une mère célibataire.
P_HH_U5_Sec_Edu_W	Proportion de ménages comptant des enfants âgés de moins de 5 ans et des femmes âgées de moins de 49 ans ayant suivi au moins un enseignement secondaire.
P_HH_2 ⁺ _U5	Proportion de ménages ayant au moins deux enfants âgés de moins de cinq ans.
P_Ru_HH_U5	Proportion de ménages ruraux ayant des enfants de moins de cinq ans.
P_Ru_HH_2 ⁺ _U5	Proportion de ménages ruraux ayant au moins deux enfants âgés de moins de cinq ans.
P_Ur_HH_2 ⁺ _U5	Proportion de ménages urbains ayant au moins deux enfants âgés de moins de cinq ans.

Tableau A.2

Effets fixes et aléatoires des modèles de F-H spécifiques à une année d'enquête pour ANC0

Année d'enquête	Transformation	Effets fixes	Effets aléatoires	Données de recensement
1994	Non	1 + Division + P_HH_High_Edu_W + P_Ur_HH_2 ⁺ _U5	OOA : Niveau du district Ordonnée à l'origine aléatoire	1991
1997	Non	1 + Division + P_MW_Sec_Edu + P_HH_Sec_Edu_Head	OOA	1991
2000	Non	1 + Division + P_Ru_HH_U5 + P_W	OOA	1991
2004	Non	1 + Division + P_HH_U5_Sec_Edu_W + P_MW	OOA	2001
2007	Non	1 + Division + P_U5 + P_Ru_HH_Size_4 ⁺	OOA	2001
2011	TPRC	1 + Division + sqrt(P_Ru_HH_U5) + sqrt(P_MW)	OOA	2011
2014	TPRC	1 + Division + sqrt(P_Ur_HH_2 ⁺ _U5) + sqrt(P_Ru_HH_Elec)	OOA	2011

Tableau A.3

Effets fixes et aléatoires des modèles de F-H spécifiques à une année d'enquête pour ANC4

Année d'enquête	Transformation	Effets fixes	Effets aléatoires	Données de recensement
1994	TPRC	1 + Division + sqrt(P_Ru_U5) + sqrt(P_HH_2 ⁺ _U5) + sqrt(P_Ur_HH_2 ⁺ _U5)	OOA : Niveau du district Ordonnée à l'origine aléatoire	1991
1997	TPRC	1 + Division + sqrt(P_HH_U5_Sec_Edu_W) + log(P_HH_Sec_Edu_Head)	OOA	1991
2000	TPRC	1 + Khulna + Region + sqrt(P_MW _{prim_Edu}) + sqrt(P_HH_W _{lli_Edu})	OOA	1991
2004	TPRC	1 + Division + sqrt(P_HH_U5_Prim_Edu_W) + sqrt(P_W)	OOA	2001
2007	TPRC	1 + Rangpur + Region + sqrt(P_U5) + sqrt(P_Ru_HH_Size_4 ⁺)	OOA	2001
2011	TPRC	1 + Rangpur + Chittagong + sqrt((P_HH_U5_Sec_Edu_W) + sqrt(P_W)	OOA	2011
2014	TPRC	1 + Rangpur + Chittagong + Rajshahi + Region + sqrt(P_W) + sqrt(P_Ru_HH_Sing_Mot)	OOA	2011

Bibliographie

- BBS (2020). *SDG Tracker: Bangladesh Development Mirror*. https://www.sdg.gov.bd/page/thirty_nine_plus_one_indicator/5#1, [En ligne; consulté le 18 décembre 2020].
- BBS et UNICEF (2019). *Progotir Pathay Bangladesh: Bangladesh Multiple Indicator Cluster Survey 2019*. Rapport technique, Bangladesh Bureau of Statistics (BBS).
- Besag, J., et Kooperberg, C. (1995). On conditional and intrinsic autoregression. *Biometrika*, 82(4), 733-746.
- Bollineni-Balabay, O., van den Brakel, J., Palm, F. et Boonstra, H.J. (2017). Multilevel hierarchical Bayesian versus state space approach in time series small area estimation: The Dutch Travel Survey. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 180(4), 1281-1308.
- Boonstra, H.J. (2021). *mcmcscsae: Markov Chain Monte Carlo Small Area Estimation*. R package version 0.6.0.
- Boonstra, H.J., et van den Brakel, J. (2019). [Estimation du niveau et de la variation du chômage au moyen de modèles de séries chronologiques structurels](https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2019003/article/00005-fra.pdf). *Techniques d'enquête*, 45, 3, 421-454. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2019003/article/00005-fra.pdf>.
- Boonstra, H.J., et van den Brakel, J. (2022). Multilevel time series models for small area estimation at different frequencies and domain levels. *Annals of Applied Statistics*, 16, 4, 2314-2338, DOI: 10.1214/21-AOAS1592.
- Boonstra, H.J., van den Brakel, J. et Das, S. (2021). Multilevel time series modelling of mobility trends in the Netherlands for small domains. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 184(3), 985-1007.
- Carvalho, C.M., Polson, N.G. et Scott, J.G. (2010). The horseshoe estimator for sparse signals. *Biometrika*, 97(2), 465-480.
- Das, S., Kumar, B. et Kawsar, L.A. (2020). Disaggregated level child morbidity in Bangladesh: An application of small area estimation method. *PLoS ONE*, 15(5), e0220164.
- Das, S., van den Brakel, J., Boonstra, H.J. et Haslett, S. (2021). Multilevel time series modelling of antenatal care coverage in Bangladesh at disaggregated administrative levels. Document de discussion, Statistics Netherlands.

- Datta, G.S., Lahiri, P. et Maiti, T. (2002). Empirical Bayes estimation of median income of four-person families by state using time series and cross-sectional data. *Journal of Statistical Planning and Inference*, 102(1), 83-97.
- 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.
- DHS (2021). *The DHS Program: Demographic and Health Surveys*. <https://dhsprogram.com/>, consulté le 01/11/2021.
- 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.
- 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.
- Ghosh, M., Nangia, N. et Kim, D.H. (1996). Estimation of median income of four-person families: A Bayesian time series approach. *Journal of the American Statistical Association*, 91(436), 1423-1431.
- Haslett, S., et Jones, G. (2004). *Local Estimation of Poverty and Malnutrition in Bangladesh*. Rapport technique, Bangladesh Bureau of Statistics and United Nations World Food Programme.
- Haslett, S., Jones, G. et Isidro, M. (2014). *Local Estimation of Poverty and Malnutrition in Bangladesh*. Rapport technique, United Nations World Food Programme and International Fund for Agricultural Development.

- Hossain, M.J., Das, S. et Chandra, H. (2020). Disaggregate level estimates and spatial mapping of food insecurity in Bangladesh by linking survey and census data. *PLoS ONE*, 15(4), e0230906.
- Marhuenda, Y., Molina, I. et Morales, D. (2013). Small area estimation with spatio-temporal Fay–Herriot models. *Computational Statistics & Data Analysis*, 58, 308-325.
- Mrisho, M., Obrist, B., Schellenberg, J.A., Haws, R.A., Mushi, A.K., Mshinda, H., Tanner, M. et Schellenberg, D. (2009). The use of antenatal and postnatal care: Perspectives and experiences of women and health care providers in rural Southern Tanzania. *BMC Pregnancy and Childbirth*, 9(1), 10.
- NIPORT (2015a). *Bangladesh Demographic and Health Survey 2014: Key Indicators*. Rapport technique, National Institute of Population Research and Training.
- NIPORT (2015b). *Bangladesh Demographic and Health Survey 2014: Policy Briefs*. Rapport technique, National Institute of Population Research and Training.
- NIPORT (2016). *Bangladesh Demographic and Health Survey 2014*. Rapport technique, National Institute of Population Research and Training.
- 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.
- OMS, UNICEF et autres (2014). *Trends in Maternal Mortality: 1990 to 2013: Estimates by WHO, UNICEF, UNFPA, the World Bank and the United Nations Population Division*. Rapport technique de l'Organisation mondiale de la santé.
- Park, T., et Casella, G. (2008). The Bayesian Lasso. *Journal of the American Statistical Association*, 103(482), 681-686.
- Pfeffermann, D., et Burck, L. (1990). [Estimation robuste pour petits domaines par la combinaison de données chronologiques et transversales](https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/1990002/article/14534-fra.pdf). *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. (2006). Small area estimation with state-space models subject to benchmark constraints. *Journal of the American Statistical Association*, 101, 1387-1397.
- Pratesi, M., et Salvati, N. (2008). Small area estimation: The EBLUP estimator based on spatially correlated random area effects. *Statistical Methods and Applications*, 17(1), 113-141.

- R Core Team (2015). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienne, Autriche.
- Rahman, M.S., Mohiuddin, H., Kafy, A.-A., Sheel, P.K. et Di, L. (2019). Classification of cities in Bangladesh based on remote sensing derived spatial characteristics. *Journal of Urban Management*, 8(2), 206-224.
- 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, 56.
- Sakia, R.M. (1992). The Box-Cox transformation technique: A review. *The Statistician*, 41, 169-178.
- 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.
- Tibshirani, R. (1996). Regression shrinkage and selection via the Lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, 267-288.
- Watanabe, S. (2010). Asymptotic equivalence of Bayes cross validation and widely applicable information criterion in singular learning theory. *Journal of Machine Learning Research*, 11, 3571-3594.
- Watanabe, S. (2013). A widely applicable Bayesian information criterion. *Journal of Machine Learning Research*, 14, 867-897.
- West, M. (1984). Outlier models and prior distributions in bayesian linear regression. *Journal of the Royal Statistical Society, Series B (Methodological)*, 431-439.
- Wolter, K. (2007). *Introduction to Variance Estimation*. Springer.
- You, Y. (2008). [Une approche intégrée de modélisation de l'estimation du taux de chômage pour les régions infraprovinciales au Canada](https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2008001/article/10614-fra.pdf). *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 Rao, J.N.K. (2000). [Estimation bayésienne hiérarchique des moyennes pour petites régions à l'aide de modèles à plusieurs niveaux](https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2000002/article/5537-fra.pdf). *Techniques d'enquête*, 26, 2, 197-206. Article accessible à l'adresse <https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2000002/article/5537-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](https://www150.statcan.gc.ca/n1/fr/pub/12-001-x/2003001/article/6602-fra.pdf). *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>.