

Techniques d'enquête

Répartir le fardeau de réponse dans les enquêtes-entreprises à Statistique Pays-Bas : évaluation des méthodes de coordination d'échantillons ciblant des entreprises ayant un lourd fardeau

par Marc J.E. Smeets et Jonas Klingwort

Date de diffusion : le 29 juin 2026



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 la ministre responsable de Statistique Canada

© Sa Majesté le Roi du chef du Canada, représenté par la ministre de l'Industrie, 2026

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.

Répartir le fardeau de réponse dans les enquêtes-entreprises à Statistique Pays-Bas : évaluation des méthodes de coordination d'échantillons ciblant des entreprises ayant un lourd fardeau

Marc J.E. Smeets et Jonas Klingwort¹

Résumé

Les instituts nationaux de statistique exploitent des systèmes de coordination d'échantillons pour répartir le fardeau de réponse dans les enquêtes-entreprises. Malgré la coordination d'échantillons appliquée et la surveillance du fardeau de réponse, certaines entreprises pourraient tout de même être grandement échantillonnées dans une courte période. Cela pourrait entraîner une pointe du fardeau de réponse pour certaines entreprises, ce qui influerait sur les taux de réponse et la qualité des réponses. Dans le présent article, on propose une nouvelle méthode de coordination d'échantillons fondée sur l'échantillonnage de Poisson spatialement corrélé adapté et visant les entreprises ayant un lourd fardeau de réponse. Les effets sur le fardeau de réponse sont évalués dans deux études par simulation et sont comparés à ceux d'une approche stratifiée, d'une méthode pragmatique dans laquelle les fractions d'échantillonnage sont ajustées manuellement et d'une méthode de référence qui ignore le fardeau de réponse. Les simulations sont réalisées à partir de scénarios concrets et de données de Statistique Pays-Bas. La première étude par simulation se rapporte à une situation pratique dans laquelle un échantillon donné est ajusté dans le but d'éviter l'échantillonnage des entreprises dont le fardeau de réponse est à son point culminant. La seconde étude par simulation vise à analyser les effets à plus long terme des différentes méthodes de coordination d'échantillons et est axée sur la réduction et la répartition du fardeau de réponse. Les avantages et les inconvénients des différentes méthodes sont expliqués et exposés en détail, puis des recommandations sur l'application de ces méthodes par les instituts nationaux de statistique et d'autres organismes d'enquête sont formulées.

Mots-clés : Coordination d'échantillons; échantillonnage de Pareto; échantillonnage de Poisson; échantillonnage de Poisson spatialement corrélé adapté; fardeau de réponse.

1. Introduction

Dans le contexte des statistiques officielles, on a souvent besoin de réduire le fardeau de réponse (perçu) lié aux enquêtes-entreprises (Bradburn, 1978; Bavdaž, Giesen, Černe, Loöfgren et Raymond-Blaess, 2015; Giesen, Vella, Brady, Brown, Ravindra et Vaasen-Otten, 2018; Erikson, Giesen, Houben et Saraiva, 2023). Même si le fardeau de réponse est un concept difficile à définir, puisqu'il peut renfermer à la fois des facteurs objectifs, comme le temps nécessaire pour répondre au questionnaire, et des facteurs subjectifs, comme ce qui est perçu comme un fardeau par les répondants (Yan, Fricker et Tsai, 2019; Bottone, Modugno et Neri, 2021; Yan et Williams, 2022), des données empiriques indiquent qu'une augmentation du fardeau perçu est associée à une plus grande attrition, à une réponse partielle, à une moindre qualité des données et à de plus faibles taux de réponse (Lorenc, Kloek, Abrahamsson et Eckman, 2013; Bottone et coll., 2021).

Les instituts nationaux de statistique (INS) font appel à deux grandes stratégies pour contrôler le fardeau de réponse. Ils réduisent d'abord le fardeau de réponse total en envoyant, par exemple, moins de

1. Marc J.E. Smeets, Statistics Netherlands, Pays-Bas. Courriel : mje.smeets@cbs.nl; Jonas Klingwort, Statistics Netherlands, Pays-Bas. Courriel : j.klingwort@cbs.nl.

questionnaires, puis en le dispersant de façon à ce qu'il soit réparti entre les entreprises données aussi uniformément que possible. Pour répartir le fardeau de réponse, les INS appliquent la coordination d'échantillons (Ernst, 1999; Giesen et coll., 2018; Erikson et coll., 2023) et utilisent un système de coordination d'échantillons. La coordination d'échantillons vise à contrôler la taille du chevauchement entre les échantillons prélevés de façon successive, par opposition au cas où les échantillons sont prélevés de manière indépendante. Une coordination positive vise à tenter d'optimiser le chevauchement afin d'obtenir des estimations plus précises de la variation. Les méthodes de coordination négative servent à réduire au minimum le chevauchement pour réduire le nombre d'enquêtes dans lesquelles une entreprise particulière est sélectionnée dans une période donnée et, par conséquent, pour répartir le fardeau de réponse (Bradburn, 1978). Matei et Smith (2023) donnent un aperçu des méthodes de coordination d'échantillons et examinent aussi les systèmes de coordination d'échantillons utilisés par les INS. Dans le reste du présent article, nous référons à la coordination d'échantillons négative lorsque nous employons le terme « coordination d'échantillons ».

Statistique Pays-Bas exploite un système de coordination d'échantillons appelé Système de fardeau d'enquête (SFE) (Smeets et Boonstra, 2018). Ce système applique la coordination d'échantillons parmi les groupes d'enquêtes, laquelle repose sur des nombres aléatoires permanents (NAP) et tient compte du fardeau de réponse cumulé des entreprises en fonction du nombre d'enquêtes pour lesquelles une entreprise a été sélectionnée. De façon plus précise, les entreprises sont classées en incrémentant, d'abord selon le fardeau de réponse cumulé, puis selon les NAP. Les premières entreprises de ce classement sont ensuite sélectionnées pour l'échantillon. Malgré la coordination d'échantillons appliquée, certaines entreprises sont encore grandement échantillonnées chaque année. Cela s'explique principalement par les grandes fractions d'échantillonnage et la stratification détaillée, mais aussi par le fait que les fonctions du SFE ne sont pas utilisées de façon optimale dans toutes les enquêtes. Il en résulte une pointe de fardeau de réponse pour certaines entreprises, ce qui peut causer des problèmes, surtout pour les petites entreprises. La méthode de coordination d'échantillons dont il est question dans le présent article n'est pas implantée dans le SFE. L'article vise à examiner si cette méthode peut se révéler une amélioration utile du SFE existant qui pourrait être implantée dans une prochaine version.

En réaction aux plaintes de la communauté des entreprises néerlandaises (Strategische commissie betere regelgeving bedrijven, 2020), Statistique Pays-Bas a amorcé en mai 2021 un projet visant à éviter, dans la mesure du possible, que des entreprises soient souvent échantillonnées de façon disproportionnée. Ce projet porte principalement sur les petites entreprises de 0 à 19 travailleurs. Une étude pilote qui s'est échelonnée de mai 2021 à mai 2022 a permis de surveiller la fréquence à laquelle une entreprise avait été échantillonnée au cours des 12 mois précédents (Gommans, Burgerjon et Paulissen, 2022). Dans le cadre de cette étude pilote, les entreprises comptant de 0 à 9 travailleurs ont été étiquetées en tant que « points chauds » si elles avaient été sélectionnées pour plus de 3 enquêtes différentes au cours des 12 mois précédents. Les entreprises comptant de 10 à 19 travailleurs ont été étiquetées en tant que « points chauds » si elles avaient été sélectionnées pour plus de 4 enquêtes au cours des 12 mois précédents. En règle générale, ces points chauds

peuvent être considérés comme étant des unités non désirées. Les unités désirées seraient habituellement celles ayant un petit fardeau de réponse cumulé et les unités non désirées, celles ayant un lourd fardeau de réponse cumulé (Matei, Smith, Smeets et Klingwort, 2023). Durant l'étude pilote, des méthodes pragmatiques ont été appliquées afin de contrôler le fardeau de réponse en réduisant le nombre de points chauds ou en les dispersant autant que possible (voir la sous-section 2.3). Statistique Pays-Bas a appliqué des méthodes pragmatiques parce que les politiciens le lui ont demandé (Strategische commissie betere regelgeving bedrijven, 2020). Ces méthodes représentent une situation pratique nécessitant des solutions pragmatiques en raison du manque de temps pour mettre en œuvre une méthode d'échantillonnage solide. Le présent article apporte une contribution par l'introduction d'une nouvelle méthode de coordination d'échantillons, de même que par la comparaison des méthodes pragmatiques qui visent précisément la réduction du nombre de points chauds (ou le nombre d'unités non désirées en général) et des méthodes solides sur le plan méthodologique qui visent à répartir le fardeau de réponse. Nous avons tenu compte des analyses et des résultats de la présente étude, qui est d'une importance pratique, pour d'autres organismes nationaux de statistique qui pourraient faire face à des situations comparables.

Matei et coll. (2023) ont mis au point une méthode de coordination d'échantillons au moyen de l'échantillonnage de Poisson spatialement corrélé adapté (PSCA) et celle-ci permet de faire un double contrôle du fardeau de réponse pour les entreprises étiquetées comme étant des unités non désirées. Dans ce contexte, il s'agissait d'entreprises que les INS veulent épargner le plus possible de l'échantillonnage en raison, par exemple, d'un lourd fardeau de réponse cumulé. Cette méthode fait appel à des NAP et repose sur un échantillonnage de Poisson corrélé (Grafström, 2012). Dans Matei et coll. (2023), cette méthode est appliquée dans une étude par simulation fondée sur un scénario inspiré des enquêtes-entreprises néerlandaises; on y montre qu'il est possible de réduire la variance du nombre de points chauds sélectionnés dans plusieurs échantillons comparativement à un échantillonnage de Poisson avec des NAP (Brewer, Early et Joyce, 1972) et à un échantillonnage ordonné de Pareto (Rosén, 1997a et 1997b).

Au moyen de l'échantillonnage de PSCA, il est possible de prélever des échantillons pour éviter d'y retrouver un regroupement d'unités non désirées. Le recours à l'échantillonnage de PSCA dans une méthode de coordination d'échantillons pourrait éviter la sélection d'un grand nombre d'entreprises ayant un lourd fardeau de réponse au cours d'une période donnée. La méthode de coordination d'échantillons proposée dans le présent article vise à modifier le statut de point chaud des entreprises et à appliquer l'échantillonnage de PSCA aux points chauds sans faire appel aux NAP. Nous comparons cette méthode et une approche stratifiée aux méthodes pragmatiques appliquées dans l'étude pilote, évaluons les méthodes en utilisant une simulation de Monte Carlo et analysons leurs effets à plus long terme sur le nombre de points chauds échantillonnés. Nous considérons que la comparaison avec les méthodes de l'étude pilote est particulièrement pertinente, puisqu'elles sont simples à mettre en œuvre pour les INS ou les organismes d'enquête; cependant, elles pourraient entraîner des problèmes, comme nous allons le démontrer.

Nous étudions la possibilité d'appliquer la méthode proposée, fondée sur l'échantillonnage de PSCA, en tant que solution de rechange aux méthodes utilisées dans l'étude pilote. Dans la première étude par

simulation, nous examinons la réduction du nombre de points chauds dans un échantillon donné. Dans la seconde étude par simulation, nous analysons les effets à plus long terme des méthodes étudiées sur le nombre de points chauds échantillonnés lorsqu'une réduction et une répartition sont prises en considération, et nous vérifions si les méthodes respectent les probabilités d'inclusion ciblées du plan d'échantillonnage sous-jacent. Les données de l'enquête néerlandaise sur la recherche et le développement menée en 2020 sont utilisées dans les deux études par simulation.

L'article est structuré de la manière suivante : à la section 2, une vue d'ensemble des méthodes d'échantillonnage prises en considération est présentée et la manière d'utiliser ces méthodes pour la coordination d'échantillons est décrite. Les études par simulation sont présentées à la section 3. Les résultats sont fournis à la section 4. On trouve, à la section 5, une discussion et la conclusion.

2. Méthodes d'échantillonnage et coordination d'échantillons

Dans la présente section, nous décrirons brièvement les échantillonnages de Poisson, de Pareto et de PSCA, nous expliquerons la manière de les utiliser pour la coordination d'échantillons, nous proposerons une nouvelle méthode de coordination d'échantillons et présenterons les méthodes utilisées dans l'étude pilote.

2.1 Plans d'échantillonnage de Poisson, de Pareto et de PSCA

Un échantillon de Poisson provenant d'une population de N unités est prélevé comme suit (Brewer et coll., 1972) : des nombres aléatoires sont générés $r_1, \dots, r_k, \dots, r_N \stackrel{\text{iid}}{\sim} \text{Unif}(0, 1)$ et l'unité k dans l'échantillon S est sélectionnée si $r_k < \pi_k$. Dans ce cas-ci, $\pi_k = P(k \in S)$ est la probabilité d'inclusion de l'unité, ou de l'entreprise, k . Il convient de mentionner que la taille de l'échantillon n_S d'un échantillon de Poisson est aléatoire et que la taille escomptée de l'échantillon est $E(n_S) = \sum_{k=1}^N \pi_k$.

Il est parfois préférable d'opter pour des plans d'échantillonnage à taille fixe. L'échantillonnage de Pareto (Rosén, 1997b) constitue un bon exemple de ce genre de plan d'échantillonnage, et il a aussi servi dans l'étude pilote et pour l'approche stratifiée (voir ci-dessous). Pour les nombres aléatoires donnés $r_1, \dots, r_k, \dots, r_N \stackrel{\text{iid}}{\sim} \text{Unif}(0, 1)$ et les probabilités d'inclusion $\pi_1, \dots, \pi_k, \dots, \pi_N$, un échantillon de Pareto de taille n est obtenu en sélectionnant les n unités ayant les plus petites valeurs de

$$\rho_k = \frac{r_k / (1 - r_k)}{\pi_k / (1 - \pi_k)}.$$

L'échantillonnage de PSCA est un cas spécial d'échantillonnage de Poisson corrélé (Bondesson et Thorburn, 2008). L'échantillonnage de Poisson corrélé est une méthode séquentielle d'énumération utilisée pour prélever un échantillon aléatoire S à partir d'une population U , avec des probabilités d'inclusion ciblées π_k pour $k \in U$. Examinons la probabilité de sélection $\pi_k^{(\ell-1)}$, pour $1 \leq \ell \leq N$, qui permet à l'unité ℓ d'être sélectionnée dans S à la ℓ^{e} itération de l'algorithme suivant (voir l'étape 3b) :

1. Poser $\pi_k^{(0)} = \pi_k$ pour toutes les $k \in U$.
2. Poser $\ell = 1$.
3. Alors que $\ell \leq N$,
 - a) générer r_ℓ séparément, à partir de la distribution $\text{Unif}(0, 1)$;
 - b) si $r_\ell < \pi_\ell^{(\ell-1)}$, alors $I_\ell = 1$ (ℓ est sélectionnée dans S), autrement $I_\ell = 0$ (ℓ n'est pas sélectionnée dans S), où I_ℓ est la variable indicatrice associée à l'unité $\ell \in U$;
 - c) mettre à jour les probabilités de sélection pour les unités restantes $i = \ell + 1, \dots, N$ selon

$$\pi_i^{(\ell)} = \pi_i^{(\ell-1)} - (I_\ell - \pi_\ell^{(\ell-1)}) w_\ell^{(i)},$$

où $w_\ell^{(i)}$ sont des facteurs de pondération donnés par l'unité ℓ aux autres unités $i = \ell + 1, \dots, N$;

- d) augmenter ℓ de 1.

Le choix de $w_\ell^{(i)}$ à l'étape 3.c) fournit différents plans d'échantillonnage; par exemple, l'échantillonnage de Poisson est obtenu si tous les $w_\ell^{(i)} = 0$ pour $\ell \in U$. Bondesson et Thorburn (2008) établissent des conditions sur le choix de $w_\ell^{(i)}$, de sorte que $0 \leq \pi_i^{(\ell)} \leq 1$ et des échantillons à taille fixe sont obtenus. Ils montrent également que selon ces conditions, les probabilités d'inclusion ciblées sont respectées.

Grafström (2012) a appliqué l'échantillonnage de Poisson corrélé à une population dans laquelle chaque unité était associée à des coordonnées géographiques. L'idée sous-jacente est que des unités contiguës pourraient fournir des renseignements semblables et qu'un échantillon contient plus de renseignements s'il évite le regroupement d'unités contiguës. En utilisant, à l'étape 3.c), des facteurs de pondération $w_\ell^{(i)}$ qui remplissent les conditions de Bondesson et Thorburn (2008), de sorte que l'unité ℓ donne un facteur de pondération maximal à l'unité la plus rapprochée de ℓ en distance (euclidienne), parmi les unités $i = \ell + 1, \dots, N$, on obtient l'échantillonnage de Poisson spatialement corrélé (PSC). Cette méthode procure une couverture géographique de toute la population observée en dispersant les unités échantillonnées dans la région géographique.

Au lieu d'appliquer l'échantillonnage de PSC à des coordonnées géographiques, il est aussi possible de l'appliquer à d'autres caractéristiques des unités, par exemple le fardeau de réponse cumulé ou des facteurs de pondération fondés sur le temps requis pour répondre au questionnaire, ce qui donne lieu à des échantillons bien répartis en ce qui a trait à ces caractéristiques. Il convient de mentionner que ces caractéristiques ne doivent pas dépendre d'indicateurs de réponse, car cela pourrait entraîner une sélectivité dans les échantillons prélevés, c'est-à-dire que des groupes sélectifs de répondants pourraient devenir surreprésentés dans l'échantillon. Matei et coll. (2023) appliquent un échantillonnage de PSC, qu'ils appellent « échantillonnage de PSCA », à une mesure du fardeau de réponse cumulé.

Dans ce cas-ci, nous introduisons une variable binaire qui indique si une entreprise est un point chaud ou non. Recourir à cette variable binaire en tant que mesure du fardeau de réponse cumulé mène à un cas spécial d'échantillonnage de PSCA dans lequel des échantillons spatialement équilibrés sont prélevés par

rapport à l'espace généré par cette variable. La variable binaire pour l'unité k a une valeur de 1 si k est un point chaud et de 0 dans les autres cas. Il est possible de prélever des échantillons de PSCA à l'aide de la fonction `scps()` dans le paquet `R BalancedSampling` (Grafström, Lisic et Prentius, 2024).

Nous avons également examiné une approche stratifiée de l'échantillonnage de PSCA pour les plans d'échantillonnage stratifiés ayant des probabilités d'inclusion égales au sein des strates. Supposons que $1, \dots, h, \dots, H$ désigne les strates de la population U , où une strate h affiche une taille d'échantillon n_h et une taille de population N_h . Les probabilités d'inclusion sont obtenues à l'aide de fractions d'échantillonnage de la strate, c'est-à-dire $\pi_k = f_h = n_h / N_h$ pour les unités k dans la strate h . En définissant une sous-strate distincte pour les unités qui sont des points chauds et celles qui n'en sont pas de la strate h et en prélevant des échantillons aléatoires dans les deux sous-strates ayant aussi une fraction d'échantillonnage f_h , on obtient un échantillon spatialement équilibré. Nous employons l'échantillonnage de Pareto pour prélever les échantillons dans les sous-strates.

2.2 Coordination d'échantillons

Les méthodes d'échantillonnage de la sous-section 2.1 peuvent servir à la coordination d'échantillons lorsque des NAP communs sont utilisés (Ohlsson, 1995). L'idée est que les entreprises de la population conservent les nombres aléatoires attribués $r_1, \dots, r_k, \dots, r_N \stackrel{\text{iid}}{\sim} \text{Unif}(0, 1)$ dans le processus de sélection au fil du temps. On parvient alors à une coordination d'échantillons négative en utilisant r_k et $1 - r_k$ tour à tour en tant que nombres aléatoires pour $k \in U$ lors du prélèvement d'échantillons successifs. Par exemple, l'échantillonnage de Poisson avec des NAP (Brewer et coll., 1972) fonctionne comme suit : pour des probabilités d'inclusion données $\pi_1, \dots, \pi_k, \dots, \pi_N$, l'unité k est sélectionnée dans le premier échantillon si $r_k < \pi_k$, est sélectionnée dans l'échantillon suivant si $1 - r_k < \pi_k$ et ainsi de suite.

De la même manière, des méthodes de coordination d'échantillons sont créées à partir des autres méthodes d'échantillonnage de la sous-section 2.1. La combinaison de l'échantillonnage de Pareto avec des NAP donne un échantillonnage ordonné de Pareto (Rosén, 1997a et 1997b). Matei et coll. (2023) appliquent la coordination d'échantillons en utilisant l'échantillonnage de PSCA en association avec des NAP. Les NAP r_k et $1 - r_k$ servent ensuite à l'étape 3.b) de l'algorithme de la sous-section 2.1, et l'étape 3.a) est omise. Cette méthode est présentée comme étant la stratégie de double contrôle ciblé. Elle permet un double contrôle du fardeau de réponse des unités non désirées, d'abord par l'application d'une coordination d'échantillons au moyen des NAP afin de réduire le plus possible le chevauchement entre les échantillons prélevés dans des enquêtes successives, puis par le prélèvement d'échantillons spatialement équilibrés pour éviter le regroupement d'unités non désirées dans l'échantillon.

Nous proposons une autre méthode de coordination d'échantillons basée sur l'échantillonnage de PSCA dans laquelle nous n'utilisons pas de NAP. L'idée est la suivante : tenir à jour le statut de point chaud des entreprises au fil du temps et faire appel à l'échantillonnage de PSCA pour prélever des échantillons spatialement équilibrés par rapport au statut de point chaud, comme cela est décrit à la sous-section 2.1. Nous appelons cette méthode « stratégie d'échantillonnage de PSCA pour les points chauds ». Grâce à cette

méthode, le fardeau de réponse est réparti de manière plus réaliste pour les entreprises parce qu'elle mènera à de plus longues périodes au cours desquelles les entreprises ne seront pas contactées. Par opposition, dans le cas des méthodes qui reposent sur des NAP décrites plus haut, les entreprises sont sélectionnées ou non une fois sur deux dans des échantillons successifs. Pour cette raison, nous n'envisageons pas les méthodes de coordination d'échantillons fondées sur les NAP, comme la stratégie de double contrôle ciblé, dans les études par simulation.

2.3 Méthodes pragmatiques appliquées dans l'étude pilote

Le service de collecte des données de Statistique Pays-Bas prélève des échantillons pour des enquêtes-entreprises néerlandaises. Dans la majorité des enquêtes, le service utilise le système de coordination d'échantillons (le SFE). Durant l'étude pilote, le groupe responsable de la collecte des données a surveillé les unités des échantillons prélevés qui devenaient des points chauds. Dans la mesure du possible et en collaboration avec les services de statistique, les points chauds ont été retirés manuellement des échantillons avant le début du travail sur le terrain, à l'aide des méthodes suivantes (Gommans et coll., 2022) :

1. Retirer manuellement les points chauds de l'échantillon prélevé. Cela a permis d'envoyer moins de questionnaires et d'obtenir un échantillon de plus petite taille que la taille d'échantillon ciblée.
2. Remplacer manuellement les points chauds par d'autres unités. De cette manière, la taille d'échantillon ciblée est maintenue, dans la mesure du possible. Les stratégies suivantes ont été appliquées :
 - a) Examiner un échantillon existant qui contient des points chauds. D'après le nombre de points chauds, il est possible de calculer le nombre escompté de points chauds pour un nouvel échantillon. En ajoutant le nombre escompté de points chauds à la taille d'échantillon requise, prélever de nouveau un échantillon plus grand, puis retirer manuellement les points chauds et toute unité supplémentaire de cet échantillon pour arriver à la taille d'échantillon requise. Reprendre ces étapes de manière itérative afin de remplacer le plus de points chauds possible. Cela nécessitera possiblement quelques étapes d'itération avant l'obtention de l'échantillon définitif.
 - b) Retirer manuellement tous les points chauds de la base de sondage et prélever un échantillon ayant initialement la taille d'échantillon ciblée à partir des unités restantes de la base de sondage.
 - c) Remplacer manuellement les points chauds présents dans un panel rotatif par d'autres unités durant le renouvellement du panel.

Il convient de mentionner que ces méthodes pragmatiques permettent d'ajuster efficacement les probabilités d'inclusion des unités individuelles. Cela signifie que les probabilités d'inclusion réalisées dans l'échantillon définitif ne sont plus égales aux probabilités d'inclusion ciblées dans le plan d'échantillonnage. Dans ce cas, l'échantillon n'est plus aléatoire, et l'utilisation d'un estimateur fondé sur le plan introduit un risque d'estimations biaisées. Ce biais peut être partiellement corrigé par le calibrage de la réponse observée selon les totaux de la population.

Les stratégies appliquées dans l'étude pilote visent à contrôler le fardeau de réponse en réduisant (la méthode 1 ci-dessus) ou en étendant (les méthodes 2.a), 2.b) et 2.c ci-dessus) le nombre de points chauds échantillonnés. On entend ici par répartir le fait d'attribuer la taille d'échantillon totale de façon différente à toutes les unités de la population afin de remplacer les unités qui sont des points chauds par des unités qui n'en sont pas.

Le présent article porte sur l'application de la coordination d'échantillons aux enquêtes transversales. Nous voulons comparer l'échantillonnage de PSCA et l'utilisation de sous-strates aux stratégies de l'étude pilote; ainsi, seules les méthodes 2.a) et 2.b) sont pertinentes pour la présente étude.

Les méthodes 2.a) et 2.b) sont deux méthodes différentes qui permettent d'obtenir un échantillon dans lequel les points chauds sont remplacés par d'autres unités, dans la mesure du possible. Dans la méthode 2.a), cela s'effectue en augmentant (de façon itérative) les fractions d'échantillonnage des unités qui ne sont pas des points chauds et en prélevant des échantillons provisoires jusqu'à l'obtention de l'échantillon souhaité. Dans la méthode 2.b), cela s'effectue en procédant à l'échantillonnage seulement une fois à partir de la base de sondage réduite, mais avec les probabilités d'inclusion ciblées initialement. Dans les deux méthodes, les fractions d'échantillonnage des unités qui ne sont pas des points chauds sont par conséquent accrues durant le processus d'échantillonnage. Les méthodes 2.a) et 2.b) sont considérées comme étant des méthodes semblables dans le présent article, et nous appliquerons uniquement la méthode 2.b) dans les études par simulation de la section 3. Nous faisons référence à la méthode de coordination d'échantillons dans laquelle le statut de point chaud est mis à jour au fil du temps, et la méthode 2.b) est appliquée sous forme de « fractions croissantes ».

3. Aperçu des études par simulation

La présente section donne un aperçu de deux études par simulation. Le but consiste à examiner la possibilité que l'échantillonnage de PSCA soit utilisé en tant que solution de rechange à la méthode pragmatique des fractions croissantes (section 2). La première étude par simulation porte sur la réduction du nombre de points chauds dans un échantillon donné prélevé par le SFE. La seconde étude par simulation est axée sur les effets à plus long terme sur les points chauds échantillonnés et vise à vérifier si toutes les méthodes respectent les fractions d'échantillonnage ciblées dans le plan d'échantillonnage. Cette étude par simulation fournit des renseignements sur la réduction et la variation du nombre de points chauds au fil du temps.

Dans les deux études par simulation, les données de l'Enquête néerlandaise sur la recherche et le développement menée en 2020 sont utilisées. Il s'agit d'une enquête annuelle qui repose sur un plan d'échantillonnage stratifié ayant des probabilités d'inclusion égales au sein des strates et dans laquelle les entreprises sont les unités d'échantillonnage (Statistique Pays-Bas, 2022). Cela sous-entend que les méthodes examinées peuvent être appliquées à chaque strate séparément. Les strates sont définies par une combinaison de classification des industries selon la Standaard Bedrijfsindeling (SBI; classification type des industries) et la catégorie de taille des entreprises. La SBI est fondée sur la Nomenclature statistique des activités économiques

dans les communautés européennes (la classification type des industries en Europe) et la catégorie de taille, sur le nombre de travailleurs. La base de sondage consiste en un volet d'échantillonnage, des strates à tirage complet et des strates à tirage nul. Lors de la stratification, 99 combinaisons de la SBI et 6 catégories de taille sont utilisées. L'enquête sur la recherche et le développement repose sur un échantillon pour les catégories de taille 4 à 9 (10 travailleurs ou plus). Puisque les points chauds sont définis dans les catégories de taille 0 à 4 (section 1), seuls les points chauds de l'échantillonnage pour la catégorie de taille 4 peuvent être pris en considération dans les simulations.

Pour les deux études par simulation, nous avons fait appel au plan d'échantillonnage qui avait servi dans le cadre de l'Enquête sur la recherche et le développement de 2020. À ce moment-là, Statistique Pays-Bas n'a rien fait pour retirer ou remplacer les points chauds de l'échantillon; seul le statut de point chaud a été surveillé. Par conséquent, l'échantillon qui servira dans les études par simulation constitue un scénario réaliste se produisant dans la pratique. Nous tenons uniquement compte du volet d'échantillonnage, puisque le fait de répartir le fardeau de réponse n'a aucun effet sur les échantillons quand la fraction d'échantillonnage est de 0 ou de 1 (strates à tirage nul ou à tirage complet).

Supposons que U est la population de l'enquête sur la recherche et le développement qui compte des entreprises de la catégorie de taille 4 présentes dans le volet d'échantillonnage de la base de sondage de l'Enquête sur la recherche et le développement de 2020, et que S est l'échantillon sélectionné à partir de U que le SFE a prélevé pour l'Enquête sur la recherche et le développement de 2020. La taille de U est $N = 33\,044$, et l'échantillon S se compose de $n = 926$ unités. Il y a $N_{PC} = 211$ points chauds dans U et $n_{PC} = 108$ points chauds dans S . Le tableau A.1 du document de travail d'accompagnement de Smeets et Klingwort (2023) montre les $H = 83$ strates, de même que N_h , n_h , $f_h = n_h / N_h$, le nombre de points chauds dans la population $N_{PC,h}$, la fraction de points chauds de la population $F_{PC,h} = N_{PC,h} / N_h$, le nombre escompté de points chauds échantillonnés $E(n_{PC,h}) = n_h F_{PC,h}$ et la variance du nombre de points chauds échantillonnés $V(n_{PC,h}) = n_h F_{PC,h} (1 - F_{PC,h}) (N_h - n_h) / (N_h - 1)$. Il convient de mentionner que $n_{PC,h}$ est distribué de manière hypergéométrique selon des paramètres N_h , $N_{PC,h}$ et n_h , que suivent l'espérance et la variance de $n_{PC,h}$.

3.1 Première étude par simulation : une situation pratique

La première étude par simulation traite de la situation dans laquelle les différentes méthodes sont appliquées à un échantillon donné, l'échantillon initial, dans le but d'y réduire le nombre de points chauds. L'échantillon initial, censé avoir été prélevé à l'aide du système de coordination d'échantillons, est considéré comme étant fixe dans cette étude par simulation. Ensuite, dans chaque exécution, un second échantillon est obtenu à l'aide des méthodes suivantes, respectivement :

1. *Ignorance des points chauds* : prendre l'échantillon initial et conserver tous les points chauds qu'il contient.

2. *Fractions croissantes* : remplacer manuellement les points chauds de l'échantillon initial par d'autres unités en augmentant les fractions d'échantillonnage d'unités qui ne sont pas des points chauds (méthode 2.b); voir la sous-section 2.3).
3. *Échantillonnage de PSCA* : appliquer l'échantillonnage de PSCA en tenant compte de la variable binaire (1 en cas de point chaud, 0 dans les autres cas) dans le processus d'échantillonnage (voir la sous-section 2.1).
4. *Utilisation de sous-strates* : utiliser deux sous-strates distinctes pour les points chauds et les autres unités (sous-section 2.1).

Dans la méthode d'ignorance des points chauds, l'échantillon initial n'est pas modifié et, par conséquent, le nombre de points chauds qui s'y trouvent n'est pas réduit. Il s'agit de notre méthode de référence.

Nous examinons seulement les strates qui contiennent au moins un point chaud dans la population parce que dans le cas contraire, toutes les méthodes produisent des échantillons sans points chauds. Nous examinons alors les strates de la catégorie de taille 4 avec $0 < f_h < 1$ et $N_{PC,h} > 0$. Supposons que U' est la population qui répond à ces trois exigences. Disons que $S' \subset S$ est la partie de S contenant les unités sélectionnées à partir des strates qui répondent à ces trois exigences, où S est l'échantillon donné prélevé à l'aide du SFE. Un échantillon S' sert d'échantillon initial. La population U' se compose de $H = 48$ strates et est de taille $N' = 17\,773$, et S' comprend $n' = 598$ unités et $n'_{PC} = 108$ points chauds.

Dans cette étude par simulation, la coordination d'échantillons est appliquée de manière différente, comme le décrit la sous-section 2.2. Le statut de point chaud n'est pas modifié, puisque l'échantillon initial et le second échantillon renvoient à la même période et que le but consiste à remplacer l'échantillon initial par le second échantillon. La coordination d'échantillons est appliquée en générant des nombres aléatoires qui mènent à l'échantillon initial lorsque l'on ignore les points chauds. Les autres méthodes reposent sur ces nombres aléatoires pour prélever le second échantillon. Les méthodes sont mises en œuvre dans l'algorithme ci-dessous. L'échantillonnage de Pareto (Rosén, 1997b) est décrit à la sous-section 2.1.

Algorithme 1.

1. Générer des nombres aléatoires uniformes $0 < r_k < 1$ pour toutes les $k \in U'$, de sorte que $r_k < \pi_k$ si $k \in S'$, et $r_k \geq \pi_k$ dans les autres cas. Calculer

$$\rho_k = \frac{r_k / (1 - r_k)}{\pi_k / (1 - \pi_k)} \text{ pour toutes les } k \in U'.$$

2. Appliquer les méthodes d'échantillonnage :

- *Ignorance des points chauds* : prélever un échantillon de Pareto S_1 avec un r_k et une π_k donnés en sélectionnant les n_h unités ayant les plus petites valeurs de ρ_k .

- *Fractions croissantes* : prélever un échantillon de Pareto S_2 avec un r_k et une taille d'échantillon n_h donnés, où les $N_{PC,h}$ points chauds de la strate h sont retirés de la population. L'échantillon est obtenu en sélectionnant les n_h unités de la population réduite ayant les plus petites valeurs de ρ_k . Si $n_h > (N_h - N_{PC,h})$, il faut alors appliquer « l'ignorance des points chauds » pour qu'aucun point chaud ne soit retiré de l'échantillon dans la strate h .
- *Échantillonnage de PSCA* : prélever un échantillon de PSCA S_3 avec un r_k et une π_k donnés, et le statut de point chaud.
- *Utilisation de sous-strates* : prélever des échantillons de Pareto distincts dans les sous-strates d'unités qui sont des points chauds et d'unités qui ne sont pas des points chauds avec un r_k donné, et supposer que S_4 correspond à la combinaison de ces échantillons. La taille d'échantillon $n_{h,1}$ dans la sous-strate de points chauds de la strate h est obtenue en arrondissant aléatoirement $f_h N_{PC,h}$, et la taille d'échantillon $n_{h,2}$ dans la sous-strate d'unités qui ne sont pas des points chauds est $n_{h,2} = n_h - n_{h,1}$. Un arrondissement aléatoire est appliqué selon des probabilités proportionnelles à $f_h N_{PC,h}$ et à $1 - f_h N_{PC,h}$ d'arrondissement vers le haut et le bas, respectivement. L'échantillon de la sous-strate de points chauds est obtenu en sélectionnant les $n_{h,1}$ unités ayant les plus petites valeurs de ρ_k de la sous-strate de points chauds et, pour les unités qui ne sont pas des points chauds, les $n_{h,2}$ unités ayant les plus petites valeurs de ρ_k de la sous-strate d'unités qui ne sont pas des points chauds sont sélectionnées.

Dans l'algorithme, toutes les méthodes sont appliquées par strate, ce qui donne un échantillon par strate. Les échantillons S_1 , S_2 , S_3 et S_4 sont ensuite obtenus en combinant ces échantillons par strate. Il convient de mentionner que les nombres aléatoires de l'étape 1 sont générés de telle sorte que $S_1 = S'$ dans chaque exécution de simulation. La méthode « fractions croissantes » n'est pas appliquée si $n_h > (N_h - N_{PC,h})$. Dans ce cas, le nombre de points chauds dans la strate est d'au moins $n_h - (N_h - N_{PC,h})$ et peut être réduit à ce minimum. Pour des raisons pratiques, nous avons décidé de ne pas le faire.

Dans l'étude par simulation, nous prélevons des échantillons $S_1^{(j)}$, $S_2^{(j)}$, $S_3^{(j)}$ et $S_4^{(j)}$ en reprenant les étapes 1 et 2 à chaque exécution $j = 1, \dots, R$. Au total, nous examinons $R = 1\,000$ exécutions, et chacune d'elles représente un seul prélèvement d'échantillon. Pour chaque échantillon $S_m^{(j)}$, prélevé à l'aide de la méthode m lors de l'exécution j , la taille de l'échantillon n et le nombre de points chauds échantillonnés n_{PC} sont calculés pour l'échantillon en entier et pour l'échantillon tiré de la strate h . Nous écrivons

$$\begin{aligned} n_m^{(j)} &= n(S_m^{(j)}), & n_{PC,m}^{(j)} &= n_{PC}(S_m^{(j)}), \\ n_{h,m}^{(j)} &= n_h(S_m^{(j)}), & n_{PC,h,m}^{(j)} &= n_{PC,h}(S_m^{(j)}), \end{aligned}$$

pour $m = 1, \dots, 4$, $h = 1, \dots, H$ et $j = 1, \dots, R$.

Dans l'étude par simulation, les tailles d'échantillon obtenues sont vérifiées et la réduction du nombre de points chauds échantillonnés à l'aide des différentes méthodes est examinée. À cette fin, les distributions du nombre de points chauds échantillonnés lors de l'étude par simulation sont analysées. Pour l'échantillon en entier, les tableaux statistiques créés présentent le nombre d'échantillons dans l'étude par simulation ayant un certain nombre de points chauds échantillonnés. Pour les strates, les distributions du nombre de points chauds échantillonnés $n_{PC,h,m}^{(j)}$, pour $j=1,\dots,R$, sont examinées par strate. Les résultats sont présentés à la section 4.

3.2 Seconde étude par simulation : effets à plus long terme

La seconde étude par simulation vise à analyser les effets à plus long terme sur le nombre de points chauds échantillonnés et à vérifier si toutes les unités sont prélevées selon les probabilités d'inclusion ciblées. Les quatre mêmes méthodes que dans la première étude par simulation sont examinées (ignorance des points chauds, fractions croissantes, échantillonnage de PSCA et utilisation de sous-strates). Nous appliquons maintenant la coordination d'échantillons décrite à la sous-section 2.2, où le statut de point chaud est mis à jour au fil du temps. La méthode de référence, à savoir l'ignorance des points chauds, équivaut désormais à un échantillonnage indépendant.

Nous simulons une série de tirages au fil du temps, où les périodes représentent les mois $t=1,\dots,T$. Dans chaque mois t , on détermine quelles unités de la population sont des points chauds, et les méthodes sont appliquées de manière à éviter le plus possible un échantillonnage des points chauds. Pour cette étude par simulation, nous prenons en considération toutes les strates de la catégorie de taille 4 avec $0 < f_h < 1$. Ainsi, les échantillons sont prélevés à partir de la population U décrite à la section 3.

Cette étude par simulation ne suit pas le même principe d'échantillon initial donné. Dans ce cas-ci, de nouveaux échantillons sont sélectionnés dans la première période et les suivantes. Les fractions d'échantillonnage de l'échantillon S définissent les probabilités d'inclusion dans le processus d'échantillonnage, où S est l'échantillon donné prélevé à l'aide du SFE. Les méthodes sont mises en œuvre dans l'algorithme ci-dessous, où seule la première étape (la génération des nombres aléatoires) diffère de l'algorithme 1.

Algorithme 2.

1. Générer des nombres aléatoires $r_1, \dots, r_k, \dots, r_N \stackrel{\text{iid}}{\sim} \text{Unif}(0,1)$ et calculer

$$\rho_k = \frac{r_k / (1 - r_k)}{\pi_k / (1 - \pi_k)} \text{ pour toutes les } k \in U.$$

2. Prélever des échantillons S_1, S_2, S_3 et S_4 en appliquant l'étape 2 de l'algorithme 1 pour un r_k , une π_k et une ρ_k donnés.

Dans l'algorithme 2, les méthodes de l'étape 2 sont aussi appliquées par strate. Les échantillons généraux S_1, S_2, S_3 et S_4 sont obtenus en combinant les échantillons prélevés à partir des strates. Dans cette étude

par simulation, le statut de point chaud n'est plus fixe, il est mis à jour tous les mois. Nous utilisons la même définition que celle de l'étude pilote, à savoir qu'une unité de la catégorie de taille 4 est un point chaud lorsqu'elle a été prélevée dans plus de 4 enquêtes au cours des 12 derniers mois.

Pour chaque méthode, nous simulons une série longue de 25 ans, donc $T = 300$ mois. Nous sélectionnons des échantillons $S_1^{(t)}$, $S_2^{(t)}$, $S_3^{(t)}$ et $S_4^{(t)}$ en reprenant les étapes 1 et 2 à chaque mois $t = 1, \dots, T$. Pour chaque échantillon $S_m^{(t)}$, sélectionné à l'aide de la méthode m dans un mois t , la taille de l'échantillon n et le nombre de points chauds échantillonnés n_{PC} sont calculés pour l'échantillon en entier et pour les échantillons par strate h .

$$\begin{aligned} n_m^{(t)} &= n(S_m^{(t)}), & n_{PC,m}^{(t)} &= n_{PC}(S_m^{(t)}), \\ n_{h,m}^{(t)} &= n_h(S_m^{(t)}), & n_{PC,h,m}^{(t)} &= n_{PC,h}(S_m^{(t)}), \end{aligned}$$

pour $m = 1, \dots, 4$, $h = 1, \dots, H$ et $t = 1, \dots, T$.

Afin d'étudier les effets à plus long terme des méthodes d'échantillonnage sur le nombre de points chauds échantillonnés, nous observons la croissance du nombre de points chauds échantillonnés dans la série chronologique pour toute la population et pour les strates séparément, ainsi $n_{PC,m}^{(t)}$ et $n_{PC,h,m}^{(t)}$, pour $t = 1, \dots, T$ et la méthode $m = 1, \dots, 4$.

Nous calculons enfin les fractions d'échantillonnage réalisées pour les sous-populations de points chauds et d'autres unités par strate. Les fractions d'échantillonnage réalisées sont calculées en établissant une moyenne au fil du temps de la façon suivante :

$$\begin{aligned} \hat{f}_{PC,h,m} &= \sum_{t=1}^T n_{PC,h,m}^{(t)} / \sum_{t=1}^T N_{PC,h}^{(t)}, \\ \hat{f}_{pPC,h,m} &= \sum_{t=1}^T n_{pPC,h,m}^{(t)} / \sum_{t=1}^T N_{pPC,h}^{(t)}. \end{aligned}$$

Dans ce cas-ci, $N_{PC,h}^{(t)}$ et $N_{pPC,h}^{(t)}$ sont, respectivement, le nombre de points chauds et d'unités qui ne sont pas des points chauds dans la population au cours du mois t et au sein de la strate h , et $n_{pPC,h,m}^{(t)}$ est le nombre d'unités échantillonnées qui ne sont pas des points chauds au cours du mois t et au sein de la strate h selon la méthode m . Nous calculons des estimations pour les marges de 95 % des fractions d'échantillonnage réalisées en supposant la distribution normale de la façon suivante

$$2 \sqrt{\frac{f_h(1-f_h)}{\sum_{t=1}^T N_{PC,h}^{(t)}}} \quad \text{et} \quad 2 \sqrt{\frac{f_h(1-f_h)}{\sum_{t=1}^T N_{pPC,h}^{(t)}}},$$

pour les sous-populations de points chauds et d'unités qui ne sont pas des points chauds, respectivement. Dans ce cas-ci, $f_h = n_h / N_h$ correspond à la fraction d'échantillonnage ciblée de la strate h . Nous évaluons si les fractions d'échantillonnage ciblées se situent dans les intervalles de confiance à 95 % selon ces marges. Il convient de mentionner que ces marges ne sont qu'une approche et ne rendent pas compte de la nature

discrète des nombres. De plus, ces marges prennent en considération seulement la variation au fil du temps, et non au cours d'un mois donné. Les résultats sont présentés à la section 4.

4. Résultats

La présente section porte sur les résultats des deux études par simulation. Les figures et les tableaux correspondants y sont présentés.

4.1 Première étude par simulation : une situation pratique

Par définition, toutes les méthodes examinées visent à sélectionner des échantillons dont la taille est fixe (sous-section 3.1). Les tailles d'échantillon globales réalisées sont par conséquent égales aux tailles d'échantillon ciblées pour toutes les méthodes et dans toutes les exécutions de simulation.

Le tableau 4.1 montre la distribution du nombre de points chauds échantillonnés pour toutes les méthodes examinées. Par définition, l'ignorance des points chauds ne réduit pas le nombre de points chauds et donne lieu à la présence de 108 points chauds dans chaque échantillon. Augmenter les fractions d'échantillonnage mène au plus petit nombre de points chauds échantillonnés pour chaque exécution de simulation, soit 36. Le nombre de points chauds ne varie pas, puisque l'échantillon initial et les étiquettes de point chaud des unités ne changent pas. L'échantillonnage de PSCA et l'utilisation de sous-strates réduisent le nombre de points chauds dans l'échantillon initial. L'utilisation de sous-strates donne lieu à moins de points chauds (avec le mode 55) que l'échantillonnage de PSCA (avec le mode 72), mais la fourchette s'élargit au fil des exécutions de simulation. Même si la méthode pragmatique de fractions croissantes permet d'obtenir la plus grande réduction, cette méthode présente des inconvénients, que nous aborderons dans la seconde étude par simulation.

La figure 4.1 montre le nombre moyen de points chauds échantillonnés par strate, ajusté au fil des 1 000 exécutions de simulation. La couleur indique la fraction de points chauds $F_{PC,h}$ dans la population. En augmentant les fractions d'échantillonnage, le nombre de points chauds échantillonnés est réduit à 0 dans la majorité des strates. Lors de l'utilisation de sous-strates, le nombre moyen de points chauds échantillonnés dans 26 des 48 strates est près de 0. Avec l'échantillonnage de PSCA, seules 11 strates présentent ce résultat, principalement dans la partie inférieure des strates affichées. En règle générale, nous pouvons affirmer que l'échantillonnage de PSCA donne lieu à moins de points chauds échantillonnés que l'ignorance des points chauds, mais tout de même à plus de points chauds que l'utilisation de sous-strates et les fractions croissantes.

Tableau 4.1

Nombre de points chauds échantillonnés dans 1 000 échantillons de la première étude par simulation (les sommes d'une ligne égalent 1 000)

Nombre de points chauds échantillonnés																	
	36	53	68	69	70	71	72	73		74	75	76					108
Ignorance																1 000	
Fractions croissantes	1 000																
PSCA				3	18	88			223	275	221			131	35	6	
a) Les méthodes « ignorance des points chauds », « fractions croissantes » et « échantillonnage de PSCA ».																	
Nombre de points chauds échantillonnés																	
	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61	62	
Sous-strates	1	3	7	26	60	97	134	145	172	123	102	67	31	20	7	5	
b) Les méthodes « utilisation de sous-strates ».																	

Note : La première colonne contient les méthodes appliquées et les colonnes suivantes, le nombre de points chauds échantillonnés. Les cellules du tableau montrent le nombre d'échantillons sélectionnés qui contiennent le nombre de points chauds échantillonnés de la colonne correspondante. Par exemple, un échantillonnage de Poisson spatialement corrélé adapté (PSCA) permet de sélectionner 3 échantillons contenant chacun 68 points chauds.

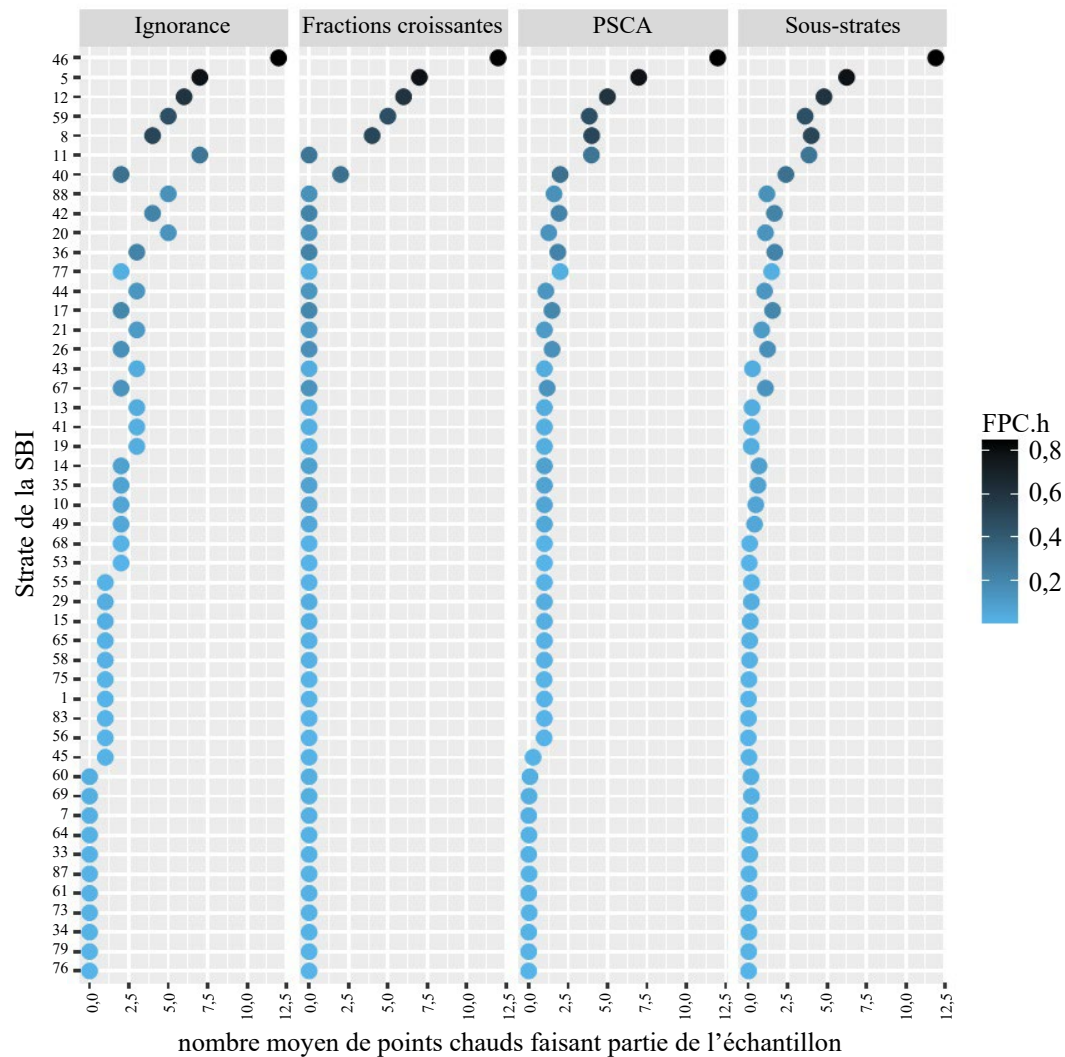
La figure 4.2 a) affiche la réduction moyenne du nombre de points chauds échantillonnés comparativement à la méthode d'ignorance des points chauds. La couleur indique l'espérance du nombre de points chauds échantillonnés $E(n_{PC,h}) = n_h F_{PC,h}$ (section 3). Les plus grandes réductions sont obtenues dans les strates 11, 20 et 88 de la SBI, et ce, pour les trois méthodes. Dans ces strates, l'échantillon initial contient un nombre (relativement) important de points chauds (voir la figure 4.1), et il s'y trouve suffisamment d'unités qui ne sont pas des points chauds dans la population pour y remplacer les points chauds. Pour ces strates, nous avons $N_h - N_{PC,h} \geq n_h$ ou, de façon équivalente, $f_h + F_{PC,h} \leq 1$. Dans la strate 46 de la SBI, le nombre de points chauds dans l'échantillon initial n'est pas réduit par les trois méthodes d'échantillonnage malgré la valeur élevée de $E(n_{PC,h})$. Dans cette strate, nous avons $f_h + F_{PC,h} > 1$, ce qui sous-entend qu'il manque d'espace dans la population pour remplacer les points chauds.

La figure 4.2 b) compare l'échantillonnage de PSCA avec l'utilisation de sous-strates et montre la réduction moyenne du nombre de points chauds échantillonnés comparativement à la méthode de fractions croissantes. La couleur indique l'espérance du nombre de points chauds échantillonnés $E(n_{PC,h}) = n_h F_{PC,h}$. Il s'avère que l'échantillonnage de PSCA mène à moins de points chauds seulement dans les strates 12 et 59 de la SBI. Dans les deux strates, le nombre de points chauds de l'échantillon provisoire est relativement important, mais il manque d'espace dans la population pour remplacer les points chauds par des unités qui ne sont pas des points chauds.

L'utilisation de sous-strates et l'échantillonnage de PSCA entraînent des résultats comparables. Dans la partie centrale des strates affichées dans la figure 4.2 b), l'utilisation de sous-strates donne lieu à une réduction moyenne de près de 0 point chaud, tandis que dans le cas de l'échantillonnage de PSCA, 1 point chaud de plus est sélectionné par rapport à la méthode de fractions croissantes. Dans ces strates, les fractions d'échantillonnage sont petites et la population contient seulement un peu de points chauds, de sorte que $E(n_{PC,h})$ se rapproche de 0. Pour l'utilisation de sous-strates, des échantillons de la strate de points chauds ayant une taille d'échantillon égale au nombre arrondi de façon aléatoire de $E(n_{PC,h})$ (section 3, algorithme 1),

qui est principalement de 0 dans ces strates, sont sélectionnés. Pour l'échantillonnage de PSCA, des échantillons sont sélectionnés dans la strate complète avec la fraction d'échantillonnage ciblée. L'arrondissement est aussi appliqué à l'échantillonnage de PSCA, mais l'effet sur les fractions d'échantillonnage réalisées s'avère plus petit que dans l'approche des sous-strates. Cela s'explique par le fait que l'échantillonnage est appliqué au niveau de l'unité dans l'échantillonnage de PSCA et au niveau de la sous-strate dans l'approche des sous-strates.

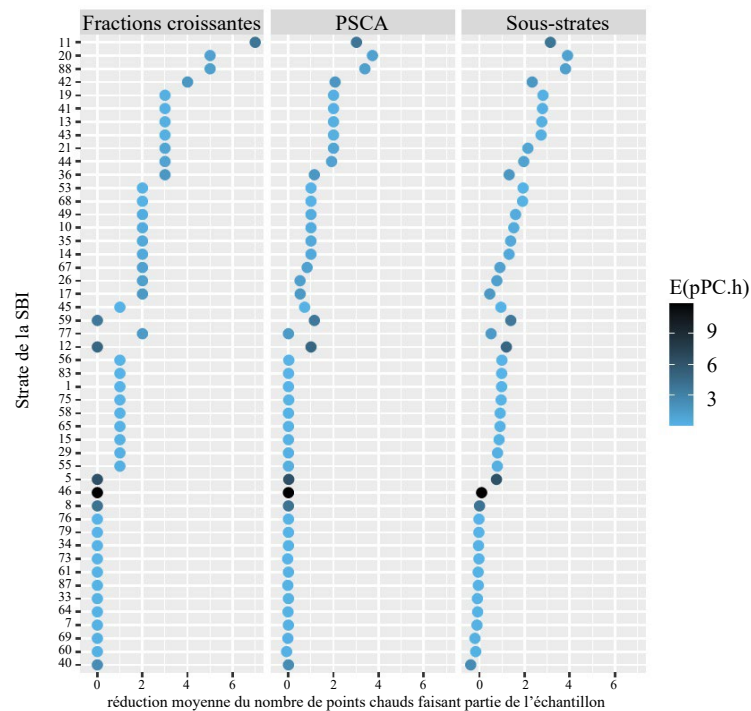
Figure 4.1 Nombre moyen de points chauds échantillonnés par strate dans la première étude par simulation



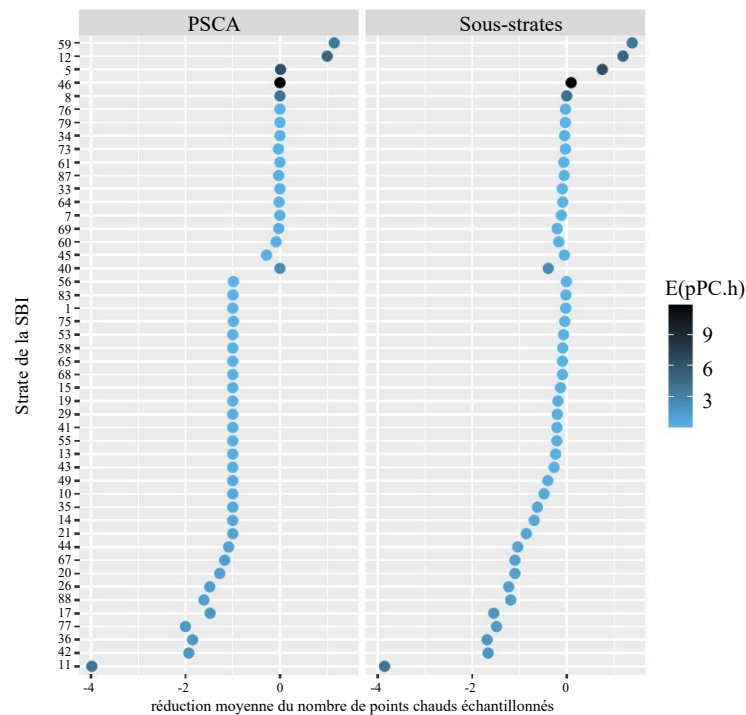
Notes :

- L'axe des x montre le nombre moyen de points chauds échantillonnés et l'axe des y , la classification SBI à 2 chiffres. La couleur indique la fraction de points chauds $F_{PC,h}$ de la population dans la strate h .
- Standaard Bedrijfsindeling (classification type des industries) (SBI).

Figure 4.2 Réduction moyenne des points chauds échantillonnés par strate dans la première étude par simulation



a) Réduction moyenne du nombre de points chauds échantillonnés comparativement à l'ignorance des points chauds.



b) Réduction moyenne du nombre de points chauds échantillonnés comparativement aux fractions croissantes.

Notes :

- L'axe des x montre la réduction moyenne des points chauds échantillonnés et l'axe des y , la classification SBI à 2 chiffres. La couleur indique l'espérance du nombre de points chauds échantillonnés $E(n_{PC,h}) = n_h F_{PC,h}$ dans la strate h .
- Échantillonnage de Poisson spatialement corrélé adapté (PSCA); Standaard Bedrijfsindeling (classification type des industries) (SBI).

En règle générale, nous pouvons affirmer que si la population compte suffisamment d'espace pour remplacer les points chauds, les fractions croissantes donnent lieu à moins de points chauds échantillonnés que l'échantillonnage de PSCA et l'utilisation de sous-strates. Dans ce cas, il est possible d'obtenir une grande réduction lorsque le nombre de points chauds présents dans l'échantillon initial est (relativement) important. L'utilisation de sous-strates et l'échantillonnage de PSCA montrent des résultats comparables. Grâce à un arrondissement, l'utilisation de sous-strates donne lieu à légèrement moins de points chauds dans les strates avec une petite valeur de $E(n_{PC,h})$.

4.2 Seconde étude par simulation : effets à plus long terme

La figure 4.3 a) montre la croissance au fil du temps du nombre total de points chauds échantillonnés pour les quatre méthodes, dans toutes les strates ensemble. Les 12 premiers mois ne sont pas présentés parce que, selon la définition (section 1), les points chauds apparaissent seulement après 12 mois. Le tableau 4.2 a) résume les distributions du nombre de points chauds échantillonnés pour les mois 13 à 300. Le premier élément à remarquer est que le nombre de points chauds générés par l'échantillonnage de PSCA et l'utilisation de sous-strates est au même niveau que celui pour l'ignorance des points chauds avec une moyenne de 185. Une nette réduction d'environ 20 points chauds échantillonnés apparaît lorsque les fractions sont augmentées. L'ignorance des points chauds entraîne le plus grand écart-type, ce qui sous-entend qu'il est possible d'étendre plus efficacement le fardeau de réponse à l'aide des fractions croissantes, de l'échantillonnage de PSCA et de l'utilisation de sous-strates. La meilleure répartition du fardeau de réponse est obtenue à l'aide de l'échantillonnage de PSCA.

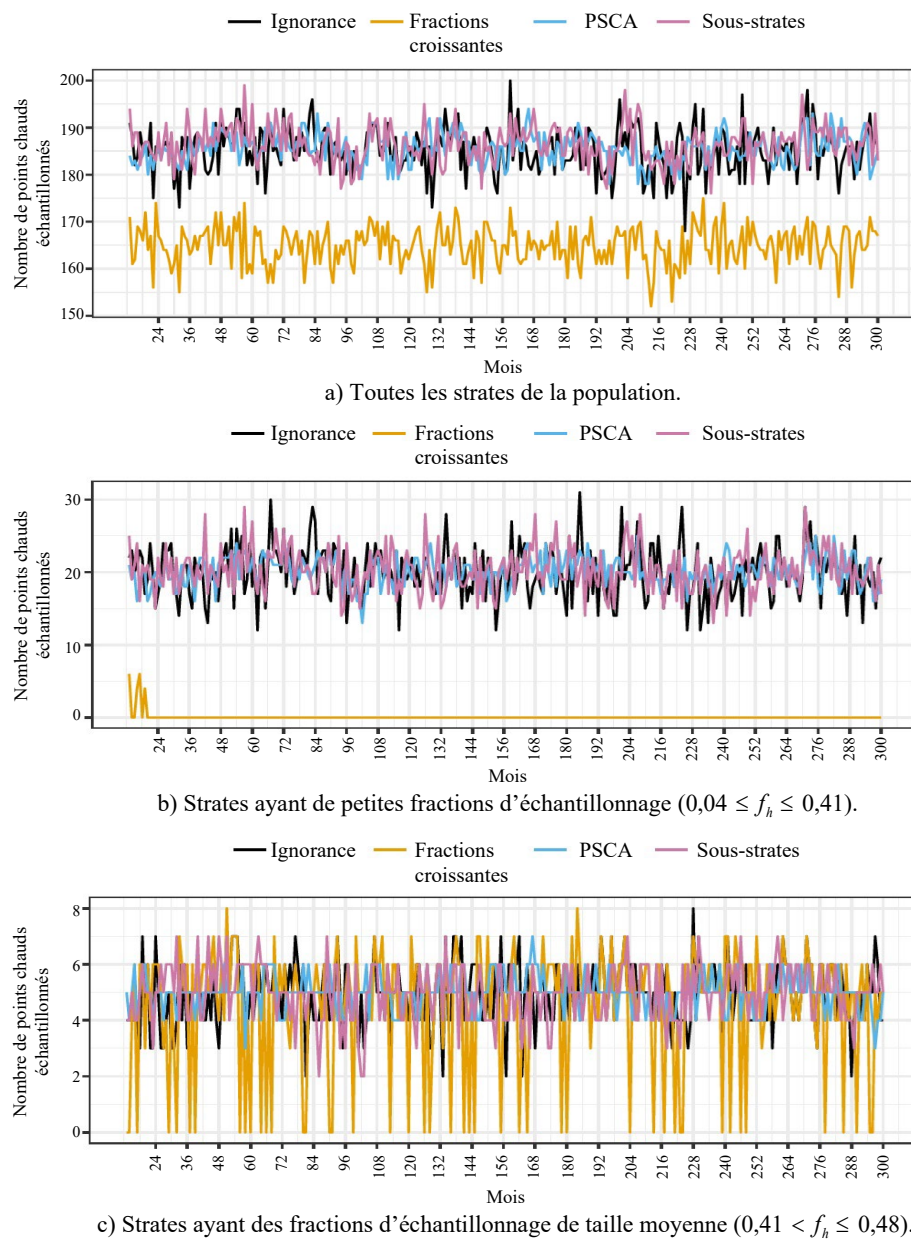
Les figures 4.3 b) et 4.3 c), et les tableaux 4.2 b) et 4.2 c) montrent les résultats dans des sélections de strates ayant, respectivement, des fractions d'échantillonnage de petite taille ($0,04 \leq f_h \leq 0,41$) et de taille moyenne ($0,41 < f_h \leq 0,48$). Pour les strates ayant d'autres fractions d'échantillonnage, les résultats des méthodes sont semblables et, par conséquent, ne sont pas présentés. Lorsque les fractions d'échantillonnage sont très faibles ($0 < f_h < 0,04$), aucun point chaud n'est créé et le nombre de points chauds échantillonnés est toujours de 0 pour chaque méthode. Lorsque les fractions d'échantillonnage sont élevées ($0,48 \leq f_h < 1$), le nombre de points chauds échantillonnés est au même niveau élevé d'environ 160 pour toutes les méthodes. Dans le cas des fractions d'échantillonnage élevées, l'échantillonnage de PSCA permet d'obtenir le plus petit écart-type du nombre de points chauds.

Dans les strates ayant de petites fractions d'échantillonnage ($0,04 \leq f_h \leq 0,41$), la population a assez de place pour que les points chauds soient remplacés chaque mois. La plus forte réduction est par conséquent obtenue en augmentant les fractions, où le nombre de points chauds échantillonnés chaque mois est égal à 0. Le nombre de points chauds échantillonnés à l'aide de l'échantillonnage de PSCA et de l'utilisation de sous-strates se situe au même niveau que lorsque les points chauds sont ignorés, alors ces méthodes ne réduisent pas le nombre de points chauds. Comparativement à l'ignorance des points chauds et à l'utilisation de sous-strates, l'échantillonnage de PSCA donne lieu à la meilleure répartition du fardeau de réponse, c'est-à-dire le plus petit écart-type du nombre de points chauds.

Dans les strates ayant des fractions d'échantillonnage de taille moyenne ($0,41 < f_h \leq 0,48$), la méthode de fractions croissantes montre une tendance intéressante. Le nombre de points chauds échantillonnés est

sensiblement au même niveau pour toutes les méthodes, tandis que les fractions croissantes procurent le plus grand écart-type. En remplaçant les points chauds par des unités qui ne sont pas des points chauds, de plus en plus d'unités de la population deviennent des points chauds. À un certain point, il ne reste plus suffisamment de place pour remplacer les points chauds. Cela entraîne la pire répartition du fardeau de réponse et aucune réduction du nombre de points chauds échantillonnés. De même, dans ces strates, la répartition optimale du fardeau de réponse est obtenue à l'aide de l'échantillonnage de PSCA.

Figure 4.3 Croissance au fil du temps du nombre de points chauds échantillonnés dans la seconde étude par simulation



Note : Échantillonnage de Poisson spatialement corrélé adapté (PSCA).

Tableau 4.2

Résumé du nombre de points chauds échantillonnés pour les mois 13 à 300 dans la seconde étude par simulation, pour toutes les strates de la population et par fraction d'échantillonnage

Méthode	min.	moyenne	max.	écart-type
Ignorance des points chauds	168	185,2	200	4,82
Fractions croissantes	152	164,9	175	4,14
Échantillonnage de PSCA	178	185,6	194	3,25
Utilisation de sous-strates	176	186,4	199	4,06

a) Toutes les strates.

Méthode	min.	moyenne	max.	écart-type
Ignorance des points chauds	12	19,96	31	3,53
Fractions croissantes	0	0,07	6	0,60
Échantillonnage de PSCA	13	20,17	25	2,04
Utilisation de sous-strates	13	20,29	29	3,00

b) Strates ayant de petites fractions d'échantillonnage ($0,04 \leq f_h \leq 0,41$).

Méthode	min.	moyenne	max.	écart-type
Ignorance des points chauds	2	4,92	8	1,14
Fractions croissantes	0	4,51	8	2,19
Échantillonnage de PSCA	3	5,00	7	0,64
Utilisation de sous-strates	2	4,98	7	0,99

c) Strates ayant des fractions d'échantillonnage de taille moyenne ($0,41 < f_h \leq 0,48$).

Note : Poisson spatialement corrélé adapté (PSCA).

Pour le résultat définitif, nous examinons les fractions d'échantillonnage réalisées dans la seconde étude par simulation. Par définition (section 3), les fractions d'échantillonnage réalisées $\hat{f}_{h,m}$ par strate h sont égales aux fractions d'échantillonnage ciblées f_h pour toutes les méthodes m . La figure 4.4 présente les fractions d'échantillonnage réalisées $\hat{f}_{PC,h,m}$ dans la sous-population de points chauds et les intervalles de confiance à 95 % (sous-section 3.2). Les strates de l'axe des x sont classées selon les fractions d'échantillonnage ciblées f_h . Les fractions d'échantillonnage ciblées sont représentées par des points verts. Dans les strates ayant de très petites fractions d'échantillonnage ($0 < f_h < 0,04$), aucun point chaud n'a été créé lors de la simulation. Par conséquent, il est impossible de calculer les fractions d'échantillonnage ou les marges.

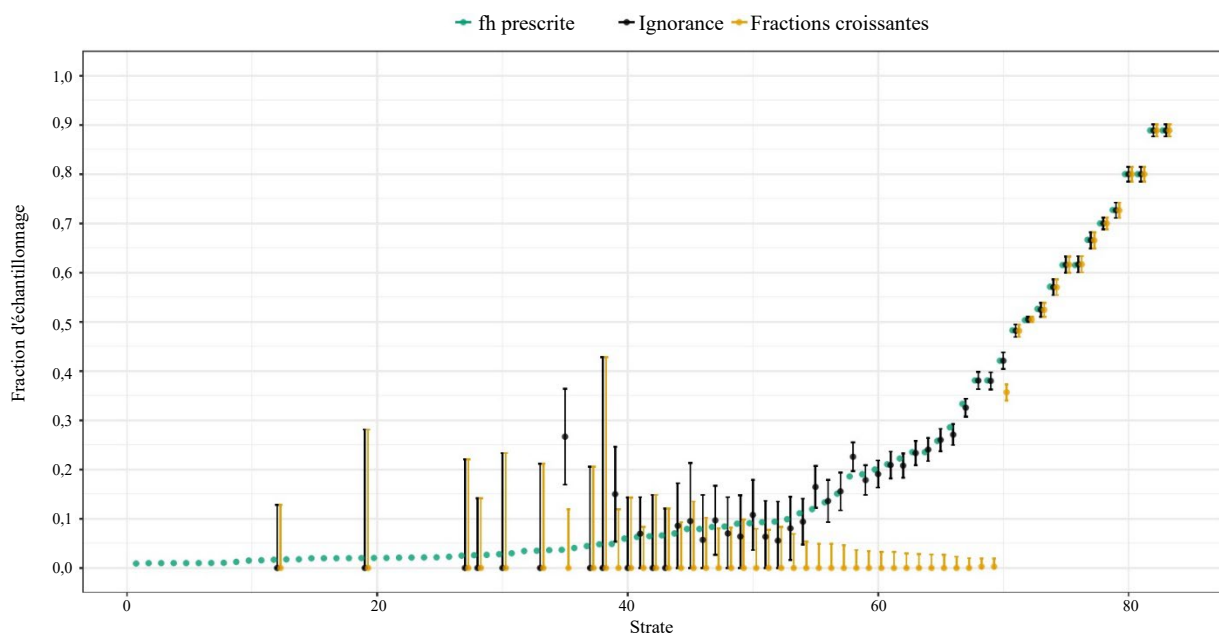
Lorsque les points chauds sont ignorés, les intervalles de confiance à 95 % ne couvrent pas les fractions d'échantillonnage ciblées dans quatre strates (indices 35, 39, 55 et 58). Les fractions réalisées sont trop petites dans 23 strates (indices 47, 48 et 50 à 70) pour permettre une augmentation des fractions. Elles correspondent principalement aux strates ayant des fractions d'échantillonnage de taille moyenne ($0,08 \leq f_h \leq 0,42$). En ce qui concerne l'échantillonnage de PSCA et l'utilisation de sous-strates, seulement une strate ayant une fraction f_h se situe à l'extérieur de l'intervalle de confiance à 95 % (indices 35 et 12, respectivement), ce qui est sans doute attribuable à la chance.

Les résultats obtenus dans la sous-population d'unités qui ne sont pas des points chauds sont semblables et ne sont pas présentés. Une augmentation des fractions d'échantillonnage donne maintenant lieu à des

fractions d'échantillonnage trop élevées, principalement dans les strates ayant des fractions d'échantillonnage de taille moyenne ($0,19 \leq f_h \leq 0,42$).

Nous concluons, pour la seconde étude par simulation, que dans une fourchette précise de petites fractions d'échantillonnage, il est possible d'obtenir une réduction du nombre de points chauds échantillonnés en augmentant les fractions, mais il y a un risque d'introduire une sélectivité dans les échantillons prélevés. Pour les valeurs plus élevées de fractions d'échantillonnage, aucune réduction du nombre de points chauds échantillonnés n'est obtenue par cette méthode, mais il existe un risque d'aggraver la répartition du fardeau de réponse. En règle générale, l'échantillonnage de PSCA et l'utilisation de sous-strates ne réduisent pas le nombre escompté de points chauds échantillonnés, mais donnent lieu à une répartition plus uniforme du fardeau de réponse sans introduire de sélectivité dans les échantillons prélevés. L'échantillonnage de PSCA et l'utilisation de sous-strates procurent des résultats semblables, mais l'échantillonnage de PSCA permet d'obtenir une répartition plus uniforme dans toutes les situations examinées.

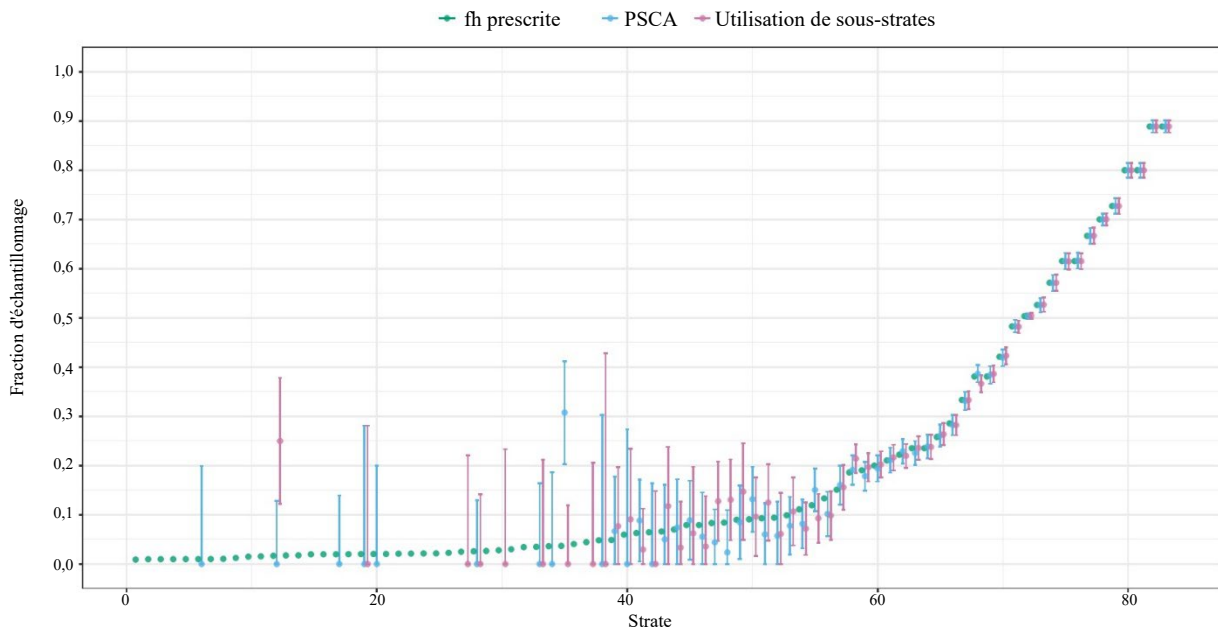
Figure 4.4 Fractions d'échantillonnage réalisées $\hat{f}_{PC,h,m}$ et intervalles de confiance à 95 % dans la seconde étude par simulation, pour les sous-populations de points chauds par strate h et méthode m , et les fractions d'échantillonnage ciblées f_h par strate h



a) Ignorance des points chauds et fractions croissantes.

Note : Poisson spatialement corrélé adapté (PSCA).

Figure 4.4 (suite) Fractions d'échantillonnage réalisées $\hat{f}_{PC,h,m}$ et intervalles de confiance à 95 % dans la seconde étude par simulation, pour les sous-populations de points chauds par strate h et méthode m , et les fractions d'échantillonnage ciblées f_h par strate h



b) Échantillonnage de PSCA et utilisation de sous-strates.

Note : Poisson spatialement corrélé adapté (PSCA).

5. Discussion et conclusion

Le présent article présente une nouvelle méthode de coordination d'échantillons fondée sur l'échantillonnage de Poisson spatialement corrélé adapté (PSCA), axée sur les petites entreprises ayant un lourd fardeau de réponse (les points chauds) et ne faisant pas appel à des nombres aléatoires permanents. Cette méthode (la stratégie d'échantillonnage de PSCA pour les points chauds) est comparée, dans le cadre de deux études par simulation, avec une approche stratifiée (l'utilisation de sous-strates) et une méthode pragmatique dans laquelle des fractions d'échantillonnage sont ajustées manuellement (les fractions croissantes). Toutes les méthodes visent à mieux répartir le fardeau de réponse élevé des petites entreprises. Quels que soient les résultats des simulations dont il est question ci-dessous, nous insistons sur la différence dans l'approche inférentielle entre la méthode pragmatique et la méthode de PSCA. La stratégie d'échantillonnage de PSCA pour les points chauds respecte les probabilités d'inclusion ciblées du plan d'échantillonnage et permet une inférence statistique directe. La méthode pragmatique, en particulier, occasionne des problèmes de sélectivité qu'il faut régler avant de tirer des conclusions statistiques. Dans ce cas-ci, aucune probabilité d'inclusion n'est attribuée aux points chauds, et les probabilités d'inclusion ciblées des autres unités ne sont plus respectées non plus. Les simulations reposent sur l'Enquête sur la recherche et le développement de Statistique Pays-Bas menée en 2020.

La première étude par simulation vise à examiner une situation pratique que rencontrent habituellement les instituts nationaux de statistique (INS) ou d'autres organismes d'enquête. Les méthodes sont appliquées à un échantillon initial donné dans le but d'y remplacer les unités non désirées (les points chauds) par autant d'unités désirées que possible (les unités qui ne sont pas des points chauds). L'étude par simulation montre que dans les strates qui comptent suffisamment d'unités qui ne sont pas des points chauds dans la population, c'est-à-dire lorsque la fraction d'échantillonnage ne dépasse pas la fraction d'unités qui ne sont pas des points chauds dans la population, il est possible de remplacer le plus grand nombre de points chauds en augmentant les fractions, comparativement à l'échantillonnage de PSCA et à l'utilisation de sous-strates. Dans ce cas, si la population compte suffisamment d'unités pour que les points chauds soient remplacés, il est possible d'obtenir une importante réduction proportionnelle du nombre de points chauds.

La seconde étude par simulation vise à analyser les effets à plus long terme des méthodes examinées. Une série de tirages mensuels d'une durée de 25 ans est simulée et dans celle-ci, le statut de point chaud est mis à jour chaque mois pour toutes les entreprises de la population. Cette étude par simulation montre que dans une fourchette précise de petites fractions d'échantillonnage, il est possible d'obtenir une réduction du nombre de points chauds échantillonnés en augmentant les fractions. Cependant, il existe un risque d'introduire une sélectivité dans les échantillons prélevés, c'est-à-dire que des groupes sélectifs d'unités deviennent surreprésentés ou sous-représentés dans l'échantillon, ce qui peut causer des estimations biaisées. Pour d'autres valeurs de fractions d'échantillonnage, aucune réduction du nombre de points chauds échantillonnés n'est obtenue en augmentant les fractions, et il existe un risque d'aggraver la répartition du fardeau de réponse. Cela s'explique par le fait qu'en remplaçant les points chauds par des unités qui ne sont pas des points chauds, plus d'unités de la population deviennent des points chauds. À un certain point, il ne reste plus suffisamment de place pour remplacer les points chauds échantillonnés par des unités qui ne sont pas des points chauds de la population. L'échantillonnage de PSCA et l'utilisation de sous-strates ne réduisent pas le nombre escompté de points chauds échantillonnés, mais permettent d'avoir une répartition plus uniforme du fardeau de réponse sans introduire de sélectivité dans les échantillons prélevés. Dans cette étude par simulation, l'échantillonnage de PSCA et l'utilisation de sous-strates procurent des résultats semblables en ce qui a trait au nombre de points chauds échantillonnés, mais l'échantillonnage de PSCA donne lieu à la meilleure répartition dans toutes les situations examinées.

En règle générale, nous concluons que la méthode pragmatique d'augmentation des fractions d'échantillonnage peut réduire substantiellement le nombre d'unités non désirées dans un échantillon donné, lorsque la population compte suffisamment d'autres unités. Cependant, à long terme, rien ne garantit que le nombre de points chauds demeurera à un faible niveau avec l'application de cette méthode. De plus, il existe un risque de créer une sélectivité dans les échantillons et de l'accroître en appliquant une coordination d'échantillons. Il est par conséquent conseillé d'utiliser cette méthode seulement pour les strates ayant de petites fractions d'échantillonnage, afin de remplacer seulement quelques points chauds par d'autres unités et de limiter le risque d'introduire un biais. La stratégie d'échantillonnage de PSCA pour les points chauds montre qu'elle permet effectivement de réduire la variance du nombre de points chauds échantillonnés et d'éviter la présence d'un grand nombre de points chauds dans l'échantillon. L'utilisation de sous-strates a donné des

résultats semblables. Bien que la stratégie d'échantillonnage de PSCA soit plus généralement applicable et qu'elle permette d'éviter les problèmes d'arrondissement comparativement à la stratification, la méthode consistant à utiliser des sous-strates pourrait être plus facile à implanter dans un système existant de coordination d'échantillons utilisé par un INS ou un organisme d'enquête. Ainsi, dans le cadre d'enquêtes-entreprises reposant sur un plan d'échantillonnage stratifié, l'utilisation de sous-strates se révélerait une bonne solution de rechange. L'utilisation de sous-strates pourrait présenter l'inconvénient, dans le cas d'une stratification détaillée, de ne plus permettre de réaliser des probabilités d'inclusion ciblées en raison de l'arrondissement. Nous n'avons pas rencontré ce problème dans les études par simulation.

Nous concluons le présent article par quelques remarques et recommandations. Précisons d'abord que les études par simulation présentées reposent sur une enquête-entreprise menée par Statistique Pays-Bas. Il est toutefois possible de généraliser les résultats à d'autres enquêtes-entreprises reposant sur un plan d'échantillonnage stratifié et lorsque des échantillons successifs sont prélevés pour des enquêtes ayant différentes stratifications. Dans le dernier cas, une stratification de base pourrait se définir comme étant un croisement des strates d'enquêtes distinctes. Ensuite, lorsqu'il faut procéder à une réduction structurelle du nombre d'unités non désirées, comme les points chauds, les fractions d'échantillonnage des enquêtes doivent être ajustées. Il est en outre conseillé de mettre régulièrement à jour les attributions et les stratifications des enquêtes afin de pouvoir utiliser la taille d'échantillon requise le plus efficacement possible. Enfin, à Statistique Pays-Bas, le statut de point chaud des entreprises sert à mesurer le fardeau de réponse dans les études par simulation. Le statut de point chaud repose sur le nombre d'enquêtes dans lesquelles une entreprise a été sélectionnée au cours des 12 derniers mois. D'autres mesures du fardeau de réponse, comme les facteurs de pondération fondés sur le temps nécessaire pour répondre au questionnaire, ne mèneraient fort probablement pas à d'autres conclusions.

Avis de non-responsabilité

Les points de vue exprimés dans le présent article sont ceux des auteurs et ne reflètent pas nécessairement les politiques de Statistique Pays-Bas.

Remerciements

Les auteurs tiennent à remercier deux examinateurs anonymes et un rédacteur associé pour leur lecture attentive du texte et les commentaires judicieux qu'ils ont formulés sur ses anciennes versions.

Bibliographie

Bavdaž, M., Giesen, D., Černe, S.K., Loöfgren, T. et Raymond-Blaess, V. (2015). Response burden in official business surveys: Measurement and reduction practices of national statistical institutes. *Journal of Official Statistics*, 31(4), 559-588.

- Bondesson, L. et Thorburn, D. (2008). A list sequential sampling method suitable for real-time sampling. *Scandinavian Journal of Statistics*, 35, 466-483.
- Bottone, M., Modugno, L. et Neri, A. (2021). Response burden and data quality in business surveys. *Journal of Official Statistics*, 37(4), 811-836.
- Bradburn, N. (1978). Respondent burden. *Proceedings of the Survey Research Methods Section*, American Statistical Association, 35, 35-40. Alexandria, Virginie.
- Brewer, K., Early, I. et Joyce, S. (1972). Selecting several samples from a single population. *Australian Journal of Statistics*, 3, 231-239.
- Erikson, J., Giesen, D., Houben, L. et Saraiva, P. (2023). Managing response burden for official statistics business surveys – Experiences and recent developments at Statistics Netherlands, Statistics Portugal, and Statistics Sweden. Dans *Advances in Business Statistics, Methods and Data Collection*, (Éds., G. Snijkers, M. Bavdaž, S. Bender, J. Jones, S. MacFeely, J.W. Sakshaug, K.J. Thompson et A. van Delden), Hoboken: John Wiley & Sons, Inc., 193-223.
- Ernst, I. (1999). The maximization and minimization of sample overlap problems: A half century of results. *Proceedings of the International Statistical Institute*, 52, 168-182. Helsinki, Finlande.
- Giesen, D., Vella, M., Brady, C.F., Brown, P., Ravindra, D. et Vaasen-Otten, A. (2018). Response burden management for establishment surveys at four national statistical institutes. *Journal of Official Statistics*, 34(2), 397-418.
- Gommans, G., Burgerjon, L. et Paulissen, R. (2022). Hotspots. Internal CBS memo, Statistics Netherlands.
- Grafström, A. (2012). Spatially correlated Poisson sampling. *Journal of Statistical Planning and Inference*, 142(1), 139-147.
- Grafström, A., Lisic, J. et Prentius, W. (2024). *BalancedSampling: Balanced and Spatially Balanced Sampling*. R package version 2.1.1. Accessible à l'adresse : <https://CRAN.R-project.org/package=BalancedSampling>.
- Lorenc, B., Kloek, W., Abrahamsson, L. et Eckman, S. (2013). Description of actual response burden and response behaviour with swedish register. Dans *Comparative Report on Integration of Case Study Results Related to Reduction of Response Burden and Motivation of Businesses for Accurate Reporting*, (Éds., D. Giesen, M. Bavdaž et I. Bolko). *BLUE-ETS deliverable 8.1*, 11-27.

- Matei, A. et Smith, P.A. (2023). Sample coordination methods and systems for establishment surveys. Dans *Advances in Business Statistics, Methods and Data Collection*, (Éds., G. Snijkers, M. Bavdaž, S. Bender, J. Jones, S. MacFeely, J.W. Sakshaug, K.J. Thompson et A. van Delden), Chichester, Angleterre, Royaume-Uni: John Wiley & Sons, Inc., 637-657.
- Matei, A., Smith, P.A., Smeets, M.J.E. et Klingwort, J. (2023). [Double contrôle ciblé du fardeau dans de multiples enquêtes](https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2023002/article/00010-fra.pdf). *Techniques d'enquête*, 49(2), 391-416. Accessible à l'adresse : <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2023002/article/00010-fra.pdf>.
- Ohlsson, E. (1995). Coordination of samples using permanent random numbers. Dans *Business Survey Methods*, (Éds., B.G. Cox, D. Binder, B. Chinnapa, A. Christianson, M. Colledge et P. Kott), New York: John Wiley & Sons, Inc., États-Unis, chapitre 9, 153-169.
- Rosén, B. (1997a). Asymptotic theory for order sampling. *Journal of Statistical Planning and Inference*, 62, 135-158.
- Rosén, B. (1997b). On sampling with probability proportional to size. *Journal of Statistical Planning and Inference*, 62, 159-191.
- Smeets, M. et Boonstra, H.J. (2018). Sampling coordination of business surveys at Statistics Netherlands. Dans *The Unit Problem and Other Current Topics in Business Survey Methodology*, (Éds., B. Lorenc, P. Smith, M. Bavdaž, G. Haraldsen, D. Nedyalkova, L.-C. Zhang et T. Zimmermann), Cambridge Scholars, Newcastle-upon-Tyne, 127-137.
- Smeets, M. et Klingwort, J. (2023). Sampling methods to better spread the response burden for small businesses. Document de discussion du CBS, Statistics Netherlands. Accessible à l'adresse : <https://www.cbs.nl/en-gb/background/2023/38/sampling-methods-to-better-spread-the-response-burden-for-small-businesses>.
- Statistique Pays-Bas (2022). Onderzoeksomschrijvingen: Research and development. Accessible à l'adresse : <https://www.cbs.nl/nl-nl/onze-diensten/methoden/onderzoeksomschrijvingen/korte-onderzoeksomschrijvingen/research-development>.
- Strategische commissie betere regelgeving bedrijven (2020). Onevenredige enquêtebelasting van specifieke ondernemers in het MKB. Note interne du CBS, Statistics Netherlands.
- Yan, T., Fricker, S. et Tsai, S. (2019). Response burden: What is it and what predicts it? Dans *Advances in Questionnaire Design, Development, Evaluation and Testing*, (Éds., P. Betty, D. Collins, L. Kaye, J.L. Padilla, G. Willis et A. Wilmot), Hoboken: John Wiley & Sons, Inc., 193-212.
- Yan, T. et Williams, D. (2022). Response burden – Review and conceptual framework. *Journal of Official Statistics*, 38(4), 939-961.