



La déclaration RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Health Data) : directives pour la communication des études réalisées à partir de données de santé collectées en routine

Eric Benchimol, Liam Smeeth, Astrid Guttman, Katie Harron, David Moher, Irene Petersen, Henrik Sørensen, Jean-Marie Januel, Erik von Elm, Sinéad Langan

► To cite this version:

Eric Benchimol, Liam Smeeth, Astrid Guttman, Katie Harron, David Moher, et al.. La déclaration RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Health Data) : directives pour la communication des études réalisées à partir de données de santé collectées en routine. CMAJ, 2019, 191 (8), pp.E216-E230. 10.1503/cmaj.181309 . hal-03703553

HAL Id: hal-03703553

<https://ehesp.hal.science/hal-03703553>

Submitted on 24 Jun 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

La déclaration RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Health Data) : directives pour la communication des études réalisées à partir de données de santé collectées en routine

Eric I. Benchimol MD PhD, Liam Smeeth MBChB PhD, Astrid Guttman MDCM MSc, Katie Harron PhD, David Moher PhD, Irene Petersen PhD, Henrik T. Sørensen MD PhD, Jean-Marie Januel PhD MPH, Erik von Elm MD MSc, Sinéad M. Langan MSc PhD; pour le comité de travail RECORD*

■ Citation : *CMAJ* 2019 February 25;191:E216-30. doi: 10.1503/cmaj.181309

La disponibilité croissante des données générées par les prestataires de soins de santé ainsi que par la surveillance de l'incidence des maladies et des indicateurs de l'évolution des maladies a transformé le caractère de la recherche en santé. Les données de santé collectées en routine sont définies comme étant des données collectées sans que des questions de recherche spécifiques ne soient développées de manière a priori¹. Ces données représentent pourtant de nombreuses ressources pour la recherche (p. ex., registres des maladies), la gestion clinique (p. ex., bases de données de soins primaires), la planification du système de santé (p. ex., données médico-administratives), la documentation clinique des soins (p. ex., dossiers médicaux électroniques) ou la surveillance épidémiologique (p. ex., registres des cancers et données des rapports de santé publique). Ces données, recueillies dans divers contextes de soins de santé et de lieux géographiques, offrent des opportunités pour développer des recherches innovantes, efficaces et rentables pour éclairer les décisions cliniques en médecine, pour la planification des services de santé et en santé publique². À l'échelle internationale, les gouvernements et les instances de financement de la santé ont priorisé l'utilisation de ces données pour améliorer les soins offerts aux patients, transformer la recherche en santé et augmenter l'efficacité des soins³.

Malgré le fait que l'explosion de la disponibilité de ces données puisse offrir des opportunités importantes pour répondre à des questions de recherche urgentes, elles posent cependant des défis à ceux qui mènent et évaluent la recherche et mettent en œuvre des décisions issues des conclusions à partir d'études utilisant ces données. Les très nombreuses sources de données de santé collectées en routine et l'expansion rapide du champ de ces données rendent difficile l'identification de leurs avantages et de

leurs limites, mais aussi des biais associés aux sources de données individuelles. Un rapport incomplet ou inadéquat des études observationnelles basées sur des données collectées de manière routinière exacerbe les défis liés à l'utilisation de ces données. Une analyse systématique d'un échantillon d'études utilisant des sources de données collectées de façon routinière a identifié de nombreux domaines où le rapport est incomplet ou peu clair⁴. Les lacunes comprennent des informations inadéquates ou manquantes concernant le codage des facteurs d'expositions et des résultats, ainsi que dans les détails concernant la concordance des informations entre 2 sources de données différentes. Deux revues systématiques ont montré également des résultats médiocres concernant la validation des données de routine^{5,6}, qui peuvent masquer différents types de biais, entraver les efforts pour réaliser des méta-analyses à partir de ces données, et conduire à des conclusions erronées. Les définitions des populations qui sont pertinentes dans les études utilisant des données collectées de manière routinière sont présentées dans l'encadré 1.

Des directives pour la communication des études réalisées à partir de données de santé collectées en routine ont été élaborées pour guider un processus qui couvre un large éventail de points clés allant de la conception au contexte de telles études, afin d'établir la qualité de ces dernières^{8,9}. En particulier, la déclaration STROBE (Strengthening the Reporting of Observational Studies in Epidemiology) a été élaborée pour améliorer la transparence dans les études observationnelles^{10,11} et a été largement adoptée et approuvée par la plupart des journaux scientifiques médicaux. Il a été démontré que la mise en œuvre de la déclaration STROBE dans le processus éditorial améliorait la qualité des études publiées^{12,13}. La plupart des études conduites à partir de données collectées en routine répondent aux

caractéristiques des études observationnelles, et par conséquent, les directives proposées par la déclaration STROBE sont pertinentes et applicables. Cependant, étant donné que la déclaration STROBE a été conçue pour être appliquée à toutes les études observationnelles, les problèmes spécifiques liés aux données collectées de façon routinière ne sont pas abordés dans celle-ci. Un groupe international de scientifiques ayant un intérêt particulier pour l'utilisation des données de santé collectées en routine et des représentants du groupe STROBE se sont rencontrés après le « Primary Care Database Symposium » à Londres en 2012 pour parler de STROBE dans le cadre d'études utilisant des données de santé collectées en routine^{14,15}. Les scientifiques ont identifié d'importantes lacunes dans STROBE pour évaluer les études utilisant ce type de données et ont suggéré qu'une exten-

sion de STROBE était nécessaire. Ainsi, l'initiative RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Data) a été établie sur la base d'un processus collaboratif international ayant abouti à une extension de STROBE pour explorer et adresser les problèmes spécifiques à la communication des études réalisées à partir de données collectées en routine. L'initiative RECORD a impliqué plus de 100 parties prenantes au niveau international comprenant des chercheurs, des rédacteurs de journaux et des utilisateurs de ces données, y compris ceux qui utilisent les résultats des études réalisées à partir de données de routine pour guider leur prise de décision. La méthodologie utilisée pour élaborer les directives RECORD a été détaillée par Nicholls et collaborateurs¹⁶. Les directives RECORD sont fondées sur des méthodes reconnues pour élaborer des lignes directrices dans le rapport qualitatif d'études¹⁷. En bref, les parties prenantes ont été interrogées à 2 reprises pour établir et classer par ordre de priorité les thèmes à inclure dans les directives RECORD. Un comité de travail s'est ensuite réuni en présentiel pour rédiger les directives finales. Les intervenants du groupe de travail ont examiné chaque directive et fait des commentaires. La liste finale des items d'évaluation et le présent document d'explication ont été rédigés par les membres du comité de pilotage de RECORD, avec examen et approbation par le groupe de travail. Les membres du comité de pilotage de STROBE ont été impliqués dans la création de RECORD.

Conformément à l'approche STROBE, les directrices RECORD ne sont pas conçues pour recommander des méthodes de recherche, mais pour améliorer la présentation de la recherche afin d'assurer les lecteurs, les pairs évaluateurs, les rédacteurs de journaux et toutes autres personnes d'un moyen d'évaluation de la validité interne et externe des études utilisant des données de santé collectées en routine. En améliorant la qualité des informations rapportées dans de telles études, nous avons cherché à réduire les ambiguïtés et à respecter les principes du processus scientifique : la découverte, la transparence et la reproductibilité¹⁸.

Encadré 1 : Définitions des termes de la population (population source, population de la base de données et population étudiée)

Il existe 3 niveaux dans la hiérarchie des populations qui sont pertinents dans les études utilisant des données collectées de manière routinière et qui seront mentionnés tout au long du document. Ces populations comprennent la population source, qui représente celle à partir de laquelle la population de la base de données est dérivée et donc de laquelle les chercheurs relèvent des déductions, la population de la base de données, qui est dérivée de la population source et comprend des personnes avec des données incluses dans la source de données, et la population étudiée, identifiée à partir de la population de la base de données par les chercheurs à l'aide de codes et d'algorithmes (figure 1)⁷. Par exemple, dans le cas du « Clinical Practice Research Datalink » (CPRD), la population source comprend toutes les personnes qui fréquentent des médecins généralistes au Royaume-Uni. La population de la base de données comprend les personnes incluses dans le CPRD, alors que la population étudiée comprend celles choisies du CPRD en utilisant des codes et des algorithmes qui seront décrits dans l'étude spécifique.



Figure 1 : Hiérarchie des populations dans les études menées en utilisant des données collectées de manière routinière.

Ce travail a été approuvé par le Comité d'éthique de la recherche du Centre hospitalier pour enfants de l'est de l'Ontario.

Les éléments de la liste de contrôle RECORD

La liste des items RECORD est fournie dans le tableau 1. Comme RECORD est fondé sur une extension de STROBE, les énoncés spécifiques à RECORD sont présentés dans une deuxième colonne à celle de STROBE et organisés pour chaque section d'un manuscrit. Nous conseillons aux auteurs d'aborder de manière adéquate chaque item de la liste, mais ne fixons pas d'ordre ou d'emplacement précis dans le manuscrit. Nous incluons ci-dessous un texte explicatif pour chaque item de la liste de contrôle RECORD, selon les sections d'un manuscrit. Dans le cas où aucun élément supplémentaire n'est requis en plus de ceux élaborés dans STROBE dans les études utilisant des données de santé collectées en routine, les critères d'évaluation sont ceux fournis par l'item STROBE respectif.

Titre et résumé

- Item RECORD 1.1: Le type de données utilisées devrait être spécifié dans le titre ou le résumé. Lorsque cela est possible, le nom des bases de données utilisées devrait être mentionné.
- Item RECORD 1.2: Le cas échéant, les contextes géographique et temporel dans lesquels l'étude a été réalisée devraient être indiqués dans le titre ou le résumé.
- Item RECORD 1.3: Si un chaînage entre différentes bases de données a été réalisé pour l'étude, cela devrait être clairement indiqué dans le titre ou le résumé.

Exemples

Les articles énumérés ci-dessous fournissent 2 exemples d'une bonne documentation de cette section :

1. « Perforations and haemorrhages after colonoscopy in 2010: a study based on comprehensive French health insurance data (SNIIRAM) »¹⁹.
2. « The Dutch hospital standardised mortality ratio (HSMR) method and cardiac surgery: benchmarking in a national cohort using hospital administration data versus a clinical database »²⁰.

Explication

Comme il n'existe pas de termes « MeSH » (Medical Subject Heading) permettant l'identification des études réalisées à partir de données de santé collectées en routine, il est important de pouvoir identifier ce type d'étude. Compte tenu de la grande diversité des types de données, il ne suffit pas de documenter le fait que des données de routine ont été utilisées. Le type de données de routine devrait être spécifié dans le titre et/ou le résumé. Les exemples de type de données comprennent entre autres les données médico-administratives, les autres données administratives (assurance, registres de naissance/décès et emploi), les registres concernant des maladies, les données de soins primaires, les données issues des dossiers médicaux électroniques ainsi que les registres de population. Le fait de nommer les bases de données utilisées est impor-

tant mais ne remplace pas l'indication du type de source de données dans le titre ou le résumé.

La région géographique et la période de production des données (contexte temporel) sont des critères déjà présents dans STROBE. Nous suggérons que ces informations soient également documentées dans les sections du titre ou du résumé des manuscrits utilisant la liste de contrôle RECORD. La présentation des informations fondamentales concernant la région géographique et la période de production des données devraient évidemment respecter les limites du nombre de mots et tenir compte des questions liées à la confidentialité des données. La région devrait être identifiée au moins par le plus haut niveau géographique pour définir la population étudiée (p. ex., pays, état, province ou région, département ou canton).

De plus, si un chaînage entre plusieurs bases de données a été réalisé, cela devrait être mentionné dans le titre ou le résumé. Les exemples de termes permettant d'indiquer ce chaînage comprennent par exemple « utilisant plusieurs bases de données médico-administratives chaînées » ou « [nom de la base de données] chaînée à [nom de la base de données] ». Les mots « chaîné » ou « chaînage » fournissent suffisamment d'informations dans le titre ou le résumé; les détails supplémentaires sur la méthodologie du chaînage devraient être fournis dans la section Méthodes du manuscrit.

Introduction

Aucun critère spécifique aux directives RECORD n'est nécessaire en plus de ceux établis dans STROBE. Les directives de STROBE recommandent que des objectifs précis, y compris des hypothèses prédéfinies, soient énoncés dans la section Introduction. L'énoncé des objectifs de recherche spécifiques est essentiel pour la réplication et la traduction de toute recherche observationnelle. Pour les études utilisant des données collectées en routine, les auteurs devraient clarifier si les analyses étaient exploratoires dans le but d'identifier de nouvelles relations potentielles dans les données collectées (p. ex., études génératrices d'hypothèses^{21,22}) ou confirmatives dans le but de tester une ou plusieurs hypothèses²³. Les auteurs devraient indiquer si leurs hypothèses ont été établies avant ou après l'analyse des données. Ils devraient également indiquer clairement s'il existait un protocole d'étude, comment on peut y accéder le cas échéant, et si l'étude a été enregistrée dans un registre d'étude accessible au public. Étant donné que les points forts et les limites des méthodes utilisées pour effectuer des recherches avec des données collectées en routine peuvent être controversées, une description claire des objectifs de l'étude est essentielle^{23,24}. Il est néanmoins insuffisant de qualifier simplement une étude de descriptive sans préciser si elle vise à établir ou à examiner une hypothèse.

Méthodes (contexte)

Aucun critère RECORD supplémentaire à ceux déjà présents dans les directives STROBE n'est nécessaire pour évaluer la description du contexte, des lieux et des dates pertinentes, y compris les périodes de recrutement de la population d'étude, d'exposition, de suivi et de collecte de données. Les auteurs

Tableau 1 (partie 1 sur 2) : La déclaration RECORD : liste de contrôle, étendue de la déclaration STROBE, des éléments qui devraient être rapportés dans les études observationnelles utilisant des données de santé collectées de manière routinière*

Section du manuscrit	N° d'item	Item STROBE	Item RECORD
Titre et résumé			
	1	(a) Indiquer dans le titre ou le résumé le plan d'étude avec un terme couramment utilisé. (b) Fournir un résumé informatif et équilibré de ce qui a été fait et ce qui a été trouvé.	RECORD 1.1 : Le type de données utilisées devrait être spécifié dans le titre ou le résumé. Lorsque cela est possible, le nom des bases de données utilisées devrait être mentionné. RECORD 1.2 : Le cas échéant, les contextes géographique et temporel dans lesquels l'étude a été réalisée devraient être indiqués dans le titre ou le résumé. RECORD 1.3 : Si un chaînage entre différentes bases de données a été réalisé pour l'étude, cela devrait être clairement indiqué dans le titre ou le résumé.
Introduction			
Contexte/ justification	2	Expliquer le contexte scientifique et la logique sous-jacente dans l'étude rapportée.	
Objectifs	3	Énoncer les objectifs spécifiques, y compris toutes les hypothèses préétablies.	
Méthodes			
Plan d'étude	4	Présenter les éléments clés du plan d'étude au début du document.	
Contexte	5	Décrire le contexte, les lieux et les dates pertinentes, y compris les périodes de recrutement, d'exposition, de suivi et de collecte de données.	
Participants	6	(a) <i>Étude de cohorte</i> — Donner les critères d'admissibilité ainsi que les sources et les méthodes de sélection des participants. Décrire les méthodes de suivi. <i>Étude cas-témoins</i> — Donner les critères d'admissibilité, ainsi que les sources et les méthodes de détermination des cas et de sélection des témoins. Donner la justification du choix des cas et des témoins. <i>Étude transversale</i> — Donner les critères d'admissibilité, ainsi que les sources et les méthodes de sélection des participants. (b) <i>Étude de cohorte</i> — Pour les études appariées, indiquer les critères d'appariement et le nombre de participants exposés et non exposés. <i>Étude cas-témoins</i> — Pour les études appariées, donner les critères d'appariement et le nombre de témoins par cas.	RECORD 6.1 : Les méthodes de sélection de la population étudiée (telles que les codes ou les algorithmes utilisés pour identifier les participants) devraient être énumérées en détail. Si cela n'est pas possible, une explication devrait être fournie. RECORD 6.2 : Toute étude de validation des codes ou des algorithmes utilisés pour sélectionner la population devrait être mise en référence. Si la validation a été effectuée pour cette étude en particulier et n'a pas été publiée par ailleurs, le détail des méthodes utilisées et des résultats de cette validation des données devraient être fournis dans le manuscrit. RECORD 6.3 : Si l'étude impliquait le chaînage de plusieurs bases de données, les auteurs devraient envisager l'utilisation d'un diagramme de flux ou d'une autre représentation graphique pour décrire le processus de chaînage des données, y compris le nombre de personnes avec les données chaînées à l'issue de chaque étape du chaînage.
Variables	7	Définir clairement tous les critères d'évaluation, les expositions, les prédictors, les facteurs confondants potentiels et les facteurs d'influence. Donner les critères de diagnostic, le cas échéant.	RECORD 7.1 : Une liste complète des codes et des algorithmes utilisés pour classer les expositions, les résultats, les facteurs confondants et les modificateurs d'effets devrait être fournie. Si ces codes ou algorithmes ne peuvent pas être documentés, une explication devrait être fournie.
Sources de données/mesures	8	Pour chaque variable d'intérêt, donner les sources de données et les détails des méthodes d'évaluation (mesure). Décrire la comparabilité des méthodes d'évaluation s'il y a plus d'un groupe.	
Biais	9	Décrire les efforts déployés pour adresser les sources potentielles de biais.	
Taille de l'étude	10	Décrire comment la taille de l'échantillon dans l'étude a été déterminée.	
Variables quantitatives	11	Expliquer comment les variables quantitatives ont été traitées dans les analyses. Décrire, s'il y a lieu, quels groupements ont été choisis et pourquoi.	
Méthodes statistiques	12	(a) Décrire toutes les méthodes statistiques, y compris celles utilisées pour contrôler les facteurs confondants. (b) Décrire toutes les méthodes utilisées pour examiner les sous-groupes et les interactions. (c) Expliquer comment les données manquantes ont été adressées. (d) <i>Étude de cohorte</i> — Le cas échéant, expliquer comment la perte de suivi a été adressée. <i>Étude cas-témoins</i> — Le cas échéant, expliquer comment l'appariement des cas et des témoins a été réalisé. <i>Étude transversale</i> — Le cas échéant, décrire les méthodes d'analyse tenant compte de la stratégie d'échantillonnage. (e) Décrire toutes les analyses de sensibilité.	

Tableau 1 (partie 2 sur 2) : La déclaration RECORD : liste de contrôle, étendue de la déclaration STROBE, des éléments qui devraient être rapportés dans les études observationnelles utilisant des données de santé collectées de manière routinière

Section du manuscrit	N° d'item	Item STROBE	Item RECORD
Accès aux données et méthodes de traitement	SO		RECORD 12.1 : Les auteurs devraient décrire à quel degré les chercheurs ont eu accès à la base de données utilisée pour créer la population étudiée. RECORD 12.2 : Les auteurs devraient fournir des informations sur les méthodes de traitement des données utilisées dans l'étude.
Chaînage des données	SO		RECORD 12.3 : Indiquer si l'étude a inclus un chaînage de données au niveau de la personne, au niveau de l'établissement ou d'autres données dans 2 bases de données ou plus. Les méthodes de chaînage et d'évaluation de la qualité de celui-ci devraient être fournies.
Résultats			
Participants	13	(a) Indiquer le nombre de personnes à chaque étape de l'étude (p. ex., les nombres potentiellement admissibles, examinés pour l'admissibilité, confirmés admissibles, inclus dans l'étude, ayant complété le suivi et dont les données ont été analysées). (b) Indiquer les raisons de la non-participation à chaque étape. (c) Envisager l'utilisation d'un diagramme de flux.	RECORD 13.1 : Décrire en détail la sélection des personnes incluses dans l'étude (c.-à-d., la sélection de la population étudiée), y compris le filtrage basé sur la qualité des données, la disponibilité des données et le chaînage. La sélection des personnes incluses peut être décrite dans le texte et/ou au moyen d'un diagramme de flux.
Données descriptives	14	(a) Indiquer les caractéristiques des participants à l'étude (p. ex., démographiques, cliniques, sociales) et les informations sur les expositions et les facteurs confondants potentiels. (b) Indiquer le nombre de participants avec des données manquantes pour chaque variable d'intérêt. (c) <i>Étude de cohorte</i> — Résumer le temps de suivi (p. ex., la moyenne et le temps total).	
Données de résultats	15	<i>Étude de cohorte</i> — Rapporter le nombre d'événements ou les indicateurs mesurés au fil du temps. <i>Étude cas-témoins</i> — Rapporter le nombre de participants dans chaque catégorie d'exposition, ou les indicateurs du niveau d'exposition mesurés. <i>Études transversales</i> — Rapporter le nombre d'événements ou les indicateurs mesurés.	
Résultats principaux	16	(a) Indiquer les estimations non ajustées et, le cas échéant, les estimations ajustées pour les facteurs confondants et leur précision (p. ex., intervalle de confiance à 95%). Indiquer clairement pour quels facteurs confondants l'ajustement a été fait et la raison pour laquelle ils ont été inclus. (b) Mentionner les limites des intervalles quand des variables continues ont été catégorisées. (c) Le cas échéant, envisager de traduire les estimations du risque relatif en risque absolu pour une période de temps significative.	
Autres analyses	17	Indiquer les autres analyses faites — par exemple, analyses de sous-groupes et d'interactions, et analyses de sensibilité.	
Discussion			
Résultats clés	18	Résumer les résultats clés en faisant référence aux objectifs de l'étude.	
Limites	19	Discuter des limites de l'étude, en tenant compte des sources de biais potentiels ou d'imprécisions. Discuter de la tendance et de l'ampleur de tout biais potentiel.	RECORD 19.1 : Discuter des implications de l'utilisation de données qui n'ont pas été créées ou collectées pour répondre aux questions de recherche spécifiques. Inclure une discussion des biais de classification erronée, des facteurs confondants non mesurés, des données manquantes et de l'évolution de l'admissibilité au fil du temps, en ce qui a trait à l'étude en question.
Interprétation	20	Fournir une interprétation globale prudente des résultats compte tenu des objectifs, des limites de l'étude, de la multiplicité des analyses, des résultats d'études similaires et tout autre élément pertinent.	
Généralisabilité	21	Discuter de la généralisabilité (validité externe) des résultats de l'étude.	
Autres informations			
Financement	22	Indiquer la source de financement et le rôle des financeurs dans la présente étude et, le cas échéant, dans l'étude originale sur laquelle est fondé le présent article.	
Accessibilité du protocole, des données brutes et du code de programmation	SO		RECORD 22.1 : Les auteurs devraient fournir des informations sur la façon d'accéder à toute information supplémentaire telle que le protocole d'étude, les données brutes ou le code de programmation.

Nota : RECORD = Reporting of Studies Conducted Using Observational Routinely Collected Health Data, SO = sans objet, STROBE = Strengthening the Reporting of Observational Studies in Epidemiology.

*La liste de contrôle est protégée par la licence Creative Commons Attribution (CC BY).

devraient noter qu'au-delà du type de base de données indiqué dans le titre et/ou le résumé, des informations devraient être fournies afin de permettre aux lecteurs de comprendre le contenu et la validité de la base de données et les raisons originales pour lesquelles les données ont été collectées. Par exemple, un dossier de santé électronique peut être utilisé par des spécialistes ou des médecins de soins primaires, pour des soins ambulatoires ou hospitaliers, ou par des médecins séniors ou des étudiants en médecine. Les utilisateurs peuvent être spécialement formés pour une saisie de données exhaustive et reproductible, mais il se peut aussi qu'ils n'aient reçus aucune formation particulière²⁵. Les auteurs devraient également décrire comment la population issue de la base de données est comparable avec la population source, notamment en signalant les critères de sélection potentiels, afin que les lecteurs puissent déterminer si les résultats peuvent être appliqués à la population source.

Méthodes (participants)

- Item RECORD 6.1 : Les méthodes de sélection de la population étudiée (telles que les codes ou les algorithmes utilisés pour identifier les participants) devraient être énumérées en détail. Si cela n'est pas possible, une explication devrait être fournie.
- Item RECORD 6.2 : Toute étude de validation des codes ou des algorithmes utilisés pour sélectionner la population devrait être mise en référence. Si la validation a été effectuée pour cette étude en particulier et n'a pas été publiée par ailleurs, le détail des méthodes utilisées et des résultats de cette validation des données devraient être fournis dans le manuscrit.
- Item RECORD 6.3 : Si l'étude impliquait le chaînage de plusieurs bases de données, les auteurs devraient envisager l'utilisation d'un diagramme de flux ou d'une autre représentation graphique pour décrire le processus de chaînage des données, y compris le nombre de personnes avec les données chaînées à l'issue de chaque étape du chaînage.

Exemples

Item RECORD 6.1 : L'extrait suivant fournit un exemple de bonne documentation²⁶ :

The OCCC [Ontario Crohn's and Colitis Cohort] uses validated algorithms to identify patients with IBD based on age group. Each of these algorithms was validated in Ontario, in the specific age group to which it was applied, in multiple cohorts, medical practice types, and regions. For children younger than 18 years, the algorithm was defined by whether they underwent colonoscopy or sigmoidoscopy. If they had undergone endoscopy, children required 4 outpatient physician contacts or 2 hospitalizations for IBD within 3 years. If they had not undergone endoscopy, children required 7 outpatient physician contacts or 3 hospitalizations for IBD within 3 years. ... This algorithm correctly identified children with IBD with a sensitivity of....

Cet article fait référence à 2 études de validation antérieures d'algorithmes visant à identifier des patients atteints d'une maladie inflammatoire de l'intestin à des âges différents, y compris des mesures de précision du diagnostic. Item RECORD 6.2 :

1. Ducharme et collaborateurs²⁷ ont décrit en détail la validation des codes pour identifier les enfants atteints d'intussusception, puis ont utilisé ces codes validés pour décrire les caractéristiques

épidémiologiques de la maladie. Les codes impliqués dans l'étude de validation sont décrits dans la figure 2 de l'article.

2. Benchimol et collaborateurs²⁶ n'ont pas conduit de travail de validation ; cependant, ils ont mis en référence le travail de validation mené précédemment. La précision des codes diagnostiques utilisés pour développer l'algorithme d'identification a été décrit en détail.

Item RECORD 6.3 : Les exemples des figures 2, 3 et 4 sur le site Web de RECORD illustrent plusieurs façons pour décrire le processus de chaînage entre des bases de données :

1. Figure 2 : Diagramme de Venn pour illustrer un processus de chaînage (reproduit avec la permission de Herrett et collaborateurs²⁸ sur notre site Web [www.record-statement.org/images/figure2.jpg]).
2. Figure 3 : Diagramme de flux mixte et diagramme de Venn illustrant un processus de chaînage (reproduit avec la permission de van Herk-Sukel et collaborateurs²⁹ sur notre site Web [www.record-statement.org/images/figure3.jpg]).
3. Figure 4 : Diagramme de chaînage combiné avec un diagramme de flux des participants (reproduit avec la permission de Fosbøl et collaborateurs³⁰ sur notre site Web [www.record-statement.org/images/figure4.jpg]).

Explication

Items RECORD 6.1 et 6.2 : Il est essentiel de documenter la validité des codes/des algorithmes d'identification utilisés pour sélectionner la population étudiée pour permettre l'évaluation de la transparence dans la présentation d'une recherche observationnelle menée à partir de données de santé collectées en routine. En outre, la documentation des codes/des algorithmes permet à d'autres chercheurs de procéder à une validation externe ou interne.

Les méthodes utilisées pour identifier les participants devraient être explicitement et clairement énoncées, notamment si l'identification est basée sur l'utilisation de codes uniques ou d'algorithmes (combinaisons de codes), le chaînage entre différentes bases de données ou des champs de textes non structurés.

Comme dans beaucoup d'autres études épidémiologiques, le risque de biais de classification dans les études utilisant des données de santé collectées en routine peut menacer la validité des résultats de l'étude³¹. Bien que le risque d'erreurs de classification soit amplifié dans des études utilisant des bases de données contenant de larges populations, de telles études offrent l'opportunité d'étudier des maladies rares ou peu communes³². La validation des méthodes d'identification a été de plus en plus mise en relief comme étant essentielle pour les études utilisant des données de santé collectées en routine, en particulier pour les codes de maladies dans les études utilisant des données médico-administratives collectées à des fins de facturation³³. Les études de validation externes impliquent généralement la comparaison des codes ou des algorithmes utilisés pour identifier les populations étudiées à une norme utilisée comme la référence. Les standards de référence les plus courantes sont les dossiers médicaux, les enquêtes auprès des patients ou des praticiens et les registres cliniques^{5,34}. De plus, une validation interne des bases

de données peut être entreprise pour comparer les sources de données qui se chevauchent au sein d'une même base de données³⁵. Les mesures de précision, y compris la sensibilité, la spécificité, les valeurs prédictives positives et négatives, ou les coefficients κ , sont similaires à celles rapportées dans les études de tests diagnostiques^{5,34}.

Ainsi, pour les études observationnelles utilisant des données de santé collectées en routine, nous recommandons que les détails de la validation externe ou interne des codes/des algorithmes d'identification soient présentés dans la section Méthodes du manuscrit. Si une ou plusieurs études de validation ont déjà été réalisées, celles-ci doivent être citées dans les références. Si de telles études de validation n'ont pas été réalisées, cela devrait être explicitement indiqué. En outre, une brève discussion sur la précision des méthodes d'identification (en utilisant des termes communs d'exactitude diagnostique) et sur leur fonctionnement dans les sous-populations étudiées devrait être incluse. Si un travail de validation a été réalisé dans le cadre de l'étude observationnelle en question, nous suggérons aux auteurs d'utiliser les directives publiées pour les études de validation⁵. Il est important d'indiquer si la validation s'est produite dans une population source ou une population issue d'une base de données différente de celle sélectionnée pour l'étude en question, car les codes peuvent fonctionner différemment dans différentes populations ou bases de données³⁶. En outre, s'il existe des problèmes connus liés à la norme utilisée comme référence à laquelle les données ont été comparées, comme par exemple une incomplétude ou une imprécision, ces problèmes devraient être documentés et discutés comme étant une limite de l'étude. Les auteurs devraient discuter des implications de l'utilisation des codes/des algorithmes sélectionnés pour identifier les populations étudiées et les résultats, du risque d'erreurs de classification et des impacts potentiels concernant les résultats de l'étude. Il est particulièrement important d'examiner les implications du recours à une étude de validation menée dans une population différente de celle examinée.

Item RECORD 6.3: Un diagramme de flux ou une autre présentation graphique peut fournir des informations utiles sur le processus de chaînage et peut simplifier une description potentiellement longue. De telles illustrations peuvent fournir des données clés telles que des informations sur la proportion et les caractéristiques des participants qui ont pu être chaînés et ceux qui n'ont pas pu l'être. Les lecteurs devraient être en mesure d'établir la proportion des populations issues des bases de données qui ont été chaînées avec succès ainsi que la représentativité de la population étudiée qui en résulte. Les organigrammes de chaînage peuvent être soit des diagrammes autonomes (p. ex., diagramme de Venn ou diagramme de flux), soit combinés avec un diagramme de flux de participants tel que recommandé par les directives de STROBE. Comme les graphiques peuvent être fournis dans de nombreux formats, nous ne recommandons pas de modèle spécifique.

Méthodes (variables)

- Item RECORD 7.1: Une liste complète des codes et des algorithmes utilisés pour classer les expositions, les résultats, les

facteurs confondants et les modificateurs d'effets devrait être fournie. Si ces codes ou algorithmes ne peuvent pas être documentés, une explication devrait être fournie.

Exemples

1. Hardelid et collaborateurs³⁷ ont fourni tous les codes dans le tableau S1 de leur annexe 1 dans le supplément de données 2.
2. Murray et collaborateurs³⁸ ont fourni tous les codes pour les groupes à risque dans leur annexe S1.

Explication

Tout comme pour les codes/les algorithmes utilisés pour identifier la population étudiée, les codes/les algorithmes permettant de classer les expositions, les résultats, les facteurs confondants ou les modificateurs d'effets exposent la recherche à un risque potentiel de biais de classification. Afin de permettre la réplique, l'évaluation et les comparaisons avec d'autres études, nous recommandons qu'une liste de tous les codes diagnostiques, de procédures, de médicaments ou d'autres codes utilisés pour réaliser l'étude soit fournie dans le manuscrit, dans une annexe en ligne et/ou sur un site Web externe. Pour les données de routine comprenant des résultats d'enquêtes, les questions de l'enquête devraient être accompagnées de la formulation précise fournie aux participants. Compte tenu du risque de biais de classification dans toutes les recherches, y compris dans les recherches menées à partir de données de santé collectée en routine³¹, les auteurs devraient fournir suffisamment de détails pour rendre leur recherche reproductible et le risque de biais visible. Les études de validation peuvent être décrites dans le manuscrit ou citées dans les références ou encore accessible en ligne. Comme indiqué ci-dessus, les auteurs devraient indiquer si l'étude de validation a été réalisée dans une population source ou avec une base de données différente de celle examinée dans l'étude en question.

Nous reconnaissons que dans certaines situations, les chercheurs peuvent être empêchés de fournir des listes de codes et des algorithmes utilisés dans une publication, car cette information peut être considérée comme exclusive ou protégée par le droit d'auteur, la propriété intellectuelle ou d'autres lois. Par exemple, certains indices d'ajustement sur des comorbidités ont été créés par des entreprises à but lucratif et vendus à des chercheurs dans le cadre de recherches universitaires^{39,40}. Dans ces situations, les auteurs peuvent avoir fait appel à des fournisseurs de données ou à des tiers pour collecter, traiter et/ou chaîner les données. Les auteurs devraient fournir une explication détaillée concernant leur incapacité à fournir des listes de codes ou autres détails sur la façon dont les personnes ou les conditions ont été identifiées. Ils devraient également œuvrer pour inclure les coordonnées du groupe détenant des droits de propriété sur ces listes. En outre, les auteurs devraient indiquer comment leur incapacité à fournir cette information peut avoir un impact en termes de réplique et d'évaluation de la recherche. Idéalement, les tiers devraient fournir des informations détaillées sur la manière dont les données ont été collectées, traitées ou chaînées. Une

communication améliorée entre les fournisseurs et les utilisateurs de données pourrait être mutuellement bénéfique.

Certains ont trouvé qu'il était nécessaire que les listes de codes représentent la propriété intellectuelle des chercheurs. La publication de listes de codes par d'autres chercheurs que ceux qui les ont développés pourrait priver ces derniers de leur propriété intellectuelle et de la reconnaissance pour avoir créé de telles listes. Nous estimons que cette opinion n'est pas compatible avec la norme scientifique définie par la transparence qui permet la reproduction de travaux de recherche. Par conséquent, en dehors de ceux qui sont protégés par la loi ou par contrat, nous recommandons que les listes complètes de codes soient publiées.

Considérant le nombre de mots et les restrictions d'espace dans de nombreux journaux et la longueur potentielle de ces listes de codes/d'algorithmes, nous reconnaissons que la publication dans un article de journal en format imprimé peut être impossible. Les informations détaillées peuvent être reportées dans le texte, des tableaux publiés, ainsi que des compléments en ligne sur le site Web du journal en annexe, entretenus en permanence par les auteurs ou d'autres personnes, ou déposés dans un référentiel de données tiers (p. ex., Dryad ou Figshare). Les sections de texte et de références du manuscrit doivent fournir des informations détaillées sur la façon d'accéder aux listes de codes. Les référentiels de code tels que ClinicalCodes.org (université de Manchester) sont très prometteurs pour la documentation et la transparence des codes utilisés dans la recherche basée sur les données de santé⁴¹. Si les listes de codes sont publiées dans des suppléments en ligne sur le site Web du journal ou sur un site Web externe fourni par les auteurs, le lien devrait être publié dans l'article du journal. La publication sur un site Web de journal ou sur PubMed Central (www.ncbi.nlm.nih.gov/pmc/) augmente la probabilité que le supplément soit disponible tant que le journal est opérationnel. Si la publication sur un site Web externe privé ou institutionnel est la seule option, nous recommandons que ces listes restent disponibles pendant au moins 10 ans après la publication de l'article. Si l'adresse URL est modifiée, une redirection automatique à partir de l'ancienne adresse Web est requise. Ces mesures permettront aux futurs lecteurs de l'article d'avoir accès aux listes de codes complètes.

En plus des listes de codes fournies dans l'article (ou dans une annexe en ligne), les auteurs doivent inclure une réflexion sur la possibilité que le choix des codes/des algorithmes utilisés dans l'étude conduisent à des biais. Un tel biais pourrait inclure un biais de classification erronée, un biais de constatation et un biais dû à des données manquantes. Si des analyses de sensibilité ont été effectuées en se basant sur de différents ensembles de codes/d'algorithmes, ceux-ci doivent également être décrits et évalués. Une discussion sur les biais potentiels pourrait également être liée à d'autres parties des listes de contrôle RECORD et STROBE, telles que la sélection des participants étudiés et la validation ou non-validation des codes.

Méthodes (méthodes statistiques)

L'accès aux données et les méthodes de traitement

- Item RECORD 12.1 : Les auteurs devraient décrire à quel

degré les chercheurs ont eu accès à la base de données utilisée pour créer la population étudiée.

- Item RECORD 12.2 : Les auteurs devraient fournir des informations sur les méthodes de traitement des données utilisées dans l'étude.

Le chaînage des données

- Item RECORD 12.3 : Indiquer si l'étude a inclus un chaînage de données au niveau de la personne, au niveau de l'établissement ou d'autres données dans 2 bases de données ou plus. Les méthodes de chaînage et d'évaluation de la qualité de celui-ci devraient être fournies.

Exemples

Item RECORD 12.1 : Les articles suivants décrivent l'accès à un sous-ensemble de la « General Practice Research Database » (GPRD) du Royaume-Uni :

- « The GPRD restricts its data sets to 100,000 individuals for projects funded through the Medical Research Council licence agreement. This restriction mandated a case-control rather than cohort design to ensure we identified sufficient cases of cancer for each particular symptom »⁴².
- « A random sample from the General Practice Research Database ... was obtained under a Medical Research Council licence for academic institutions »⁴³.

Item RECORD 12.2 : Voici un exemple de description de méthodes de traitement de données⁴⁴ :

Completeness of common identifiers for linking varied between datasets and by time (identifiers were more complete in recent years). For LabBase2, completeness of identifiers varied by unit (figure 2). For PICA Net [Paediatric Intensive Care Audit Network], date of birth and hospital number were 100% complete, and the majority of other identifiers were >98% complete, with the exception of NHS [National Health Service] number (85% complete). For both datasets, cleaning and data preparation were undertaken: NHS or hospital numbers such as "Unknown" or "9999999999" were set to null; generic names (e.g., "Baby," "Twin 1," "Infant Of") were set to null; multiple variables were created for multiple surname and first names; postcodes beginning "ZZ" (indicating no UK postcode) were set to null.

Item RECORD 12.3 : Les extraits d'articles suivants sont de bons exemples de bons comptes rendus sur le niveau de chaînage de données, les techniques et méthodes de chaînage utilisées et les méthodes utilisées pour évaluer la qualité du chaînage :

- « We linked live birth and fetal death certificates into chronological chains of events that, excluding induced abortions and ectopic pregnancies, constituted the reproductive experience of individual women »⁴⁵.
- Deux articles contiennent d'excellentes descriptions de chaînage qui a été entrepris spécifiquement pour l'étude signalée^{44,45}. Dans l'article de Harron et collaborateurs⁴⁴, une explication détaillée sur le chaînage est fournie avec une représentation graphique du processus d'appariement. En outre, les méthodes de calcul de la probabilité de chaînage sont décrites :

Match probabilities $P(M|agreement\ pattern)$ were calculated to estimate the probability of a match given agreement on a joint set of identifiers. This avoided the assumption of independence between identifiers. Probabilities

were derived as the number of links divided by the total number of pairs for each agreement pattern (based on probable links identified in the training datasets). For example, if 378 comparison pairs agreed on date of birth and Soundex but disagreed on sex, and 312 of these were probable links, the match probability for the agreement pattern [1,1,0] was $312/378 = 0.825$.

L'article de Adams et collaborateurs⁴⁵ fournit également une explication détaillée du processus de chaînage :

The deterministic linkage consisted of phase I, which entailed six processing steps during which chains were formed and individual (previously unlinked) records were added to chains. Next followed phase n, which entailed multiple passes through the file to combine chains belonging to the same mother.

3. En revanche, si une étude se réfère à des données chaînées antérieurement, la référence faite à un document antérieur serait adéquate comme suit : « Records from both databases were linked to the municipal registries based on date of birth, gender and zip code, and were subsequently linked to each other. The linkage was performed by Statistics Netherlands and is described in previous publications »²⁰.
4. L'extrait suivant est un exemple de bon rapport sur les caractéristiques des personnes chaînées et non chaînées⁴⁶ :

For the purposes of this paper unmatched ISC [Inpatient Statistics Collection] records will be referred to as ISC residuals, unmatched MDC [Midwives Data Collection] records as MDC residuals and linked pairs as matched records. ... Selected variables that were available on both data sets were compared across three groups—ISC residuals, MDC residuals and matched records.

Explication

Items RECORD 12.1 et 12.2 : Des erreurs peuvent survenir si les analystes de données qui ne connaissent pas les nuances de la création de cohortes ou les objectifs de l'étude créent les cohortes d'étude. Par conséquent, la mesure dans laquelle les auteurs ont accès à la base de données devrait être signalée. La description des méthodes de traitement des données à différentes étapes de l'étude doit inclure celles utilisées pour dépister les données erronées et manquantes, y compris les vérifications des intervalles, les vérifications des doublons et le traitement des mesures répétées^{47,48}. D'autres méthodes à signaler pourraient inclure l'évaluation des distributions de fréquences et des tableaux croisés de données et l'exploration graphique ou l'utilisation de méthodes statistiques pour la détection des valeurs aberrantes⁴⁹. Des détails supplémentaires pourraient être fournis sur le diagnostic d'erreur, y compris les définitions de plausibilité et la gestion des erreurs dans l'analyse. Une description claire et transparente des méthodes de traitement des données est importante, car le choix des méthodes pourrait influencer sur les résultats de l'étude, la répétabilité de l'étude et la reproductibilité des résultats⁵⁰.

Item RECORD 12.3 : Pour les études utilisant le chaînage, nous suggérons de rendre compte du taux estimé de chaînage réussi, de l'utilisation du chaînage déterministe contre probabiliste, de la qualité et du type de variables utilisées pour le chaînage et des résultats de toute validation de chaînage. Si le chaînage des entrées à travers les bases de données a été réalisé spécifiquement pour l'étude, les méthodes de chaînage et d'évaluation de

la qualité du chaînage devraient être rapportées, y compris les informations à propos de la personne qui a effectué le chaînage. Si possible, des détails doivent être fournis sur les variables de blocage, l'exhaustivité des variables de chaînage, les règles de chaînage, les seuils et la révision manuelle⁴⁴. Si le chaînage a été réalisé avant l'étude (pour des études antérieures ou pour un usage général) ou si le chaînage de données a été effectué par un fournisseur externe, tel qu'un centre de chaînage de données, une référence est nécessaire décrivant la ressource de données et les méthodes du chaînage.

Les données décrivant les méthodes de chaînage et évaluant leur succès sont essentielles pour permettre au lecteur d'évaluer l'impact de toute erreur de chaînage et des biais associés⁵¹. Plus précisément, le lecteur devrait savoir si le type de chaînage utilisé était déterministe et/ou probabiliste, afin de déterminer si le chaînage pourrait être influé par de fausses appariements ou d'appariement manquantes. Le chaînage déterministe est utile lorsqu'un identifiant unique est disponible dans les différentes sources de données. Lorsqu'un tel identifiant n'est pas disponible, une description des règles de chaînage d'enregistrements appliquées (ou des clés de chaînage statistiques) est essentielle. En revanche, le chaînage probabiliste utilise plusieurs identifiants, parfois avec des poids différents, et les correspondances sont considérées comme présentes au-dessus d'un seuil spécifique. Des méthodes mixtes peuvent également être utilisées. Par exemple, un chaînage déterministe peut être utilisé pour certaines entrées, et un chaînage probabiliste peut être appliqué lorsque des identifiants uniques ne sont pas disponibles pour d'autres entrées. Le biais associé avec le chaînage se produit lorsque des associations sont présentes entre la probabilité d'erreur de chaînage (p. ex., les correspondances fausses et manquantes) et les variables d'intérêt. Par exemple, les taux de chaînage peuvent varier selon les caractéristiques du patient (p. ex., l'âge, le sexe et l'état de santé). Même de petites erreurs dans le processus de chaînage peuvent causer un biais et conduire à des résultats qui peuvent surestimer ou sous-estimer les associations étudiées⁵². Les auteurs devraient signaler une erreur de chaînage en utilisant des approches standard, y compris des comparaisons avec des normes de référence ou des ensembles de données de référence, des analyses de sensibilité et en comparant les caractéristiques des données chaînées et non chaînées⁵³. La signalisation d'une erreur de chaînage permet au lecteur de déterminer la qualité du chaînage et la possibilité de biais lié aux erreurs de chaînage.

Résultats (participants)

- Item RECORD 13.1 : Décrire en détail la sélection des personnes incluses dans l'étude (c.-à-d., la sélection de la population étudiée), y compris le filtrage basé sur la qualité des données, la disponibilité des données et le chaînage. La sélection des personnes incluses peut être décrite dans le texte et/ou au moyen d'un diagramme de flux.

Exemple

Un exemple de bon compte-rendu est donné dans l'extrait suivant⁵⁴ :

We identified 161,401 Medicare beneficiaries given a diagnosis of one or more cases of cancer of the lung and bronchus in the SEER [Surveillance, Epidemiology, and End Results] registries between 1998 and 2007. Among these patients, we identified a total of 163,379 separate diagnoses of incident lung cancer. (Some patients had two cases of primary lung cancer separated by more than a year during the study period.) Fig 1 shows the derivation of the final cohort of 46,544 patients with 46,935 cases of NSCLC [non-small cell lung cancer]

Voir la figure 5 (reproduit avec la permission de Dinan et collaborateurs⁵⁴) sur le site Web de RECORD (www.record-statement.org/images/figure5.jpg) pour un exemple d'un diagramme de flux.

Explication

Les auteurs devraient fournir une description claire de la dérivation de la/des population(s) étudiée(s) de la base de données originale des données de santé collectées de manière routinière, car les différences entre la population étudiée et la population de la base de données doivent être documentées pour permettre l'application des résultats (voir aussi l'item RECORD 6.1). Les chercheurs utilisant des sources de données collectées de manière routinière limitent souvent leur population étudiée en fonction de facteurs tels que la qualité des données disponibles. Par exemple, ils peuvent restreindre la période d'étude à un temps pendant lequel la qualité des données est jugée acceptable, ce qui entraîne l'exclusion de participants potentiels. Des études peuvent exclure des pratiques médicales avec des entrées de dossiers de santé électroniques incohérentes ou attendre que ces pratiques deviennent cohérentes^{38,55}. La population étudiée peut également être restreinte en fonction de la disponibilité de données. Par exemple, dans les études utilisant les données de l'assurance-maladie des États-Unis, les bénéficiaires actuellement inscrits dans un organisme de maintien de la santé sont souvent exclus en raison de l'absence d'entrées d'événements cliniques^{54,56}. Lors de l'utilisation de sources de données dans lesquelles l'admissibilité fluctue au fil du temps (p. ex., les bases de données d'assurance), les chercheurs doivent préciser clairement comment ils ont défini l'admissibilité et géré les modifications d'admissibilité. Si une étude utilise des données de routine chaînées, la population de l'étude est fréquemment réduite en se limitant aux personnes pour lesquelles des données chaînées sont disponibles⁵⁷. Des cohortes fortement restreintes peuvent également être utilisées pour des raisons méthodologiques afin d'éliminer certains facteurs confondants.

Ainsi, les étapes suivies pour obtenir les populations étudiées définitives, les critères d'inclusion et d'exclusion, et l'inclusion et l'exclusion des participants aux différentes étapes de la création et de l'analyse des cohortes doivent être clairement définies soit dans le texte, soit sous forme d'une représentation schématique. Les populations étudiées peuvent être dégagées à l'aide de codes et/ou d'algorithmes différents (voir l'item RECORD 6.1), ce qui peut influencer sur la population étudiée au fil du temps^{58,59}. Certaines études peuvent également avoir utilisé plusieurs définitions de cas plus ou moins sensibles/spécifiques, ce qui peut influencer sur les analyses ultérieures. La délimitation de ces étapes est importante pour évaluer la validité externe des résultats de

l'étude et, dans certaines circonstances, pour évaluer un éventuel biais de sélection. Des analyses de sensibilité peuvent être présentées pour évaluer l'impact potentiel de l'absence de données et de la représentativité de la population étudiée. Le fait de fournir des informations sur la sélection de la/des population(s) étudiée(s) à partir de la base de données initiale permet également de répliquer l'étude. Des analyses subsidiaires peuvent avoir été effectuées sur différentes populations d'étude et peuvent potentiellement être présentées dans des annexes en ligne.

Discussion (limites)

- Item RECORD 19.1 : Discuter des implications de l'utilisation de données qui n'ont pas été créées ou collectées pour répondre aux questions de recherche spécifiques. Inclure une discussion des biais de classification erronée, des facteurs confondants non mesurés, des données manquantes et de l'évolution de l'admissibilité au fil du temps, en ce qui a trait à l'étude en question.

Exemples

Les 2 articles suivants décrivent les limites associées à l'utilisation des données administratives :

Third, this study was a retrospective, claims-based analysis. Only PET [positron emission tomography] scans paid for by Medicare could be detected in the analysis. To minimize the proportion of missed claims, all analyses were limited to Medicare beneficiaries with both Medicare Part A and Part B coverage and no enrollment in managed care or Medicare Part C for the 12 months before and after diagnosis. Fourth, patients in the SEER registry are more likely to be nonwhite, to live in areas with less poverty, and to live in urban areas, which may limit the generalizability of the findings. Fifth, during the study period, disease stage was based on SEER data obtained over 4 months or until first surgery. In 2004, data collection for SEER changed to the collaborative staging system. It is unclear how our results would differ with this newer approach⁵⁴.

Despite several strengths of the SEER-Medicare data, including a comparatively large sample size, generalizability to the US population, and detailed information on prescriptions, our study was limited by the lack of laboratory data on cholesterol, triglyceride, and glucose levels that would have informed the extent of metabolic disturbances in the population ... thus having laboratory-based data could have reduced residual confounding by severity of metabolic disease. We also lacked more granular data on cancer progression, which could have confounded the association between statin use and death, given that statin treatment may be withheld or discontinued in patients with short expected survival time⁶⁰.

Explication

Les données de santé collectées de manière routinière ne sont généralement pas collectées en se basant sur une question de recherche spécifique développée a priori, et les raisons qui motivent la collecte de données peuvent varier. De nombreux domaines potentiels de biais, y compris toutes les sources habituelles de biais associées à la recherche observationnelle mais aussi plus spécifiques à la recherche observationnelle utilisant des données de routine, menacent les conclusions des chercheurs. Les auteurs devraient considérer les éléments suivants comme des sources potentielles de biais : 1) les codes ou les

algorithmes utilisés pour identifier les populations étudiées, les critères de jugement, les facteurs confondants ou les modificateurs d'effets (biais de classification erronée), 2) des variables manquantes (facteurs confondants non mesurée), 3) des données manquantes et 4) les changements d'admissibilité au fil du temps.

La logique sous-jacente à la collecte de données de routine peut influencer sur la qualité et l'applicabilité des données aux questions de recherche examinées. Par exemple, les registres utilisés pour les analyses rétrospectives peuvent mettre en œuvre un meilleur contrôle de la qualité que les organisations collectant d'autres types de données collectées de manière routinière, bien que cela puisse varier. De même, certaines données administratives sont soumises à un contrôle de qualité minutieux, alors que d'autres ne le sont pas. Les données administratives sont particulièrement susceptibles au codage abusif ou opportuniste. Par exemple, lorsque le remboursement hospitalier est basé sur la complexité de la gamme de cas, les hôpitaux pourraient maximiser le remboursement en appliquant généreusement des codes de maladie plus complexes aux dossiers des patients⁶¹. De plus, des changements dans les stratégies de codage peuvent avoir une incidence sur la validité ou la consistance des données. Par exemple, l'introduction de codes d'incitation à la facturation du fournisseur peut modifier la probabilité d'utilisation d'un code au fil du temps^{62,63}. D'autres codes peuvent être évités en raison de la stigmatisation du patient ou des pénalités au fournisseur⁶⁴. De plus, des changements dans les versions des systèmes de classification des codes (p. ex., de la Classification internationale des maladies [CIM], 9^{ième} révision à la CIM-10) peuvent altérer la validité de la constatation utilisant des données codées^{65,66}. La variation de la pratique clinique dans les hôpitaux et selon les populations peut entraîner une localisation des examens de laboratoire dans des endroits et/ou des pratiques spécifiques, ce qui peut influencer sur un algorithme diagnostique. Si l'une de ces sources potentielles de biais de classification se présente, elle devrait être considérée comme étant une limite à l'étude.

Les facteurs confondants non mesurés sont définies comme étant des facteurs associés à des variables non incluses dans les données de l'étude, ce qui entraîne un biais de confusion résiduel⁶⁷. Bien qu'ils représentent une source potentielle de biais dans toutes les recherches observationnelles, ils sont particulièrement prédominants dans les études utilisant des données collectées de façon routinière.

L'analyse pourrait nécessiter des variables qui n'ont pas été prises en compte lors de la planification des bases de données ou lors de la collecte des données. Diverses méthodes ont été proposées pour adresser cette source potentielle de biais⁶⁸⁻⁷¹, y compris les scores de propension. Cependant, les analyses de score de propension, comme les analyses de régression standard et l'appariement, ne peuvent garantir qu'un équilibre des participants à l'étude sur les variables disponibles dans les données. Un type particulier de confusion non mesurée est la confusion par indication; c'est souvent un problème lorsqu'on examine l'efficacité et l'innocuité des traitements (médicamenteux) en utilisant des données de routine. Par conséquent, le pronostic des per-

sonnes recevant le traitement (médicamenteux) peut être mieux ou pire que celui des autres, mais les données sur le pronostic et/ou la sévérité de la maladie sous-jacente peuvent ne pas être disponibles dans les données⁷². Les auteurs devraient discuter de ces questions et signaler les méthodes utilisées pour en tenir compte (si possible).

Les données manquantes sont problématiques pour toute recherche observationnelle et cela a été adressé dans l'encadré 6 de l'article explicatif STROBE¹⁰. Les données manquantes sont un problème particulier pour les données collectées de façon routinière, car les chercheurs ne peuvent pas contrôler la collecte de données⁷³. Les données manquantes peuvent entraîner un biais de sélection s'il existe des valeurs manquantes dans les variables utilisées pour définir la cohorte de l'étude ou des identifiants manquants qui entravent le chaînage des entrées, surtout si les données manquantes existent de manière non aléatoire. Les variables manquantes créent des défis similaires. Les auteurs devraient décrire les variables manquantes suspectées de causer un biais de confusion non mesuré, la raison pour laquelle ces variables manquaient, comment cela a pu influencer sur les résultats de l'étude et les méthodes utilisées pour ajuster pour les variables manquantes. Par exemple, le tabagisme a un fort effet sur la sévérité de la maladie de Crohn et a été associé à l'évolution de cette maladie. Cependant, le statut tabagique est rarement inclus dans les données administratives sur la santé. Dans une étude évaluant l'association entre le statut socio-économique et les résultats de la maladie de Crohn, le statut tabagique a été abordé comme un facteur de confusion potentiel non mesuré⁷⁴. Fréquemment, les données/les variables manquantes ne sont découvertes qu'après le début de la recherche utilisant des données de santé collectées de manière routinière, ce qui oblige les chercheurs à s'écarter de leur protocole de recherche original. Les détails de la déviation par rapport au protocole, quelle que soit sa raison, doivent toujours être indiqués. Les raisons de la déviation et les implications sur la recherche et les conclusions devraient être abordées.

Une autre limitation potentielle importante est l'évolution des pratiques de codage ou des critères d'admissibilité résultant d'une modification de la composition de la base de données, de la population étudiée ou des deux au fil du temps. La définition de la population de base de données peut changer dans un certain nombre de circonstances, comme par exemple si les cabinets de recrutement cessent de collaborer avec la base de données, changent de logiciel ou modifient les critères d'inscription dans la base de données, par exemple, un registre. La population étudiée dans les sources de données administratives (p. ex., les bases de données d'assurance) peut changer si l'admissibilité des personnes n'est pas constante, en raison de changements d'emploi, de statut de résidence ou de fournisseur de soins médicaux. Un changement dans la façon dont les entrées sont codées (p. ex., codage abusif ou changements dans les systèmes de codage, comme décrit ci-dessus) peut modifier la population de l'étude^{63,75,76}. Lorsqu'ils abordent les limites, les auteurs devraient expliquer comment ils ont adressé le changement d'admissibilité dans l'analyse afin que le lecteur puisse évaluer la possibilité d'un biais. Comme indiqué par les directives STROBE,

la discussion devrait inclure l'orientation et l'ampleur de tout biais potentiel et les efforts déployés pour remédier à ce biais.

Autres informations

- Item RECORD 22.1 : Les auteurs devraient fournir des informations sur la façon d'accéder à toute information supplémentaire telle que le protocole d'étude, les données brutes ou le code de programmation.

Exemples

1. L'article de Taljaard et collaborateurs⁷⁷ représente le protocole de recherche complet pour une étude basée sur l'Enquête sur la santé dans les collectivités canadiennes.
2. Guttmann et collaborateurs⁷⁸ invitent des demandes pour le protocole de leur étude : « Data sharing: The technical appendix, dataset creation plan/protocol, and statistical code are available from the corresponding author at [adresse de courriel] ».

Explication

Nous soutenons fortement la diffusion d'informations détaillées sur les méthodes de l'étude et les résultats. Lorsque cela est possible, nous encourageons la publication préalable ou simultanée du protocole d'étude, des résultats de données brutes et, le cas échéant, du code de programmation. Ces informations sont utiles aux évaluateurs et aux lecteurs pour évaluer la validité des résultats de l'étude. Un certain nombre d'opportunités sont à la disposition des chercheurs pour la publication ouverte de ces données, telles des documents supplémentaires en ligne, des sites Web personnels, des sites Web institutionnels, des sites de médias sociaux à vocation scientifique (p. ex., www.ResearchGate.net et www.Academia.edu), des répertoires de données (p. ex., Dryad ou Figshare) ou des sites Web gouvernementaux ouverts.⁷⁹ Nous reconnaissons que certaines organisations de recherche, sociétés, institutions ou lois peuvent interdire ou restreindre la disponibilité libre de ces informations. Bien qu'une discussion sur la propriété et l'utilisation de cette propriété intellectuelle soit hors du champ d'application des directives RECORD, la publication de ces données doit toujours être effectuée dans le respect des directives légales et éthiques de l'environnement institutionnel des chercheurs, sous la direction des rédacteurs de journaux. Cette information serait également utile à d'autres chercheurs qui souhaiteraient accéder à ces données pour reproduire ou développer la recherche décrite dans le manuscrit. Quel que soit le format ou l'étendue des informations supplémentaires disponibles, nous recommandons que la référence à l'emplacement de ces informations soit clairement indiquée dans le manuscrit.

Interprétation

Les directives RECORD sont spécifiques à la recherche observationnelle réalisée à partir de données de santé collectées de manière routinière et servent à compléter et non à remplacer les directives STROBE. Elles ont été créées pour guider les auteurs, les rédacteurs de journaux, les pairs évaluateurs et les

autres parties prenantes afin d'encourager la transparence et l'exhaustivité de la communication de recherche menée en utilisant des données de santé collectées de manière routinière. La liste de contrôle est destinée à être utilisée par tout chercheur utilisant de telles données, et nous encourageons une large diffusion auprès de toutes les parties intéressées. Nous prévoyons que l'approbation et la mise en œuvre de RECORD par les journaux amélioreront la transparence de la communication de recherche utilisant des données de santé collectées de manière routinière.

Au fur et à mesure que la disponibilité des données de santé collectées de manière routinière augmente, nous nous attendons à une plus grande implication des chercheurs des régions dans lesquelles ces données ne sont pas actuellement accessibles. Nous nous attendons à des commentaires continus et à des discussions de document RECORD, par le biais de notre site Web (www.record-statement.org) et du babillard électronique, de la part des parties intéressées, ce qui pourrait entraîner des révisions officielles à l'avenir. Grâce à cette communauté en ligne, RECORD deviendra un document vivant qui peut s'adapter aux changements sur le terrain.

La publication d'une ligne directrice et l'approbation des journaux ne sont pas suffisantes pour améliorer les rapports de recherche⁸⁰. La manière de mettre en œuvre les lignes directrices par les chercheurs, les journaux et les pairs évaluateurs est d'une importance cruciale pour que RECORD ait un impact mesurable⁸¹. Par conséquent, le babillard électronique inclura un forum de discussion sur la mise en œuvre. Nous encourageons également l'évaluation de l'impact de RECORD sur les rapports sur le terrain afin de garantir que les directives fournissent des avantages mesurables.

Limites

L'application de STROBE et de RECORD est uniquement destinée aux études de recherche observationnelles. Cependant, des données de santé collectées de manière routinière sont parfois utilisées pour des recherches menées avec d'autres plans d'étude, telles que les essais randomisés en groupe pour l'évaluation des systèmes de soins de santé. De plus, le chaînage de données provenant d'essais randomisés à des données administratives peut être utilisé pour le suivi à long terme des résultats, et les études associées ne seraient pas considérées comme des observations. À mesure que le domaine évolue, nous prévoyons étendre RECORD à d'autres plans de recherche en utilisant des méthodes rigoureusement similaires.

Bien que RECORD représente notre meilleure tentative pour refléter l'intérêt et les priorités des parties prenantes, nous reconnaissons que les méthodes utilisées pour mener des recherches utilisant des données de santé collectées de façon régulière évoluent rapidement et que la disponibilité des types de données pour ces recherches augmente. Par exemple, les applications de santé mobiles deviennent de plus en plus disponibles pour les téléphones intelligents et les technologies portables. Bien que la recherche à l'aide de ces sources de données soit actuellement limitée, nous prévoyons une croissance rapide de l'utilisation de ces données prochainement, et de nouvelles

méthodologies seront créées pour gérer cette ressource. En outre, le comité de travail a décidé de se concentrer sur les données de santé et non sur toutes les sources de données utilisées pour mener des recherches sur la santé (p. ex., données environnementales, données financières). Par conséquent, la liste de contrôle de RECORD pourrait ne pas refléter les thèmes qui deviendront importants à l'avenir, et une révision pourrait être nécessaire à un moment donné.

Nous avons déployé des efforts importants pour inclure une large représentation des parties prenantes dans la création de ces lignes directrices. Nous avons recruté des parties prenantes à travers des appels ouverts et des invitations ciblées en utilisant une variété de canaux¹⁶. Cependant, la représentation des parties prenantes provenait principalement de régions menant des recherches utilisant des données de santé collectées de manière routinière, avec seulement quelques représentants de pays en développement et de pays non anglophones. Néanmoins, nous croyons que le groupe des intervenants était représentatif de la communauté actuelle des chercheurs et des utilisateurs de connaissances générées. Bien que de nombreuses contributions aient été obtenues grâce aux enquêtes et retours du groupe de parties prenantes, la faisabilité dictait que les déclarations soient rédigées par un comité de travail plus restreint composé de 19 membres qui se rencontraient en personne¹⁷. À l'avenir, la technologie et les médias sociaux pourraient permettre une participation plus active de groupes plus importants aux réunions du comité de travail.

Conclusion

La déclaration RECORD étend les critères STROBE aux études observationnelles menées à l'aide de données de santé collectées de manière routinière. Grâce à la contribution de la communauté de la recherche et de l'édition, nous avons créé des lignes directrices pour la présentation de rapports sous la forme d'une liste de contrôle et du présent document explicatif. Il a été démontré que les lignes directrices pour la présentation des rapports améliorent le compte rendu de la recherche, ce qui permet aux consommateurs de la recherche d'être conscients des points forts, des limites et de l'exactitude des conclusions^{12,82-84}. Bien que nous prévoyions que RECORD changera avec l'évolution des méthodes de recherche sur le terrain, ces lignes directrices aideront à faciliter la communication adéquate de la recherche au cours des prochaines années. Grâce à la mise en œuvre des lignes directrices RECORD par les auteurs, les rédacteurs de journaux et les pairs évaluateurs, nous prévoyons que RECORD se traduira par la transparence, la reproductibilité et l'exhaustivité de la communication de recherche menée en utilisant des données de santé collectées de manière routinière.

Références

- Spasoff RA. *Epidemiologic methods for health policy*. New York: Oxford University Press;1999.
- Morrato EH, Elias M, Gericke CA. Using population-based routine data for evidence-based health policy decisions: lessons from three examples of setting and evaluating national health policy in Australia, the UK and the USA. *J Public Health (Oxf)* 2007;29:463-71.
- De Coster C, Quan H, Finlayson A, et al. Identifying priorities in methodological research using ICD-9-CM and ICD-10 administrative data: report from an international consortium. *BMC Health Serv Res* 2006;6:77.
- Hemkens LG, Benchimol EI, Langan SM, et al., éditeurs. Reporting of studies using routinely collected health data: systematic literature analysis [affiche]. Dans : *REWARD/EQUATOR Conference 2015* ; 28-30 septembre 2015 ; Edinburgh.
- Benchimol EI, Manuel DG, To T, et al. Development and use of reporting guidelines for assessing the quality of validation studies of health administrative data. *J Clin Epidemiol* 2011;64:821-9.
- Herrett E, Thomas SL, Schoonen WM, et al. Validation and validity of diagnoses in the General Practice Research Database: a systematic review. *Br J Clin Pharmacol* 2010;69:4-14.
- Rothman KJ, Greenland S, Lash TL. *Modern epidemiology*. 3^{ème} éd. Philadelphia: Lippincott Williams & Wilkins;2008.
- Plint AC, Moher D, Morrison A, et al. Does the CONSORT checklist improve the quality of reports of randomised controlled trials? A systematic review. *Med J Aust* 2006;185:263-7.
- Library for health research reporting. Oxford (R.-U.) : Enhancing the Quality and Transparency Of health Research (EQUATOR) Network, University of Oxford; 2015. Accessible ici : www.equator-network.org/library/ (consulté le 7 mars 2015).
- Vandenbroucke JP, von Elm E, Altman DG, et al.; STROBE Initiative. Strengthening the Reporting of Observational Studies in Epidemiology (STROBE): explanation and elaboration. *PLoS Med* 2007;4:e297.
- von Elm E, Altman DG, Egger M, et al.; STROBE Initiative. The Strengthening the Reporting of Observational Studies in Epidemiology (STROBE) statement: guidelines for reporting observational studies. *PLoS Med* 2007;4:e296.
- Sorensen AA, Wojahn RD, Manske MC, et al. Using the Strengthening the Reporting of Observational Studies in Epidemiology (STROBE) statement to assess reporting of observational trials in hand surgery. *J Hand Surg Am* 2013; 38:1584-9.e2.
- Cobo E, Cortés J, Ribera JM, et al. Effect of using reporting guidelines during peer review on quality of final manuscripts submitted to a biomedical journal: masked randomised trial. *BMJ* 2011;343:d6783.
- Benchimol EI, Langan S, Guttman A; RECORD Steering Committee. Call to RECORD: the need for complete reporting of research using routinely collected health data. *J Clin Epidemiol* 2013;66:703-5.
- Langan SM, Benchimol EI, Guttman A, et al. Setting the RECORD straight: developing a guideline for the REporting of studies Conducted using Observational Routinely collected Data. *Clin Epidemiol* 2013;5:29-31.
- Nicholls SG, Quach P, von Elm E, et al. The REporting of studies Conducted using Observational Routinely-collected health Data (RECORD) statement: methods for arriving at consensus and developing reporting guidelines. *PLoS One* 2015;10:e0125620.
- Moher D, Schulz KF, Simera I, et al. Guidance for developers of health research reporting guidelines. *PLoS Med* 2010;7:e1000217.
- Glaziov P, Altman DG, Bossuyt P, et al. Reducing waste from incomplete or unusable reports of biomedical research. *Lancet* 2014;383:267-76.
- Blotière PO, Weill A, Ricordeau P, et al. Perforations and haemorrhages after colonoscopy in 2010: a study based on comprehensive French health insurance data (SNIIRAM). *Clin Res Hepatol Gastroenterol* 2014;38:112-7.
- Siregar S, Pouw ME, Moons KG, et al. The Dutch hospital standardised mortality ratio (HSMR) method and cardiac surgery: benchmarking in a national cohort using hospital administration data versus a clinical database. *Heart* 2014;100:702-10.
- Price SD, Holman CD, Sanfilippo FM, et al. Use of case-time-control design in pharmacovigilance applications: exploration with high-risk medications and unplanned hospital admissions in the Western Australian elderly. *Pharmacoepidemiol Drug Saf* 2013;22:1159-70.
- Gross CP, Andersen MS, Krumholz HM, et al. Relation between Medicare screening reimbursement and stage at diagnosis for older patients with colon cancer. *JAMA* 2006;296:2815-22.
- Vandenbroucke JP. Observational research, randomised trials, and two views of medical science. *PLoS Med* 2008;5:e67.
- Smith GD, Ebrahim S. Data dredging, bias, or confounding. *BMJ* 2002;325: 1437-8.
- Prokosh HU, Ganslandt T. Perspectives for medical informatics. Reusing the electronic medical record for clinical research. *Methods Inf Med* 2009; 48:38-44.
- Benchimol EI, Manuel DG, Guttman A, et al. Changing age demographics of inflammatory bowel disease in Ontario, Canada: a population-based cohort study of epidemiology trends. *Inflamm Bowel Dis* 2014;20:1761-9.
- Ducharme R, Benchimol EI, Deeks SL, et al. Validation of diagnostic codes for intussusception and quantification of childhood intussusception incidence in Ontario, Canada: a population-based study. *J Pediatr* 2013;163:1073-9.e3.

28. Herrett E, Shah AD, Boggan R, et al. Completeness and diagnostic validity of recording acute myocardial infarction events in primary care, hospital care, disease registry, and national mortality records: cohort study. *BMJ* 2013;346:f2350.
29. van Herk-Sukel MP, van de Poll-Franse LV, Lemmens VE, et al. New opportunities for drug outcomes research in cancer patients: the linkage of the Eindhoven Cancer Registry and the PHARMO Record Linkage System. *Eur J Cancer* 2010;46:395-404.
30. Fosbøl EL, Granger CB, Peterson ED, et al. Prehospital system delay in ST-segment elevation myocardial infarction care: a novel linkage of emergency medicine services and in hospital registry data. *Am Heart J* 2013;165:363-70.
31. Manuel DG, Rosella LC, Stukel TA. Importance of accurately identifying disease in studies using electronic health records. *BMJ* 2010;341:c4226.
32. *Narcolepsy in association with pandemic influenza vaccination: a multi-country European epidemiological investigation* [rapport technique]. Stockholm: European Centre for Disease Prevention and Control;2012.
33. Lewis JD, Schinnar R, Bilker WB, et al. Validation studies of The Health Improvement Network (THIN) database for pharmacoepidemiology research. *Pharmacoepidemiol Drug Saf* 2007;16:393-401.
34. Sørensen HT, Sabroe S, Olsen J. A framework for evaluation of secondary data sources for epidemiological research. *Int J Epidemiol* 1996;25:435-42.
35. Baron JA, Lu-Yao G, Barrett J, et al. Internal validation of Medicare claims data. *Epidemiology* 1994;5:541-4.
36. Marston L, Carpenter JR, Walters KR, et al. Smoker, ex-smoker or non-smoker? The validity of routinely recorded smoking status in UK primary care: a cross-sectional study. *BMJ Open* 2014;4:e004958.
37. Hardelid P, Dattani N, Gilbert R; Programme Board of the Royal College of Paediatrics and Child Health; Child Death Overview Working Group. Estimating the prevalence of chronic conditions in children who die in England, Scotland and Wales: a data linkage cohort study. *BMJ Open* 2014;4:e005331.
38. Murray J, Bottle A, Sharland M, et al.; Medicines for Neonates Investigator Group. Risk factors for hospital admission with RSV bronchiolitis in England: a population-based birth cohort study. *PLoS One* 2014;9:e89186.
39. Berry JG, Hall M, Hall DE, et al. Inpatient growth and resource use in 28 children's hospitals: a longitudinal, multi-institutional study. *JAMA Pediatr* 2013;167:170-7.
40. Shahian DM, Wolf RE, Iezzoni LI, et al. Variability in the measurement of hospital-wide mortality rates. *N Engl J Med* 2010;363:2530-9.
41. Springate DA, Kontopantelis E, Ashcroft DM, et al. ClinicalCodes: an online clinical codes repository to improve the validity and reproducibility of research using electronic medical records. *PLoS One* 2014;9:e99825.
42. Dommett RM, Redaniel MT, Stevens MC, et al. Features of childhood cancer in primary care: a population-based nested case-control study. *Br J Cancer* 2012;106:982-7.
43. Tsang C, Bottle A, Majeed A, et al. Adverse events recorded in English primary care: observational study using the General Practice Research Database. *Br J Gen Pract* 2013;63:e534-42.
44. Harron K, Goldstein H, Wade A, et al. Linkage, evaluation and analysis of national electronic healthcare data: application to providing enhanced blood-stream infection surveillance in paediatric intensive care. *PLoS One* 2013;8:e85278.
45. Adams MM, Wilson HG, Casto DL, et al. Constructing reproductive histories by linking vital records. *Am J Epidemiol* 1997;145:339-48.
46. Ford JB, Roberts CL, Taylor LK. Characteristics of unmatched maternal and baby records in linked birth records and hospital discharge data. *Paediatr Perinat Epidemiol* 2006;20:329-37.
47. Weiskopf NG, Weng C. Methods and dimensions of electronic health record data quality assessment: enabling reuse for clinical research. *J Am Med Inform Assoc* 2013;20:144-51.
48. Sandall J, Murrells T, Dodwell M, et al. The efficient use of the maternity workforce and the implications for safety and quality in maternity care: a population-based, cross-sectional study. *Health Serv Deliv Res* 2014;2(38).
49. Welch C, Petersen I, Walters K, et al. Two-stage method to remove population- and individual-level outliers from longitudinal data in a primary care database. *Pharmacoepidemiol Drug Saf* 2012;21:725-32.
50. Van den Broeck J, Cunningham SA, Eckels R, et al. Data cleaning: detecting, diagnosing, and editing data abnormalities. *PLoS Med* 2005;2:e267.
51. Bohensky MA, Jolley D, Sundararajan V, et al. Data linkage: a powerful research tool with potential problems. *BMC Health Serv Res* 2010;10:346.
52. Harron K, Wade A, Muller-Pebody B, et al. Opening the black box of record linkage. *J Epidemiol Community Health* 2012;66:1198.
53. Lariscy JT. Differential record linkage by Hispanic ethnicity and age in linked mortality studies: implications for the epidemiologic paradox. *J Aging Health* 2011;23:1263-84.
54. Dinan MA, Curtis LH, Carpenter WR, et al. Variations in use of PET among Medicare beneficiaries with non-small cell lung cancer, 1998-2007. *Radiology* 2013;267:807-17.
55. Horsfall L, Walters K, Petersen I. Identifying periods of acceptable computer usage in primary care research databases. *Pharmacoepidemiol Drug Saf* 2013;22:64-9.
56. Gerber DE, Laccetti AL, Xuan L, et al. Impact of prior cancer on eligibility for lung cancer clinical trials. *J Natl Cancer Inst* 2014;106:dju302.
57. Carrara G, Scirè CA, Zambon A, et al. A validation study of a new classification algorithm to identify rheumatoid arthritis using administrative health databases: case-control and cohort diagnostic accuracy studies. Results from the RECO linkage On Rheumatic Diseases study of the Italian Society for Rheumatology. *BMJ Open* 2015;5:e006029.
58. Rait G, Walters K, Griffin M, et al. Recent trends in the incidence of recorded depression in primary care. *Br J Psychiatry* 2009;195:520-4.
59. Wijlaars LP, Nazareth I, Petersen I. Trends in depression and antidepressant prescribing in children and adolescents: a cohort study in The Health Improvement Network (THIN). *PLoS One* 2012;7:e33181.
60. Jeon CY, Pandolfi SJ, Wu B, et al. The association of statin use after cancer diagnosis with survival in pancreatic cancer patients: a SEER-Medicare analysis. *PLoS One* 2015;10:e0121783.
61. Pruitt Z, Pracht E. Upcoding emergency admissions for non-life-threatening injuries to children. *Am J Manag Care* 2013;19:917-24.
62. McLintock K, Russell AM, Alderson SL, et al. The effects of financial incentives for case finding for depression in patients with diabetes and coronary heart disease: interrupted time series analysis. *BMJ Open* 2014;4:e005178.
63. Brunt CS. CPT fee differentials and visit upcoding under Medicare Part B. *Health Econ* 2011;20:831-41.
64. Walters K, Rait G, Griffin M, et al. Recent trends in the incidence of anxiety diagnoses and symptoms in primary care. *PLoS One* 2012;7:e41670.
65. Nilson F, Bonander C, Andersson R. The effect of the transition from the ninth to the tenth revision of the International Classification of Diseases on external cause registration of injury morbidity in Sweden. *Inj Prev* 2015;21:189-94.
66. Jagai JS, Smith GS, Schmid JE, et al. Trends in gastroenteritis-associated mortality in the United States, 1985-2005: variations by ICD-9 and ICD-10 codes. *BMC Gastroenterol* 2014;14:211.
67. European Network of Centres for Pharmacoepidemiology and Pharmacovigilance. 4.2.2.6. Unmeasured confounding. Dans: *ENCEPP guide on methodological standards in pharmacoepidemiology*. Londres: European Medicines Agency;2015 [mise à date le 5 octobre 2018]. Accessible ici: www.encepp.eu/standards_and_guidances/methodologicalGuide4_2_2_6.shtml (consulté le 16 octobre 2018).
68. Toh S, García Rodríguez LA, Hernán MA. Confounding adjustment via a semi-automated high-dimensional propensity score algorithm: an application to electronic medical records. *Pharmacoepidemiol Drug Saf* 2011;20:849-57.
69. Stukel TA, Fisher ES, Wennberg DE, et al. Analysis of observational studies in the presence of treatment selection bias: effects of invasive cardiac management on AMI survival using propensity score and instrumental variable methods. *JAMA* 2007;297:278-85.
70. Austin PC. An introduction to propensity score methods for reducing the effects of confounding in observational studies. *Multivariate Behav Res* 2011;46:399-424.
71. Sterne JA, White IR, Carlin JB, et al. Multiple imputation for missing data in epidemiological and clinical research: potential and pitfalls. *BMJ* 2009;338:b2393.
72. Freemantle N, Marston L, Walters K, et al. Making inferences on treatment effects from real world data: propensity scores, confounding by indication, and other perils for the unwary in observational research. *BMJ* 2013;347:f6409.
73. Marston L, Carpenter JR, Walters KR, et al. Issues in multiple imputation of missing data for large general practice clinical databases. *Pharmacoepidemiol Drug Saf* 2010;19:618-26.
74. Benchimol EI, To T, Griffiths AM, et al. Outcomes of pediatric inflammatory bowel disease: socioeconomic status disparity in a universal-access healthcare system. *J Pediatr* 2011;158:960-7.e1-4.
75. Nassar N, Dixon G, Bourke J, et al. Autism spectrum disorders in young children: effect of changes in diagnostic practices. *Int J Epidemiol* 2009;38:1245-54.
76. Tan GH, Bhoo-Pathy N, Taib NA, et al. The Will Rogers phenomenon in the staging of breast cancer — Does it matter? *Cancer Epidemiol* 2015;39:115-7.
77. Taljaard M, Tuna M, Bennett C, et al. Cardiovascular Disease Population Risk Tool (CVDpORT): predictive algorithm for assessing CVD risk in the community setting. A study protocol. *BMJ Open* 2014;4:e006701.
78. Guttman A, Schull MJ, Vermeulen MJ, et al. Association between waiting times and short term mortality and hospital admission after departure from emergency department: population based cohort study from Ontario, Canada. *BMJ* 2011;342:d2983.

79. Nicol A, Caruso J, Archambault É. *Open data access policies and strategies in the European research area and beyond*. Montréal: Science-Metrix; 2013.
80. Fuller T, Pearson M, Peters J, et al. What affects authors' and editors' use of reporting guidelines? Findings from an online survey and qualitative interviews. *PLoS One* 2015;10:e0121585.
81. Turner L, Shamseer L, Altman DG, et al. CONSolidated Standards Of Reporting Trials (CONSORT) and the completeness of reporting of randomised controlled trials (RCTs) published in medical journals. *Cochrane Database Syst Rev* 2012; (11):MR000030.
82. Armstrong R, Waters E, Moore L, et al. Improving the reporting of public health intervention research: advancing TREND and CONSORT. *J Public Health (Oxf)* 2008;30:103-9.
83. Moher D, Cook DJ, Eastwood S, et al. Improving the quality of reports of meta-analyses of randomised controlled trials: the QUOROM statement. Quality of Reporting of Meta-analyses. *Lancet* 1999;354:1896-900.
84. Prady SL, Richmond SJ, Morton VM, et al. A systematic evaluation of the impact of STRICTA and CONSORT recommendations on quality of reporting for acupuncture trials. *PLoS One* 2008;3:e1577.

Editor's note: The original English version of this article is available at *PLoS Med* 2015;12: e1001885.

Intérêts concurrents: Liam Smeeth déclare des subventions de la Wellcome Trust, du Medical Research Council, du National Institute for Health Research, de GlaxoSmithKline, de la British Heart Foundation et du Diabetes UK en dehors du travail soumis et siège au conseil d'administration de la British Heart Foundation. Katie Harron et Sinéad Langan déclarent des subventions de Wellcome Trust pendant que le travail soit mené. Henrik Sørensen déclare que le Département d'épidémiologie clinique de l'université d'Aarhus prend partie à des études financées par des subventions de recherche accordées à l'université par diverses firmes et administrées par celle-ci. Les autres auteurs n'ont aucun conflit d'intérêt à déclarer.

La version originale de cet article a été révisée par des pairs chez *PLoS Medicine*.

Affiliations: Institut de recherche du Centre hospitalier pour enfants de l'est de l'Ontario (Benchimol); Département de pédiatrie (Benchimol), Université d'Ottawa; École d'épidémiologie et de santé publique (Benchimol, Moher), Université d'Ottawa, Ottawa, Ont.; ICES (Benchimol, Guttman), Toronto, Ont.; London School of Hygiene and Tropical Medicine (Smeeth, Harron, Langan), Londres, Royaume-Uni; Department of Paediatrics (Guttman), The Hospital for Sick Children; Institute of Health Policy, Management and Evaluation (Guttman), University of Toronto, Toronto, Ont.; Institut de recherche de l'Hôpital d'Ottawa (Moher), Ottawa, Ont.; Département de soins primaires et santé publique (Petersen), University College London, Londres, Royaume-Uni; Département d'épidémiologie clinique (Sørensen), université d'Aarhus, Aarhus, Danemark; Management des organisations de santé (EA 7348 MOS) (Januel), Institut du Management, École des hautes études en santé publique, Rennes, France; Chaire d'excellence en Management de la santé (Januel), Université Sorbonne Paris Cité, Paris, France; Cochrane Suisse (von

Elm), Institut universitaire de médecine sociale et préventive, Université de Lausanne, Lausanne, Suisse

Contributeurs: Les auteurs principaux (Erik von Elm et Sinéad Langan) ont contribué de manière égale à ce travail. Rédaction de la première version du manuscrit: Eric Benchimol et Sinéad Langan. Contribution à la rédaction du manuscrit et approbation des résultats et conclusions: Eric Benchimol, Liam Smeeth, Astrid Guttman, Katie Harron, David Moher, Irene Petersen, Henrik Sørensen, Erik von Elm, Sinéad Langan et membres du comité de travail RECORD. Rédaction de la traduction: Jean-Marie Januel. Chaque auteur a approuvé la version finale à être publiée et a consenti à être responsable de tout aspect du travail.

Financement: Ce travail a été mené avec le soutien financier des Instituts de recherche en santé du Canada (subvention 130512), des Fonds national suisse de la recherche scientifique (subvention IZ32Z0_147388/1) et du Département d'épidémiologie clinique de l'université d'Aarhus.

Avis de non-responsabilité: Les bailleurs de fonds n'ont joué aucun rôle dans la conception du travail, la collecte et l'analyse des données, la décision de publier ou la préparation du manuscrit.

Traduction: Cet article a été publié antérieurement en anglais (Benchimol EI, Smeeth L, Guttman A, et al.; RECORD Working Committee. The REporting of studies Conducted using Observational Routinely-collected health Data (RECORD) statement. *PLoS Med* 2015;12: e1001885). Il a été traduit par Jean-Paul Salameh (Institut de recherche du Centre hospitalier pour enfants de l'est de l'Ontario et École d'épidémiologie et de santé publique, Université d'Ottawa, Ottawa, Ont.) et Yara Boutros (École de traducteurs et d'interprètes de Beyrouth, Université Saint-Joseph de Beyrouth, Beyrouth, Liban).

Remerciements: Les auteurs remercient les parties prenantes qui ont participé aux enquêtes d'avoir accordé la priorité aux thèmes à inclure dans la liste de contrôle

(annexe 1, disponible sur www.cmaj.ca/lookup/suppl/doi:10.1503/cmaj.181309/-/DC1). Les auteurs sont également reconnaissants aux membres du groupe d'initiative STROBE (Strengthening the Reporting of Observational Studies in Epidemiology), qui ont guidé et soutenu la création de RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Health Data). Les auteurs remercient chaleureusement Pauline Quach et Danielle Birman, coordonnatrices de la recherche à RECORD, ainsi que Andrew Perlmutter, concepteur du site Web et administrateur de www.record-statement.org. Les auteurs sont également reconnaissants pour la contribution de toutes les parties prenantes qui ont participé à l'élaboration de ces lignes directrices.

***Membres du comité de travail RECORD (Reporting of Studies Conducted Using Observational Routinely Collected Health Data):** Douglas Altman (Centre for Statistics in Medicine, Oxford University), Nicholas de Klerk (University of Western Australia), Lars G. Hemkens (Hôpital universitaire de Bâle), David Henry (University of Toronto et Institute for Clinical Evaluative Sciences), Jean-Marie Januel (Université de Lausanne), Marie-Annick Le Pogam (Institut universitaire de médecine sociale et préventive, Centre hospitalier universitaire vaudois), Douglas Manuel (Institut de recherche de l'Hôpital d'Ottawa et Université d'Ottawa), Kirsten Patrick (rédactrice, *CMAJ*), Pablo Perel (London School of Hygiene and Tropical Medicine), Patrick S. Romano (University of California, Davis et co-rédacteur en chef, *Health Services Research*), Peter Tugwell (Université d'Ottawa et rédacteur en chef, *Journal of Clinical Epidemiology*), Joan Warren (National Institutes of Health/National Cancer Institute), Wim Weber (rédacteur en chef européen, *BMJ*) et Margaret Winker (ancienne rédactrice, *PLoS Medicine* et secrétaire, World Association of Medical Editors).

Correspondance: Eric Benchimol, ebenchimol@cheo.on.ca, Sinéad Langan, sinead.langan@lshtm.ac.uk