

UNIVERSITÉ DU QUÉBEC À MONTRÉAL

MÉCANISMES IMPLIQUÉS DANS L'ASSOCIATION ENTRE L'INDICE DE
MASSE CORPORELLE ET L'ÉPAISSEUR DE L'INTIMA MÉDIA
CAROTIDIEN

MÉMOIRE
PRÉSENTÉ
COMME EXIGENCE PARTIELLE
DE LA MAÎTRISE EN MATHÉMATIQUES

PAR
FRANCOIS FREDDY ATEBA

MAI 2023

UNIVERSITÉ DU QUÉBEC À MONTRÉAL
Service des bibliothèques

Avertissement

La diffusion de ce mémoire se fait dans le respect des droits de son auteur, qui a signé le formulaire *Autorisation de reproduire et de diffuser un travail de recherche de cycles supérieurs* (SDU-522 – Rév.04-2020). Cette autorisation stipule que «conformément à l'article 11 du Règlement no 8 des études de cycles supérieurs, [l'auteur] concède à l'Université du Québec à Montréal une licence non exclusive d'utilisation et de publication de la totalité ou d'une partie importante de [son] travail de recherche pour des fins pédagogiques et non commerciales. Plus précisément, [l'auteur] autorise l'Université du Québec à Montréal à reproduire, diffuser, prêter, distribuer ou vendre des copies de [son] travail de recherche à des fins non commerciales sur quelque support que ce soit, y compris l'Internet. Cette licence et cette autorisation n'entraînent pas une renonciation de [la] part [de l'auteur] à [ses] droits moraux ni à [ses] droits de propriété intellectuelle. Sauf entente contraire, [l'auteur] conserve la liberté de diffuser et de commercialiser ou non ce travail dont [il] possède un exemplaire.»

TABLE DES MATIÈRES

LISTE DES FIGURES	v
LISTE DES TABLEAUX	vii
RÉSUMÉ	ix
CHAPITRE I ÉTAT DE L'ART EN ANALYSE DE MÉDIATION ET INFÉRENCE CAUSALE	14
1.1 Approches de médiation traditionnelles	15
1.1.1 Le concept de médiation	15
1.1.2 La méthode de Baron et Kenny	16
1.1.3 La méthode du produit	19
1.1.4 La méthode de la différence	20
1.1.5 Équivalence entre la méthode du produit et la méthode de la différence	20
1.1.6 Limites des méthodes traditionnelles de médiation	21
1.2 Inférence causale basée sur le cadre contrefactuel	22
1.2.1 Notations préliminaires	23
1.2.2 Effet total : définitions, hypothèses et identification	24
1.2.3 Estimation de l'effet causal moyen	27
1.2.4 Effets direct et indirect : Notations, définitions et identification	31
1.2.5 Identification	36
1.2.6 Formules d'identification	38
1.2.7 Méthodes de régression pour les effets directs et indirects	40
1.2.8 Estimations	42
CHAPITRE II INTRODUCTION AUX CONCEPTS DE LA RANDO- MISATION MENDÉLIENNE	46
2.1 Introduction à la méthode des variables instrumentales	47

2.1.1	Notion d'instrument	47
2.1.2	La méthode d'instrumentation	48
2.1.3	Les hypothèses d'une variable instrumentale	49
2.2	Concept de randomisation mendélienne	52
2.2.1	Quelques concepts et définitions de génétique	53
2.3	Méthodes d'estimation en randomisation mendélienne : les méthodes 2SLS et de Wald	56
2.3.1	La méthode de Wald	57
2.3.2	La méthode <i>2SLS</i>	59
CHAPITRE III APPLICATION DES MÉTHODES DE RANDOMISA- TION MENDÉLIENNE EN ANALYSE DE MÉDIATION		64
3.1	La RM en deux étapes	66
3.2	Présentation des simulations	70
3.2.1	Équations génératrices des données de simulation	71
3.2.2	Scénarios de simulation	72
3.2.3	Les scénarios de simulation	73
3.2.4	Analyse des données simulées	75
3.3	Résultats	76
3.3.1	Commentaires sur les résultats	77
3.3.2	Les types d'erreurs détectées	77
3.3.3	Impact des erreurs détectées sur les conclusions de Carter et al. (2021)	84
CHAPITRE IV ÉTUDE DE L'EFFET CAUSAL DE L'OBÉSITÉ SUR L'ÉPAISSEUR DE L'INTIMA MÉDIA CAROTIDIEN MÉDIÉ PAR LA PROTÉINE C-RÉACTIVE		86
4.1	Description de l'application	87
4.2	Données et variables	88
4.2.1	Source des données	88
4.2.2	Détermination de l'échantillon analytique	89

4.2.3 Variables utilisées et échelles de mesures	92
4.2.4 Analyses	95
CHAPITRE V CONCLUSION	105
5.1 Rappel des objectifs de l'étude	105
5.2 Limites et forces de notre étude	107
ANNEXE A	109
ANNEXE B	111
RÉFÉRENCES	124

LISTE DES FIGURES

Figure		Page
1.1	Diagramme causal illustrant le modèle de médiation conceptuel. .	16
1.2	Diagramme causal illustrant le modèle de médiation conceptuel de Baron et Kenny (1986) et incluant les coefficients décrits dans les Équations (1.1) à (1.3).	17
1.3	Identification des effets de médiation, le chemin en rouge représente une possible violation de H_4 (confusion induite par l'exposition). .	37
2.1	Diagramme causal pour la relation entre un instrument Z , une exposition A , des facteurs de confusion non mesurés U et une réponse Y	48
2.2	Diagramme causal illustrant des cas de violation des hypothèses $H_2 - H_3$ de la méthode des IVs pour une variable instrumentale Z , une exposition A , des facteurs de confusion non mesurés U et une réponse Y	51
2.3	Diagramme causal illustrant le phénomène de la stratification dans lequel la structure populationnelle SP confond l'association entre l'instrument Z associé à l'exposition A et la réponse Y	55
2.4	Diagramme causal illustrant le phénomène de la pléiotropie verticale dans lequel l'instrument Z est associé à un autre phénotype W situé en amont de l'exposition A ; U est un facteur de confusion de la relation entre A et Y	55
2.5	Diagramme causal illustrant le phénomène de la pléiotropie horizontale dans lequel l'instrument Z associé à l'exposition A est indépendamment associé à la réponse Y directement ou indirectement par le biais d'un autre phénotype W_2	56
3.1	Diagramme causal illustrant les hypothèses causales dans la RM en deux étapes. L'hypothèse H_1 est matérialisée par la présence des flèches $Z_1 \rightarrow A$ et $Z_2 \rightarrow M$	68

4.1	Diagramme causal illustrant le modèle de médiation supposé dans l'application d'intérêt.	87
4.2	Diagramme illustrant la constitution de l'échantillon final retenu pour les analyses.	91

LISTE DES TABLEAUX

Tableau		Page
3.1	Définition des paramètres de simulation dans Carter et al. (2021).	73
3.2	Présentation des scénarios de simulation avec confusion non mesurée considérés dans notre étude de l'article de Carter et al. (2021).	74
3.3	Définition des métriques de performance des estimateurs considérés dans les simulations.	75
3.4	Fréquences d'erreurs d'arrondi, de calcul et de recopie des méthodes de la <i>RM</i> en deux étapes (<i>TSMR</i>) et du produit des coefficients standard (<i>Prod</i>) avec confusion non mesurée pour les effets de médiation et l'effet total.	78
3.5	Estimations originale et répliquée de l'effet total (<i>ET</i>) pour les méthodes <i>RM</i> en deux étapes (<i>TSMR</i>) et du produit des coefficients standard (<i>Prod</i>) avec confusion non mesurée.	80
3.6	Estimations originale et répliquée de l'effet direct (<i>ED</i>) pour les méthodes <i>RM</i> en deux étapes (<i>TSMR</i>) et du produit des coefficients standard (<i>Prod</i>) avec confusion non mesurée.	81
3.7	Estimations originale et répliquée de l'effet indirect (<i>EI</i>) pour les méthodes <i>RM</i> en deux étapes (<i>TSMR</i>) et du produit des coefficients standard (<i>Prod</i>) avec confusion non mesurée.	82
3.8	Estimations originale et répliquée de la proportion médiée (<i>PM</i>) pour les méthodes <i>RM</i> en deux étapes (<i>TSMR</i>) et du produit des coefficients standard (<i>Prod</i>) avec confusion non mesurée.	83
4.1	Description des 5 <i>SNPs</i> significatifs associés à la <i>CRP</i> (extrait de Said et al. (2022)).	94
4.2	Description de cinq <i>SNPs</i> communs significatifs associés à l' <i>IMC</i> (extraits de Locke et al. (2015)).	95
4.3	Analyse descriptive des variables de l'échantillon d'étude.	100

4.4	Forces de l'instrument <i>SNP rs1558902</i> pour l'exposition (<i>IMC</i>) et de l'instrument <i>SNP rs61946383</i> pour le médiateur (<i>CRP</i>). Associations obtenues par moindres carrés ordinaires.	101
4.5	Estimations des effets de médiation du taux de protéine C-réactive dans le sang dans l'association entre l'indice de masse corporelle et l'épaisseur de l'intima média-carotidien par les méthodes du produit des coefficients et de la randomisation mendélienne en deux étapes.	102
A.1	Données non génétiques mesurées de la cohorte globale de l' <i>ÉLCV</i> utilisées dans notre analyse.	110

RÉSUMÉ

L'analyse de médiation cherche à étudier les mécanismes qui définissent les relations entre trois entités : une variable d'exposition, une variable réponse et une ou un ensemble de variables intermédiaires appelées médiateurs. Les méthodes traditionnelles de variables non instrumentales pour l'analyse de la médiation rencontrent un certain nombre de difficultés méthodologiques y compris celles dues à la présence de biais de confusion entre l'exposition, le médiateur et la réponse. La randomisation mendélienne (*RM*) peut être utilisée dans le but d'améliorer l'inférence causale pour l'analyse de la médiation. Nous décrivons deux approches récemment utilisées pour tenter de pallier au problème de confusion non mesurée en médiation en y intégrant la *RM*. L'objectif principal de ce mémoire est de comprendre, explorer et illustrer les différences entre une méthode standard de médiation (la méthode du produit des coefficients) et la *RM* en deux étapes (*TSMR*) pour la médiation dans le contexte où il existe des variables de confusion non mesurées. L'étude réalisée a permis de faire un état de l'art en analyse de médiation et en inférence causale, d'expliquer la théorie sur les variables instrumentales et de présenter des concepts en *RM*. Une contribution apportée par ce mémoire est la réplication de certaines simulations de Carter et al. (2021) ainsi que la correction d'erreurs qui se sont glissées dans l'article original. Les deux méthodes explorées ont été évaluées par une application sur les données extraites de la cohorte de l'Étude longitudinale canadienne sur le vieillissement qui est une plateforme de recherche nationale contenant des données longitudinales dont un des principaux objectifs est de comprendre les facteurs qui déterminent la qualité du vieillissement. Plus spécifiquement, le but de l'application est d'examiner le rôle médiateur de la protéine C-réactive dans l'association entre l'indice de masse corporelle et l'épaisseur de l'intima média carotidien.

Mots-clés : médiation par inférence causale, variable instrumentale, randomisation mendélienne, réplication, simulation.

INTRODUCTION

L'analyse de médiation est une approche statistique permettant de comprendre les processus par lesquels une variable d'exposition affecte ou cause une variable réponse (MacKinnon *et al.*, 2012). En effet, l'analyse de médiation permet d'étudier les mécanismes qui définissent les relations entre trois entités : la variable d'exposition (A), la variable réponse (Y) et une ou un ensemble de variables intermédiaires appelées médiateurs (M). En bref, le ou les médiateurs servent à clarifier la nature de la relation entre la variable d'exposition (A) et la variable réponse (Y) (Baron et Kenny, 1986; MacKinnon et Luecken, 2008). En épidémiologie, l'analyse de médiation peut améliorer la compréhension étiologique des pathologies et de nombreux exemples d'application modernes existent (Richiardi *et al.*, 2013; Kaufman *et al.*, 2004; Rizzo *et al.*, 2022).

La méthode de Baron et Kenny (1986) et celles découlant de celle-ci (Hernan et Robins, 2020; VanderWeele, 2015; VanderWeele, 2016) sont utilisées couramment en recherche appliquée pour les analyses de médiation. Récemment, une plus grande attention a été portée à certaines hypothèses fortes inhérentes aux modèles de médiation, liées à la confusion, ce qui a mené au développement du cadre de travail causal en médiation (MacKinnon *et al.*, 2012; VanderWeele, 2011; VanderWeele *et al.*, 2012; VanderWeele *et al.*, 2014). Cependant, une faiblesse des modèles de médiation estimés à partir de données observationnelles est qu'ils ne gèrent pas la confusion non mesurée, c'est-à-dire qu'ils ne prennent pas en compte les variables non mesurées qui sont communes à la fois de l'exposition, du médiateur et de la réponse.

Récemment, des approches ont tenté de pallier au problème de confusion non mesurée en médiation en y intégrant la randomisation mendélienne (RM). La RM

est une approche statistique relativement récente née à la faveur des progrès dans les domaines de la génétique ainsi que de la disponibilité des données génétiques (Davey Smith et Ebrahim, 2003; Burgess *et al.*, 2012; Davey Smith et Hemani, 2014). La *RM* est une approche par variables instrumentales (*IV*) qui utilise des variants génétiques comme *IV* (Burgess *et al.*, 2015) dans le but de faire une estimation de l'effet causal d'une exposition A sur une réponse Y qui, sous certaines hypothèses, est exempte de biais dû aux variables de confusion non mesurées (Davey Smith et Ebrahim, 2003; Burgess *et al.*, 2012; Davey Smith et Hemani, 2014; Burgess *et al.*, 2015).

La *RM* peut être utilisée pour améliorer l'inférence causale en analyse de médiation afin d'estimer les effets direct et indirect de l'exposition sur la réponse en présence de variables de confusion non mesurées (Burgess *et al.*, 2015; Burgess et Thompson, 2015). Deux approches utilisées dans ce contexte sont développées par Carter *et al.* (2021). Ces auteurs proposent une application pratique de la *RM* en deux étapes (Burgess *et al.*, 2015; Relton et Davey Smith, 2012; Wu *et al.*, 2020) et de la randomisation mendélienne multivariable (*RMMV*) (Burgess *et al.*, 2017b; Sanderson, 2021) en médiation pour réduire les biais des estimateurs d'effets.

L'objectif principal de ce mémoire est de comprendre, explorer et illustrer les différences entre une méthode standard de médiation (la méthode du produit des coefficients) et la *RM* en deux étapes pour la médiation dans le contexte où il existe des variables de confusions non mesurées. Une contribution apportée par ce mémoire est la réplication de certaines simulations de Carter *et al.* (2021) ainsi que la correction d'erreurs qui se sont glissées dans l'article original.

En guise d'illustration, nous procédons à l'analyse d'un jeu de données issu de l'Étude longitudinale canadienne sur le vieillissement (*ÉLCV*). Le but de l'ana-

lyse est d'évaluer, à l'aide de modèles d'analyse de médiation causale, le rôle du médiateur putatif, la protéine C-réactive (*CRP*), comme courroie de transmission de l'effet de l'indice de masse corporelle (*IMC*) sur l'épaisseur de l'intima media carotidien (*cIMT*), cette dernière étant un marqueur de risque du développement de maladies cardiovasculaires. L'hypothèse sous-jacente est que l'adiposité affecte le *cIMT* par l'action de marqueurs de l'inflammation systémique comme le niveau de *CRP* dans le sang (Sproston *et al.*, 2018).

La structure du mémoire est la suivante. Le Chapitre I présente l'état de l'art en analyse de médiation et en inférence causale. Le Chapitre II est consacré à la théorie sur les *IV* d'une part et sur la présentation des concepts en *RM* d'autre part. Le Chapitre III porte sur les méthodes et les approches étudiées par Carter *et al.* (2021) ainsi que de leur évaluation par simulations. Le Chapitre IV est consacré à l'analyse du jeu de données réelles issu de la cohorte de l'*ÉLCV*. Le mémoire se termine par un chapitre de conclusion qui fait la synthèse de tous les chapitres et qui énonce des perspectives pour les travaux futurs.

CHAPITRE I

ÉTAT DE L'ART EN ANALYSE DE MÉDIATION ET INFÉRENCE CAUSALE

Ce chapitre donne un aperçu de l'état de l'art dans le domaine de l'analyse de médiation et de l'inférence causale. Il donne dans un survol détaillé, la théorie, les hypothèses, les limites des méthodes traditionnelles de médiation et aide à la compréhension des concepts préliminaires de l'inférence causale. Le chapitre est organisé comme suit. La première section traite de trois approches en analyse de médiation traditionnelle : l'approche de Baron et Kenny, la méthode du produit et la méthode de la différence, toutes popularisées par Baron et Kenny (1986). Cette section discute également de certaines limites des méthodes traditionnelles de médiation. La deuxième section présente une introduction à la médiation causale basée sur le cadre de travail contrefactuel et aborde également les hypothèses requises pour l'interprétation causale de différents types d'effets, soit l'effet direct contrôlé et les effets direct et indirect naturels.

Dans ce chapitre, nous portons une attention particulière au traitement du biais de confusion fait par chacune des méthodes considérées. En effet, dans certaines situations, deux variables, par exemple l'exposition et la réponse, sont expliquées

par une ou plusieurs variables communes. Ces variables sont dites de confusion car elles ont le potentiel de fausser l'interprétation de l'association entre les deux variables de départ. La difficulté à identifier les variables de confusion ou à en tenir compte de manière appropriée dans les analyses peut engendrer la présence d'un biais (erreur) dans l'estimation des effets direct, indirect et total. Finalement, nous supposons dans ce chapitre que le médiateur et la variable réponse sont tous les deux continus.

1.1 Approches de médiation traditionnelles

L'analyse de médiation est une méthode statistique permettant d'étudier l'effet d'une variable d'exposition A sur une variable réponse Y à travers une variable intermédiaire M appelée médiateur. Les fondements de l'analyse de médiation (Woodworth, 1928; Alwin et Hauser, 1975; Judd et Kenny, 1981) sont antérieurs aux travaux séminaux de Baron et Kenny (1986). L'article de Baron et Kenny (1986) a cependant eu pour effet de populariser l'analyse de médiation et de la rendre accessible à une plus large communauté de chercheurs issus de plusieurs domaines, favorisant ainsi la multiplicité de ses applications. Dans cette section, nous discutons des méthodes de médiation traditionnelles qui se déclinent en trois méthodes : la méthode dite de Baron et Kenny, la méthode du produit et la méthode de la différence.

1.1.1 Le concept de médiation

On utilise le diagramme causal de la Figure 1.1 pour schématiser le rôle d'une variable intermédiaire dans la relation causale entre une exposition et une réponse (Baron et Kenny, 1986).

En effet, ce diagramme schématise un modèle de médiation simple qui permet

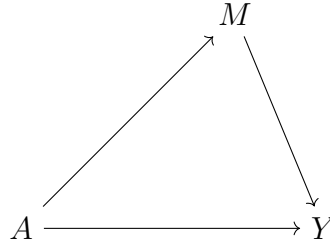


FIGURE 1.1. Diagramme causal illustrant le modèle de médiation conceptuel.

de définir conceptuellement les notions d'effets direct et indirect de l'exposition A sur la réponse Y à travers M . Le but ultime d'une analyse de médiation est la décomposition de l'effet total de l'exposition sur la réponse, en quantifiant la magnitude de l'effet direct et de l'effet indirect. L'effet indirect est défini comme l'effet de A sur Y à travers M , représenté par la chaîne $A \rightarrow M \rightarrow Y$. De façon complémentaire, l'effet direct est défini comme l'effet de A sur Y à travers des mécanismes qui n'impliquent pas le médiateur M , représenté par la flèche directe $A \rightarrow Y$. L'effet total de A sur Y peut alors être vu comme l'aggrégation des effets indirect et direct, soit la combinaison des chaînes $A \rightarrow M \rightarrow Y$ et $A \rightarrow Y$.

1.1.2 La méthode de Baron et Kenny

La méthode de Baron et Kenny permet d'établir l'existence d'un effet médié de l'exposition sur la réponse à travers un système de trois modèles de régression qui peuvent s'écrire à l'aide de la Figure 1.2. L'approche consiste à régresser la variable réponse Y sur la variable d'exposition A , le médiateur M et les facteurs de confusion ou covariables mesurées que nous désignons par C . Notons que même si le modèle originel de Baron et Kenny n'incluait pas de covariables initialement, celles-ci s'intègrent facilement dans les modèles de régression.

La méthode de Baron et Kenny émet l'hypothèse que les modèles de régression

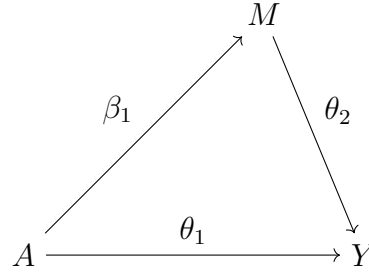


FIGURE 1.2. Diagramme causal illustrant le modèle de médiation conceptuel de Baron et Kenny (1986) et incluant les coefficients décrits dans les Équations (1.1) à (1.3).

linéaire présentés dans les Équations (1.1) à (1.3) caractérisent les données observées.

Pour établir l'existence d'une médiation, la méthode de Baron et Kenny procède en quatre étapes de vérification de conditions sur les coefficients des modèles présentés dans les Équations (1.1) à (1.3).

Ces régressions modélisent les interrelations entre (1) Y et A conditionnellement aux covariables mesurées, (2) M et A conditionnellement aux covariables mesurées, (3) Y et A conditionnellement au médiateur et aux covariables mesurées :

$$E(Y|A = a, C = c) = \zeta_0 + \zeta_1 a + \zeta_2' c \quad (1.1)$$

$$E(M|A = a, C = c) = \beta_0 + \beta_1 a + \beta_2' c \quad (1.2)$$

$$E(Y|A = a, M = m, C = c) = \theta_0 + \theta_1 a + \theta_2 m + \theta_4' c. \quad (1.3)$$

Dans la discussion qui suit, nous supposons que l'interprétation des effets s'applique au changement d'une unité dans les niveaux de l'exposition, c'est-à-dire $A = a + 1$ versus $A = a$. L'Équation (1.1) traduit la possibilité que la variable d'exposition A ait un effet total sur la variable réponse Y : l'effet total est encodé dans le paramètre ζ_1 et existe si $\zeta_1 \neq 0$. L'association entre la variable d'exposi-

tion et le médiateur est représentée par le coefficient β_1 dans l'Équation (1.2) et cette association existe lorsque $\beta_1 \neq 0$. Dans l'Équation (1.3), θ_1 est le coefficient qui encode l'effet de la variable d'exposition sur la variable réponse ajusté pour le médiateur ; cet effet existe lorsque $\theta_1 \neq 0$. L'effet du médiateur sur la variable réponse en fixant la variable d'exposition est encodé par le coefficient θ_2 , et cet effet existe lorsque $\theta_2 \neq 0$. Dans l'Équation (1.1), le coefficient ζ'_2 est un vecteur ligne de coefficients de régression pour les covariables C . L'inclusion de ce coefficient joue en fait un rôle d'ajustement pour les variables de confusion potentielles entre l'exposition A et la réponse Y ; l'inclusion du coefficient β'_2 joue un rôle analogue dans l'Équation (1.2), tout comme pour le coefficient θ'_4 dans l'Équation (1.3).

La méthode de Baron et Kenny consiste en quatre étapes de vérification de conditions. La première est l'existence d'une association entre la variable d'exposition A et la variable réponse Y ($\zeta_1 \neq 0$). Cette étape établit qu'il y a un effet de l'exposition sur la réponse qui peut être expliqué, totalement ou partiellement, à travers le médiateur. La deuxième est que la variable d'exposition A est associée avec le médiateur M ($\beta_1 \neq 0$). La troisième établit que le médiateur M affecte la variable réponse Y ($\theta_2 \neq 0$). Finalement, en considérant les trois premières conditions rencontrées, on détermine qu'il y a médiation totale de l'effet de l'exposition A sur la variable réponse Y lorsque $\theta_1 = 0$. Dans le cas où $\theta_1 \neq 0$, on établit l'existence d'une médiation partielle. L'approche de Baron et Kenny considère les coefficients estimés $\hat{\zeta}_1$, $\hat{\beta}_1$, $\hat{\theta}_1$, $\hat{\theta}_2$ et leur significativité statistique pour établir la présence d'un effet médié.

La méthode de Baron et Kenny établit l'existence de l'effet indirect sous la condition que $\beta_1 \neq 0$ et $\theta_2 \neq 0$, mais elle n'explicite pas sa définition par rapport à ces deux paramètres¹. Elle caractérise tout de même l'effet direct θ_1 grâce au

1. La première condition (étape) est en réalité non nécessaire, car dans un cas extrême il

modèle de régression (1.3). Les deux méthodes que nous présentons dans la suite permettent quant à elles d'exprimer et d'estimer l'effet direct mais aussi indirect.

1.1.3 La méthode du produit

La méthode du produit, ou méthode du produit des coefficients, est utilisée plus fréquemment dans les sciences sociales (VanderWeele, 2016). Avec la méthode du produit, seules les Équations (1.2) et (1.3) sont utilisées. L'effet direct (ED) est toujours défini comme θ_1 , qui correspond au coefficient de la variable d'exposition A dans le modèle de régression pour la réponse qui inclut le médiateur (Équation 1.3). L'effet indirect (EI) est pris comme le produit $\beta_1\theta_2$, où β_1 est le coefficient de la variable d'exposition A dans le modèle ayant le médiateur comme variable dépendante (Équation 1.2) et θ_2 est le coefficient du médiateur M dans le modèle ayant la variable réponse Y comme variable dépendante (Équation 1.3). Ainsi, dans la méthode du produit, on définit les effets direct, indirect et total comme suit.

Définition 1.1.1 (Les effets de médiation dans la méthode du produit)

Les effets de la méthode du produit sont :

- *Effet direct*, $ED = \theta_1$;
- *Effet indirect*, $EI = \beta_1\theta_2$;
- *Effet total*, $ET = ED + EI = \theta_1 + \beta_1\theta_2$.

Notons $\widehat{ED}$, $\widehat{EI}$ et $\widehat{ET}$ les estimateurs des effets direct, indirect et total. On les obtient en remplaçant les coefficients θ_1 , θ_2 et β_2 par leurs estimations comme suit :

pourrait y avoir un effet médié alors que l'effet total est nul.

- $\widehat{ED} = \hat{\theta}_1$;
- $\widehat{EI} = \hat{\beta}_1 \hat{\theta}_2$;
- $\widehat{ET} = \widehat{ED} + \widehat{EI} = \hat{\theta}_1 + \hat{\beta}_1 \hat{\theta}_2$.

1.1.4 La méthode de la différence

La méthode de la différence est fréquemment utilisée en épidémiologie et en sciences biomédicales (VanderWeele, 2016). Les coefficients utilisés pour obtenir une estimation de l'effet indirect avec la méthode de la différence sont issus des Équations (1.1) et (1.3). Rappelons que le coefficient ζ_1 dans le modèle de régression (1.1) représente l'effet total (ET) de la variable d'exposition A sur la variable réponse Y , alors que θ_1 représente l'effet direct ($ED = \theta_1$). Si ζ_1 est différent de θ_1 , alors ceci indique que le médiateur explique une portion de l'effet de la variable d'exposition A sur la variable réponse Y . L'effet indirect est donc la différence entre ces deux coefficients ($EI = \zeta_1 - \theta_1$). Les définitions données pour l'effet indirect et direct satisfont donc la relation suivante :

$$ET = ED + EI.$$

1.1.5 Équivalence entre la méthode du produit et la méthode de la différence

La méthode du produit et la méthode de la différence se basent sur les équations de régression (1.1), (1.2) et (1.3) pour formaliser la quantification des effets direct et indirect. Une question qui se pose naturellement est de savoir laquelle de ces deux méthodes choisir pour le calcul de ces effets dans diverses situations. En fait, pour une variable réponse et un médiateur continu avec des modèles de régression linéaire ajustés par les moindres carrés ordinaires, les deux approches coïncident (VanderWeele, 2016; MacKinnon *et al.*, 1995). On a donc l'équivalence suivante :

$$\zeta_1 - \theta_1 = \beta_1 \theta_2,$$

lorsque les modèles (1.1), (1.2) et (1.3) sont correctement spécifiés. La spécification correcte de ces modèles prend pour acquis qu'il n'y a pas d'interaction entre la variable d'exposition et le médiateur pour expliquer la variable réponse. Tel que nous le verrons, la possibilité d'incorporer une telle interaction dans le modèle de réponse donné dans l'Équation (1.3) est une différence fondamentale entre les approches traditionnelles et causales.

1.1.6 Limites des méthodes traditionnelles de médiation

Dans cette section, nous abordons quatre limites importantes des approches traditionnelles de médiation.

Une première limite des méthodes traditionnelles est que la définition des effets direct et indirect est indissociée des modèles spécifiés pour les données. En effet, les méthodes traditionnelles de médiation ne fournissent pas une définition générale de l'effet direct et indirect, car elles définissent les effets à partir de paramètres des modèles. Nous aimerions toutefois avoir une définition qui s'applique peu importe la spécification du modèle, c'est-à-dire une définition qui peut s'appliquer quels que soient les modèles utilisés pour décrire les données.

Une seconde limite est que la méthode de Baron et Kenny suppose à la base que l'effet total est non nul ($\zeta_1 \neq 0$). Plusieurs auteurs ont contesté cette limite (Kenny *et al.*, 1998; Judd et Kenny, 2010; Zhao *et al.*, 2010), établissant qu'en effet, on peut avoir de la médiation même si $\zeta_1 = 0$. Ce cas extrême survient dans la situation où l'effet direct est de même magnitude que l'effet indirect mais dans une direction opposée (médiation inconsistante).

Une troisième limite que partagent les trois méthodes traditionnelles d'analyse de médiation est la non prise en compte de l'interaction entre la variable d'exposition et le médiateur. Cette limite peut engendrer des biais importants dans les

estimations des effets direct et indirect lorsque cette interaction existe réellement dans le phénomène étudié (Pearl, 2001).

Une quatrième limite est que la randomisation de l'exposition ou même l'ajustement sur toutes les variables de confusion de la relation entre l'exposition et la réponse ne garantit pas l'interprétation causale des effets direct et indirect. En effet, l'article de Baron et Kenny (1986) ne mentionne pas la nécessité de considérer la confusion possible entre le médiateur et la réponse pour l'estimation sans biais des effets direct et indirect. L'absence d'ajustement pour les variables de confusion entre le médiateur et la réponse dans le modèle de médiation peut toutefois causer des biais importants dans les estimations des effets direct et indirect par les méthodes traditionnelles de médiation (VanderWeele, 2016).

1.2 Inférence causale basée sur le cadre contrefactuel

Les méthodes d'inférence causale pour l'analyse de médiation font référence à une discipline qui considère les hypothèses et les stratégies d'estimation qui permettent aux chercheurs de tirer des conclusions causales sur la base de données provenant de devis randomisés ou observationnels (Hill et Stuart, 2015). Ces méthodes sont une extension et une généralisation des approches traditionnelles de médiation. Elles permettent de faire la décomposition des effets totaux ainsi que l'identification et les estimations des effets directs et indirects (contrôlés ou naturels), même en présence de variables de confusion.

Parmi les méthodes d'inférence causale, un cadre de travail fréquemment utilisé est l'approche contrefactuelle (Holland, 1986). Elle est également appelée approche des réponses potentielles. L'approche contrefactuelle présente de nombreux avantages dont un des plus grands est de fournir un cadre précis qui formalise les règles permettant de déterminer la présence de confusion (Burgess et Thompson, 2015).

Il devient donc possible de contrôler la confusion dans les associations et d’obtenir des estimations d’effets qui sont interprétables causalement.

1.2.1 Notations préliminaires

Dans cette section, nous clarifions la notation que nous utilisons pour définir formellement l’effet de l’exposition sur la réponse, aussi appelé effet total dans le contexte de la médiation en analyse causale (Pearl, 2001; Robins et Greenland, 1992). Dans tout le texte suivant, on utilise le terme *variables potentielles* pour référer aux variables contrefactuelles. Les termes *traitement* et *exposition* sont parfois utilisées indifféremment, mais dans ce mémoire nous utilisons le terme *exposition*. La représentation des réalisations des variables aléatoires A et C est notée a et c respectivement. On note $\mathcal{A}$ l’ensemble continu des niveaux de l’exposition, c’est-à-dire le domaine des valeurs possibles de A . On note A_i le niveau de l’exposition pour l’individu i et Y_i la réponse observée pour cet individu.

Le cadre contrefactuel est basé sur l’idée de comparer la réponse qui serait observée sous différents niveaux de l’exposition. Dans cette partie du mémoire, on se concentre sur la comparaison (ou le contraste) entre a et a^* qui représentent deux niveaux de l’exposition possibles de l’ensemble $\mathcal{A}$. Le cadre contrefactuel définit pour chaque individu deux réponses potentielles : (1) Y^a , la réponse potentielle qui correspond au niveau d’exposition fixé à a et (2) Y^{a^*} , la réponse potentielle qui correspond au niveau d’exposition fixé à a^* . Pour l’individu i , celles-ci sont notées Y_i^a et $Y_i^{a^*}$. Les variables aléatoires Y_i^a et $Y_i^{a^*}$ sont des réponses potentielles, parce qu’elles ne correspondent pas nécessairement à la réponse observée Y dans la réalité.

1.2.2 Effet total : définitions, hypothèses et identification

La caractérisation de l'effet causal individuel de l'exposition A sur la réponse Y peut être approchée en étudiant le contraste $\tau^i = Y_i^a - Y_i^{a^*}$ entre les réponses qui seraient observées chez le même individu i sous les deux niveaux d'exposition fixés $A = a$ et $A = a^*$.

Définition 1.2.1 (L'effet causal individuel) *L'existence d'un effet causal pour l'individu i est établie si le contraste τ_i est non nul, c'est-à-dire si $\tau^i \neq 0$.*

Pour expliquer l'effet causal individuel, on considère l'exemple consistant à évaluer l'association entre l'obésité (A) et l'épaisseur de l'intima carotidien (Y). On conceptualise qu'il est possible de fixer l'*IMC* de l'individu i à deux valeurs (niveaux d'exposition) différentes, a et a^* , de façon simultanée dans deux réalités parallèles. On considère les valeurs $a = 34$ et $a^* = 22$ comme étant respectivement des valeurs² prises dans les intervalles de valeurs continues $[30, +\infty)$ (valeurs de l'*IMC* qualifiant un état d'obésité) et $[20, 25)$ (valeurs idéales de l'*IMC*) (Hales *et al.*, 2018). Pour cet individu i , on a donc un niveau d'obésité et un niveau de non obésité et on désire déterminer l'impact de l'obésité sur l'épaisseur de l'intima carotidien. Ainsi, on peut définir pour cet individu deux réponses potentielles, Y_i^{34} , l'épaisseur de l'intima carotidien qui serait observée on fixait l'*IMC* de cet individu à 34 et Y_i^{22} , celle qui serait observée si on fixait plutôt l'*IMC* de cet individu à 22. L'impact de l'obésité sur l'épaisseur de l'intima carotidien est étudié en comparant les deux réponses potentielles. Il y a effet causal de l'obésité pour cet individu si on observe que ses deux réponses potentielles sont différentes, c'est-à-dire que $Y_i^{34} \neq Y_i^{22}$.

2. Dans toute la suite de ce chapitre on utilisera les valeurs typiques $a = 34$ et $a^* = 22$ pour faire des illustrations. Les cas particuliers seront signalés.

Le problème fondamental de l'inférence causale (Holland, 1986) est que les deux réponses potentielles Y^a et Y^{a^*} ne peuvent pas être observées simultanément pour un même individu, de sorte que les quantités causales qui en dépendent directement comme l'effet causal individuel ne peuvent être estimées à partir des données. Pour surmonter cette difficulté, nous nous intéressons aux effets causaux d'ordre populationnel que nous sommes capables d'identifier et d'estimer grâce à une série d'hypothèses. Dans cette section, nous commençons par énoncer l'hypothèse de cohérence et l'hypothèse de stabilité; d'autres hypothèses, comme la positivité, l'ignorabilité marginale et l'ignorabilité conditionnelle, sont énoncées dans la section 1.2.3.

L'existence d'un effet causal moyen de l'exposition A , comparant le niveau a au niveau a^* , sur la réponse Y peut être déterminée en étudiant le contraste $\tau = E[Y^a] - E[Y^{a^*}]$ qui compare les moyennes de la variable réponse dans les cas hypothétiques où tous les sujets seraient assignés au niveau d'exposition $A = a$ et celui où tous les sujets seraient assignés au niveau d'exposition $A = a^*$.

Définition 1.2.2 (L'effet causal moyen) *L'existence d'un effet causal moyen est établie si le contraste τ est non nul, c'est-à-dire si $\tau \neq 0$.*

Par exemple, afin de détecter un effet causal moyen de l'obésité sur l'épaisseur de l'intima carotidien, on imagine une situation où tous les sujets seraient assignés à un niveau d'*IMC* de $a = 34$ et celle où tous les sujets seraient assignés à un niveau d'*IMC* de $a^* = 22$. Dans cette situation, si le contraste obtenu en observant le groupe sous ces conditions d'obésité différentes est non nul, alors cet effet causal est établi.

Pour établir formellement la relation entre les réponses potentielles (Y^a et Y^{a^*}) et la réponse observée (Y), on suppose que la réponse factuelle, c'est-à-dire la réponse réellement observée est la même que la réponse potentielle correspondant

au niveau d'exposition observé : c'est l'hypothèse de cohérence.

Définition 1.2.3 (L'hypothèse de cohérence) *On dit que l'hypothèse de cohérence est satisfaite si :*

$$Y_i = 1_{A_i=a^*} \times Y_i^{a^*} + 1_{A_i=a} \times Y_i^a.$$

En effet, si on observe réellement une épaisseur de l'intima carotidien (Y) pour un individu i sous un niveau d'obésité fixé $A = a$, on s'attend à ce que la réponse potentielle sous ce même niveau d'obésité lui soit égale.

Pour s'assurer de la validité de la définition de l'effet causal, on fait une autre hypothèse qui s'articule en deux points : c'est l'hypothèse de stabilité de la valeur de l'exposition pour l'unité, aussi appelée « Stable Unit Treatment Value Assumption » (SUTVA).

Définition 1.2.4 (L'hypothèse de stabilité)

1. *Il n'existe pas plusieurs versions de chaque niveau d'exposition qui peuvent donner lieu à des réponses différentes selon la version considérée.*
2. *Les réponses potentielles d'un individu i correspondant aux niveaux d'exposition fixés ne dépendent pas de l'exposition des autres individus.*

En effet, la première partie de l'hypothèse de stabilité requiert que les résultats potentiels soient « stables », au sens où une exposition bien définie correspond à une réponse potentielle qui est elle aussi bien définie. Dans notre exemple sur l'étude de la relation causale entre l'obésité et l'épaisseur de l'intima carotidien, on note qu'un individu i pourra atteindre un niveau d'*IMC* fixé de plusieurs manières différentes, par exemple en pratiquant un régime sportif intense ou un régime alimentaire particulier. Pour cet individu, l'épaisseur de l'intima carotidien observée pourrait varier en fonction de la méthode utilisée pour atteindre un

niveau d'*IMC* donné. C'est pour contourner ce problème que la première partie de l'hypothèse de stabilité est énoncée. Par ailleurs, le niveau d'obésité d'un sujet i n'est pas susceptible à influencer l'épaisseur de l'intima carotidien potentielle d'un autre sujet j . Cette situation est exprimée par la seconde partie de l'hypothèse de stabilité. Cependant, certains contextes ne se prêtent pas à cette hypothèse de stabilité comme le cas des maladies contagieuses pour lesquelles l'état d'exposition d'un individu donné peut influencer le risque de maladie des autres individus.

Les hypothèses de cohérence (Définition 1.2.3) et de stabilité (Définition 1.2.4) sont des préalables nécessaires pour bien définir les variables réponses potentielles. Les implications de ces dernières hypothèses sont utiles pour spécifier les estimateurs des effets causaux.

1.2.3 Estimation de l'effet causal moyen

Deux types de devis sont fréquemment utilisés dans les études épidémiologiques et médicales, soient les devis randomisés et les devis observationnels. Dans les devis randomisés, les personnes éligibles sont assignées au hasard à l'un des groupes d'exposition. Un groupe reçoit l'intervention (comme un nouveau médicament) tandis que le ou les autres groupe(s) témoin(s) peuvent recevoir : aucun traitement, un placebo ou une exposition conventionnelle (dans le cas où l'on veut tester un nouveau processus curatif). Dans les devis observationnels, les chercheurs observent l'effet d'une exposition ou d'une autre intervention sans essayer de la modifier (Faraoni et Schaefer, 2016).

Dans la pratique, pour des raisons éthiques ou de faisabilité, on ne peut pas toujours utiliser des devis randomisés qui sont pourtant le type d'études idéal pour établir un lien causal. En effet, les devis randomisés sont considérés comme les devis étalon-or pour établir ou investiguer l'existence d'un effet causal dans les

études épidémiologiques et médicales (Hariton et Locascio, 2018). Lorsque qu'il est impossible de faire un devis randomisé, on mise plutôt sur un devis observationnel. Par exemple, on ne pourrait pas assigner des individus à un régime alimentaire ou sportif drastique dans le but de leur faire perdre du poids pendant une courte période de temps pour modifier leur *IMC*. Pour cette exposition, bien qu'un effet bénéfique en ce qui concerne l'épaisseur de l'intima carotidien de ces individus pourrait être obtenu, une perte de poids subite et importante pourrait dégrader leur état de santé général.

L'identification de l'effet causal moyen dans un devis randomisé ou observationnel requiert que d'autres hypothèses soient rajoutées à la suite de celles déjà énoncées. Ces hypothèses sont l'hypothèse de positivité et une des deux versions de l'hypothèse d'ignorabilité (marginale ou conditionnelle). Nous énonçons tour à tour chacune de ces hypothèses dans la suite.

Définition 1.2.5 (L'hypothèse de positivité) *L'hypothèse de positivité stipule que chaque individu i a une probabilité strictement comprise entre 0 et 1 de recevoir chaque niveau a et $a^* \in \mathcal{A}$ de la variable d'exposition, c'est-à-dire*

$$0 < P(A_i = a) < 1 \text{ et } 0 < P(A_i = a^*) < 1 \text{ pour } a, a^* \in \mathcal{A}.$$

En effet, si la population étudiée comprend certains individus pour lesquels on ne peut pas observer les deux niveaux d'exposition considérés, c'est-à-dire que ces individus ont une probabilité nulle ou certaine de recevoir (ou de ne pas recevoir) un niveau d'exposition, on ne peut pas parler de l'effet de l'exposition A sur la réponse Y pour cette population en entier³. Dans les devis observationnels, on ne

3. Si on exclut ces individus, on pourrait alors estimer un effet causal sur ce sous-ensemble de la population.

contrôle pas l'assignation de l'exposition dans la population et dans cette situation, l'hypothèse de positivité n'est pas toujours satisfaite. En revanche, l'hypothèse de positivité est satisfaite par défaut dans un devis randomisé parce que chaque individu a la possibilité de recevoir chaque niveau de l'exposition étudiée.

On conceptualise qu'il soit possible pour chaque individu de la population de modifier son IMC pour atteindre les niveaux d' IMC de $a = 34$ et $a^* = 22$. En effet, que l'individu ait un IMC anormalement élevé (obésité) ou anormalement bas (maigreur), cet individu a une probabilité non nulle et non certaine de modifier son IMC pour atteindre les deux niveaux d' IMC fixés a et a^* . Dans ce cas, l'hypothèse de positivité est respectée naturellement.

Une deuxième hypothèse effectuée pour identifier l'effet causal moyen est l'hypothèse forte d'ignorabilité. En effet, cette hypothèse s'interprète comme l'absence de confusion dans la relation entre l'exposition et la réponse.

Nous utilisons dans la suite de ce mémoire la notation $U \perp\!\!\!\perp V|C$ pour indiquer que U est indépendant de V conditionnellement à C .

Définition 1.2.6 (L'hypothèse d'ignorabilité marginale) *L'hypothèse d'ignorabilité marginale entre l'exposition et les réponses potentielles est définie par :*

$$A \perp\!\!\!\perp \{Y^a, Y^{a^*}\} \text{ pour } a \text{ et } a^* \in A.$$

On note que dans un devis randomisé, l'hypothèse d'ignorabilité est satisfaite par définition, car dans ces devis on fixe le niveau d'exposition de chaque individu de la population de manière aléatoire. Cette hypothèse nous permet de procéder

dans les étapes d'identification de l'effet moyen comme suit :

$$\begin{aligned}
\tau &= E[Y^a - Y^{a^*}] \\
&= E[Y^a] - E[Y^{a^*}] \\
&= E[Y^a|A = a] - E[Y^{a^*}|A = a^*] \quad (\text{par ignorabilité}) \\
&= E[Y|A = a] - E[Y|A = a^*] \quad (\text{par cohérence}).
\end{aligned}$$

L'effet causal moyen τ devient alors estimable, car il est exprimé à partir de variables observables.

L'hypothèse d'ignorabilité (Définition 1.2.6) dans le cadre de devis randomisés nous permet d'établir les égalités suivantes :

$E[Y^a] = E[Y|A = a]$ et $E[Y^{a^*}] = E[Y|A = a^*]$. En revanche, dans le cadre d'un devis observationnel, l'hypothèse d'ignorabilité ne tient généralement plus. En effet, il est possible, quoique excessivement rare d'avoir des groupes d'exposition dans lesquels les sujets ont des caractéristiques similaires pour les niveaux d'exposition fixés $A = a$ et $A = a^*$ ($a, a^* \in \mathcal{A}$). Reprenons notre exemple dans un cadre observationnel. Les individus ayant naturellement un *IMC* à 22 peuvent différer des individus ayant un *IMC* naturellement à 34, par exemple en ce qui concerne l'âge, le diabète, ou la dyslipidémie, tous des facteurs qui peuvent influencer l'épaisseur de l'intima média. Ces facteurs sont alors considérés comme variables de confusion pour l'association entre l'*IMC* et l'épaisseur de l'intima média carotidien (Herinirina *et al.*, 2015). Dans un devis observationnel, la présence de variables de confusion font que les inégalités suivantes sont de mise : $E[Y^a] \neq E[Y|A = a]$ et $E[Y^{a^*}] \neq E[Y|A = a^*]$.

Pour pallier au problème que l'hypothèse d'ignorabilité n'est typiquement pas satisfaite dans les devis observationnels et que l'effet causal de l'exposition n'est donc pas identifiable, nous introduisons l'hypothèse d'ignorabilité conditionnelle.

Définition 1.2.7 (L'hypothèse d'ignorabilité conditionnelle) *L'hypothèse d'ignorabilité conditionnelle est définie par :*

$$A \perp\!\!\!\perp \{Y^a, Y^{a^*}\} | C.$$

L'hypothèse d'ignorabilité conditionnelle est une condition suffisante qui permet d'estimer sans biais l'effet causal en « contrôlant » pour un ensemble de variables de confusion C (Rosenbaum et Rubin, 1983). En effet, l'hypothèse d'ignorabilité conditionnelle s'interprète comme l'absence de confusion dans la relation entre l'exposition et la réponse après avoir conditionné sur C .

Cette dernière hypothèse nous permet d'exprimer l'effet causal moyen comparant les niveaux d'exposition a et a^* , comme suit :

$$\begin{aligned} \tau &= E[Y^a - Y^{a^*}] \\ &= E[Y^a] - E[Y^{a^*}] \\ &= E_C[E[Y^a|C=c]] - E_C[E[Y^{a^*}|C=c]] \quad (\text{par le théorème de l'espérance totale}) \\ &= E_C[E[Y^a=c, A=a]] - E_C[E[Y^{a^*}|C=c, A=a^*]] \quad (\text{par ignorabilité conditionnelle}) \\ &= E_C[E[Y|C=c, A=a]] - E_C[E[Y|C=c, A=a^*]] \quad (\text{par cohérence}). \end{aligned}$$

L'effet causal moyen τ dans le cadre observationnel est alors identifiable, car les quantités $E[Y|C=c, A=a]$ et $E[Y|C=c, A=a^*]$ peuvent être estimées à partir des données observées.

1.2.4 Effets direct et indirect : Notations, définitions et identification

La théorie de l'inférence causale et en particulier les notations contrefactuelles permettent, contrairement aux approches de médiation traditionnelles, de définir formellement et de généraliser les effets directs et indirects. Ces effets directs et indirects peuvent être estimés à partir de modèles de régression sous réserve que les variables de confusion soient prises en compte et que les modèles soient correcte-

ment spécifiés (Pearl, 2001; VanderWeele, 2015). Dans cette section, nous donnons les définitions de ces différents effets ainsi que leurs hypothèses d'identification.

1.2.4.1 Notations

Nous notons M^a le médiateur potentiel si l'exposition A avait été fixée à la valeur a , Y^{am} la réponse potentielle si A avait été fixée à a et M à m . Nous définissons la réponse potentielle emboîtée $Y^{aM^{a^*}}$ qui désigne la réponse potentielle qui aurait été observée si l'exposition A avait été fixée à a et le médiateur M à la valeur M^{a^*} qu'il aurait prise si l'exposition avait été fixée au niveau a^* .

Nous faisons aussi l'hypothèse de cohérence, qui stipule que si l'exposition A avait été fixée à a , les valeurs potentielles Y^a et M^a seraient respectivement égales aux valeurs observées Y et M . Nous supposons également que si A avait été fixée à a et M à m , la réponse potentielle Y^{am} serait fixée à la réponse Y observée. Notons que toutes ces variables potentielles peuvent être définies par rapport à un individu i en particulier. Le cadre formel ainsi établi (Robins et Greenland, 1992), nous procédons à la définition des effets individuels (direct contrôlé, naturel direct, naturel indirect) ainsi que leurs versions populationnelles (Pearl, 2001; Robins et Greenland, 1992).

1.2.4.2 Définitions

L'effet direct contrôlé (*CDE*) individuel de l'exposition A sur la réponse Y compare les niveaux d'exposition $A = a$ et $A = a^*$ pour une valeur fixée du médiateur M . Il mesure l'effet non médié par M de l'exposition A sur la réponse Y , c'est-à-dire l'effet de A sur Y après être intervenu pour fixer le médiateur M à la valeur m .

Définition 1.2.8 (*L'effet direct contrôlé pour les niveaux d'exposition a versus a^**)

$$CDE_i(m) = Y_i^{am} - Y_i^{a^*m}.$$

Pour illustrer⁴ la notion de CDE , on conceptualise qu'il soit possible qu'un individu puisse modifier son IMC pour les deux niveaux d' IMC $a = 34$ et $a^* = 22$. On rappelle que a et a^* sont respectivement des valeurs typiques prises dans les intervalles de valeurs $[30, +\infty)$ et $[20, 25)$. On imagine aussi pouvoir fixer pour ce même individu le niveau de *protéine C-réactive* dans le sang, représentant la valeur du médiateur, à $m = 0.6$ (valeur typique prise dans l'intervalle $[0.3, 1.0]$). Si on mesure l'épaisseur de l'intima carotidien chez ce même individu, l'effet direct contrôlé de l' IMC sur l'épaisseur de l'intima carotidien peut être défini par le contraste : $CDE(0.6) = Y^{34\ 0.6} - Y^{22\ 0.6}$.

L'effet naturel direct (NDE) de l'exposition A sur la réponse Y compare les réponses potentielles sous les niveaux d'exposition $A = a$ et $A = a^*$ après être intervenu pour fixer le médiateur M à la valeur qu'il aurait naturellement pris si l'exposition avait été fixée à a^* .

Définition 1.2.9 (*L'effet naturel direct pour les niveaux d'exposition a versus a^**)

$$NDE_i = Y_i^{aM_i^{a^*}} - Y_i^{a^*M_i^{a^*}}.$$

On explique cet effet en reprenant notre exemple utilisé dans le cadre du CDE . Considérons $a^* = 22$ comme le niveau de l' IMC de référence et $M^{a^*} = 0.6$ la valeur que prend le médiateur sous ce niveau d'exposition. Si on intervient pour fixer le médiateur à cette même valeur $M = 0.6$ et qu'il soit possible de

4. Nous avons adapté cet exemple du point technique 22.1, page 278 du livre *Causal Inference : What If* (Hernan et Robins, 2020).

mesurer l'épaisseur de l'intima carotidien pour ce même individu sous le niveau d'exposition de référence et sous un autre niveau d'exposition $a = 34$, on définit l'effet naturel direct par le contraste des réponses potentielles suivant :

$$NDE = Y^{34 \ 0.6} - Y^{22 \ 0.6}.$$

L'effet naturel indirect (NIE) individuel est défini comme le contraste après avoir fixé l'exposition au niveau $A = a$, entre la réponse potentielle d'un individu i si le médiateur avait été fixé à la valeur qu'il aurait pris à un niveau d'exposition $A = a$ et la réponse potentielle du même individu i si le médiateur avait été fixé à la valeur qu'il aurait pris si le niveau d'exposition avait été fixé à $A = a^*$.

Définition 1.2.10 (*L'effet naturel indirect pour les niveaux d'exposition a versus a^**)

$$ENI_i = Y_i^{aM_i^a} - Y_i^{aM_i^{a^*}}.$$

On poursuit notre exemple dans le cadre de l'effet naturel indirect. On se situe toujours dans notre cadre imaginaire où il est possible pour l'individu i de modifier son IMC aux valeurs $a = 34$ et $a^* = 22$. On considère $M^a = 5$ et $M^{a^*} = 0.6$ les valeurs du médiateur respectivement sous ces mêmes niveaux d'exposition a et a^* pour cet individu. L'effet naturel indirect correspond au contraste des réponses potentielles : $NIE = Y^{34 \ 5} - Y^{34 \ 0.6}$.

Il existe aussi des versions populationnelles pour les effets direct contrôlé, naturel direct, naturel indirect que nous définissons ici.

Définition 1.2.11 (*Versions populationnelles des effets de médiation*)

1. *Effet direct contrôlé* : $CDE = E[Y^{am} - Y^{a^*m}]$;
2. *Effet naturel direct* : $NDE = E[Y^{aM^{a^*}} - Y^{a^*M^{a^*}}]$;

3. *Effet naturel indirect* : $NIE = E[Y^{aM^a} - Y^{aM^{a^*}}]$;

4. *Effet total* : $TE = NDE + NIE$.

Nous pouvons montrer que l'effet total (TE) peut être effectivement décomposé en un effet naturel direct et indirect. Pour ce faire, nous devons invoquer l'hypothèse de composition qui stipule que $Y^a = Y^{aM^a}$ et $Y^{a^*} = Y^{a^*M^{a^*}}$. Cette hypothèse traduit que la réponse potentielle de Y qui se réaliserait si l'exposition A avait été fixée à a (Y^a) est fixée à la réponse potentielle de Y qui se réaliserait si l'exposition A avait été fixée à a et le médiateur fixé à la valeur qu'il aurait pris si l'exposition A avait été fixée à a (Y^{aM^a}). Dans ce cas, l'effet total peut s'écrire comme :

$$\begin{aligned}
 TE &= E[Y^a - Y^{a^*}] \\
 &= E[Y^{aM^a} - Y^{a^*M^{a^*}}] \quad (\text{par composition}) \\
 &= E[Y^{aM^a} - Y^{aM^{a^*}}] + E[Y^{aM^{a^*}} - Y^{a^*M^{a^*}}] \\
 &= NIE + NDE.
 \end{aligned} \tag{1.4}$$

On note que la décomposition de l'effet total (Équation 1.4) fonctionne pour les effets direct et indirect naturels, mais pas pour l'effet contrôlé direct. En effet, si l'on soustrait un effet direct contrôlé d'un effet total, la quantité résultante ne peut généralement pas être interprétée comme un effet indirect (Kaufman *et al.*, 2004).

Il est aussi possible de définir des effets de médiation conditionnels, c'est-à-dire des effets définis pour une strate de population partageant des valeurs communes pour les variables C .

Définition 1.2.12 (*Versions conditionnelles des effets de médiation*)

1. *Effet direct contrôlé* : $CDE(m, c) = E[Y^{am} - Y^{a^*m} | C = c]$;

2. *Effet naturel direct* : $NDE(c) = E[Y^{aM^{a^*}} - Y^{a^*M^{a^*}} | C = c]$;

3. *Effet naturel indirect* : $NIE(c) = E[Y^{aM^a} - Y^{aM^{a^*}} | C = c]$;
4. *Effet total* : $TE(c) = NDE(c) + NIE(c)$.

1.2.5 Identification

Dans cette section, nous présentons les hypothèses d'identification de l'effet direct contrôlé ainsi que des effets naturels. En particulier, les quatre hypothèses que nous énonçons ici sont utilisées pour l'identification des effets naturels direct et indirect de l'exposition A sur la réponse Y .

La première de ces hypothèses suppose que des données sont disponibles sur un ensemble de covariables C qui est suffisant pour contrôler la confusion pour la relation entre l'exposition A et la réponse Y (Pearl, 2001; VanderWeele et Vansteelandt, 2009) :

$$(H_1) \quad Y^{am} \perp\!\!\!\perp A \mid C, \text{ pour tous niveaux } A = a \text{ et } M = m.$$

La seconde hypothèse fait référence aux variables de confusion entre le médiateur et la réponse et stipule que chaque Y^{am} est indépendant de M parmi les individus partageant un même niveau d'exposition A et un ensemble de covariables⁵ C :

$$(H_2) \quad Y^{am} \perp\!\!\!\perp M \mid A, C, \text{ pour tous niveaux } A = a \text{ et } M = m.$$

Troisièmement, le recours aux contrefactuels M^a exige des hypothèses supplémentaires pour identifier l'effet de A sur M (Pearl, 2001). Cette hypothèse a été formalisée (VanderWeele et Vansteelandt, 2009) en supposant qu'un ensemble de

5. L'ensemble de covariables C considéré ici n'est pas nécessairement le même que pour l'hypothèse précédente.

covariables C est également suffisant pour contrôler la confusion pour la relation entre A et M :

$$(H_3) \quad M^a \perp\!\!\!\perp A \mid C, \text{ pour tout niveau d'exposition } A = a.$$

Enfin, le recours à des contrefactuels composites non observables, tels que $Y^{aM^{a^*}}$ avec $a \neq a^*$, nécessite une condition d'identification supplémentaire (Pearl, 2001) :

$$(H_4) \quad Y^{am} \perp\!\!\!\perp M^{a^*} \mid C, \text{ pour tous niveaux } A = a, A = a^* \text{ et } M = m.$$

L'hypothèse (H_4) est relative à l'indépendance entre les réponses potentielles et le médiateur potentiel. Elle exclut la possibilité d'une confusion médiateur-réponse induite par l'exposition (Pearl, 2001; VanderWeele et Vansteelandt, 2009).

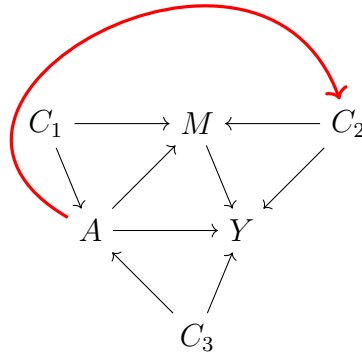


FIGURE 1.3. Identification des effets de médiation, le chemin en rouge représente une possible violation de H_4 (confusion induite par l'exposition).

La Figure 1.3 illustre la situation où les hypothèses $H_1 - H_4$ sont valides. En effet, les relations $A-M$, $M-Y$ et $A-Y$ sont ajustées pour l'ensemble de leurs confondants respectifs C_1 , C_2 et C_3 ou pour tout autre ensemble incluant minimalement ceux-ci. Il apparaît donc que les hypothèses $H_1 - H_3$ sont respectées dans ce cadre.

Pour affirmer que H_4 est aussi respectée, il faut s'assurer que C_2 n'est pas causée par l'exposition A , c'est-à-dire que le lien représenté en rouge n'existe pas.

Alors que les hypothèses $(H_1 - H_4)$ permettent d'identifier les effets naturels direct et indirect, les seules hypothèses (H_1) et (H_2) permettent d'identifier l'effet direct contrôlé. On note finalement que si l'exposition A est randomisée, alors les hypothèses (H_1) et (H_3) sont automatiquement valides sans conditionnement sur C , cependant les hypothèses (H_2) et (H_4) ne le sont généralement pas.

1.2.6 Formules d'identification

Nous avons présenté des hypothèses suffisantes qui permettent l'identification des effets contrôlé et naturels comparant les niveaux d'exposition a et a^* . Nous proposons dans cette section⁶ les formules permettant leur description⁷.

Proposition 1 (*Pearl, 2001*) : *Si les hypothèses $(H_1 - H_2)$ sont satisfaites, alors l'effet naturel direct contrôlé conditionnel aux covariables $C = c$ est identifié et donné par :*

$$E[Y^{am} - Y^{a^*m}|c] = E[Y|a, m, c] - E[Y|a^*, m, c].$$

Démonstration. (VanderWeele, 2015)

$$\begin{aligned} E[Y^{am} - Y^{a^*m}|c] &= E[Y^{am}|a, c] - E[Y^{a^*m}|a^*, c] \quad (\text{par } H_1) \\ &= E[Y^{am}|a, m, c] - E[Y^{a^*m}|a^*, m, c] \quad (\text{par } H_2) \\ &= E[Y|a, m, c] - E[Y|a^*, m, c] \quad (\text{par cohérence}). \end{aligned}$$

6. On note dans ce mémoire $g(m|a^*, c)$ la densité de M conditionnelle à $A = a^*$ et $C = c$.

7. Pour alléger la notation, dans la suite du chapitre on note par exemple $E[Y|a, m, c]$ au lieu de $E[Y|A = a, M = m, C = c]$.

Proposition 2 (Pearl, 2001) : *Si les hypothèses $(H_1 - H_4)$ sont satisfaites, alors l'effet naturel direct conditionnel aux covariables C est identifié et donné par :*

$$E \left[Y^{aM^{a^*}} - Y^{a^*M^{a^*}} | c \right] = \int_m \{ E[Y|a, m, c] - E[Y|a^*, m, c] \} g(m|a^*, c) dm.$$

Démonstration. (VanderWeele, 2015)

$$\begin{aligned} E[Y^{aM^{a^*}} | c] &= \int_m E[Y^{am} | c, M^{a^*} = m] g(M^{a^*} = m | c) dm \\ &= \int_m E[Y^{am} | c] g(M^{a^*} = m | a^*, c) dm \quad (\text{par } H_4 \text{ et } H_3) \\ &= \int_m E[Y^{am} | a, c] g(M = m | a^*, c) dm \quad (\text{par } H_1 \text{ et par cohérence}) \\ &= \int_m E[Y^{am} | a, m, c] g(m | a^*, c) dm \quad (\text{par } H_2) \\ &= \int_m E[Y | a, m, c] g(m | a^*, c) dm \quad (\text{par cohérence}). \end{aligned}$$

On note que l'égalité $E[Y^{aM^{a^*}} | c] = \int_m E[Y | a, m, c] g(m | a^*, c) dm$ est souvent appelée la formule de médiation (Pearl, 2013). En refaisant le même développement que précédemment mais en remplaçant a par a^* , on obtient :

$$E[Y^{a^*M^{a^*}} | c] = \int_m E[Y | a^*, m, c] g(m | a^*, c) dm.$$

Si on procède à la différence des deux expressions, $E[Y^{aM^{a^*}} | c]$ et $E[Y^{a^*M^{a^*}} | c]$, on obtient :

$$E[Y^{aM^{a^*}} - Y^{a^*M^{a^*}} | c] = \int_m \{ E[Y | a, m, c] - E[Y | a^*, m, c] \} g(m | a^*, c) dm.$$

Proposition 3 (Pearl, 2001) : *Si les hypothèses $(H_1 - H_4)$ sont satisfaites, alors l'effet naturel indirect conditionnel aux covariables C est identifié et donné par :*

$$E[Y^{aM^a} - Y^{aM^{a^*}} | c] = \int_m E[Y | a, m, c] \{ g(m | a, c) - g(m | a^*, c) \} dm.$$

1.2.7 Méthodes de régression pour les effets directs et indirects

Dans cette section, nous dérivons les formes paramétriques des différents effets conditionnels pour les modèles de régression linéaire en utilisant les formules des effets de médiation (Proposition 1-Proposition 3). Les modèles de régression linéaire conceptualisés dans cette section peuvent prendre en compte l'interaction potentielle entre l'exposition A et le médiateur M .

Nous supposons les modèles suivants pour la réponse Y et le médiateur M :

$$E[Y|a, m, c] = \theta_0 + \theta_1 a + \theta_2 m + \theta_3 am + \theta'_4 c. \quad (1.5)$$

$$E[M|a, c] = \beta_0 + \beta_1 a + \beta'_2 c. \quad (1.6)$$

Proposition 4 (VanderWeele et Vansteelandt, 2009) : *Si les hypothèses ($H_1 - H_4$) sont vérifiées et si les modèles de régression (Équations 1.5–1.6) pour Y et M sont correctement spécifiés, l'effet direct contrôlé et les effets naturels direct et indirect populationnels conditionnels à $C = c$ sont donnés par :*

$$\begin{aligned} CDE(m, c) &= E[Y^{am} - Y^{a^*m}|c] \\ &= (\theta_1 + \theta_3 m)(a - a^*). \end{aligned} \quad (1.7)$$

$$\begin{aligned} NDE(c) &= E[Y^{aM^{a^*}} - Y^{a^*M^{a^*}}|c] \\ &= \{\theta_1 + \theta_3(\beta_0 + \beta_1 a^* + \beta'_2 c)\}(a - a^*). \end{aligned} \quad (1.8)$$

$$\begin{aligned} NIE(c) &= E[Y^{aM^a} - Y^{aM^{a^*}}|c] \\ &= (\theta_2 \beta_1 + \theta_3 \beta_1 a)(a - a^*). \end{aligned} \quad (1.9)$$

Démonstration. (VanderWeele, 2015)

En utilisant la Proposition 1, nous pouvons dériver l'effet direct contrôlé présenté

dans l'Équation (1.7) en procédant comme suit :

$$\begin{aligned}
CDE(m, c) &= E[Y|c, a, m] - E[Y|c, a^*, m] \\
&= (\theta_0 + \theta_1 a + \theta_2 m + \theta_3 a m + \theta'_4 c) \\
&\quad - (\theta_0 + \theta_1 a^* + \theta_2 m + \theta_3 a^* m + \theta'_4 c) \\
&= (\theta_1 a + \theta_3 a m - \theta_1 a^* - \theta_3 a^* m) \\
&= \theta_1 (a - a^*) + \theta_3 m (a - a^*) \\
&= (\theta_1 + \theta_3 m) (a - a^*).
\end{aligned}$$

En utilisant la Proposition 2, nous pouvons dériver l'effet naturel direct présenté dans l'Équation (1.8) en procédant comme suit :

$$\begin{aligned}
NDE(c) &= \int_m \{E[Y|a, m, c] - E[Y|a^*, m, c]\} g(m|a^*, c) dm \\
&= \int_m \{(\theta_0 + \theta_1 a + \theta_2 m + \theta_3 a m + \theta'_4 c) \\
&\quad - (\theta_0 + \theta_1 a^* + \theta_2 m + \theta_3 a^* m + \theta'_4 c)\} g(m|c, a^*) dm \\
&= \int_m \{(\theta_1 a + \theta_2 m + \theta_3 a m) - (\theta_1 a^* + \theta_2 m + \theta_3 a^* m)\} g(m|c, a^*) dm \\
&= \{(\theta_1 a + \theta_2 E[M|a^*, c] + \theta_3 a E[M|a^*, c]) \\
&\quad - (\theta_1 a^* + \theta_2 E[M|a^*, c] + \theta_3 a^* E[M|a^*, c])\} \\
&= \{(\theta_1 a + \theta_2 (\beta_0 + \beta_1 a^* + \beta'_2 c) + \theta_3 a (\beta_0 + \beta_1 a^* + \beta'_2 c)) \\
&\quad - (\theta_1 a^* + \theta_2 (\beta_0 + \beta_1 a^* + \beta'_2 c) - \theta_3 a^* (\beta_0 + \beta_1 a^* + \beta'_2 c))\} \\
&= (\theta_1 a + \theta_3 a (\beta_0 + \beta_1 a^* + \beta'_2 c)) - (\theta_1 a^* + \theta_3 a^* (\beta_0 + \beta_1 a^* + \beta'_2 c)) \\
&= (\theta_1 + \theta_3 \beta_0 + \theta_3 \beta_1 a^* + \theta_3 \beta'_2 c) (a - a^*) \\
&= \{\theta_1 + \theta_3 (\beta_0 + \beta_1 a^* + \beta'_2 c)\} (a - a^*).
\end{aligned}$$

En utilisant la Proposition 3, nous pouvons dériver l'effet naturel indirect présenté

dans l'Équation (1.9) en procédant comme suit :

$$\begin{aligned}
NIE(c) &= \int_m E[Y|a, m, c] \{g(m|a, c) - g(m|a^*, c)\} dm \\
&= \int_m (\theta_0 + \theta_1 a + \theta_2 m + \theta_3 a m + \theta'_4 c) \{g(m|a, c) - g(m|a^*, c)\} dm \\
&= \int_m (\theta_0 + \theta_1 a + \theta_2 m + \theta_3 a m + \theta'_4 c) g(m|a, c) dm \\
&\quad - \int_m (\theta_0 + \theta_1 a + \theta_2 m + \theta_3 a m + \theta'_4 c) g(m|a^*, c) dm \\
&= (\theta_0 + \theta_1 a + \theta_2 E[M|a, c] + \theta_3 a E[M|a, c] + \theta'_4 c) \\
&\quad - (\theta_0 + \theta_1 a + \theta_2 E[M|a^*, c] + \theta_3 a E[M|a^*, c] + \theta'_4 c) \\
&= (\theta_0 + \theta_1 a + \theta_2 (\beta_0 + \beta_1 a + \beta'_2 c) + \theta_3 a (\beta_0 + \beta_1 a + \beta'_2 c) + \theta'_4 c) \\
&\quad - (\theta_0 + \theta_1 a + \theta_2 (\beta_0 + \beta_1 a^* + \beta'_2 c) + \theta_3 a (\beta_0 + \beta_1 a^* + \beta'_2 c) + \theta'_4 c) \\
&= \theta_2 \beta_1 (a - a^*) + \theta_3 \beta_1 a (a - a^*) \\
&= (\theta_2 \beta_1 + \theta_3 \beta_1 a) (a - a^*).
\end{aligned}$$

On note que lorsque le terme d'interaction θ_3 est nul, c'est-à-dire $\theta_3 = 0$, alors on retrouve les effets tels qu'obtenus précédemment dans la méthode du produit (pour $a - a^* = 1$). La méthode du produit peut donc être considérée comme un cas spécial de la médiation par inférence causale.

1.2.8 Estimations

Dans cette section, nous partons de la définition paramétrique des effets définis à la section 1.2.7 pour poser les bases de leur estimation.

On suppose que les modèles de régression (Équations (1.5)–(1.6)) ont été correctement spécifiés et que les estimateurs $\hat{\beta}$ de β et $\hat{\theta}$ de θ des coefficients de ces modèles sont définis par : $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}'_2)'$ et $\hat{\theta} = (\hat{\theta}_0, \hat{\theta}_1, \hat{\theta}_2, \hat{\theta}_3, \hat{\theta}'_4)'$.

En partant des Équations (1.7–1.9) décrivant les effets naturels direct, indirect et

total comparant les niveaux d'exposition a et a^* , on pose les estimateurs suivants :

$$\begin{aligned}\widehat{CDE}(m, c) &= (\hat{\theta}_1 + \hat{\theta}_3 m) (a - a^*), \\ \widehat{NDE}(c) &= \left\{ \hat{\theta}_1 + \hat{\theta}_3 (\hat{\beta}_0 + \hat{\beta}_1 a^* + \hat{\beta}'_2 c) \right\} (a - a^*), \\ \widehat{NIE}(c) &= (\hat{\theta}_2 \hat{\beta}_1 + \hat{\theta}_3 \hat{\beta}_1 a) (a - a^*), \\ \widehat{TE}(c) &= \widehat{NDE}(c) + \widehat{NIE}(c).\end{aligned}$$

Désignons par Σ_β et Σ_θ les matrices de covariance respectives de $\hat{\beta}$ et $\hat{\theta}$. La matrice de covariance de $(\hat{\beta}', \hat{\theta}')$ est définie par :

$$\Sigma = \begin{pmatrix} \Sigma_\beta & 0 \\ 0 & \Sigma_\theta \end{pmatrix}.$$

On note, $V[\widehat{CDE}(m, c)]$, $V[\widehat{NDE}(c)]$, $V[\widehat{NIE}(c)]$ et $V[\widehat{TE}(c)]$ la variance de $\widehat{CDE}(m, c)$, $\widehat{NDE}(c)$, $\widehat{NIE}(c)$ et $\widehat{TE}(c)$, respectivement. Ces variances peuvent être obtenues en utilisant la méthode delta, immédiatement obtenues de la forme générale suivante (VanderWeele, 2015) :

$$\Gamma' \Sigma \Gamma (|a - a^*|)^2. \quad (1.10)$$

On note ici que Γ se calcule de la manière suivante selon l'effet pour lequel on souhaite déterminer sa variance :

$$\begin{aligned}\Gamma_{CDE} &= \left(\frac{\partial}{\partial \beta'} CDE(m, c), \frac{\partial}{\partial \theta'} CDE(m, c) \right)', \\ \Gamma_{NDE} &= \left(\frac{\partial}{\partial \beta'} NDE(c), \frac{\partial}{\partial \theta'} NDE(c) \right)', \\ \Gamma_{NIE} &= \left(\frac{\partial}{\partial \beta'} NIE(c), \frac{\partial}{\partial \theta'} NIE(c) \right)'. \end{aligned}$$

Il suffit alors de remplacer Γ dans l'Équation (1.10) par $\Gamma_{CDE} = (0, 0, 0', 0, 1, 0, m, 0)'$ pour obtenir $V[\widehat{CDE}(m, c)]$, de remplacer Γ par $\Gamma_{NDE} = (\theta_3, \theta_3 a^*, \theta_3 c', 0, 1, 0, \beta_0 + \beta_1 a^* + \beta'_2 c, 0)'$ pour $V[\widehat{NDE}(c)]$ et par $\Gamma_{NIE} = (0, \theta_2 + \theta_3 a, 0', 0, 0, \beta_1, \beta_1 a, 0)'$ pour $V[\widehat{NIE}(c)]$, où $0'$ représente un vecteur constitué uniquement de zéros et de

même dimension que le vecteur de covariables C . La variance pour l'effet total s'obtient directement en utilisant la propriété d'additivité des dérivées se basant sur la décomposition de l'effet total en effet direct et indirect. Les variances sont estimées en substituant les coefficients par leurs estimations.

CHAPITRE II

INTRODUCTION AUX CONCEPTS DE LA RANDOMISATION MENDÉLIENNE

La randomisation mendélienne (RM) peut être vue comme un cas particulier de la méthode des variables instrumentales (IVs) afin d’estimer l’effet causal d’une exposition A sur une réponse Y même en présence de facteurs de confusion non mesurés. La particularité de la RM est que des instruments génétiques sont utilisés pour l’analyse (Davey Smith et Ebrahim, 2003; Smith et Ebrahim, 2004; Lawlor *et al.*, 2008). Dans ce chapitre, nous abordons la RM après avoir explicité la notion de IVs ainsi que les principaux concepts qui les sous-tendent comme l’instrumentation, les propriétés des variables instrumentales ainsi que des suppositions sur leur validité. Enfin, nous discutons de deux méthodes d’estimation d’effets causaux utilisées en RM à savoir : la méthode de Wald et la méthode des moindres carrés en deux étapes (en anglais $2SLS$ pour « two-stage least squares »).

2.1 Introduction à la méthode des variables instrumentales

Dans les devis observationnels, la confusion mesurée de l'association entre deux variables peut être contrôlée avec une variété de méthodes sophistiquées telles que l'appariement basé sur le score de propension, la stratification et l'utilisation des modèles de régression multiples (Zhang *et al.*, 2018). Cependant, ces méthodes ne tiennent pas compte des facteurs de confusion non mesurés. La méthode des *IVs* est un outil populaire en statistique pour surmonter le problème des facteurs de confusion non mesurés (Angrist et Imbens, 1995).

La technique des *IVs* est basée sur le principe qu'il existe une variable, l'instrument, qui est associée à l'exposition et dont l'effet sur la réponse est uniquement dû à l'exposition. Les variables instrumentales résolvent alors le problème des confondants exposition-réponse non mesurés en utilisant seulement une partie de la variabilité de l'exposition, plus précisément une partie qui n'est pas corrélée avec les facteurs de confusion non mesurés, pour estimer l'effet de l'exposition sur la réponse (Baiocchi *et al.*, 2014).

Cette méthode a été largement utilisée dans de nombreux domaines en dehors de la statistique dont l'économie (Angrist et Krueger, 2001), la génomique et l'épidémiologie (Vegte *et al.*, 2020; Ludl et Michoel, 2021), la sociologie, la psychologie (Bollmann *et al.*, 2019) et les sciences politiques (Sovey et Green, 2011). Elle fait l'objet de cette section.

2.1.1 Notion d'instrument

Un instrument est une variable qui permet d'inférer la relation causale entre deux autres variables, par exemple une exposition A et une réponse Y . D'une manière plus simple, une *IV* (Z) est une quantité mesurable qui est associée à une ex-

position d'intérêt et qui lui est antécédante en mesure. Une *IV* n'est associée à aucun facteur de confusion pouvant affecter la relation exposition-réponse. Elle n'est pas non plus associée à la réponse, sauf via la voie causale hypothétique à travers l'exposition d'intérêt (Burgess et Thompson, 2015).

2.1.2 La méthode d'instrumentation

L'instrumentation est une méthode permettant l'estimation d'un effet causal entre une exposition A et une réponse Y lorsqu'il existe une (ou plusieurs) variable(s) de confusion aussi bien mesurée(s), mal mesurée(s) que non mesurée(s). Dans le cas où la ou les variables de confusion sont non mesurées, on les note U (Zhang *et al.*, 2018). Dans cette méthode, pour estimer l'effet causal de A sur Y libre de tout biais de confusion dû à U , on utilise un instrument Z , qui est corrélé avec l'exposition A mais indépendant des variables de confusion U .

Le diagramme causal de la Figure 2.1 présente une situation dans laquelle on peut estimer le lien causal entre A et Y en partant de l'intuition liée à la méthode de l'instrumentation. En effet, la relation entre A et Y est confondue par la variable U et Z est un instrument pour A . Il est à noter que l'association entre Z et A (lien $Z \rightarrow A$) n'est pas nécessairement causale, mais celle-ci est souvent représentée comme telle pour des raisons de simplification.

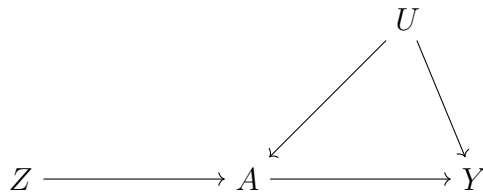


FIGURE 2.1. Diagramme causal pour la relation entre un instrument Z , une exposition A , des facteurs de confusion non mesurés U et une réponse Y .

Par exemple, imaginons qu'on s'intéresse à étudier la relation causale entre l'obé-

sité mesurée par l'*IMC* (A) et l'épaisseur de l'intima média carotidien, notée *cIMT* (Y), dans une population d'adolescents fréquentant une école secondaire située à proximité (P) d'un restaurant rapide à petit prix. Dans cet exemple, la variable P peut être considérée comme un instrument (Z) parce qu'elle est antécédente (en mesure) et supposément associée à l'*IMC* (Bahadoran *et al.*, 2015). Pour qu'une variable Z soit qualifiée d'instrument valide, il faut qu'elle satisfasse certaines hypothèses ou conditions que nous énonçons dans la suite.

2.1.3 Les hypothèses d'une variable instrumentale

Dans la littérature, les hypothèses qui dictent les conditions sous lesquelles l'instrumentation mène à une estimation d'effet causal valide ont connues de diverses et légères modifications autant en nombre que dans leurs formulations (Angrist et Imbens, 1995; Baiocchi *et al.*, 2014; Hernan et Robins, 2020), mais trois hypothèses générales peuvent être identifiées et sont communément acceptées par bon nombre d'auteurs.

Définition 2.1.1 (Validité d'un instrument) *Pour être utilisée comme IV pour l'association entre une exposition A et une réponse Y , une variable Z doit satisfaire les trois hypothèses suivantes :*

- *L'hypothèse de pertinence (H_1), soit que Z est corrélée avec A ;*
- *L'hypothèse d'indépendance (H_2), soit que l'instrument Z ne doit être associé à aucun facteur de confusion pouvant affecter Y ;*
- *L'hypothèse d'exclusion-restriction (H_3), soit que tout lien dirigé (séquence de flèches) de Z allant vers Y passe uniquement par A .*

La Figure 2.2 illustre des violations de H_2 et H_3 . Les trajets en rouge représentent des violations des hypothèses fondamentales. Les trois hypothèses fondamentales

($H_1 - H_3$) qui définissent une *IV* peuvent être interprétées comme suit. L'hypothèse H_1 traduit que Z utilisé en tant que *IV* est associé à l'exposition A , ce qui est encodé par la présence de la flèche $Z \rightarrow A$ dans le diagramme présenté en Figure 2.2. Dans une situation où le lien $Z \rightarrow A$ ne serait pas représenté, on se trouverait dans une situation de violation de l'hypothèse H_1 où l'association entre l'instrument potentiel Z et la réponse n'est pas établie. En pratique, on souhaite avoir un instrument qui a une forte association avec A , c'est-à-dire qui explique une grande proportion de la variabilité de A .

L'hypothèse H_2 traduit qu'il n'existe pas de facteurs de confusion non ou mal mesurés pour la relation entre Z et la réponse Y . Dans le diagramme de la Figure 2.2, le lien $U_2 \rightarrow Z$ induit une situation de violation de l'hypothèse H_2 dans la mesure où U_2 confond l'association entre Z et Y . En effet, il existe ici une cause commune non mesurée (U_2) entre l'instrument potentiel Z et la réponse Y .

L'hypothèse H_3 traduit que Z n'a d'effet causal sur la réponse Y qu'à travers A , c'est-à-dire qu'il n'existe pas d'autres mécanismes d'effet de Z vers Y ne passant pas par A . Une violation de l'hypothèse H_3 est représentée dans le diagramme de la Figure 2.2. En effet, le lien $Z \rightarrow Y$ qui ne passe pas par l'exposition A matérialise une situation de violation de l'hypothèse H_3 dans laquelle l'instrument potentiel Z affecte la réponse Y autrement qu'à travers l'exposition A . On remarque également que l'hypothèse H_3 serait aussi violée si Z allait à U_2 ou si Z allait à Y via d'autres variables.

L'un des plus grands défis de la méthode des *IVs* est de trouver des instruments valides, c'est-à-dire des instruments qui respectent l'ensemble des trois hypothèses des *IVs* que sont H_1 , H_2 et H_3 . Lorsqu'un instrument potentiel Z ne vérifie pas toutes ces hypothèses, l'inférence résultante est généralement biaisée. L'hypothèse H_1 est la seule parmi les hypothèses $H_1 - H_3$ qu'il est possible de vérifier à partir

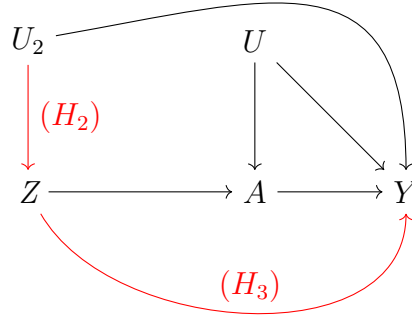


FIGURE 2.2. Diagramme causal illustrant des cas de violation des hypothèses $H_2 - H_3$ de la méthode des *IVs* pour une variable instrumentale Z , une exposition A , des facteurs de confusion non mesurés U et une réponse Y .

des données en utilisant des tests statistiques. Pour la vérifier, il suffit d'établir l'existence d'une association entre l'instrument potentiel Z et l'exposition A . On rappelle que l'instrument potentiel Z n'a pas forcément besoin d'être une cause de l'exposition A , il suffit de lui être associé (Angrist et Imbens, 1995; Savin, 1990).

L'association $(Z \rightarrow A)$ est évaluée en utilisant les statistiques F de Fisher ou le coefficient de détermination R^2 (Baiocchi *et al.*, 2014). Ces dernières sont obtenues suite à l'estimation de la régression de A sur Z .

La capacité de la méthode des *IVs* à produire des estimateurs d'effets causaux valides et efficaces dépend de la force de l'association entre l'exposition A et l'instrument Z . On observe que lorsque l'instrument est faible, c'est-à-dire que Z n'est que faiblement associé à l'exposition, son utilisation mène à des estimateurs imprécis et amplifie tout biais dû à une violation de H_2 ou H_3 (Ertefaie *et al.*, 2018). Pour ces raisons, il est conseillé d'utiliser un instrument fortement associé à l'exposition.

Une *IV* forte, lorsqu'elle est disponible, est toujours préférable. Dans de nombreuses situations pratiques, cependant, les *IVs* fortes ne sont pas toujours disponibles. La combinaison de plusieurs instruments faibles ou forts est une approche intuitive pour améliorer la force globale d'un instrument ainsi construit. De nombreux travaux ont proposé de renforcer des *IVs* faibles existants pour réduire significativement la magnitude des biais des estimations causales (Heng *et al.*, 2019).

Contrairement à l'hypothèse H_1 qui est vérifiable, l'absence de connaissance de l'ensemble des confondants fait qu'il est impossible de prouver que les hypothèses H_2 et H_3 sont respectées à partir des données (Hernan et Robins, 2006; Baiocchi *et al.*, 2014). En effet, cette difficulté s'explique par le fait que les données d'intérêt ne contiennent pas toujours de manière exhaustive toutes les variables impliquées dans le phénomène étudié, mais aussi qu'on ne connaît pas nécessairement toutes les variables impliquées. Dans la pratique, pour juger de la plausibilité du respect des hypothèses H_2 et H_3 , une bonne connaissance (littérature, données ou sources diverses) de la nature des phénomènes étudiés doit être appliquée. L'idée est ici d'écarter la possibilité de l'existence de toute cause commune entre Z et Y et de tout effet direct de l'instrument Z sur la réponse Y . Il est à souligner qu'il existe plusieurs techniques pour tenter d'invalidier H_2 et H_3 mais elles ne font pas l'objet de ce mémoire.

2.2 Concept de randomisation mendélienne

Afin de rendre la compréhension de la méthode de *RM* plus aisée, nous abordons au préalable quelques concepts de génétique sur lesquelles elle repose. Dans cette section, les concepts génétiques ne sont abordés que de manière sommaire. Des explications plus approfondies sont apportées par plusieurs auteurs (Lawlor *et al.*,

2008; Davey Smith et Ebrahim, 2003; Burgess *et al.*, 2017b; Burgess et Thompson, 2015) dont nous nous sommes inspirés pour aborder les notions présentées ici.

2.2.1 Quelques concepts et définitions de génétique

Pour contextualiser l'utilisation des concepts d'inférence causale en génétique, nous définissons de façon sommaire les concepts suivants : *SNP* (pour « single nucleotide polymorphism »), les phénotypes, les allèles, le déséquilibre de liaison, la pléiotropie et la stratification de la population.

Un *SNP* est une mutation génétique à un endroit spécifique du génome et qui encode parfois un phénotype (un trait observable) comme la couleur des yeux ou le risque élevé de surpoids. Certains phénotypes sont dits « complexes » dans le sens qu'ils sont influencés par plusieurs *SNPs*.

Les *SNPs* sont transmis à la conception et l'information qu'ils contiennent (les allèles) vient des deux parents. Ainsi, un *SNP* est habituellement codé en trois groupes formés des allèles A et a : soit AA , Aa et aa . Pour chaque paire, un des allèles vient d'un parent et l'autre de l'autre parent. Si A est l'allèle de risque, alors il est associé avec une plus grande valeur du phénotype, par exemple un plus grand risque d'obésité. La personne qui a deux allèles de risque (AA) est considérée plus à risque que celle qui ne possède qu'un seul allèle de risque (Aa) ou aucun (aa). Les groupes Aa et aA sont considérés comme identiques.

Tous les *SNPs* ne sont pas causaux, c'est-à-dire que certains *SNPs* ne causent pas de changements de phénotype connus. Cependant, en raison des règles gouvernant leur passation des parents aux enfants, les *SNPs* sont transmis en groupes. On observe ainsi une forte corrélation entre certains *SNPs* qui sont côte-à-côte sur le génome. On parle de déséquilibre de liaison, souvent noté LD , pour représenter la

corrélation entre deux *SNPs*. Donc un *SNP* qui n'est pas causalement associé à un phénotype peut quand même être fortement associé à celui-ci s'il est en déséquilibre de liaison fort avec le *SNP* causal.

Les *SNPs* représentent une source potentielle d'*IVs*, en particulier ceux dont les fonctions biologiques ont été identifiées, qui sont associés à l'exposition et qui respectent les hypothèses H_2 et H_3 . En effet, les *SNPs* ont plusieurs propriétés avantageuses les prédisposant à servir d'instrument(s) de l'exposition. Par exemple, les allèles hérités le sont à la conception et sont donc antécédants à la plupart des exposition d'intérêt, ils ne changent donc pas non plus avec le temps. Certaines expositions peuvent cependant faire exception car l'expression d'un *SNP* peut être modifiée par d'autres phénotypes. On note aussi que l'hérédité aléatoire des *SNPs* rend la distribution des génotypes indépendante des facteurs socio-économiques et du mode de vie (Davey Smith et Ebrahim, 2003). Pour comprendre l'architecture génétique des phénotypes, des études d'association sont généralement réalisées. Ces études sont ensuite utilisées pour identifier des instruments probables pour des expositions bien déterminées.

La stratification de la population (Diagramme de la Figure 2.3) se produit lorsqu'il existe des sous-groupes de population qui connaissent à la fois différentes distributions de phénotypes et ont différentes fréquences d'allèles d'intérêt. Cela peut entraîner des associations fallacieuses (confondues) entre le génotype et la réponse dans l'ensemble de la population étudiée. Elle ne devient un problème que lorsqu'elle est un facteur de confusion entre l'instrument Z et la réponse Y car elle viole alors l'hypothèse H_2 (Hodgkin, 2002) matérialisée par l'existence du lien en rouge dans le Diagramme de la Figure 2.3. Pour pouvoir surmonter le problème de confusion, on peut par exemple prendre une population homogène en terme de composition génétique.

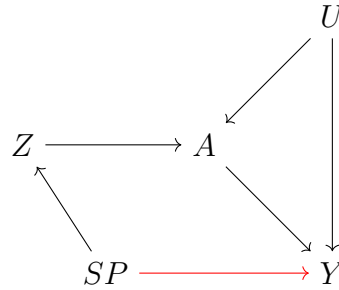


FIGURE 2.3. Diagramme causal illustrant le phénomène de la stratification dans lequel la structure populationnelle SP confond l’association entre l’instrument Z associé à l’exposition A et la réponse Y .

On observe souvent le phénomène biologique appelé pléiotropie qui stipule qu’un SNP peut influencer plus d’un phénotype à la fois. Il existe deux types de pléiotropie : la pléiotropie verticale et la pléiotropie horizontale.

Dans la pléiotropie verticale (Diagramme de la Figure 2.4), le SNP influence la réponse uniquement à travers des chemins causaux passant par l’exposition, mais dont ceux-ci impliquent des phénotypes intermédiaires antécédents à l’exposition.

Ce type de pléiotropie représente le fondement même de la MR car elle repose sur le fait que le SNP n’influence la réponse que via l’exposition.

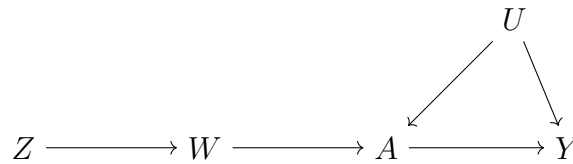


FIGURE 2.4. Diagramme causal illustrant le phénomène de la pléiotropie verticale dans lequel l’instrument Z est associé à un autre phénotype W situé en amont de l’exposition A ; U est un facteur de confusion de la relation entre A et Y .

La pléiotropie horizontale (Diagramme de la Figure 2.5) se produit lorsque le *SNP* a un effet sur le phénotype réponse (en passant par W_2) en dehors de son effet sur l'exposition. Ce type de pléiotropie conduit à la violation de l'hypothèse H_3 (existence du lien en rouge dans le diagramme de la Figure 2.5), ce qui a pour conséquence possible d'entraîner un biais grave lors des estimations (Verbanck *et al.*, 2018; Hemani *et al.*, 2018).

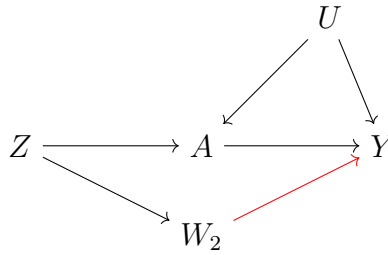


FIGURE 2.5. Diagramme causal illustrant le phénomène de la pléiotropie horizontale dans lequel l'instrument Z associé à l'exposition A est indépendamment associé à la réponse Y directement ou indirectement par le biais d'un autre phénotype W_2 .

2.3 Méthodes d'estimation en randomisation mendélienne : les méthodes 2SLS et de Wald

Dans cette section, nous présentons deux estimateurs communément utilisés en *RM*. Nous considérons uniquement le cas où autant l'exposition que la réponse sont continues. De plus, on suppose que les relations instrument-exposition et instrument-réponse sont linéaires. Nous commençons par présenter la méthode d'estimation de Wald ; par après, nous présentons la méthode d'estimation des moindres carrés en deux étapes (abrégée en anglais *2SLS*).

2.3.1 La méthode de Wald

La méthode du rapport des coefficients, ou méthode de Wald (Wald et Wolfowitz, 1940), est un des moyens permettant d'estimer l'effet causal de l'exposition A sur la réponse Y . La méthode de Wald utilise une seule IV ; si plusieurs IVs sont disponibles, les estimations causales peuvent être faites séparément et combinées ensuite. On note qu'en RM , un SNP peut être représenté comme une variable aléatoire encodant le nombre d'allèles de risque qui peut prendre les valeurs 0, 1 ou 2, pour des raisons statistiques ou biologiques (voir Burgess et Thompson (2015a) pour plus de détails).

On considère le modèle de régression structurel suivant :

$$Y = \zeta_0 + \zeta_1 A + \epsilon, \quad (2.1)$$

où ζ_1 encode l'effet causal de A sur Y et ϵ est une variable aléatoire qui représente toutes les autres causes de Y .

L'estimateur IV de l'effet causal A sur Y par la méthode de Wald est donné par le ratio (Lousdal, 2018) suivant :

$$\hat{\tau}_{Wald} = \frac{\widehat{cov}(Y, Z)}{\widehat{cov}(A, Z)}, \quad (2.2)$$

avec $\widehat{cov}(Y, Z) = (n-1)^{-1} \sum_{i=1}^n (Y_i - \bar{Y})(Z_i - \bar{Z})$, où n est la taille échantillonnale et où la covariance estimée $\widehat{cov}(A, Z)$ se définit de manière analogue.

Nous pouvons considérer la méthode du ratio comme traduisant la variation de la réponse Y suite à une variation d'une unité de l'exposition A .

Proposition 2.3.1 *L'estimateur du ratio de Wald, $\hat{\tau}_{Wald}$, est un estimateur convergent pour ζ_1 .*

Démonstration.

Pour établir la convergence, on évalue la limite en probabilité de $\hat{\tau}_{Wald}$. Partant de l'Équation (2.2), on écrit :

$$\begin{aligned}
 plim_{n \rightarrow \infty} \hat{\tau}_{Wald} &= plim_{n \rightarrow \infty} \frac{\widehat{cov}(Y, Z)}{\widehat{cov}(A, Z)} \\
 &= \frac{plim_{n \rightarrow \infty} \widehat{cov}(Y, Z)}{plim_{n \rightarrow \infty} \widehat{cov}(A, Z)} \\
 &= \frac{cov(Y, Z)}{cov(A, Z)} \quad (\text{en supposant } cov(A, Z) \neq 0 \text{ par } H_1) \\
 &= \frac{cov(\zeta_0 + \zeta_1 A + \epsilon, Z)}{cov(A, Z)} \quad (\text{par Équation 2.1}) \\
 &= \frac{cov(\zeta_0, Z) + cov(\zeta_1 A, Z) + cov(\epsilon, Z)}{cov(A, Z)} \\
 &= \frac{\zeta_1 cov(A, Z) + cov(\epsilon, Z)}{cov(A, Z)} \\
 &= \zeta_1 + \frac{cov(\epsilon, Z)}{cov(A, Z)} \\
 &= \zeta_1 \quad (\text{car } cov(\epsilon, Z) = 0 \text{ par } H_2 \text{ et } H_3).
 \end{aligned}$$

Nous avons donc établi que l'estimateur $\hat{\tau}_{Wald}$ est un estimateur convergent de l'effet de A sur Y . La variance de l'estimateur de Wald (Teumer, 2018), que nous notons $V(\hat{\tau}_{Wald})$, est approximée par la méthode delta comme suit :

$$\widehat{V}(\hat{\tau}_{Wald}) \cong \frac{\widehat{V}(\hat{\gamma}_1)}{\hat{\gamma}_2^2} + \frac{\hat{\gamma}_1^2}{\hat{\gamma}_2^4} \widehat{V}(\hat{\gamma}_2) - 2 \frac{\hat{\gamma}_1}{\hat{\gamma}_2^3} \widehat{cov}(\hat{\gamma}_1, \hat{\gamma}_2), \quad (2.3)$$

où $\hat{\gamma}_1 = \widehat{cov}(Y, Z)$ et $\hat{\gamma}_2 = \widehat{cov}(A, Z)$. L'expression de variance donnée en Équation (2.3) rend explicite le fait qu'une IV faible (γ_2 petit et du fait de sa présence au dénominateur) augmente la variance de l'estimateur de l'effet de l'exposition sur la réponse.

2.3.2 La méthode *2SLS*

Dans cette section, nous présentons la méthode *2SLS* qui est une méthode largement utilisée en instrumentation. Nous en faisons la présentation dans le contexte particulier de la génétique.

La méthode *2SLS* est préférée à l'estimateur du ratio de Wald dans le cadre des réponses continues. Cette méthode permet d'ajuster sur des variables potentiellement confondantes mesurées ou des facteurs de risque et d'utiliser de multiples *IVs* simultanément contrairement à l'estimateur du ratio de Wald. En présence d'une *IV* unique, l'estimateur obtenu par la méthode *2SLS* est le même (numériquement) que l'estimateur du ratio de Wald (Burgess et Thompson, 2015).

La méthode *2SLS* consiste en deux régressions consécutives. Dans la première étape de la méthode, l'exposition A est régressée sur le(s) $IV(s)$. Dans la deuxième étape de la méthode, la réponse Y est régressée sur les valeurs ajustées pour l'exposition de la régression obtenue lors de la première étape. Dans un premier temps, nous présentons la méthode *2SLS* dans la situation où il n'y a qu'une seule variable instrumentale Z . La présentation est inspirée de Greene (2003) et Burgess et al. (2017a).

Considérons la régression linéaire simple, exprimée sous forme matricielle, découlant du modèle donné en Équation (2.1)

$$Y = \mathbf{A}\zeta + \epsilon,$$

où Y est le vecteur des réponses individuelles de dimension $n \times 1$, $\mathbf{A}$ est la matrice de design de dimension $n \times 2$ constituée d'un vecteur unitaire et du vecteur des expositions individuelles A , $\zeta = (\zeta_0, \zeta_1)'$ et ϵ est un vecteur d'erreurs de dimension $n \times 1$.

Il est bien connu que l'estimateur par les moindres carrés de ζ est donné par

$$\hat{\zeta}_{OLS} = (\mathbf{A}'\mathbf{A})^{-1}\mathbf{A}'Y.$$

La limite en probabilité de $\hat{\zeta}_{OLS}$ est

$$plim_{n \rightarrow \infty} \hat{\zeta}_{OLS} = \zeta + plim_{n \rightarrow \infty} \frac{(\mathbf{A}'\mathbf{A})^{-1}}{n} \times \underbrace{plim_{n \rightarrow \infty} \frac{(\mathbf{A}'\epsilon)}{n}}_{(*)}.$$

Lorsque ϵ contient des causes de Y qui sont également des causes de A (i.e. qu'il existe des confondants), le terme $(*)$ est non nul et $\hat{\zeta}_{OLS}$ ne converge pas en probabilité vers ζ .

Soit $\mathbf{Z}$ la matrice de design de dimension $n \times 2$ constituée d'un vecteur unitaire et du vecteur associé aux valeurs individuelles de l'instrument Z . Nous rappelons que l'instrument Z est sélectionné tel que $Cov(Z, \epsilon) = 0$ par les hypothèses H_2 et H_3 . Considérons

$$\hat{\zeta}_{2SLS} = (\mathbf{Z}'\mathbf{A})^{-1}\mathbf{Z}'Y.$$

Puisque $Y = \mathbf{A}\zeta + \epsilon$, alors $\hat{\zeta}_{2SLS}$ est un estimateur convergent pour ζ :

$$\begin{aligned} plim_{n \rightarrow \infty} \hat{\zeta}_{2SLS} &= \zeta + plim_{n \rightarrow \infty} \frac{(\mathbf{Z}'\mathbf{A})^{-1}}{n} \times \underbrace{plim_{n \rightarrow \infty} \frac{(\mathbf{Z}'\epsilon)}{n}}_{=0} \\ &= \zeta. \end{aligned}$$

Nous supposons maintenant que nous avons $K \geq 1$ IVs valides disponibles, $Z = (Z_1, Z_2, \dots, Z_K)$, avec

$$A = \mathbf{Z}\alpha + \epsilon, \tag{2.4}$$

où $\mathbf{Z}$ est la matrice de design correspondante. Dans ce cas, la matrice $\mathbf{Z}'\mathbf{A}$ est de dimension $(K+1) \times 2$. Si $K > 1$, cette matrice possède un rang de $2 < K+1$ et n'admet pas d'inverse. Ainsi, $\hat{\zeta}_{2SLS}$ n'est pas défini.

Pour résoudre ce problème, on considère la projection des colonnes de $\mathbf{A}$ dans l'espace des colonnes de $\mathbf{Z}$:

$$\begin{aligned}\hat{\mathbf{A}} &= \mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{A} \\ &= \mathbf{P}_Z\mathbf{A},\end{aligned}\tag{2.5}$$

où la matrice de projection $\mathbf{P}_Z = \mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'$ est symétrique ($\mathbf{P}_Z' = \mathbf{P}_Z$) et idempotente ($\mathbf{P}_Z\mathbf{P}_Z = \mathbf{P}_Z$).

Un résultat intéressant est le suivant. Considérant l'expression pour $\hat{\mathbf{A}}$ dans (2.5), si la j^e colonne de $\mathbf{A}$ est l'une des colonnes de $\mathbf{Z}$, disons la l^e , alors la j^e colonne de $(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{A}$ est la l^e colonne d'une matrice identité de dimension $(K+1) \times (K+1)$. Ce résultat signifie que la j^e colonne de $\mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{A}$ est la l^e colonne de $\mathbf{Z}$ et donc également la j^e colonne de $\mathbf{A}$. Dans notre cas, la seule variable partagée entre $\mathbf{A}$ et $\mathbf{Z}$ est le vecteur unitaire correspondant à l'ordonnée à l'origine.

On remarque que la matrice $\hat{\mathbf{A}}$ donnée en (2.5) correspond aux valeurs ajustées pour $\mathbf{A}$ par la méthode des moindres carrés appliqué à (2.4), c'est-à-dire, $\hat{\mathbf{A}}$ est constituée d'un vecteur unitaire et du vecteur des valeurs prédites pour A , c'est-à-dire $\hat{A} = \mathbf{Z}\hat{\alpha}$. Finalement, l'estimateur *2SLS* défini généralement pour des instruments multiples, $\hat{\zeta}_{2SLS,M}$, est donné par l'estimateur par moindres carrés de la régression de la réponse sur les valeurs prédites de l'exposition :

$$Y = \hat{\mathbf{A}}\zeta + \epsilon,$$

En d'autres termes,

$$\begin{aligned}\hat{\zeta}_{2SLS,M} &= (\hat{\mathbf{A}}'\hat{\mathbf{A}})^{-1}\hat{\mathbf{A}}'Y \\ &= [(\mathbf{P}_Z\mathbf{A})'(\mathbf{P}_Z\mathbf{A})]^{-1}(\mathbf{P}_Z\mathbf{A})'Y && (\text{par Équation 2.5}) \\ &= [(\mathbf{A}'\mathbf{P}_Z)(\mathbf{P}_Z\mathbf{A})]^{-1}(\mathbf{A}'\mathbf{P}_Z')Y \\ &= [(\mathbf{A}'\mathbf{P}_Z\mathbf{P}_Z\mathbf{A})]^{-1}\mathbf{A}'\mathbf{P}_ZY \\ &= (\mathbf{A}'\mathbf{P}_Z\mathbf{A})^{-1}\mathbf{A}'\mathbf{P}_ZY.\end{aligned}$$

Alors $\hat{\zeta}_{2SLS,M}$ est un estimateur convergent pour ζ :

$$\begin{aligned}
 plim_{n \rightarrow \infty} \hat{\zeta}_{2SLS,M} &= (\mathbf{A}'\mathbf{P}_Z\mathbf{A})^{-1}\mathbf{A}'\mathbf{P}_Z\mathbf{Y} \\
 &= (\mathbf{A}'\mathbf{P}_Z\mathbf{A})^{-1}\mathbf{A}'\mathbf{P}_Z(\mathbf{A}\zeta + \epsilon) \\
 &= \zeta + plim_{n \rightarrow \infty} \frac{(\mathbf{A}'\mathbf{P}_Z\mathbf{A})^{-1}}{n} \times \underbrace{plim_{n \rightarrow \infty} \frac{(\mathbf{A}'\mathbf{P}_Z\epsilon)}{n}}_{=0} \\
 &= \zeta,
 \end{aligned}$$

puisque toute combinaison linéaire de Z est aussi non corrélée avec ϵ .

L'estimateur de la variance asymptotique de l'estimateur $2SLS$ (Burgess *et al.*, 2017a) est donné par :

$$\begin{aligned}
 \hat{V}(\hat{\zeta}_{2SLS,M}) &= \hat{\sigma}^2(A'Z(Z'Z)^{-1}Z'A)^{-1} \\
 &= \hat{\sigma}^2(\hat{A}'\hat{A})^{-1},
 \end{aligned} \tag{2.6}$$

où $\hat{\sigma}^2$ est l'estimation de la variance de la régression de Y sur $\hat{A}$ et non pas celle de Y sur A . Il est à préciser qu'une correction est nécessaire lors de l'estimation de la variance car il faut s'assurer que la variance estimée lors de la régression de la seconde étape tienne compte de la variabilité générée lors de l'étape précédente.

CHAPITRE III

APPLICATION DES MÉTHODES DE RANDOMISATION MENDÉLIENNE EN ANALYSE DE MÉDIATION

La *RM* englobe une variété d’approches qui utilisent des *IVs* génétiques pour estimer l’effet causal total d’une exposition sur une réponse d’intérêt. Une de ses forces est qu’elle rend possible les inférences même en présence de facteurs de confusion non mesurés (voir Chapitre II pour plus de détails). L’analyse de médiation, quant à elle, regroupe un ensemble de méthodes permettant de décomposer l’effet total d’une exposition sur une réponse d’intérêt en un effet direct et un effet indirect par le biais d’une variable médiatrice. Il est à noter qu’en présence de facteurs de confusion non mesurés, les méthodes de médiation produisent des estimateurs généralement biaisés des effets de médiation ainsi que de l’effet total (voir Chapitre I pour plus de détails). Pour tenter d’obtenir des estimateurs des effets de médiation à faible biais ou idéalement non biaisés dans ce contexte, Carter et al. (2021) proposent l’utilisation de deux approches traditionnellement utilisées en *RM*. Il s’agit de la *RM* en deux étapes (en anglais, « Two-Step Mendelian Randomization », *TSMR*) (Relton et Davey Smith, 2012) et de la randomisation

mendélienne multivariée (en anglais, « Multivariable Mendelian Randomization », *MVMR*) (Burgess et Thompson, 2015). L'intérêt de l'article de Carter et al. (2021) est qu'il fait une application pratique de ces deux approches permettant d'intégrer la *RM* en médiation pour gérer les problèmes de variables de confusion non mesurées pour les associations entre l'exposition et la réponse, l'exposition et le médiateur, et le médiateur et la réponse. Les simulations présentées dans Carter et al. (2021) couvrent le cas de réponses tout autant continues que binaires. Dans le cadre de nos travaux, on s'intéresse uniquement à la *RM* en deux étapes appliquée à la méthode de médiation du produit des coefficients. La présentation est limitée au cas où M et Y sont continus. Le but des simulations dans Carter et al. (2021) est de comparer la magnitude des biais éventuels des effets de médiation obtenus par l'approche de *RM* en deux étapes avec la méthode du produit des coefficients standard en présence de variables de confusion non mesurées.

Il est important de souligner que l'article de Carter et al. (2021) ne rencontre pas toutes les bonnes pratiques en ce qui a trait à la façon dont les simulations sont présentées. Ainsi, une des contributions de ce chapitre est d'apporter des précisions quant à la génération, l'analyse et l'interprétation des données de simulation. De plus, lors de l'étude de l'article de Carter et al. (2021), nous avons observé quelques inexactitudes dans les résultats de simulation rapportés. Certaines d'entre elles ont notamment été détectées lorsque nous avons reproduit les simulations en reprenant le code de Carter et al. (2021) disponible sur le GitHub (simulations : <https://github.com/eleanorsanderson/MediationMR>). La Section 3.3.1 de ce chapitre se consacre à l'identification et à la correction de ces erreurs pour le contexte qui nous intéresse. On précise que dans la suite, en ce qui concerne les résultats répliqués, lorsqu'il y a correction, les valeurs originales sont entre crochets.

3.1 La RM en deux étapes

La méthode de *RM* en deux étapes (Relton et Davey Smith, 2012) a été introduite dans le but de déterminer si une association exposition-réponse est médiée par un ou plusieurs autres facteurs ou variables intermédiaires, en estimant séparément l'effet de l'exposition sur la variable médiatrice et l'effet de la variable médiatrice sur la réponse. Dans un contexte d'analyse génomique, Relton et Smith (2012) la présentent comme une stratégie pour investiguer les relations causales entre une exposition (facteur environnemental, génétique ou autre), la méthylation de l'*ADN* (le médiateur) et une réponse (pathologie) d'intérêt. C'est une approche dont la mise en œuvre nécessite, pour sa première étape, l'identification d'une *IV* génétique (*SNP*) de l'exposition qui est liée à la méthylation de l'*ADN* pour évaluer la relation entre l'exposition et le médiateur. Dans la deuxième étape, cette méthode nécessite l'identification d'une *IV* génétique (*SNP*) de la méthylation pour évaluer la relation entre ce médiateur de méthylation et la réponse. En résumé, Relton et Smith (2012) infèrent sur la présence ou l'absence d'un effet médié par l'estimation séparée des deux associations exposition-médiateur et médiateur-réponse décrites ci-haut sans toutefois quantifier la magnitude de l'effet indirect.

Carter et al. (2021) ont proposé l'application de la *RM* en deux étapes pour estimer sans biais les effets total, direct, indirect ainsi que la proportion médiée en présence de facteurs de confusion non mesurés. Carter et al. (2021) opérationnalisent la *RM* en deux étapes telle que décrite par Relton et Smith (2012) en se servant de la méthode du produit des coefficients (voir Section 1.1.3 du Chapitre I pour un rappel de cette méthode). En effet, le processus d'estimation des effets se fait également en deux étapes obtenues en instrumentant à la fois sur l'exposition et sur le médiateur par deux variables instrumentales. Dans chacune des deux

étapes, la méthode du *2SLS* (voir Section 2.3.2 du Chapitre II) est utilisée afin d'estimer sans biais les coefficients d'un des deux modèles utilisés dans la méthode du produit pour l'estimation de l'effet indirect. À l'issue des deux étapes de la *RM* sont calculées : l'effet causal de l'exposition sur le médiateur et l'effet causal du médiateur sur la réponse. Ces deux estimations peuvent ensuite être multipliées pour estimer l'effet indirect par la méthode du produit des coefficients de l'exposition sur la réponse. L'effet direct de l'exposition sur la réponse est également estimable sans biais dans ce cadre de la *RM*.

La Figure 3.1, inspirée de Carter et al. (2021), présente une situation propice à l'application de la *RM* en deux étapes. L'exposition est instrumentée par Z_1 et le médiateur par Z_2 . Ces instruments doivent chacun vérifier les hypothèses $H_1 - H_3$ pour les *IVs* (voir Section 2.1.3 du Chapitre II) pour que les inférences faites par cette méthode restent valides.

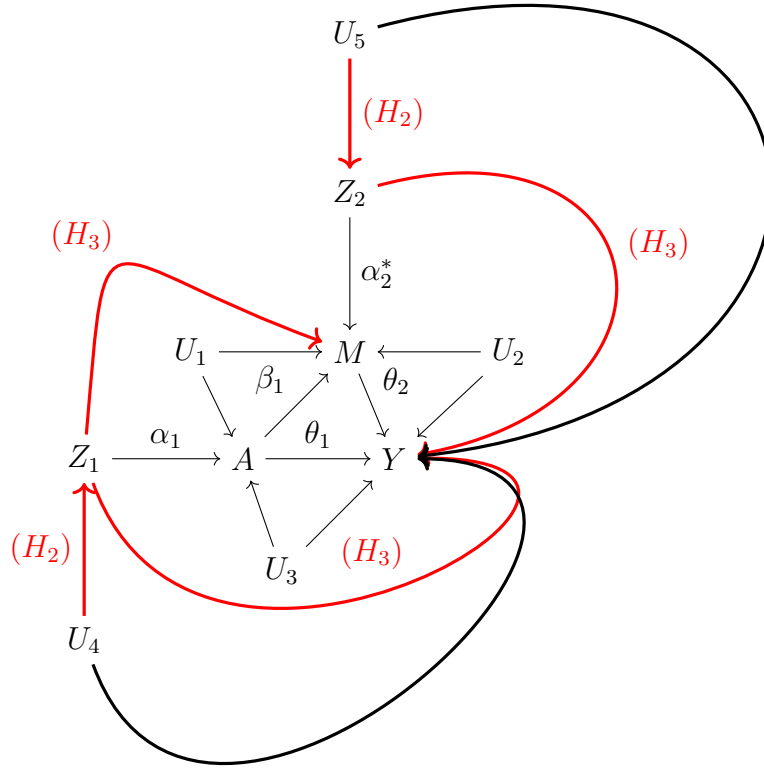


FIGURE 3.1. Diagramme causal illustrant les hypothèses causales dans la *RM* en deux étapes. L'hypothèse H_1 est matérialisée par la présence des flèches $Z_1 \rightarrow A$ et $Z_2 \rightarrow M$.

Les notions abordées dans la suite de ce paragraphe sont un rappel sommaire des notions présentées dans le Chapitre II. Sur la base du diagramme de la Figure 3.1, nous présentons les hypothèses $H_1 - H_3$ pour Z_1 , le principe restant similaire pour Z_2 . Si on suppose une association entre Z_1 et A au préalable, l'hypothèse

H_1 est vérifiée ; de plus il est souhaité que Z_1 soit un instrument fort. L'hypothèse H_2 traduit qu'il n'existe pas de facteurs de confusion non ou mal mesurés pour la relation entre Z_1 et la réponse Y . Dans la Figure 3.1, le lien $U_4 \rightarrow Z_1$ (en rouge) induit une violation de l'hypothèse H_2 dans la mesure où U_4 confond l'association entre Z_1 et Y . L'hypothèse H_3 traduit que Z_1 n'a d'effet causal sur la réponse Y qu'à travers A . En effet, le lien $Z_1 \rightarrow M$ (en rouge), qui ne passe pas par l'exposition A , matérialise une violation de l'hypothèse H_3 dans laquelle l'instrument Z_1 affecte la réponse Y autrement qu'à travers l'exposition A (voir Section 2.1.3 du Chapitre II pour plus de détails sur les hypothèses $H_1 - H_3$).

Avant de passer à l'explicitation des modèles de régression qui sous-tendent la méthode *RM* en deux étapes, on émet des hypothèses de linéarité pour les relations instrument-exposition, exposition-médiateur, exposition-réponse et on précise qu'aucune interaction statistique n'existe entre l'exposition et le médiateur en lien avec la réponse (Burgess *et al.*, 2015).

Dans la suite, pour simplifier la présentation nous utilisons une seule *IV* génétique pour chacun de A et M , mais plusieurs *IVs* génétiques peuvent en fait être utilisées comme instruments de l'exposition, chose restant valable pour le médiateur.

Dans la première étape de la méthode (voir Équations (3.1)-(3.2)), on instrumente sur l'exposition A en se servant d'une *IV* génétique Z_1 qui lui est associée. L'Équation (3.1) encode la relation entre l'exposition et Z_1

$$A = \alpha_0 + \alpha_1 Z_1 + \epsilon. \quad (3.1)$$

En utilisant l'idée de la méthode *2SLS* présentée à la Section 2.3.2 du Chapitre II, on utilise les valeurs ajustées issues de l'Équation (3.1) dans l'Équation (3.2) qui lie le médiateur à l'exposition. Cette étape permet d'estimer sans biais le paramètre β_1 qui encode l'effet de l'exposition sur le médiateur

$$M = \beta_0 + \beta_1 \hat{A} + \epsilon. \quad (3.2)$$

Dans la deuxième étape de la méthode (voir Équations (3.3)-(3.5)), on instrumente le médiateur M avec les IVs génétiques Z_1 (instrument de l'exposition) et Z_2 (instrument du médiateur).

Cette étape de la RM permet d'estimer sans biais le paramètre θ_1 qui encode l'effet direct de l'exposition sur la réponse, ainsi que θ_2 , l'effet du médiateur sur la réponse :

$$A = \alpha_0 + \alpha_1 Z_1 + \alpha_2 Z_2 + \epsilon \quad (3.3)$$

$$M = \alpha_0^* + \alpha_1^* Z_1 + \alpha_2^* Z_2 + \epsilon \quad (3.4)$$

$$Y = \theta_0 + \theta_1 \widehat{A} + \theta_2 \widehat{M} + \epsilon. \quad (3.5)$$

On note que les valeurs prédites $\widehat{A}$ et $\widehat{M}$ utilisées dans l'Équation (3.5) sont obtenues en ajustant séparément les modèles donnés dans les Équations (3.3) et (3.4). On précise que cette manière d'instrumenter le médiateur est une particularité de Carter et al. (2021), car l'application stricte de la méthode de RM en deux étapes ne requiert d'instrumenter l'exposition que par Z_1 et le médiateur que par Z_2 .

Alors que l'estimateur de l'effet direct de l'exposition sur la réponse obtenu par la RM est donné par $\hat{\theta}_1$, l'effet indirect est quant à lui défini comme étant le produit des estimateurs des coefficients issus des deux étapes : $\hat{\beta}_1 \hat{\theta}_2$. En ce qui concerne l'intervalle de confiance de cet effet indirect, il peut être obtenu grâce à la méthode du bootstrap (MacKinnon *et al.*, 2004).

3.2 Présentation des simulations

Dans cette section, nous reprenons quelques scénarios de simulations présentés dans Carter et al. (2021). Le but de ces simulations est d'évaluer le niveau des biais des effets de médiation obtenus par la méthode de RM en deux étapes en comparaison avec les méthodes de médiation standard dans un contexte marqué par la présence de facteurs de confusion non mesurés. Plus particulièrement, nous

présentons les équations génératrices des données de simulation, les paramètres de simulation et l'analyse des données des simulations.

3.2.1 Équations génératrices des données de simulation

La génération des données est conçue sur la base du diagramme causal de la Figure 3.1 sans les liens en rouge, le diagramme respectant alors les hypothèses H_1 à H_3 pour les deux instruments. Les IVs Z_1 de l'exposition et Z_2 du médiateur sont indépendantes et normalement distribuées de moyenne nulle et de variance 1 : $Z_1 \sim \mathcal{N}(0, 1)$ et $Z_2 \sim \mathcal{N}(0, 1)$.

La génération des variables A , M et Y se fait sur la base des Équations (3.6)-(3.8) :

$$A = \alpha_0 + \alpha_1 Z_1 + \epsilon_A, \quad (3.6)$$

$$M = \beta_0 + \beta_1 A + \beta_2 Z_2 + \epsilon_M, \quad (3.7)$$

$$Y = \theta_0 + \theta_1 A + \theta_2 M + \epsilon_Y. \quad (3.8)$$

On note qu'en lieu et place de simuler les facteurs de confusion non mesurés $U = (U_1, U_2, U_3)$ (se référant au diagramme de la Figure 3.1), on simule au préalable des erreurs corrélées $\epsilon = (\epsilon_A, \epsilon_M, \epsilon_Y)$ qui elles-mêmes intègrent la confusion non mesurée de façon implicite. Les erreurs sont simulées selon une normale multivariée $(\epsilon_A, \epsilon_M, \epsilon_Y) \sim \mathcal{N}(\mu, \Sigma)$ où

$$\mu = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \quad \text{et} \quad \Sigma = \begin{pmatrix} \sigma_A^2 & \sigma_{AM} & \sigma_{AY} \\ \sigma_{MA} & \sigma_M^2 & \sigma_{MY} \\ \sigma_{YA} & \sigma_{YM} & \sigma_Y^2 \end{pmatrix}.$$

Deux spécifications pour Σ ont été étudiées par Carter et al. (2021). Une première

spécification pose

$$\Sigma = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}. \quad (3.9)$$

Elle matérialise une situation dans laquelle on est en présence d'une variable de confusion non mesurée qui est une cause commune de A , M et Y simultanément.

La deuxième spécification matérialise une absence de confusion non mesurée pour chacun des liens $A - M$, $M - Y$ et $A - Y$. Dans ce cas, les éléments non diagonaux sont tous égaux à 0 et on a :

$$\Sigma = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}. \quad (3.10)$$

Dans la suite, on s'intéresse uniquement à la situation où on est en présence de facteurs de confusion non mesurés. En effet, en absence de confusion non mesurée, les estimateurs d'effets sont en principe non biaisés, car l'exposition, le médiateur ainsi que la réponse ne sont pas sujets aux erreurs de mesure dans les scénarios de simulation que nous avons considérés.

3.2.2 Scénarios de simulation

Les paramètres précisant les différents scénarios de simulation sont formés de coefficients ou de combinaison de coefficients issus des Équations (3.7) et (3.8). Dans Carter et al. (2021), la valeur des paramètres suivants est précisée : l'effet total (ET) et la proportion médiée (PM). Ainsi, les valeurs de l'effet indirect et de l'effet direct ne sont spécifiées qu'indirectement :

$$EI = ET \times PM$$

$$ED = ET - EI.$$

On peut ainsi par la suite faire correspondre les paramètres constituant ces effets en utilisant les définitions du Tableau 3.1.

TABLEAU 3.1. Définition des paramètres de simulation dans Carter et al. (2021).

Paramètres d'entrée	
Effet direct (ED)	θ_1
Effet de l'exposition sur la réponse	β_1
Effet du médiateur sur la réponse	θ_2
Effet total (ET)	$\theta_1 + \beta_1\theta_2$
Proportion médiée (PM)	$\frac{\beta_1\theta_2}{\theta_1 + \beta_1\theta_2}$

3.2.3 Les scénarios de simulation

Les simulations qui considèrent de la confusion non mesurée se basent sur quatorze scénarios qui correspondent à (1) l'absence de médiation, (2) la médiation inconsistante (voir Section 1.1.6 du Chapitre I) et (3) la variation de la PM pour chaque valeur de l' ET (voir Tableau 3.2 pour les détails concernant nos simulations d'intérêt). On précise que dans chaque scénario, l'effet du médiateur sur la réponse (θ_2) est fixé à 0.2.

Lorsque $ET = ED + EI = 0$, alors $\theta_1 + \beta_1\theta_2 = 0$. La résolution de cette dernière équation pour trouver les valeurs de θ_1 et β_1 nous laisse une infinité de possibilités pour le choix de θ_1 et β_1 sous la condition que $\frac{\theta_1}{\beta_1} = -0.2$. Carter et al. (2021) ont défini par exemple un scénario où $ET = 0$ et $PM = 0.05$; cependant au vu de la définition de la PM , il apparaît que ce scénario ainsi que tous ceux incluant un $ET = 0$ sont mal définis car la PM n'est en réalité pas calculable dans ce cas. La raison pour laquelle nous utilisons le terme “À Définir” au lieu du terme “Non

définie” dans le Tableau 3.2 pour ces paramètres est que la possibilité de leur fixer des valeurs existe pour peu que la condition posée pour leur ratio est remplie. Même si Carter et al. (2021) ont intégré ces scénarios dans leurs simulations et obtenus des résultats, nous les considérons dans la suite comme étant des scénarios invalides (voir Tableau 3.2) car il est erroné de calculer des biais pour l’*ET*, l’*ED*, l’*EI* ainsi que la *PM* quand ils sont mal ou pas définis.

Pour chacun des scénarios considérés, 1000 échantillons de taille $n = 5000$ ont été générés.

TABLEAU 3.2. Présentation des scénarios de simulation avec confusion non mesurée considérés dans notre étude de l’article de Carter et al. (2021).

	θ_1	β_1	$\beta_1\theta_2$ (<i>EI</i>)	$\theta_1 + \beta_1\theta_2$ (<i>ET</i>)	$\frac{\beta_1\theta_2}{\theta_1 + \beta_1\theta_2}$ (<i>PM</i>)
<i>Absence de médiation</i>	0.5	0	0	0.5	0
<i>Médiation inconsistante</i>	0.75	-1.25	-0.25	0.5	-0.5
<i>Variation de la PM</i>	À définir	À définir	À définir	0	Non définie [0.05]
	À définir	À définir	À définir		Non définie [0.25]
	À définir	À définir	À définir		Non définie [0.75]
	0.190	0.050	0.010	0.2	0.05
	0.150	0.250	0.050		0.25
	0.050	0.750	0.150		0.75
	0.475	0.125	0.025	0.5	0.05
	0.375	0.625	0.125		0.25
	0.125	1.875	0.375		0.75
	0.950	0.250	0.050	1	0.05
	0.750	1.250	0.250		0.25
	0.250	3.750	0.750		0.75

3.2.4 Analyse des données simulées

Dans cette section, nous présentons les métriques de performance des estimateurs considérés dans les simulations de Carter et al. (2021). Pour chacune des deux méthodes ciblées, nous calculons le biais, le biais relatif et l'erreur standard (pour les estimateurs des quantités d'intérêt (effet direct, effet indirect, effet total et proportion médiée). Ce calcul est effectué pour les quatorze scénarios considérés (voir Tableau 3.3 pour les détails sur la définition des métriques considérées). On définit un estimateur $\hat{\phi}$ d'un paramètre générique ϕ pour illustrer le calcul des valeurs moyennes des estimations, des biais et des erreurs standards pour tous les effets considérés. Dans le Tableau 3.3, il suffit donc de remplacer l'estimateur ($\hat{\phi}$) par la formule de celui pour lequel on souhaite faire les calculs précédemment cités. Par exemple, pour le cas de l'*ED*, il suffit de remplacer $\hat{\phi}$ par $\hat{\theta}$ et ϕ par θ dans les formules, un raisonnement similaire peut être fait pour l'*EI*, l'*ET* et la *PM*.

TABLEAU 3.3. Définition des métriques de performance des estimateurs considérés dans les simulations.

Paramètres de sortie (estimations)	
Effet estimé moyen ($\bar{\hat{\phi}}$)	$\sum_{j=1}^{1000} \frac{\hat{\phi}_j}{1000}$
Erreur standard $\sqrt{\widehat{var}(\hat{\phi})}$	$\sqrt{\frac{\sum_{j=1}^{1000} (\hat{\phi}_j - \bar{\hat{\phi}})^2}{1000 - 1}}$
Biais	$\bar{\hat{\phi}} - \phi$
Biais relatif	$\frac{\bar{\hat{\phi}} - \phi}{\phi}$

3.3 Résultats

Dans cette section, nous présentons les résultats pour les scénarios de simulation décrits précédemment (Section 3.2.3). Il est à noter que dans la suite nous réorganisons la présentation des résultats en combinant dans un même tableau (correspondant à un effet spécifique) les résultats obtenus par la méthode du produit des coefficients standard et par la *RM* en deux étapes.

Dans la présentation de ces résultats, nous appliquons un arrondi à 2 chiffres après la virgule (par la fonction *round()* de R) sur les résultats obtenus par Carter et al. (2021). L'objectif est d'obtenir des tableaux lisibles et ainsi de pouvoir comparer plus aisément les résultats des répliqués (voir code en Annexe B pour plus de détails) et ceux obtenus par Carter et al. (2021). Notons que les arrondis à 2 chiffres après la virgule ne sont appliqués que sur les résultats finaux des calculs de biais ou de variance, et non sur les résultats intermédiaires. Notre approche est à l'opposé de celle de Carter et al. (2021) qui effectuent les calculs de métriques sur Excel en arrondissant les résultats intermédiaires. En effet, les arrondissements dans les calculs intermédiaires peuvent artificiellement grandir ou réduire les valeurs estimées des effets.

Les résultats sont organisés de la manière suivante. Dans un premier temps, nous recensons les types d'erreurs présents dans les Tableaux 3.5-3.8. Nous parlons ensuite des différences observées entre les résultats des Tableaux 3.5-3.8 répliqués et ceux de Carter et al. (2021). Enfin, nous abordons la performance globale des méthodes de la *RM* en deux étapes et du produit des coefficients standard dans les résultats des Tableaux 3.5-3.8 répliqués.

3.3.1 Commentaires sur les résultats

Nous commentons quelques différences observées entre les tableaux de Carter et al. (2021) et ceux obtenus par nos répliques. Étant donné que nous utilisons le code original de Carter et al. (2021), les erreurs ou imprécisions possibles surviennent lors des calculs des métriques de biais, à l'exception des erreurs de recopie, d'arrondi ou/et de troncature.

3.3.2 Les types d'erreurs détectées

La réplique des résultats sur la base du code original de Carter et al. (2021) présentés dans les Tableaux 3.5-3.8 nous a permis de mettre en évidence trois types d'erreurs (arrondi, calcul et recopie) qui peuvent expliquer les écarts constatés d'avec les résultats proposés par Carter et al. (2021). Le Tableau 3.4 nous présente la fréquence des erreurs observées, par effet, métrique de performance et méthode. Les erreurs d'arrondi, de calcul et recopie sont présentes sur l'ensemble des Tableaux 3.5-3.8. Pour chacune des méthodes considérées, on note la prédominance d'erreurs dans les calculs de biais, en particulier pour la méthode du produit des coefficients standard.

TABLEAU 3.4. Fréquences d'erreurs d'arrondi, de calcul et de recopie des méthodes de la *RM* en deux étapes (*TSMR*) et du produit des coefficients standard (*Prod*) avec confusion non mesurée pour les effets de médiation et l'effet total.

	<i>TSMR</i>			<i>Prod</i>		
	<i>Arrondi</i>	<i>Calcul</i>	<i>Recopie</i>	<i>Arrondi</i>	<i>Calcul</i>	<i>Recopie</i>
<i>ET moyen</i>						
$\widehat{var}(\hat{\theta}_1 + \hat{\beta}_1\hat{\theta}_2)^{1/2}$						
<i>Biais</i>					9	
<i>Biais relatif</i>	1				9	
<i>ED moyen</i>	1			1		
$\widehat{var}(\hat{\theta}_1)^{1/2}$						
<i>Biais</i>						
<i>Biais relatif</i>	1				8	
<i>EI moyen</i>	1					
$\widehat{var}(\hat{\beta}_1\hat{\theta}_2)^{1/2}$						
<i>Biais</i>		1			1	
<i>Biais relatif</i>	2	1		1	1	
<i>PM moyenne</i>	1		9			
$\widehat{var}\left(\frac{\hat{\beta}_1\hat{\theta}_2}{\hat{\theta}_1 + \hat{\beta}_1\hat{\theta}_2}\right)^{1/2}$	1					
<i>Biais</i>		1				
<i>Biais relatif</i>	2	1			10	

Pour marquer la distinction entre les trois types d'erreurs observées dans les Tableaux 3.5-3.8, nous utilisons quatre couleurs distinctes, à savoir : le noir, le bleu, le violet et le rouge. Les valeurs en noir sont les valeurs de cellules de tableaux qui sont fidèles aux résultats de Carter et al. (2021). Les valeurs en bleu marquent des erreurs supposées être dues aux arrondis numériques, celles en violet indiquent des erreurs supposées être dues à la recopie (des données par les auteurs) et enfin

celles en rouge suggèrent des erreurs de calcul dans les biais des estimateurs des effets.

Les résultats répliqués pour l'*ET* présentent quelques différences par rapport à ceux obtenus par Carter et al. (2021). Le biais et le biais relatif calculés par la méthode du produit des coefficients standard diminue de façon marquée comparativement aux valeurs rapportées dans Carter et al. (2021) (voir Tableau 3.5 pour plus de détails). Pour l'*ED*, les résultats de la réplication pour la méthode de *RM* en deux étapes sont restés quasi identiques aux résultats de Carter et al. (2021). Bien que la même tendance soit observée en ce qui concerne la méthode du produit des coefficients standard, on note quelques différences sur certains biais, mais elles ne sont pas importantes (voir Tableau 3.6 pour plus de détails). Les résultats présentés par Carter et al. (2021) donnaient des estimations de biais pour l'*EI* de valeurs nulles pour la *RM* en deux étapes. Nous avons été en mesure de répliquer ces résultats. Les biais obtenus par la méthode du produit des coefficients standard restent globalement les mêmes après réplication sauf pour un scénario ($PM = -0.5$ et $ET = 0.5$) où le biais de la méthode est significativement réduit (voir Tableau 3.7 pour plus de détails). En ce qui concerne la *PM*, on note que les biais dans les résultats présentés par Carter et al (2021) restent les mêmes globalement pour la méthode du produit des coefficients standard (voir Tableau 3.8).

TABLEAU 3.5. Estimations originale et répliquée de l'effet total (ET) pour les méthodes RM en deux étapes ($TSMR$) et du produit des coefficients standard ($Prod$) avec confusion non mesurée.

	PM	ET	ET moyen (ES)	$Biais$	$Biais$ relatif
$TSMR$	0.00	0.5	0.50 (0.02)	0.00	0.00
$Prod$			1.10 (0.01)	0.60	1.20
$TSMR$	-0.50	0.5	0.50 (0.02)	0.00	0.00 [0.01]
$Prod$			1.10 (0.01)	0.60	1.20
$TSMR$	0.05	0.0	Non définie [0.00 (0.02)]	Non définie [0.00]	Non définie
$Prod$			Non définie [0.60 (0.01)]	Non définie [0.60]	Non définie
$TSMR$		0.2	0.20 (0.02)	0.00	0.00
$Prod$			0.80 (0.01)	0.60 [0.80]	3.00 [4.00]
$TSMR$		0.5	0.50 (0.02)	0.00	0.00
$Prod$			1.10 (0.01)	0.60 [0.90]	1.20 [1.80]
$TSMR$		1.0	1.00 (0.02)	0.00	0.00
$Prod$			1.60 (0.01)	0.60 [1.10]	0.60 [1.10]
$TSMR$	0.25	0.0	Non définie [0.00 (0.02)]	Non définie [0.00]	Non définie
$Prod$			Non définie [0.60 (0.01)]	Non définie [0.60]	Non définie
$TSMR$		0.2	0.20 (0.02)	0.00	0.00
$Prod$			0.80 (0.01)	0.60 [0.80]	3.00 [4.00]
$TSMR$		0.5	0.50 (0.02)	0.00	0.00
$Prod$			1.10 (0.01)	0.60 [0.90]	1.20 [1.80]
$TSMR$		1.0	1.00 (0.02)	0.00	0.00
$Prod$			1.60 (0.01)	0.60 [1.10]	0.60 [1.10]
$TSMR$	0.75	0.0	Non définie [0.00 (0.02)]	Non définie [0.00]	Non définie
$Prod$			Non définie [0.60 (0.01)]	Non définie [0.60]	Non définie
$TSMR$		0.2	0.20 (0.02)	0.00	0.00
$Prod$			0.80 (0.01)	0.60 [0.80]	3.00 [4.00]
$TSMR$		0.5	0.50 (0.02)	0.00	0.00
$Prod$			1.10 (0.01)	0.60 [0.90]	1.20 [1.80]
$TSMR$		1.0	1.00 (0.02)	0.00	0.00
$Prod$			1.60 (0.01)	0.60 [1.10]	0.60 [1.10]

TABLEAU 3.6. Estimations originale et répliquée de l'effet direct (ED) pour les méthodes RM en deux étapes ($TSMR$) et du produit des coefficients standard ($Prod$) avec confusion non mesurée.

	PM	ET	EI	ED	ED moyen (ES)	$Biais$	$Biais$ relatif
$TSMR$	0.00	0.5	0.00	0.50	0.50 (0.01)	0.00	0.00
$Prod$					0.83 (0.01)	0.33	0.67
$TSMR$	-0.50	0.5	-0.25	0.75	0.75 (0.02)	0.00	0.00
$Prod$					1.50 (0.01)	0.75	1.00 [1.5]
$TSMR$	0.05	0.0	À définir [0.00]	À définir [0.00]	Non définie [0.00 (0.01)]	Non définie [0.00]	Non définie
$Prod$					Non définie [0.33 (0.01)]	Non définie [0.33]	Non définie
$TSMR$		0.2	0.01	0.19	0.19 (0.02)	0.00	0.00
$Prod$					0.51 (0.01)	0.32	1.67 [1.58]
$TSMR$		0.5	0.03	0.48	0.48 (0.01)	0.00	0.00
$Prod$					0.77 (0.01)	0.30	0.62 [0.58]
$TSMR$		1.0	0.05	0.95	0.95 (0.014)	0.00	0.00
$Prod$					1.20 (0.01)	0.25	0.25
$TSMR$	0.25	0.0	À définir [0.00]	À définir [0.00]	Non définie [0.00 (0.01)]	Non définie [0.00]	Non définie
$Prod$					Non définie [0.33 (0.01)]	Non définie [0.33]	Non définie
$TSMR$		0.2	0.05	0.15	0.15 (0.01)	0.00	0.00
$Prod$					0.40 (0.01)	0.25	1.66 [1.25]
$TSMR$		0.5	0.12	0.38	0.38 (0.02)	0.00	0.00
$Prod$					0.50 (0.01)	0.12	0.25
$TSMR$		1.0	0.25	0.75	0.75 (0.02)	0.00	0.00
$Prod$					0.67 (0.01)	-0.08	-0.11 [-0.08]
$TSMR$	0.75	0.0	À définir [0.00]	À définir [0.00]	Non définie [0.00 (0.01)]	Non définie [0.00]	Non définie
$Prod$					Non définie [0.33 (0.01)]	Non définie [0.33]	Non définie
$TSMR$		0.2	0.15	0.05	0.05 (0.02)	0.00	0.00
$Prod$					0.13 (0.01)	0.08	1.66 [0.42]
$TSMR$		0.5	0.38	0.12	0.12 [0.13] (0.03)	0.00	0.00 [0.01]
$Prod$					-0.17 (0.017)	-0.29	-2.42 [-0.58]
$TSMR$		1.0	0.75	0.25	0.25 (0.06)	0.00	0.00
$Prod$					-0.67 (0.03)	-0.92	-3.67 [-0.92]

TABLEAU 3.7. Estimations originale et répliquée de l'effet indirect (EI) pour les méthodes RM en deux étapes ($TSMR$) et du produit des coefficients standard ($Prod$) avec confusion non mesurée.

	PM	ET	EI	EI moyen (ES)	$Biais$	$Biais$ relatif
$TSMR$	0.00	0.5	0.00	0.00 (0.00)	0.00	NA
$Prod$				0.27 (0.01)	0.27	NA
$TSMR$	-0.50	0.5	-0.25	-0.25 (0.02)	0.00	0.00
$Prod$				-0.40 (0.01)	-0.15 [-1.15]	0.60 [4.60]
$TSMR$	0.05	0	À définir [0.00]	Non définie [0.00 (0.00)]	Non définie [0.00]	Non définie
$Prod$				Non définie [0.27 (0.01)]	Non définie [0.27]	Non définie
$TSMR$		0.2	0.01	0.01 (0.00)	-0.09 [0.00]	-9.00 [0.01]
$Prod$				0.29 (0.01)	0.28	28.00
$TSMR$		0.5	0.03	0.03 (0.00)	0.00	0.00 [-0.01]
$Prod$				0.33 (0.01)	0.30	12.20 [12.32]
$TSMR$		1	0.05	0.05 (0.00)	0.00	0.00
$Prod$				0.40 (0.01)	0.35	7.00
$TSMR$	0.25	0	À définir [0.00]	Non définie [0.00 (0.00)]	Non définie [0.00]	Non définie
$Prod$				Non définie [0.27 (0.01)]	Non définie [0.27]	Non définie
$TSMR$		0.2	0.05	0.05 (0.00)	0.00	0.00 [-0.01]
$Prod$				0.40 (0.01)	0.35	7.00
$TSMR$		0.5	0.12	0.12 (0.01)	0.00	0.00
$Prod$				0.60 (0.01)	0.48	3.84
$TSMR$		1	0.25	0.25 (0.02)	0.00	0.00
$Prod$				0.93 (0.01)	0.68	2.73
$TSMR$	0.75	0	À définir [0.00]	Non définie [0.00 (0.00)]	Non définie [0.00]	Non définie
$Prod$				Non définie [0.27 (0.01)]	Non définie [0.27]	Non définie
$TSMR$		0.2	0.15	0.15 (0.01)	0.00	0.00
$Prod$				0.67 (0.01)	0.52	3.44
$TSMR$		0.5	0.38	0.37 [0.38] (0.03)	0.00	0.00
$Prod$				1.27 (0.02)	0.90	2.38
$TSMR$		1	0.75	0.75 (0.06)	0.00	0.00
$Prod$				2.27 (0.03)	1.52	2.03 [2.02]

TABLEAU 3.8. Estimations originale et répliquée de la proportion médiée (PM) pour les méthodes RM en deux étapes ($TSMR$) et du produit des coefficients standard ($Prod$) avec confusion non mesurée.

	PM	ET	PM moyenne (ES)	$Biais$	$Biais$ relatif
$TSMR$	0	0.5	0.00 (0.01) [(0.00)]	0.00	NA
$Prod$			0.24 (0.01)	0.24	NA [0.12]
$TSMR$	-0.5	0.5	-0.50 (0.04)	0.00	0.00
$Prod$			-0.36 (0.01)	0.14	-0.27 [0.07]
$TSMR$	0.05	0	Non définie [0.00 (0.16)]	Non définie [0.11]	Non définie [2.28]
$Prod$			Non définie [0.45 (0.01)]	Non définie [0.39]	Non définie [0.20]
$TSMR$		0.2	0.05 (0.02) [0.004 (0.049)]	0.00	0.00 [-0.01]
$Prod$			0.37 (0.01)	0.32	6.34 [0.16]
$TSMR$		0.5	0.05 (0.01) [0.00 (0.05)]	0.00	0.00 [-0.01]
$Prod$			0.30 (0.01)	0.25	5.06 [0.13]
$TSMR$		1	0.05 (0.00) [0.00 (0.05)]	0.00	0.00
$Prod$			0.25 (0.00)	0.20	4.00 [0.10]
$TSMR$	0.25	0	Non définie [0.00 (0.11)]	Non définie [-0.14]	Non définie [-0.56]
$Prod$			Non définie [0.44 (0.01)]	Non définie [0.19]	Non définie [0.10]
$TSMR$		0.2	0.25 (0.02) [0.00 (0.25)]	0.00	0.00
$Prod$			0.50 (0.01)	0.25	1.00 [0.12]
$TSMR$		0.5	0.25 (0.02) [0.01 (0.25)]	0.00	0.00
$Prod$			0.55 (0.01)	0.30	1.18 [0.15]
$TSMR$		1	0.25 (0.02) [0.02 (0.25)]	0.00	0.00
$Prod$			0.58 (0.01)	0.33	1.33 [0.17]
$TSMR$	0.75	0	Non définie [0.00 (-0.05)]	Non définie [0.80]	Non définie [-1.06]
$Prod$			Non définie [0.45 (0.01)]	Non définie [-0.31]	Non définie [-0.15]
$TSMR$		0.2	0.75 (0.07) [0.01 (0.75)]	0.00	0.00 [0.01]
$Prod$			0.83 (0.01)	0.08	0.11 [0.04]
$TSMR$		0.5	0.75 (0.06) [0.03 (0.75)]	0.00	0.00
$Prod$			1.15 (0.02)	0.40	0.53 [0.20]
$TSMR$		1	0.75 (0.06) [0.06 (0.75)]	0.00	0.00
$Prod$			1.42 (0.02)	0.67	0.89 [0.33]

3.3.3 Impact des erreurs détectées sur les conclusions de Carter et al. (2021)

Les résultats obtenus lors de la réplication des simulations ne diffèrent pas énormément des résultats présentés par Carter et al. (2021) bien que certaines erreurs de calcul, de recopie, d'arrondi aient été mises en évidence. En effet, la *RM* en deux étapes a mené à des estimations de l'*ET*, l'*ED* et l'*EI* non biaisées contrairement à la méthode du produit des coefficients standard. De plus, la *RM* en deux étapes a estimé la *PM* avec des biais plus faibles comparés à ceux obtenus par la méthode du produit des coefficients standard. En définitive, les résultats et conclusions de Carter et al. (2021) conservent toute leur pertinence.

CHAPITRE IV

ÉTUDE DE L'EFFET CAUSAL DE L'OBÉSITÉ SUR L'ÉPAISSEUR DE L'INTIMA MÉDIA CAROTIDIEN MÉDIÉ PAR LA PROTÉINE C-RÉACTIVE

Le but de ce chapitre est de faire une application sur des données réelles de la *RM* dans le cadre de la médiation en présence supposée de confusion non mesurée. Nous partons du cadre méthodologique posé dans l'article de Carter et al. (2021) pour appliquer la méthode du produit des coefficients standard et la méthode de *RM* en deux étapes. Nous expliquons d'abord le contexte de notre application ; ainsi, nous présentons l'Étude longitudinale canadienne sur le vieillissement (*ÉLCV*) de laquelle proviennent les données utilisées (Raina *et al.*, 2019). Par la suite, nous expliquons la procédure de sélection et de validation des *SNPs* que nous avons utilisés lors de l'instrumentation respective de l'exposition et du médiateur. La dernière partie de ce chapitre se concentre sur la description et l'estimation des modèles ainsi que l'analyse des résultats obtenus.

4.1 Description de l'application

Plusieurs études ont suggéré que l'adiposité (mesurée par l'indice de masse corporelle (*IMC*)) ainsi que l'inflammation dans le corps d'une personne (mesurée par la quantité de *CRP* dans le sang) sont des facteurs de risque pour l'artériosclérose (Ellulu *et al.*, 2017). L'épaisseur de l'intima média carotidien (*cIMT*) est un marqueur du risque d'artériosclérose très utile en recherche car elle est mesurée de façon quantitative. Cependant, les mécanismes exacts expliquant les associations entre l'*IMC*, la *CRP* et la *cIMT* sont complexes et restent encore à être examinés attentivement.

Le but de notre application est d'examiner le rôle médiateur de la *CRP* dans l'association entre l'*IMC* (exposition) et la *cIMT* (réponse) (voir le diagramme de la Figure 4.1). Nous utilisons deux méthodes, à savoir la méthode du produit des coefficients dans le cadre de l'inférence causale et la méthode de RM en deux étapes (TSMR) comme dans l'article de Carter et al (2021). Cette dernière utilise deux instruments génétiques, un pour l'*IMC* et l'autre pour la *CRP* en vue d'étudier les liens causaux entre l'*IMC*, la *CRP* et la *cIMT*.

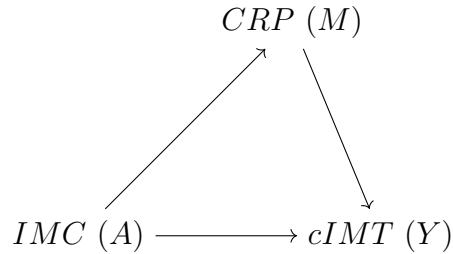


FIGURE 4.1. Diagramme causal illustrant le modèle de médiation supposé dans l'application d'intérêt.

4.2 Données et variables

4.2.1 Source des données

L'*ÉLCV* est une plateforme de recherche nationale contenant des données longitudinales dont un des principaux objectifs est de comprendre les facteurs qui déterminent la qualité du vieillissement. De l'information détaillée sur l'*ÉLCV* ainsi que ses variables peut être trouvée grâce à un portail disponible via l'URL <https://datapreview.clsa-elcv.ca/datasets>.

À la base, l'*ÉLCV* est une cohorte de $n = 51\,338$ participants des deux sexes dont l'âge varie de 45 à 85 ans à l'enrôlement. Trois des critères fondamentaux pour inclure les participants à l'*ÉLCV* sont qu'ils soient capables de lire et de parler français ou anglais, qu'ils ne vivent pas dans des établissements de soins de longue durée et qu'ils ne soient pas affectés par des troubles cognitifs. Il est important de souligner que les données de l'*ÉLCV* sont soumises à un contrôle de qualité dont l'un des critères est que chaque participant est issu d'une famille distincte assurant ainsi l'indépendance entre les individus.

L'*ÉLCV* est composée de deux cohortes complémentaires qui peuvent être étudiées séparément ou ensemble : (1) une cohorte de suivi qui intègre $n = 21\,241$ participants et (2) une cohorte *globale* de $n = 30\,097$ participants. Dans ce mémoire, nous ne nous intéressons qu'à la cohorte *globale*. Dans cette cohorte, les participants sont sélectionnés dans un rayon de 25 à 50 km du site de collecte de données le plus proche de leur lieu de résidence (parmi les 11 répartis dans 7 provinces). Il est important de noter que l'*ÉLCV* comporte une multitude de variables autant non génétiques que génétiques en dehors de celles présentées dans ce mémoire. On note par ailleurs que les données biologiques, cliniques et génétiques ont été obtenues grâce à des méthodes et des dispositifs spécialisés dans

les centres de collecte les plus proches des zones de résidence des participants. Notre accès aux données ainsi que leur utilisation ont été approuvés par le comité d'éthique du Centre de recherche hospitalier de l'Université de Montréal.

Le choix des participants à l'*ÉLCV* s'est fait sur la base de dossiers administratifs tels que les régimes provinciaux d'assurance maladie, des listes d'annuaires téléphoniques, de la composition aléatoire de numéros de téléphone, de bases de recensement et d'enquêtes sur la population active. Lors de leur entrée dans la cohorte entre 2010 et 2015, les participants étaient interrogés en personne à leur domicile. Après qu'ils eurent fourni leur consentement éclairé, on récoltait des données sur leur statut socio-démographique, leur habitudes de vie, leur état de santé global, et bien d'autres données autant biologiques que non biologiques.

Les données génétiques des participants ont été obtenues à partir d'échantillons de sang génotypés grâce à la plateforme Affymetrix (*UK Biobank Axiom*) puis imputées à l'aide des données du *Haplotype Consortium* (39.2 millions de *SNPs*).

4.2.2 Détermination de l'échantillon analytique

Dans le cadre de cette analyse, nous utilisons de manière complémentaire les données phénotypiques et génétiques de la cohorte globale de l'*ÉLCV*. Plus précisément, l'ensemble des variables est mesuré au début de l'étude ("baseline") sauf pour la mesure de la *cIMT* qui est prise lors du premier suivi. On note que notre échantillon est restreint aux participants d'origine caucasienne pour contrôler la stratification de la population et ainsi aider à satisfaire l'hypothèse d'indépendance (H_2) des variables instrumentales (voir Section 2.1.3 du Chapitre II).

Pour obtenir notre échantillon final, nous commençons par sélectionner les participants de la cohorte *globale* ($n = 30\,097$) dont les identifiants apparaissent à la fois dans le fichier de données phénotypiques ($n = 30\,097$) d'une part et gé-

nétiques ($n = 26\,622$) d'autre part, sans tenir compte du fait que les données sur ces variables y soient complètement observées ou non. Cette première étape conduit à obtenir une taille d'échantillon de $n = 26\,622$ participants. On note que l'échantillon obtenu à cette étape est lui aussi constitué de variables dont les données ne sont pas complètement observées. On note que le fichier de données contenant les mesures de *cIMT* des $n = 19\,844$ participants ne fait pas partie de la base de données initiales qui contient les données phénotypiques et génétiques de la cohorte globale. Nous poursuivons le processus en faisant l'intersection de l'échantillon des $n = 26\,622$ participants précédemment obtenus et du fichier de données contenant les mesures de *cIMT* des $n = 19\,844$ participants. Il en résulte un ensemble de $n = 17\,660$ participants contenant des données initiales et les mesures de *cIMT*.

On termine le processus en éliminant tous les participants possédant au moins une valeur manquante au niveau de l'exposition (l'*IMC*), du médiateur (la *CRP*), de la réponse (la *cIMT*) et d'une ou plusieurs variables d'ajustement. Ces variables d'ajustement sont : l'âge, le niveau de scolarité, le sexe, le statut marital, le revenu annuel total du ménage. On obtient au final un échantillon de $n = 13\,942$ participants ayant des données complètement observées que nous utilisons dans notre analyse (voir diagramme dans la Figure 4.2).

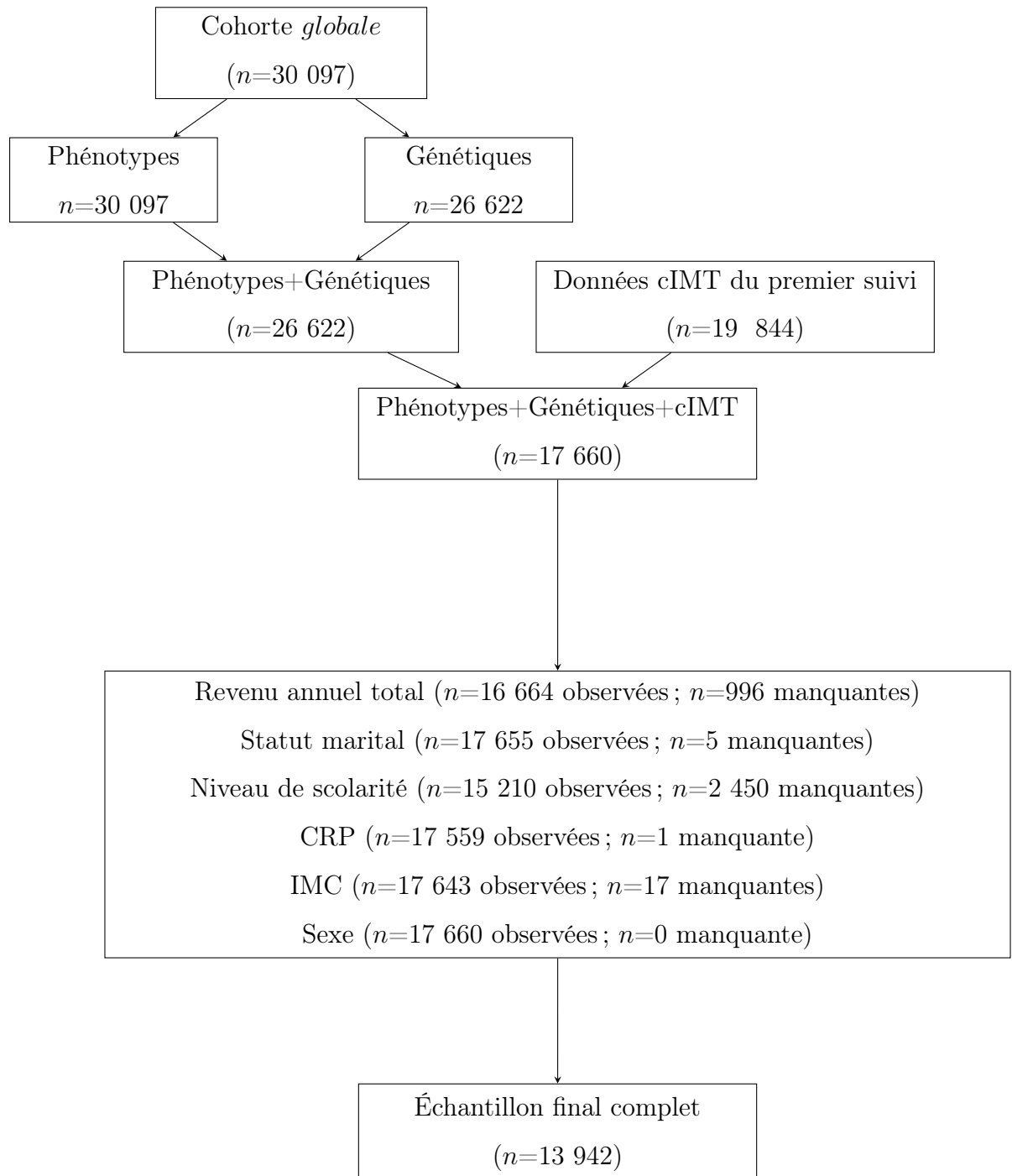


FIGURE 4.2. Diagramme illustrant la constitution de l'échantillon final retenu pour les analyses.

4.2.3 Variables utilisées et échelles de mesures

Cette section présente les variables considérées dans l'étude ainsi que la codification utilisée pour les analyses. Le sommaire des variables et de leur codification sont présentés dans le Tableau A.1.

4.2.3.1 Les variables non génétiques

Nous présentons tour à tour les variables que nous avons considérées, à savoir les variables d'exposition, médiateur et réponse et les variables d'ajustement. L'*IMC* (l'exposition) a été calculée sur la base des données de taille et de poids des participants selon la formule poids (*kg*) divisé par le carré de la taille (*m*). La *CRP* (le médiateur) est le taux mesuré (en *mg/l*) de la protéine C-réactive dans le sang. La *cIMT* (la réponse) a été calculée comme la valeur moyenne (*mm*) des trois valeurs moyennes latérale, antérieure et postérieure respectivement faites sur les parois éloignées de la carotide commune gauche et droite, des bifurcations et des artères carotidiennes lors de trois cycles cardiaques consécutifs. Ensuite viennent les covariables comme l'*âge* à l'enrôlement, le *niveau de scolarité*, le *sexe*, le *statut marital*, le *revenu annuel total du ménage* obtenu avant impôts et déductions au cours des 12 derniers mois (voir Tableau A.1 en Annexe A pour plus de détails sur ces variables).

4.2.3.2 Les variables génétiques

Dans ce qui suit, nous présentons quelques détails sur la sélection des deux instruments génétiques pour notre échantillon. En effet, un instrument a dû être sélectionné pour instrumenter l'*IMC* et un autre pour instrumenter la *CRP*. La sélection des *SNPs* parmi les *SNPs* candidats à l'instrumentation vise à identifier

des *SNPs* avec une fréquence d'allèle de risque supérieur à 5% qui ont une forte association avec le phénotype d'intérêt (*IMC* ou *CRP*). La force d'association est établie en considérant la valeur du coefficient estimé de la régression du phénotype sur le *SNP* ciblé et la *valeur-p* associée. Concrètement, nous voulons un coefficient ayant simultanément une grande valeur et une petite *valeur-p* correspondante.

Pour sélectionner chacun des instruments, nous avons considéré des *SNPs* candidats indépendants mentionnés dans les études d'associations pangénomiques (en abrégé *GWAS*) les plus complètes et récentes. Nous avons ensuite vérifié si les *SNPs* identifiés des *GWAS* se trouvaient dans la base de données de l'*ÉLCV*. Dans le cas où un *SNP* cible n'était pas disponible dans l'*ÉLCV*, nous avons remplacé celui-ci avec un *SNP* proxy en déséquilibre de liaison élevé avec le *SNP* cible.

Nous avons utilisé les travaux de Said et al. (2022) qui se sont basés sur deux populations d'ascendance caucasienne respectivement constituées de $n = 427\,367$ individus de l'*UK biobank* et $n = 575\,531$ individus de l'étude d'association à l'échelle du génome sur la *CRP* nommée *Cohorts for Heart and Aging Research in Genomic Epidemiology* (*CHARGE*). Said et al. (2022) ont identifié 266 *SNPs* indépendants dans 42 ensembles de gènes. Le Tableau 4.1 présente la liste des cinq premiers *SNPs* candidats (classés par ordre décroissant de coefficient d'association estimé) en relation avec la *CRP* identifiés par Said et al. (2022) que nous avons considérés pour sélectionner le *SNP* instrumentant la *CRP* dans notre analyse.

TABLEAU 4.1. Description des 5 *SNPs* significatifs associés à la *CRP* (extrait de Said et al. (2022)).

<i>Variant</i>	<i>Allèle d'effet</i>	$\hat{\beta}$	<i>ES</i> ¹	<i>valeur-p</i>	<i>MAF</i> ²
rs17616063	A	0.134	0.004	3.91×10^{-205}	0.079
rs11208685	A	0.097	0.003	2.25×10^{-320}	0.208
rs75460349	A	0.095	0.007	3.40×10^{-45}	0.025
rs61946383	A	0.082	0.002	4.94×10^{-324}	0.390
rs1260326	T	0.076	0.002	2.70×10^{-303}	0.411

¹ *ES* = Erreur standard ; ² *MAF* est l'acronyme anglais pour Minor allele frequency ; représente la fréquence allélique, c'est-à-dire la fréquence à laquelle se trouve l'allèle mineur d'un variant dans une population.

Nous avons sélectionné le *SNP* *rs61946383* comme instrument pour la *CRP* car ce *SNP* est celui dont l'allèle mineur est rencontré le plus fréquemment dans la population (*MAF* = 0.390), a la plus petite *valeur-p* ($\leq 4.94 \times 10^{-324}$) et possède un coefficient $\hat{\beta} = 0.082$ suffisamment grand.

Nous avons utilisé un raisonnement similaire pour identifier le *SNP* instrumentant l'*IMC*. En effet, selon nos recherches et investigations, plusieurs *GWAS* ont montré que les gènes du *FTO* sont associés au phénotype de l'obésité (mesuré par l'*IMC*) dans les populations humaines en général et dans les populations caucasiennes en particulier (Wardle *et al.*, 2008; Hales *et al.*, 2018). La sélection des *SNPs* candidats pour l'*IMC* est présentée dans le Tableau 4.2 montrant une liste de cinq *SNPs* sélectionnés de l'étude de Locke et al. (2015) basée sur $n = 339\,224$ individus caucasiens des deux sexes.

TABLEAU 4.2. Description de cinq *SNPs* communs significatifs associés à l'*IMC* (extraits de Locke et al. (2015)).

<i>Variant</i>	<i>Allèle d'effet</i>	$\hat{\beta}$	<i>ES</i> ¹	<i>valeur-p</i>	<i>MAF</i> ²
rs1558902	A	0.082	0.003	7.51×10^{-153}	0.415
rs6567160	C	0.056	0.004	3.93×10^{-53}	0.236
rs13021737	G	0.060	0.004	1.11×10^{-50}	0.828
rs10938397	G	0.040	0.003	3.21×10^{-38}	0.434
rs543874	G	0.048	0.004	2.62×10^{-35}	0.193

¹ *ES* = Erreur standard ; ² *MAF* est l'acronyme anglais pour Minor allele frequency ; représente la fréquence allélique, c'est-à-dire la fréquence à laquelle se trouve l'allèle mineur d'un variant dans une population.

Sur la base des critères de sélection du *SNP* cible préalablement définis et de sa disponibilité dans la base de données *globale*, nous sélectionnons le *SNP* *rs1558902* pour instrumenter l'*IMC*.

4.2.4 Analyses

Nous débutons par une analyse descriptive de notre échantillon extrait de la cohorte de l'*ÉLCV*. Par la suite, nous mettons en oeuvre la méthode du produit des coefficients, qui est équivalente à l'approche causale dans le cas des modèles simples considérés ici, ainsi que l'approche de *RM* en deux étapes telle qu'appliquée par Carter et al. (2021) pour inférer sur le rôle médiateur de la *CRP* dans l'association entre l'*IMC* et la *cIMT*.

Pour mettre en oeuvre l'approche de *RM* en deux étapes, nous avons déterminé au

préalable la force des instruments pour l'*IMC* et la *CRP* respectivement. Ensuite, pour chacune des deux méthodes, nous avons estimé l'effet direct, l'effet indirect, l'effet total ainsi que la proportion médiée pour un changement des valeurs de l'*IMC* de $a^* = 24$ à $a = 31$. Ces valeurs spécifiques de l'*IMC* correspondent à un *IMC* d'une personne avec poids normal et à un *IMC* d'une personne obèse, respectivement (Nuttall, 2015). Les modèles ont été ajustés ou non sur les covariables dans le but de comprendre leur rôle potentiellement confondant dans la relation entre l'*IMC* et la *cIMT* médiée par la *CRP*.

Dans nos analyses, nous appliquons une transformation logarithmique sur la variable *CRP* alors que toutes les autres variables restent non transformées. Dans toute la suite de la présentation, la mention de la variable *CRP* renvoie à sa forme logarithmique telle que définie dans le Tableau A.1. Le traitement des données et l'analyse sont faits à l'aide du logiciel R version 4.0.5.

4.2.4.1 Application aux données

Nous précisons que dans le cadre de cette application, nous avons effectué les estimations par la méthode du produit des coefficients et de la *TSMR* des divers effets de médiation dans deux situations particulières en ce qui concerne l'ajustement sur les variables potentiellement confondantes, soit des analyses sans et avec ajustement. Le but de considérer ces deux situations était d'évaluer l'impact de l'ajustement sur les effets estimés. À l'exclusion des covariables âge et sexe, il n'est habituellement pas adéquat d'ajuster la *RM* par de multiples autres variables puisque l'un des attraits principaux de cette méthode est qu'elle surmonte le problème lié à la présence de confusion non mesurée induite pour les autres variables omises dans le modèle. En effet, il est recommandé d'inclure uniquement comme covariables l'âge, le sexe, les principales composantes génomiques et les covariables

techniques (telles que le centre de recrutement), car un ajustement supplémentaire peut biaiser les estimations (Burgess *et al.*, 2019). C'est par exemple le cas si l'ajustement concerne une variable qui est sur le chemin causal menant d'une *SNP* à la réponse (un médiateur), ou si l'ajustement se fait sur un collisionneur (en anglais collider). On note qu'un collisionneur est une variable causée par au moins deux autres variables (par exemple l'exposition et la réponse). Dans le cas présent, nous avons ajusté la *RM* pour les covariables comme l'âge à l'enrôlement, le niveau de scolarité, le sexe, le statut marital et le revenu annuel total du ménage pour pouvoir la comparer avec la méthode du produit des coefficients.

Avant de passer à l'explicitation des modèles de régression que nous avons utilisés pour chacune des deux méthodes considérées, on précise que nous avons investigué si les hypothèses de linéarité pour les diverses relations sont respectées. Il s'agit notamment des relations instrument-exposition, instrument-médiateur, exposition-médiateur, médiateur-réponse et exposition-réponse. À cette fin, nous avons procédé à une inspection visuelle de la linéarité. À l'observation, nous n'avons pas remarqué un problème grave de non linéarité bien que la linéarité n'eût été parfaitement respectée dans chaque cas. Toutefois, puisque la variable âge semblait poser un problème de linéarité dans les associations qui l'impliquaient, nous avons corrigé en lui adjoignant un terme quadratique.

En ce qui concerne la méthode du produit des coefficients, les équations de régression (4.1)-(4.2) ont été considérées pour effectuer les estimations des divers effets de médiation à savoir l'*ET*, l'*EI*, l'*ED* :

$$\begin{aligned} CRP = & \beta_0 + \beta_1 IMC + \beta_2 Sexe + \beta_3 Age + \beta_4 Age^2 + \beta_5 Revenu \\ & + \beta_6 Education + \beta_7 Marital + \epsilon; \end{aligned} \quad (4.1)$$

$$\begin{aligned} cIMT = & \eta_0 + \eta_1 IMC + \eta_2 CRP + \eta_3 Sexe + \eta_4 Age + \eta_5 Age^2 + \eta_6 Revenu \\ & + \eta_6 Education + \eta_7 Marital + \epsilon. \end{aligned} \quad (4.2)$$

On conserve les mêmes équations (4.1)-(4.2) pour la version non ajustée de cette méthode mais en y enlevant toutes les covariables d'ajustement

En ce qui concerne la méthode *TSMR*, nous l'avons appliquée comme Carter et al. (2021) pour effectuer les estimations des divers effets de médiation (l'*ET*, l'*EI* et l'*ED*) et ensuite pouvoir déduire la proportion mediée.

Dans la première étape de la méthode (voir Équations (4.3)-(4.4)), on instrumente sur l'*IMC* en se servant du *SNP* *rs1558902* comme *IV*. Cette étape permet d'estimer sans biais le paramètre β_1 qui encode l'effet de l'*IMC* sur la *CRP* :

$$\begin{aligned} IMC = & \alpha_0 + \alpha_1 rs1558902 + \alpha_2 Sexe + \alpha_3 Age + \alpha_4 Age^2 + \alpha_5 Revenu \\ & + \alpha_6 Education + \alpha_7 Marital + \epsilon; \end{aligned} \quad (4.3)$$

$$\begin{aligned} CRP = & \beta_0 + \beta_1 \widehat{IMC} + \beta_2 Sexe + \beta_3 Age + \beta_4 Age^2 + \beta_5 Revenu \\ & + \beta_6 Education + \beta_7 Marital + \epsilon. \end{aligned} \quad (4.4)$$

Dans la seconde étape de la méthode (voir Équation (4.3) et Équations (4.5)-(4.6)), on instrumente la *CRP* en se servant du *SNP* *rs61946383*. Le but est d'estimer sans biais le paramètre θ_1 qui encode l'effet direct de l'*IMC* sur la *cIMT*, ainsi que θ_2 , l'effet de la *CRP* sur la *cIMT* :

$$\begin{aligned} CRP = & \alpha_0^* + \alpha_1^* rs1558902 + \alpha_2^* rs61946383 + \alpha_3^* Sexe + \alpha_4^* Age \\ & + \alpha_5^* Age^2 + \alpha_6^* Revenu + \alpha_7^* Education + \alpha_8^* Marital + \epsilon; \end{aligned} \quad (4.5)$$

$$\begin{aligned} cIMT = & \theta_0 + \theta_1 \widehat{IMC} + \theta_2 \widehat{CRP} + \theta_3 Sexe + \theta_4 Age + \theta_5 Age^2 \\ & + \theta_6 Revenu + \theta_7 Education + \theta_8 Marital + \epsilon. \end{aligned} \quad (4.6)$$

4.2.4.2 Résultats

Les statistiques descriptives sont présentées sous forme de moyenne (écart-type) pour les variables continues, alors que les variables catégorielles sont présentées

sous forme de pourcentage. L'âge moyen des $n = 13\,942$ participants à l'enrôlement était de 61.6 ans (9.67). Un équilibre de sexe avec presque qu'autant de femmes (48.07%) que d'hommes était observé. La majorité des participants étaient en couple (73%). Le revenu annuel total des ménages de 79% des participants était supérieur à 50 000 \$. Le niveau de scolarité de plus de 79% des participants était supérieur ou égal aux études secondaires. L' IMC moyen était de 27.74 (5.14). Le taux moyen de protéine C-réactive dans le sang des participants était de 0.94 (0.61). Enfin, la $cIMT$ moyenne des participants était de 0.77 (0.17) (voir Tableau 4.3 pour plus de détails).

TABLEAU 4.3. Analyse descriptive des variables de l'échantillon d'étude.

	<i>Échantillon total</i> ¹		<i>Poids normal</i>	<i>Embonpoint</i>	<i>Obésité</i>
<i>Participants, n</i>	<i>n</i> = 13 942		<i>n</i> = 4440	<i>n</i> = 5762	<i>n</i> = 3740
<i>Âge à l'enrôlement</i> ² , années	61.6 (9.67)		61.16 (10.01)	62.23 (9.81)	61.12 (8.97)
<i>Sexe, n (%)</i>					
Femmes	6702 (48%)		2604 (58.6%)	2342 (40.6 %)	1756 (47%)
<i>Statut marital, n (%)</i>					
En couple	10 160 (73%)		3207 (72.2 %)	4336 (75.3 %)	2617 (70 %)
<i>Revenu annuel total du ménage, n (%)</i>					
≥ 50 000 \$	10952 (79%)		3536 (79.6 %)	4581 (79.5 %)	2835 (75.8)
<i>Niveau d'éducation, n (%)</i>					
≥ Secondaire	11 115 (79.72%)		3711 (83.6 %)	4578 (79.5)	2826 (75.6%)
<i>Indice de masse corporelle (IMC), kg/m²</i>	27.74 (5.14)				
Poids normal			22.70 (1.72)		
Embonpoint				27.35 (1.40)	
Obésité					34.32 (4.27)
<i>Taux de protéine C-réactive (CRP), mg/l</i>	0.94 (0.61)		0.71 (0.52)	0.92 (0.55)	1.27 (0.66)
<i>Épaisseur de l'intima média carotidien (cIMT), mm</i>	0.77 (0.17)		0.74 (0.16)	0.78 (0.16)	0.79 (0.17)

¹ Les statistiques descriptives présentées dans ce tableau sont faites dans un premier temps pour l'échantillon entier et ensuite de façon stratifiée selon trois intervalles de l'exposition (*IMC*) : [0 ; 25) poids normal, [25 ; 30) embonpoint et [30.0 ; +∞) obésité (voir <https://www.cdc.gov/healthyweight/assessing/bmi/adultgmi/index.html>) ;

² Les résultats pour les variables continues sont présentés sous la forme : moyenne (écart-type).

Pour pouvoir utiliser la méthode *TSMR* dans notre analyse de médiation, un préalable est de vérifier l'existence d'une association entre l'instrument de l'exposition *SNP rs1558902* et cette dernière (*IMC*) ainsi que l'existence d'une association entre l'instrument du médiateur *SNP rs61946383* et ce dernier (*CRP*). Les instruments respectifs de l'exposition et du médiateur ont des valeurs de statistique $F > 10$ (voir Tableau 4.4) et sont donc considérés comme forts (Baiocchi *et al.*, 2014).

TABLEAU 4.4. Forces de l'instrument *SNP rs1558902* pour l'exposition (*IMC*) et de l'instrument *SNP rs61946383* pour le médiateur (*CRP*). Associations obtenues par moindres carrés ordinaires.

		<i>Effet (ES^1)</i>	<i>Statistique-F</i>
<i>IV²IMC</i>	<i>SNP rs1558902</i>	0.55 (0.06)	76.93
<i>IV CRP</i>	<i>SNP rs61946383</i>	0.05 (0.01)	45.68

¹ *ES* = Erreur standard ; ² *IV* = Variable instrumentale.

Nous présentons les résultats de manière simultanée pour deux situations, celle dans laquelle aucun ajustement n'est fait sur les covariables que sont l'âge, le niveau de scolarité, le sexe, le statut marital et le revenu annuel total du ménage et celle dans laquelle un ajustement est fait. Nous présentons aussi les estimés des effets dans ce même ordre selon les effets total, indirect et direct. Nous précisons que, malgré le fait d'avoir présenté le calcul des variances des estimateurs par la méthode delta, nous utilisons la méthode du bootstrap pour obtenir les intervalles de confiance. En effet la méthode du bootstrap est plus robuste aux données bruitées et aux erreurs de spécification des modèles (Hole, 2007), de plus pour la réplication que nous faisons, c'est la méthode du bootstrap qui a été choisie par ses auteurs (Carter *et al.*, 2021).

TABLEAU 4.5. Estimations des effets de médiation du taux de protéine C-réactive dans le sang dans l'association entre l'indice de masse corporelle et l'épaisseur de l'intima média-carotidien par les méthodes du produit des coefficients et de la randomisation mendélienne en deux étapes.

	Méthode du produit des coefficients			Méthode TSMR ¹		
	Modèles non ajustés					
	<i>Effets</i> ²	<i>ES</i> ³	<i>IC</i> 95% ⁴	<i>Effets</i>	<i>ES</i>	<i>IC</i> 95%
<i>ET</i>	0.027	0.002	[0.023 ; 0.031]	0.027	0.026	[-0.024 ; 0.079]
<i>EI</i>	0.005	0.0009	[0.003 ; 0.006]	0.012	0.018	[-0.023 ; 0.047]
<i>ED</i>	0.022	0.0022	[0.018 ; 0.026]	0.015	0.030	[-0.043 ; 0.074]
<i>PM</i>	0.175	0.0353	[0.114 ; 0.254]	0.436	15.212	[-0.536 ; 1.407]
	Modèles ajustés					
<i>ET</i>	0.028	0.0018	[0.024 ; 0.032]	0.008	0.023	[-0.038 ; 0.054]
<i>EI</i>	0.002	0.0008	[0 ; 0.003]	0.003	0.016	[-0.029 ; 0.034]
<i>ED</i>	0.026	0.002	[0.025 ; 0.03]	0.005	0.027	[-0.048 ; 0.06]
<i>PM</i>	0.06	0.0293	[0 ; 0.11]	0.312	13.225	[-0.6 ; 1.22]

¹ *TSMR* = Méthode de randomisation mendélienne en deux étapes; ² Dans les méthodes du produit des coefficients et de *TSMR*, les effets de médiation (effet total (*ET*), effet direct (*ED*) et effet indirect (*EI*) ainsi que la proportion médiée (*PM*)) sont estimés en considérant deux valeurs spécifiques de l'*IMC* dont $a^* = 24$ (correspondant à un *IMC* d'une personne normale) et $a = 31$ (correspondant à un *IMC* d'une personne obèse); ³ *IC* = Intervalles de confiance obtenus par bootstrap percentile avec $B = 2\,000$ réplifications; ⁴ *ES*=Erreur standard.

Les méthodes du produit des coefficients et la *TSMR* pour les **modèles non ajustés** montrent que les effets totaux ponctuels sont les mêmes ($ET = 0.027$). La *TSMR* présente toutefois des intervalles de confiance plus larges comparés à ceux obtenus par la méthode du produit ($IC = [0.023; 0.031]$ et $[-0.024; 0.079]$ respectivement). La méthode du produit des coefficients mène à un effet direct de l'*IMC* sur l'épaisseur de l'intima média carotidien (*cIMT*) ($ED = 0.022$ $[0.018; 0.026]$), plus important que celui de la *TSMR* (voir Tableau 4.5). On note que dans cette situation, la *TSMR* présente un intervalle de confiance plus large incluant

0, par conséquent elle ne détecte aucun effet significatif de l'*IMC* sur la *cIMT* (voir Tableau 4.5). On note aussi que la *RM* présente un effet indirect ($EI = 0.012 [-0.023; 0.047]$) plus fort et donc une plus grande proportion médiée ($PM = 0.436 [-0.536; 1.407]$).

En ce qui concerne **les modèles ajustés**, l'effet total obtenu par la méthode du produit reste sensiblement le même alors que le même effet obtenu par la *TSMR* diminue de façon marquée ($ET = 0.008 [-0.038; 0.054]$). La méthode du produit présente un effet indirect ($EI = 0.002 [0; 0.003]$) plus faible que lorsque le modèle est non ajusté, mais il reste significatif. L'effet indirect de la *TSMR* ($EI = 0.003 [-0.029; 0.034]$) diminue également et reste non significatif en comparaison au modèle ajusté. La proportion médiée reste importante pour la *TSMR* ($PM = 0.312 [-0.6; 1.22]$) mais diminue de façon marquée pour la méthode du produit ($PM = 0.06 [0; 0.11]$). Il est important de mentionner que tous les résultats ajustés de la *TSMR* sont non significatifs. En effet, les intervalles de confiance obtenus par la *TSMR* sont bien plus larges comparés à ceux de la méthode du produit (voir Tableau 4.5).

CHAPITRE V

CONCLUSION

5.1 Rappel des objectifs de l'étude

L'analyse de médiation est une méthode statistique qui permet d'étudier les mécanismes qui définissent les relations entre la variable d'exposition, la variable réponse et une ou plusieurs variables médiatrices. Elle aide à clarifier la nature de la relation entre la variable d'exposition et la variable réponse (Baron et Kenny, 1986; MacKinnon et Luecken, 2008) et permet ainsi d'améliorer la compréhension étiologique des problèmes en sciences de la santé (Richiardi *et al.*, 2013; Rizzo *et al.*, 2022).

En dehors du fait qu'elle nécessite des hypothèses fortes pour sa mise en oeuvre, une faiblesse de l'analyse de médiation causale standard est qu'elle ne prend pas en compte les confondants non mesurés, c'est-à-dire les variables non mesurées qui sont les causes communes de l'exposition, du médiateur et de la réponse. Cette limite conduit souvent à des estimations d'effets de médiation biaisés et rend ainsi difficile une interprétation causale de l'effet total et des effets naturels direct et indirect. Il était donc souhaitable que soit proposée une approche alternative d'inférence causale en médiation permettant de surmonter les limites précédem-

ment mentionnées. Cette approche est la randomisation mendélienne en analyse de médiation. La randomisation mendélienne utilise des variants génétiques comme variables instrumentales pour un ou plusieurs phénotypes dans le but de faire une estimation de l'effet causal d'une exposition sur une réponse qui, sous certaines hypothèses, est exempte de biais dûs à des variables de confusion non mesurées.

Dans ce mémoire, il était question de faire un état de l'art en analyse de médiation et en inférence causale, d'expliquer la théorie sur les *IV*, de présenter des concepts en *RM*. Le reste des objectifs portait sur la compréhension et l'évaluation, par réplication des simulations, des méthodes et approches étudiées par Carter et al. (2021). Le point final était d'appliquer ces méthodes sur le jeu de données réelles issu de la cohorte de l'*ÉLCV*.

D'un point de vue théorique, nous avons au préalable, étudié l'approche de Baron et Kenny, la méthode du produit et la méthode de la différence. Par la suite, nous avons fait une introduction à la médiation causale basée sur le cadre de travail contrefactuel et discuté des hypothèses requises pour l'interprétation causale des différents types d'effets de médiation. Nous avons ensuite abordé la *RM* après avoir explicité la notion de *IV* ainsi que les principaux concepts qui les sous-tendent comme l'instrumentation, les propriétés des variables instrumentales ainsi que des suppositions sur leur validité. Nous avons aussi discuté de deux méthodes d'estimation d'effets causaux utilisées en *RM* à savoir : la méthode de Wald et la méthode *2SLS*.

Enfin, nous avons répliqué le code de Carter et al. (2021) et corrigé quelques résultats qui semblaient incorrects. Dans ce cadre, nous avons pu détecter et corriger pour les effets total, direct et indirect, ainsi que la proportion médiée et leurs erreurs standard respectives, au total, 10 erreurs d'arrondis numériques, 43 erreurs de calcul et 9 erreurs de recopie des résultats. Les résultats obtenus lors de la ré-

plication des simulations ne différaient pas qualitativement des résultats présentés par Carter et al. (2021). Ainsi, la *TSMR* a mené à des estimations d'effets peu ou non biaisées contrairement à la méthode du produit des coefficients standard. Les résultats et conclusions de Carter et al. (2021) sont donc restés pertinents.

L'application des méthodes du produit et de la *TSMR* abordées au Chapitre III aux données de la cohorte de l'*ÉLCV* a permis d'illustrer la démarche de Carter et al. (2021). En effet, dans une situation où les hypothèses sous-jacentes à ces deux modèles sont respectivement satisfaites, il est possible d'estimer des effets causaux en présence de facteurs confusion non mesurés pour la *TSMR*. Les deux méthodes ont souligné le rôle de la *CRP* comme médiateur putatif de la relation entre l'*IMC* et la *cIMT*. Cependant, des différences entre les résultats des deux méthodes sont apparues, principalement au niveau de l'étendue des intervalles de confiance. En effet la *TSMR* a mené à des intervalles de confiance beaucoup plus larges que la méthode du produit. La différence de magnitudes entre les effets ajustés et non ajustés était aussi plus importante pour la *TSMR* comparativement à la méthode du produit.

Les différences observées dans les estimations des effets de la *TSMR*, ainsi que dans l'amplitude des intervalles de confiance pourraient être liés à une possible perte de puissance de cette dernière. En effet, il est connu que l'instrumentation mène à une perte de précision de l'estimation de l'association entre l'exposition et réponse (Burgess, 2014).

5.2 Limites et forces de notre étude

Les travaux effectués dans ce mémoire présentent certaines limites dans le sens où nous n'avons pas mené d'études de simulations basés sur des scénarios autres que ceux considérés par Carter et al. (2021). Par exemple, nous aurions pu considérer

d'autres types de structure de corrélation afin de valider plus complètement les postulats faits par Carter et al. (2021). Pour l'analyse de données réelles, une limite réside dans le fait que nous n'avons pas utilisé des scores de risque polygéniques pour améliorer la qualité des estimations des effets par la *TSMR*. Pour la méthode du produit, nous n'avons pas choisi de définir un ensemble de covariables mesurés approprié pour chaque relation spécifique exposition-médiateur, médiateur-réponse et exposition-réponse ; au lieu de cela, nous avons ajusté pour chacune des relations par le même ensemble de covariables.

Une autre limite de ce travail est l'absence de mesures de *cIMT* au premier suivi (Follow-up 1) qui nous a obligé à considérer à la place celles de la baseline. Cette difficulté aurait biaisé les conclusions de notre étude du fait de l'intervalle de temps important qui sépare la baseline du premier suivi. En effet, le recrutement et la collecte des données de référence ont été achevés en 2015, et le premier suivi a été achevé à la mi-2018 (Raina *et al.*, 2019).

Malgré ces limites, notre étude possède également quelques points forts. Elle décrit certaines des méthodes classiques et des méthodes instrumentales utilisées en analyse de médiation. Elle corrige certaines erreurs des résultats de Carter et al. (2021) et fait ainsi la confirmation des postulats qu'ils font sur la performance des méthodes de randomisation mendéliennes en analyse de médiation. Elle calcule des intervalles de confiance des estimés des effets par la même méthode (bootstrap). Enfin, notre étude a évalué le comportement empirique de l'approche de Carter et al. (2021) en analyse de médiation en l'appliquant à un jeu de données réel.

ANNEXE A

TABLEAU A.1. Données non génétiques mesurées de la cohorte globale de l'ÉLCV utilisées dans notre analyse.

	Modes de prélèvement et unité de mesure				
	Auto-rapportées	Mesurées	Unité	Catégories de l'ÉLCV	Modalités utilisées dans notre analyse
<i>Mesures socio-démographiques</i>					
Âge	Oui	-	Années	-	-
Niveau de scolarité	Oui	-	-	9	≤ Secondaire, > Secondaire
Sexe	Oui	-	-	2	Femme, Homme
Statut marital	Oui	-	-	6	Célibataire, En couple
Revenu annuel total du ménage	Oui	-	\$	7	< 50 000, ≥ 50 000
<i>Mesures physiques</i>					
Statut d'obésité (mesuré par l'IMC)	-	Oui	kg/m ²	-	-
Épaisseur de l'intima-média carotidienne	-	Oui	mm	-	-
<i>Mesures chimiques et biologiques</i>					
Protéine C-réactive à haute sensibilité	-	Oui	mg/l	-	$\log(CRP + 1)$

ANNEXE B

```
library("AER")
library("CMAverse")

function_no_adjustment=function(database,noboots,astar,a){
  #-----
  # Calcul des effets de mediation par la methode
  #du produit 2SMR : Crude
  #-----

  # Effet total
  effet.total.model=ivreg(cIMT ~ IMC | SNP1, data =database)
  ET=effet.total.model$coefficients[2]

  # Effet de A (IMC) sur le mediateur M (CRP)
  effet.indirect.model.1 <- ivreg(CRP ~ IMC | SNP1, data = database)

  # Effect of M on Y
  effet.indirect.model.2 <- ivreg(cIMT ~ CRP + IMC
  | SNP2 + SNP1, data = database)

  # Effet Indirect
  EI <- effet.indirect.model.1$coefficients[2]
  *effet.indirect.model.2$coefficients[2]

  # Effet direct
```

```

ED <- effet.indirect.model.2$coefficients[3]

# Proportion mediee
PM <- EI/ ET

results.effect=abs(a-astar)*c(ET,EI,ED,1/abs(a-astar)*PM)

#-----
#Code pour le calcul des moyennes des estimateurs et
#des intervalles de confiance percentiles Bootstrap
#-----

boot.total.result<- rep(0, noboots)
boot.indirect.result<- rep(0, noboots)
boot.direct.result<- rep(0, noboots)
boot.proportion.results<- rep(0, noboots)
set.seed(100)
for(b in 1:noboots){
  sample_idx <- sample(1:nrow(database),
    size=nrow(database), replace = T)
  boot.data=clsa[sample_idx,]
  # Effet total
  boot.total=ivreg(cIMT ~ IMC | SNP1, data = boot.data)
  boot.total.result[b]=boot.total$coefficients[2]
  # Effet de A (IMC) sur le mediateur M (CRP)
  boot.expo.media <- ivreg(CRP ~ IMC | SNP1,
    data = boot.data)
  # Effect of M on Y
  boot.media.rep <- ivreg(cIMT ~ CRP + IMC |
    SNP2 + SNP1, data = boot.data)
  # Effet Indirect
  boot.indirect.result[b] <- boot.expo.media$coefficients[2]

```

```

*boot.media.rep$coefficients[2]
# Effet direct
boot.direct.result[b] <- boot.media.rep$coefficients[3]
# Proportion mediee
boot.proportion.results[b] <- boot.indirect.result[b]
/ boot.total.result[b]
}
#-----
# Percentile bootstrap 95% confidence interval
#-----
# Percentile effet total
boot.total.lci=quantile(boot.total.result,.025)
boot.total.uci=quantile(boot.total.result,.975)
# Percentile indirect
boot.indirect.result.lci=quantile(boot.indirect.result,.025)
boot.indirect.result.uci=quantile(boot.indirect.result,.975)
# Percentile direct
boot.direct.result.lci=quantile(boot.direct.result,.025)
boot.direct.result.uci=quantile(boot.direct.result,.975)
# Percentile proportion mediee
boot.proportion.results.lci=quantile(boot.proportion.results,.025)
boot.proportion.results.uci=quantile(boot.proportion.results,.975)
#-----
# bootstrap 95% confidence interval normal
#-----
n=nrow(database)
#-----
# normal total

```

```

s=sd(boot.total.result)
error <- qnorm(0.975)*s
boot.total.lci.normal=ET-error
boot.total.uci.normal=ET+error
IC.total=c(boot.total.lci.normal,boot.total.uci.normal)

# normal indirect
s=sd(boot.indirect.result)
error <- qnorm(0.975)*s
boot.indirect.result.lci.normal=EI-error
boot.indirect.result.uci.normal=EI+error
IC.indirect=c(boot.indirect.result.lci.normal,
boot.indirect.result.uci.normal)

# normal direct
s=sd(boot.direct.result)
error <- qnorm(0.975)*s
boot.direct.result.lci.normal=ED-error
boot.direct.result.uci.normal=ED+error
IC.direct=c(boot.direct.result.lci.normal,
boot.direct.result.uci.normal)

# proportion direct
s=sqrt(PM*(1-PM))
error <- qnorm(0.975)*s
boot.proportion.results.lci.normal=PM-error
boot.proportion.results.uci.normal=PM+error
IC.proportion=c(boot.proportion.results.lci.normal,
boot.proportion.results.uci.normal)

#-----
# Resultats

```

```

    results_boot=abs(a-astar)*rbind(c(boot.total.lci,
boot.total.uci,boot.total.lci.normal,
boot.total.uci.normal),
c(boot.indirect.result.lci,
boot.indirect.result.uci,
boot.indirect.result.lci.normal
,boot.indirect.result.uci.normal)
,c(boot.direct.result.lci,boot.direct.result.uci,
boot.direct.result.lci.normal,
boot.direct.result.uci.normal),
1/abs(a-astar)*c(boot.proportion.results.lci,
boot.proportion.results.uci,
boot.proportion.results.lci.normal,
boot.proportion.results.uci.normal))
results_final=cbind(results.effect,results_boot)
rownames(results_final) <- c("ET", "EI","ED","PM")
colnames(results_final) <- c("Effets","LCI percentile",
"UCI percentile", "LCI normal", "UCI normal")
    return(results_final)
}

write.csv2(function_no_adjustment(clsa,2000,24,31),
file="no_adjustment_F.csv")
#-----
# Methode du produit des coefficients : Crude
#-----
#--- delta with normal IC

est_boot_sajdelta <- cmest(data = clsa5, model = "rb",

```

```

outcome = "cIMT", exposure = "IMC",mediator = "CRP",
EMint = FALSE,mreg = list("linear"),
yreg = "linear",mval=list(1), astar = 24, a = 31,
estimation = "paramfunc",
inference = "delta")

summary(est_boot_sajdelta)

#--- Boot with percentiles IC

est_boot_sajp <- cmest(data = clsa, model = "rb",
outcome = "cIMT", exposure = "IMC",mediator = "CRP",
EMint = FALSE,mreg = list("linear"), yreg = "linear",
mval=list(1),astar = 24, a = 31, estimation = "imputation",
inference = "bootstrap",boot.ci.type="per",nboot = 2000)
summary(est_boot_sajp)

#--- Boot with BCa IC
est_boot_sajbca <- cmest(data = clsa, model = "rb",
outcome = "cIMT",exposure = "IMC",mediator = "CRP",
EMint = FALSE,mreg = list("linear"), yreg = "linear",
mval=list(1), astar = 24, a = 31, estimation = "imputation",
inference = "bootstrap",boot.ci.type="bca",nboot = 2000)
summary(est_boot_sajbca)

#-----
library("AER")
library("CMAverse")
#-----

```

```

function_adjustment_quadra_age_cat_sexe_cat_all=
function(database,noboots,astar,a){
#-----
# Calcul des effets de mediation par la methode
# du produit 2SMR : Ajustement sur age, sexe
#-----
# Effet total
effet.total.model=ivreg(cIMT ~ IMC+Sexe+Age+
Age_quadra+Revenu+Education | SNP1+Sexe+Age+Age_quadra+
Revenu+Education, data = database)
ET=effet.total.model$coefficients[2]
# Effet de A (BMI) sur le mediateur M (CRP)
effet.indirect.model.1 <- ivreg(CRP ~ IMC+Sexe+Age+
Age_quadra+Revenu+Education| SNP1+Sexe+
Age+Age_quadra+Revenu+Education,data = database)
# Effect of M on Y
effet.indirect.model.2 <- ivreg(cIMT ~ CRP +
IMC+Sexe+Age+
Age_quadra+Revenu +Education | SNP2 + SNP1+Sexe+
Age+Age_quadra
+Revenu+Education, data = database)
#summary(CRP_cIMT)
# Effet Indirect
EI <- effet.indirect.model.1$coefficients[2]
*effet.indirect.model.2$coefficients[2]
# Effet direct
ED <- effet.indirect.model.2$coefficients[3]
# Proportion mediee

```

```

PM <- EI/ ET

results.effect=abs(a-astar)*c(ET,EI,ED,1/abs(a-astar)*PM)

#-----
# Code pour le calcul des moyennes des estimateurs et
# des intervalles de confiance percentiles Bootstrap
#-----

boot.total.result<- rep(0, noboots)
boot.indirect.result<- rep(0, noboots)
boot.direct.result<- rep(0, noboots)
boot.proportion.results<- rep(0, noboots)
set.seed(100)
for(b in 1:noboots){
  sample_idx <- sample(1:nrow(database), size=nrow(database)
    , replace = T)
  boot.data=database[sample_idx,]
  # Effet total
  boot.total=ivreg(cIMT ~ IMC+Sexe+Age+Age_quadra+Revenu+
    Education | SNP1+Sexe+Age+Age_quadra+Revenu+Education
    #, data = boot.data)
  boot.total.result[b]=boot.total$coefficients[2]
  # Effet de A (BMI) sur le mediateur M (CRP)
  boot.expo.media <- ivreg(CRP ~ IMC+Sexe+Age+
    Age_quadra+Revenu+Education | SNP1+Sexe+Age+
    Age_quadra+Revenu+Education , data = boot.data)
  # Effect de M sur Y
  boot.media.rep <- ivreg(cIMT ~ CRP + IMC+Sexe+Age
    +Age_quadra+Revenu +Education | SNP2 + SNP1+Sexe+
    Age+Age_quadra+Revenu

```

```

+Education , data = boot.data)

# Effet Indirect

boot.indirect.result[b] <- boot.expo.media$coefficients[2]
*boot.media.rep$coefficients[2]

# Effet direct

boot.direct.result[b] <- boot.media.rep$coefficients[3]

# Proportion mediee

boot.proportion.results[b] <- boot.indirect.result[b]
/ boot.total.result[b]
}

#-----
# Percentile bootstrap 95% confidence interval
#-----

# Percentile effet total

boot.total.lci=quantile(boot.total.result,.025)
boot.total.uci=quantile(boot.total.result,.975)

# Percentile indirect

boot.indirect.result.lci=quantile(boot.indirect.result,.025)
boot.indirect.result.uci=quantile(boot.indirect.result,.975)

# Percentile direct

boot.direct.result.lci=quantile(boot.direct.result,.025)
boot.direct.result.uci=quantile(boot.direct.result,.975)

# Percentile proportion mediee

boot.proportion.results.lci=quantile(boot.proportion.results,.025)
boot.proportion.results.uci=quantile(boot.proportion.results,.975)

#-----
# bootstrap 95% confidence interval normal
#-----

```

```

n=nrow(database)

#-----
# normal total
s=sd(boot.total.result)
error <- qnorm(0.975)*s
boot.total.lci.normal=ET-error
boot.total.uci.normal=ET+error
IC.total=c(boot.total.lci.normal,
boot.total.uci.normal)
# normal indirect
s=sd(boot.indirect.result)
error <- qnorm(0.975)*s
boot.indirect.result.lci.normal=EI-error
boot.indirect.result.uci.normal=EI+error
IC.indirect=c(boot.indirect.result.lci.normal,
boot.indirect.result.uci.normal)
# normal direct
s=sd(boot.direct.result)
error <- qnorm(0.975)*s
boot.direct.result.lci.normal=ED-error
boot.direct.result.uci.normal=ED+error
IC.direct=c(boot.direct.result.lci.normal,
boot.direct.result.uci.normal)
# proportion direct
s=sqrt(PM*(1-PM))
error <- qnorm(0.975)*s
boot.proportion.results.lci.normal=PM-error
boot.proportion.results.uci.normal=PM+error

```

```

    IC.proportion=c(boot.proportion.results.lci.normal,
                    boot.proportion.results.uci.normal)

#-----

# Resultats
results_boot=abs(a-astar)*rbind(c(boot.total.lci,
boot.total.uci,boot.total.lci.normal,
boot.total.uci.normal)
,c(boot.indirect.result.lci,
boot.indirect.result.uci,
boot.indirect.result.lci.normal,
boot.indirect.result.uci.normal),
c(boot.direct.result.lci,
boot.direct.result.uci,
boot.direct.result.lci.normal,
boot.direct.result.uci.normal),
1/abs(a-astar)*c(boot.proportion.results.lci,
boot.proportion.results.uci,
boot.proportion.results.lci.normal,
boot.proportion.results.uci.normal))
results_final=cbind(results.effect,results_boot)
rownames(results_final) <- c("ET", "EI","ED","PM")
colnames(results_final) <- c("Effets","LCI percentile",
"UCI percentile","LCI normal", "UCI normal")
return(results_final)
}

write.csv2(function_adjustment_quadra_age_cat_sexe_cat_all
(clsa,2000,24,31),file="adjustment_F.csv")

#-----

```

```

# Ajustement Age
#-----

# Methode du produit des coefficients : Ajustement sur Sexe, Age,
# Age_quadra, Revenu, Education et statut marital
#-----

#--- delta with normal IC
est_boot_sajdelta <- cmest(data = clsa, model = "rb",
outcome = "cIMT",exposure = "IMC",mediator = "CRP",
EMint = FALSE, basec = c("Sexe", "Age", "Age_quadra",
"Revenu", "Education"),
mreg = list("linear"), yreg = "linear",
mval=list(1), astar = 24, a = 31,
estimation = "paramfunc", inference = "delta")
summary(est_boot_sajdelta)

#--- Boot with percentiles IC
est_boot_sajp <- cmest(data = clsa, model = "rb",
outcome = "cIMT",exposure = "BMI",
mediator = "CRP", EMint = FALSE,
basec = c("Sexe", "Age", "Age_quadra",
"Revenu", "Education"),
mreg = list("linear"), yreg = "linear",
mval=list(1), astar = 24, a = 31,
estimation = "imputation",
inference = "bootstrap",boot.ci.type="per"
,nboot = 2000)
summary(est_boot_sajp)

```

```

#--- Boot with BCa IC

est_boot_sajbca <- cmest(data = clsa, model = "rb",
outcome = "cIMT", exposure = "BMI", basec = c("Sexe", "Age",
"Age_quadra", "Revenu", "Education"),
mediator = "CRP", EMint = FALSE,
mreg = list("linear"), yreg = "linear"
,mval=list(1), astar = 24, a = 31,
estimation = "imputation",
inference = "bootstrap",boot.ci.type="bca",
nboot = 2000)
summary(est_boot_sajbca)

#-----

```

RÉFÉRENCES

- Alwin, D. F. et Hauser, R. M. (1975). The decomposition of effects in path analysis. *American sociological review*, 37–47.
- Angrist, J. D. et Imbens, G. W. (1995). Two-stage least squares estimation of average causal effects in models with variable treatment intensity. *Journal of the american statistical association*, 90(430), 431–442.
- Angrist, J. D. et Krueger, A. B. (2001). Instrumental variables and the search for identification : from supply and demand to natural experiments. *Journal of economic perspectives*, 15(4), 69–85.
- Bahadoran, Z., Mirmiran, P. et Azizi, F. (2015). Fast food pattern and cardiometabolic disorders : a review of current studies. *Health promotion perspectives*, 5(4), 231.
- Baiocchi, M., Cheng, J. et Small, D. S. (2014). Instrumental variable methods for causal inference. *Statistics in medicine*, 33(13), 2297–2340.
- Baron, R. et Kenny, D. (1986). The moderator-mediator variable distinction in social psychological research : conceptual, strategic, and statistical considerations. *Journal of personality and social psychology*, 51, 1173–1182.
- Bollmann, G., Rouzinov, S., Berchtold, A. et Rossier, J. (2019). Illustrating instrumental variable regressions using the career adaptability–job satisfaction relationship. *Frontiers in psychology*, 10, 1481.
- Burgess, S. (2014). Sample size and power calculations in Mendelian randomization with a single instrumental variable and a binary outcome. *International journal of epidemiology*, 43(3), 922–929.
- Burgess, S., Butterworth, A., Malarstig, A. et Thompson, S. G. (2012). Use of Mendelian randomisation to assess potential benefit of clinical intervention. *British medical journal*, 345.
- Burgess, S., Daniel, R. M., Butterworth, A. S., Thompson, S. G. et Consortium, E.-I. (2015). Network Mendelian randomization : using genetic variants as instrumental variables to investigate mediation in causal pathways. *International journal of epidemiology*, 44(2), 484–495.

- Burgess, S., Small, D. S. et Thompson, S. G. (2017a). A review of instrumental variable estimators for Mendelian randomization. *Statistical methods in medical research*, 26(5), 2333–2355.
- Burgess, S., Smith, G. D., Davies, N. M., Dudbridge, F., Gill, D., Glymour, M. M., Hartwig, F. P., Holmes, M. V., Minelli, C., Relton, C. L. *et al.* (2019). Guidelines for performing Mendelian randomization investigations. *Wellcome open research*, 4.
- Burgess, S., Thompson, D. J., Rees, J. M., Day, F. R., Perry, J. R. et Ong, K. K. (2017b). Dissecting causal pathways using Mendelian randomization with summarized genetic data : application to age at menarche and risk of breast cancer. *Genetics*, 207(2), 481–487.
- Burgess, S. et Thompson, S. G. (2015). Multivariable Mendelian randomization : the use of pleiotropic genetic variants to estimate causal effects. *American journal of epidemiology*, 181(4), 251–260.
- Carter, A. R., Sanderson, E., Hammerton, G., Richmond, R. C., Smith, G. D., Heron, J., Taylor, A. E., Davies, N. M. et Howe, L. D. (2021). Mendelian randomisation for mediation analysis : current methods and challenges for implementation. *European journal of epidemiology*, 36(5), 465–478.
- Davey Smith, G. et Ebrahim, S. (2003). ‘Mendelian randomization’ : can genetic epidemiology contribute to understanding environmental determinants of disease? *International journal of epidemiology*, 32(1), 1–22.
- Davey Smith, G. et Hemani, G. (2014). Mendelian randomization : genetic anchors for causal inference in epidemiological studies. *Human molecular genetics*, 23(R1), R89–R98.
- Ellulu, M., Patimah, H., Khaza’ai, H., Rahmat, A. et Abed, Y. (2017). Obesity and inflammation : the linking mechanism and the complications. *Archives of medical science*, 13(4), 851–63.
- Ertefaie, A., Small, D. S. et Rosenbaum, P. R. (2018). Quantitative evaluation of the trade-off of strengthened instruments and sample size in observational studies. *Journal of the american statistical association*, 113(523), 1122–1134.
- Faraoni, D. et Schaefer, S. T. (2016). Randomized controlled trials vs. observational studies : why not just live together? *BioMed central anesthesiology*, 16(1), 1–4.
- Hales, C. M., Fryar, C. D., Carroll, M. D., Freedman, D. S. et Ogden, C. L.

- (2018). Trends in obesity and severe obesity prevalence in United states youth and adults by sex and age, 2007-2008 to 2015-2016. *Journal of the american medical association*, 319(16), 1723–1725.
- Hariton, E. et Locascio, J. J. (2018). Randomised controlled trials—the gold standard for effectiveness research. *BJOG : an international journal of obstetrics and gynaecology*, 125(13), 1716.
- Hemani, G., Bowden, J. et Davey Smith, G. (2018). Evaluating the potential role of pleiotropy in Mendelian randomization studies. *Human molecular genetics*, 27(R2), R195–R208.
- Heng, S., Zhang, B., Han, X., Lorch, S. A. et Small, D. S. (2019). Instrumental variables : to strengthen or not to strengthen ? *arXiv preprint arXiv :1911.09171*.
- Herinirina, N. F., Rajaonarison, L., Herijoelison, A. R. et Ahmad, A. (2015). Thickness of carotid intima-media and cardiovascular risk factors. *The pan african medical journal*, 21, 153–153.
- Hernan, M. A. et Robins, J. M. (2006). Instruments for causal inference : an epidemiologist’s dream ? *Epidemiology*, 360–372.
- Hernan, M. A. et Robins, J. M. (2020). *Causal Inference : what if*. CRC press.
- Hill, J. et Stuart, E. (2015). Causal inference : overview. *International encyclopedia of the social & behavioral sciences : second edition*, 255–260.
- Hodgkin, J. (2002). Seven types of pleiotropy. *International journal of developmental biology*, 42(3), 501–505.
- Hole, A. R. (2007). A comparison of approaches to estimating confidence intervals for willingness to pay measures. *Health economics*, 16(8), 827–840.
- Holland, P. W. (1986). Statistics and causal inference. *Journal of the american statistical association*, 81(396), 945–960.
- Judd, C. M. et Kenny, D. A. (1981). Process analysis : estimating mediation in treatment evaluations. *Evaluation review*, 5(5), 602–619.
- Judd, C. M. et Kenny, D. A. (2010). Data analysis in social psychology : recent and recurring issues. *Handbook of social psychology*, 1, 115–139.
- Kaufman, J. S., MacLehose, R. F. et Kaufman, S. (2004). A further critique of the analytic strategy of adjusting for covariates to identify biologic mediation. *Epidemiologic perspectives & innovations*, 1(1), 1–13.

- Kenny, D. A., Kashy, D. A., Bolger, N., Gilbert, D., Fiske, S. T. et Lindzey, G. (1998). The handbook of social psychology. *Data analysis in social psychology*, 233–265.
- Lawlor, D. A., Harbord, R. M., Sterne, J. A., Timpson, N. et Davey Smith, G. (2008). Mendelian randomization : using genes as instruments for making causal inferences in epidemiology. *Statistics in medicine*, 27(8), 1133–1163.
- Lousdal, M. L. (2018). An introduction to instrumental variable assumptions, validation and estimation. *Emerging themes in epidemiology*, 15(1), 1–7.
- Ludl, A.-A. et Michoel, T. (2021). Comparison between instrumental variable and mediation-based methods for reconstructing causal gene networks in yeast. *Molecular omics*, 17(2), 241–251.
- MacKinnon, D. P., Cheong, J. et Pirlott, A. G. (2012). *Statistical mediation analysis*. American psychological association.
- MacKinnon, D. P., Lockwood, C. M. et Williams, J. (2004). Confidence limits for the indirect effect : distribution of the product and resampling methods. *Multivariate behavioral research*, 39(1), 99–128.
- MacKinnon, D. P. et Luecken, L. J. (2008). How and for whom ? Mediation and moderation in health psychology. *Health psychology*, 27(2S), S99.
- MacKinnon, D. P., Warsi, G. et Dwyer, J. H. (1995). A simulation study of mediated effect measures. *Multivariate behavioral research*, 30(1), 41–62.
- Nuttall, F. Q. (2015). Body mass index : obesity, bmi, and health : a critical review. *Nutrition today*, 50(3), 117.
- Pearl, J. (2001). *Direct and indirect effects*. Morgan Kaufmann publishers Inc.
- Pearl, J. (2013). Direct and indirect effects. *arXiv preprint arXiv :1301.2300*.
- Raina, P., Wolfson, C., Kirkland, S., Griffith, L. E., Balion, C., Cossette, B., Dionne, I., Hofer, S., Hogan, D., Van Den Heuvel, E. et al. (2019). Cohort profile : the canadian longitudinal study on aging (clsa). *International journal of epidemiology*, 48(6), 1752–1753j.
- Relton, C. L. et Davey Smith, G. (2012). Two-step epigenetic Mendelian randomization : a strategy for establishing the causal role of epigenetic processes in pathways to disease. *International journal of epidemiology*, 41(1), 161–176.
- Richiardi, L., Bellocco, R. et Zugna, D. (2013). Mediation analysis in epidemiology : methods, interpretation and bias. *International journal of*

epidemiology, 42(5), 1511–1519.

Rizzo, R., Cashin, A. G., Bagg, M. K., Gustin, S. M., Lee, H. et McAuley, J. H. (2022). A systematic review of the reporting quality of observational studies that use mediation analyses. *Prevention science*, 1–12.

Robins, J. M. et Greenland, S. (1992). Identifiability and exchangeability for direct and indirect effects. *Epidemiology*, 143–155.

Rosenbaum, P. R. et Rubin, D. B. (1983). The central role of the propensity score in observational studies for causal effects. *Biometrika*, 70(1), 41–55.

Sanderson, E. (2021). Multivariable Mendelian randomization and mediation. *Cold spring harbor perspectives in medicine*, 11(2), 1–18.

Savin, N. (1990). Two-stage least squares and the k-class estimator. In *Econometrics* 265–275. Springer.

Smith, G. D. et Ebrahim, S. (2004). Mendelian randomization : prospects, potentials, and limitations. *International journal of epidemiology*, 33(1), 30–42.

Sovey, A. J. et Green, D. P. (2011). Instrumental variables estimation in political science : a readers' guide. *American journal of political science*, 55(1), 188–200.

Sproston, N., Ashworth, J. *et al.* (2018). Role of c-reactive protein at sites of inflammation and infection. *front immunol.* 2018 ; 9 : 754. URL : <https://www.frontiersin.org/articles/10.3389/fimmu>.

Teumer, A. (2018). Common methods for performing Mendelian randomization. *Frontiers in cardiovascular medicine*, 5, 51.

VanderWeele, T. (2015). *Explanation in causal inference : methods for mediation and interaction*. Oxford university press.

VanderWeele, T., Valeri, L. et Ogburn, E. (2012). The role of misclassification and measurement error in mediation analyses. *Epidemiology*, 23, 561–564.

VanderWeele, T. J. (2011). Controlled direct and mediated effects : definition, identification and bounds. *Scandinavian journal of statistics*, 38(3), 551–563.

VanderWeele, T. J. (2016). Mediation analysis : a practitioner's guide. *Annual review of public health*, 37(1), 17–32.

VanderWeele, T. J. et Vansteelandt, S. (2009). Conceptual issues concerning mediation, interventions and composition. *Statistics and its interface*, 2(4),

457–468.

VanderWeele, T. J., Vansteelandt, S. et Robins, J. M. (2014). Effect decomposition in the presence of an exposure-induced mediator-outcome confounder. *Epidemiology*, 25(2), 300–306.

Vegte, Y. J., Said, M. A., Rienstra, M., Harst, P., Verweij, N. *et al.* (2020). Genome-wide association studies and Mendelian randomization analyses for leisure sedentary behaviours. *Nature communications*, 11(1), 1–10.

Verbanck, M., Chen, C.-y., Neale, B. et Do, R. (2018). Detection of widespread horizontal pleiotropy in causal relationships inferred from Mendelian randomization between complex traits and diseases. *Nature genetics*, 50(5), 693–698.

Wald, A. et Wolfowitz, J. (1940). Annals of mathematics. *Stat*, 11, 147.

Wardle, J., Carnell, S., Haworth, C. M., Farooqi, I. S., O’Rahilly, S. et Plomin, R. (2008). Obesity associated genetic variation in *fto* is associated with diminished satiety. *The journal of clinical endocrinology & metabolism*, 93(9), 3640–3643.

Woodworth, R. S. (1928). *How emotions are identified and classified*. Clark university press.

Wu, F., Huang, Y., Hu, J. et Shao, Z. (2020). Mendelian randomization study of inflammatory bowel disease and bone mineral density. *BioMed central medicine*, 18(1), 1–19.

Zhang, Z., Uddin, M. J., Cheng, J. et Huang, T. (2018). Instrumental variable analysis in the presence of unmeasured confounding. *Annals of translational medicine*, 6(10).

Zhao, X., Lynch Jr, J. G. et Chen, Q. (2010). Reconsidering Baron and Kenny : myths and truths about mediation analysis. *Journal of consumer research*, 37(2), 197–206.