

DIGITALES ARCHIV

ZBW – Leibniz-Informationszentrum Wirtschaft
ZBW – Leibniz Information Centre for Economics

Fourrey, Kévin

Thesis

Décomposition des indices d'inégalité et impact des politiques publiques

Provided in Cooperation with:

ZBW OAS

Reference: Fourrey, Kévin (2019). Décomposition des indices d'inégalité et impact des politiques publiques. Caen.

This Version is available at:

<http://hdl.handle.net/11159/3587>

Kontakt/Contact

ZBW – Leibniz-Informationszentrum Wirtschaft/Leibniz Information Centre for Economics
Düsternbrooker Weg 120
24105 Kiel (Germany)
E-Mail: [rights\[at\]zbw.eu](mailto:rights[at]zbw.eu)
<https://www.zbw.eu/>

Standard-Nutzungsbedingungen:

Dieses Dokument darf zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden. Sie dürfen dieses Dokument nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, aufführen, vertreiben oder anderweitig nutzen. Sofern für das Dokument eine Open-Content-Lizenz verwendet wurde, so gelten abweichend von diesen Nutzungsbedingungen die in der Lizenz gewährten Nutzungsrechte.

<https://savearchive.zbw.eu/termsfuse>

Terms of use:

This document may be saved and copied for your personal and scholarly purposes. You are not to copy it for public or commercial purposes, to exhibit the document in public, to perform, distribute or otherwise use the document in public. If the document is made available under a Creative Commons Licence you may exercise further usage rights as specified in the licence.

Décomposition des indices d'inégalité et impact des politiques publiques

Kevin Fourrey

► To cite this version:

Kevin Fourrey. Décomposition des indices d'inégalité et impact des politiques publiques. Economies et finances. Normandie Université, 2019. Français. NNT : 2019NORMC019 . tel-02375678

HAL Id: tel-02375678

<https://tel.archives-ouvertes.fr/tel-02375678>

Submitted on 22 Nov 2019

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

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



Normandie Université

THÈSE

Pour obtenir le diplôme de doctorat

Spécialité SCIENCES ECONOMIQUES

Préparée au sein de l'Université de Caen Normandie

Décomposition des indices d'inégalité et impact des politiques publiques

**Présentée et soutenue par
Kevin FOURREY**

**Thèse soutenue publiquement le 30/09/2019
devant le jury composé de**

M. STEPHANE MUSSARD	Professeur des universités, Université de Nîmes	Rapporteur du jury
Mme AGNIESZKA RUSINOWSKA	Directeur de recherche au CNRS, Université Paris 1 Panthéon-Sorbonne	Rapporteur du jury
M. NICOLAS GRAVEL	Professeur des universités, Aix-Marseille Université	Membre du jury
Mme DOMINIQUE MEURS	Professeur des universités, Université Paris-Nanterre	Membre du jury
M. FREDERIC CHANTREUIL	Maître de conférences HDR, Université Caen Normandie	Directeur de thèse
Mme ISABELLE LEBON	Professeur des universités, Université Caen Normandie	Co-directeur de thèse

Thèse dirigée par FREDERIC CHANTREUIL et ISABELLE LEBON, Centre de recherche en économie et management (Rennes)



UNIVERSITÉ
CAEN
NORMANDIE



À ma regrettée grand-mère, Yvonne

À ma douce nièce Alice et mon tendre filleul Nolan

Remerciements

Je tiens en premier lieu à remercier très chaleureusement mes directeurs de thèse, Isabelle Lebon et Frédéric Chantreuil, qui ont été d'un soutien sans faille tout du long de ma thèse. Je les remercie pour tout le temps qu'ils m'ont accordé et les nombreux moments d'échanges enrichissants que nous avons eu le plaisir de partager. Cette thèse n'aurait pu se faire sans leurs encouragements et leur bienveillance. Je tiens à remercier très cordialement tous les membres du jury de ma soutenance de thèse, la Pr. Dominique Meurs qui préside cette prestigieuse assemblée, Agnieszka Rusinowska et le Pr. Stéphane Mussard qui ont pris le temps de faire des rapports et Pr. Nicolas Gravel qui a accepté de bon gré de discuter de ma thèse.

Je remercie grandement Vincent Merlin, directeur de l'école doctorale Économie-Gestion Normandie, qui m'a notamment permis de réaliser cette thèse dans de bonnes conditions. Je suis reconnaissant également du soutien apporté par Sébastien Courtin, en tant que membre de mon comité de suivi individuel puis en tant que co-auteur. Par la même occasion, je remercie Clémence Chrisitn et Olivier Dagnelie pour leurs conseils aiguisés et pour avoir accepté de me suivre au cours de cette thèse en tant que membres de mon comité de suivi.

La valorisation de mon travail grâce à la publication d'articles et de présentations à des séminaires n'auraient pu se faire sans le concours financier et l'appui technique du laboratoire CREM et de la MRSH, je remercie notamment Anne-Marie et ma tendre Carole. Par ailleurs, je tiens à montrer ma reconnaissance à l'égard de l'UFR SEGGAT, l'UFR des Sciences et l'IUT de l'Université Caen Normandie pour m'avoir permis d'enseigner au sein de leur composante.

Tous mes remerciements vont également à mes camarades doctorants, en particulier vers Malia qui m'a accompagné tout au long de ces années passées en thèse. Je remercie Charline pour m'avoir donné l'occasion d'être organisateur de deux événements de médiation scientifique, ainsi que de m'avoir accompagné à la rédaction d'un bilan sur la situation des doctorants normands. Je salue à leur tour mes collègues économistes : Eva, Étienne et mon futur co-auteur Rodrigue, et mes autres collègues de sciences humaines et sociales : Vika, Vicky, Annie, Hang et Min. Leurs conseils et leur aide ont toujours été les bienvenus.

Je suis également reconnaissant du soutien apporté par mes amis, et notamment la plus fidèle d'entre-eux, Claire. Ainsi que Florian, Maé, Charles, Léa, Katy, Ketty, Cécile, Alexis, Jérémy,

Mathilde, Lucile, Sophie, Nina, Simon, Nathan, et tout ceux dont il m'est donné la chance d'être ami. Sans oublier les membres de la troupe de théâtre "Sur les planches en Suisse Normande" sur qui je peux compter à tout instant.

Naturellement, je remercie ma mère Marie-Claude, mon père Michel, mes sœurs Amélie et Noémie, mon frère Jean-Baptiste, mon beau-frère Gaëtan et ma belle-sœur Noémie pour leur soutien quotidien pendant ces longues années d'études.

Table des matières

Remerciements	3
Introduction générale	7
1 Les indices d'inégalité	11
1.1 Les propriétés des indices d'inégalité	12
1.2 La variance	15
1.3 La courbe de Lorenz et l'indice de Gini	16
1.3.1 La courbe de Lorenz	16
1.3.2 L'indice de Gini	17
1.4 Les indices d'inégalité de Theil	21
1.4.1 L'indice de Theil	21
1.4.2 L'indice second de Theil	22
1.4.3 Les indices généralisés de Theil	23
1.5 Les indices d'Atkinson	25
1.5.1 Une approche normative	25
1.5.2 Les indices d'Atkinson	26
1.6 Fonction d'Évaluation Sociale	29
1.7 Discussion et conclusion	31
2 Les décompositions des indices d'inégalité	35
2.1 Introduction	36
2.2 La décomposition par sous-populations	37
2.2.1 Décomposition agrégative, additive et cohérente	37
2.2.2 Décomposition additive et cohérente des indices d'entropie	38
2.2.3 Décomposition additive de l'indice de Gini	40
2.2.4 La décomposition faible des α – Gini	43
2.3 La décomposition par facteurs	46

2.3.1	Décomposition d'un indice d'inégalité par facteurs	46
2.3.2	La décomposition par facteurs de l'indice de Gini	48
2.3.3	La décomposition tridimensionnelle de la variation de l'indice de Gini . . .	48
2.4	La décomposition multiple basée sur des méthodes statistiques : les régressions RIF	49
2.4.1	Effet de composition et effet structurel	50
2.4.2	Les régressions RIF	57
2.5	La décomposition de Shapley des indices d'inégalité	61
2.5.1	Introduction au concept de Shapley	62
2.5.2	Décomposition d'un modèle hiérarchique	64
2.5.3	Normalisation des distributions de facteurs équivalents	69
2.5.4	Incidence de l'indice d'inégalité sur la décomposition	71
2.6	Discussion et conclusion	72
3	L'importance du genre à l'inégalité salariale dans la fonction publique : une analyse régionale	75
3.1	Introduction	76
3.2	Mesure de l'intensité des inégalités liées au genre dans la fonction publique, des disparités importantes entre les régions	79
3.2.1	Méthodologie et données	79
3.2.2	Comparaison interrégionale de l'intensité des inégalités liées au genre sur l'ensemble de la fonction publique	81
3.3	Quel lien avec les spécificités régionales?	83
3.3.1	Spécificités dans la structure de la fonction publique	83
3.3.2	Un parallèle avec des caractéristiques socio-démographiques et le secteur privé	87
3.4	Intensité régionale des inégalités liées au genre par catégorie et par versant	89
3.5	Discussion et conclusion	97
3.6	Annexes	99
4	Contribution marginale pure, interactions et l'importance d'un attribut	103
4.1	Introduction	104
4.2	Jeu inégalitaire	105
4.3	L'importance d'un attribut	107
4.3.1	Contribution marginale	107
4.3.2	Les interactions entre les attributs	108
4.4	Discussion et conclusion	112
4.5	Annexe	113

5	Mesurer l'impact d'une politique égalisatrice sur l'inégalité : une application à l'écart de rémunération entre les femmes et les hommes en France	119
5.1	Introduction	120
5.2	Les attributs, leur importance et leur impact	121
5.2.1	Les distributions associées aux attributs	121
5.2.2	L'importance de l'attribut genre	122
5.2.3	L'impact d'une politique égalisatrice efficace	124
5.3	L'impact de la suppression de l'écart de salaire entre les hommes et les femmes, une application sur données françaises	128
5.3.1	Les données	128
5.3.2	L'importance de l'attribut genre	129
5.3.3	L'impact du genre	130
5.3.4	L'interprétation des interactions	130
5.3.5	Un effet de composition sur l'importance du genre	132
5.4	Discussion et conclusion	134
5.5	Annexe : Distribution des observations	135
	Conclusion générale	139
	Bibliographie	143

Introduction Générale

Une société tend vers la prospérité lorsqu'elle permet à ses membres de vivre avec un niveau de bien-être qu'ils estiment satisfaisant. Il peut être aisément admis que le bien-être d'un individu dépend en partie de la richesse qu'il détient. Que penser alors d'une société où des individus sont pauvres et d'autres sont riches ? Ces inégalités de richesses peuvent-elles être acceptées par la société et ne pas être un frein à sa prospérité ? Pour pouvoir discuter de l'acceptation des inégalités, il nous faut d'abord introduire le concept de justice sociale. La notion de justice permet de différencier les inégalités justes des inégalités injustes, et donc inéquitables. Cependant, la différenciation entre ce qui est juste de ce qui est injuste dépend du contrat social partagé entre les membres d'une communauté. Ainsi, il est nécessaire de considérer les différentes théories de la justice sociale, dont chacune regroupe « l'ensemble des principes qui régissent la définition et la répartition équitable des droits et des devoirs entre les membres de la société » (Arnsperger et Van Parijs (2000)). Ces mêmes auteurs définissent les « quatre points cardinaux » des théories de l'éthique économique et sociale à partir desquelles élaborer des principes caractérisant des institutions justes : l'utilitarisme, le libéralisme, le marxisme et l'égalitarisme libéral.

L'utilitarisme (Bentham (1789)) correspond à la recherche du plus grand bonheur pour le plus grand nombre. Il est régi selon trois principes : le conséquentialisme (les institutions sont jugées sur le résultat), le welfarisme (le résultat est apprécié en bien-être) et l'universalisme (tous les individus ont le même poids). L'objectif utilitariste est de maximiser l'utilité totale, quelle que soit sa distribution entre les membres de la société, sachant qu'aucun intérêt individuel ne peut l'emporter sur l'intérêt collectif. Autrement dit, une distribution des richesses selon le critère utilitariste est juste si elle permet une maximisation de l'utilité totale.

Cette philosophie se différencie du libéralisme (Locke (1690), Hayek (1960)) selon lequel une société juste est une société libre. Ainsi « la dignité fondamentale de chaque individu réside dans la liberté de choix qui ne peut être bafouée au nom d'un impératif collectif » (Michaud (2006) p.157). Par conséquent, peu importe la distribution des richesses, si cette distribution est le résultat de choix individuels libres et non-contraints alors elle est juste. Tandis que l'utilitarisme développe une notion de justice d'une distribution de richesses par rapport au résultat, le libéralisme l'apprécie selon les règles qui déterminent cette distribution.

Pour les tenants de l'économie marxiste (Marx (1867), Roemer (1988)), la valeur d'un bien peut être déterminée par la quantité de travail nécessaire à sa production. Tout en admettant que deux biens peuvent avoir des valeurs différentes même s'ils ont exigé la même quotité horaire de travail. Cela peut notamment s'expliquer par la différence dans la qualité du travail utilisé pour la production de ce bien, en termes d'intensité ou de savoir-faire. Quoiqu'il en soit, la différence entre la valeur ajoutée d'un bien et le montant de ce bien versé au travailleur correspond à la plus-value accaparée par la « classe bourgeoise » (propriétaire des moyens de productions). Cependant, le marxisme exige que tous les membres d'une société soient égaux, dans le sens où aucune personne ne domine une autre. Ainsi, une société égalitaire est une société sans classe sociale (ou une classe unique : les travailleurs). La distribution des richesses est alors jugée juste si chacun est libre d'exercer le métier qu'il souhaite et si la totalité de la valeur créée est distribuée aux travailleurs, l'unique classe présente dans la société.

L'égalitarisme libéral (Rawls (1971) ; Sen (1985, 1992)) fait la synthèse des idéaux de liberté et d'égalité. Dans cette optique, John Rawls définit des biens premiers, « des biens utiles quel que soit le projet de vie rationnel » (Rawls (1971)). Parmi ces biens premiers, il distingue les biens premiers naturels (santé/talents) des biens premiers sociaux (libertés et droits fondamentaux). Ces biens premiers sont déterminés à partir d'une situation fictive du voile d'ignorance, où chaque individu est amené à définir les règles de la société sans connaître sa position dans la hiérarchie sociale. Dans cette situation, les individus définissent un système social sans savoir s'ils en retireront un bénéfice, leur jugement est alors déterminé par une exigence d'impartialité et d'équité. Ainsi il en ressort un contrat social qui est composé d'une égale liberté pour tous, d'une égalité des chances et du principe de différence. Ce dernier justifiant les inégalités sur les biens premiers sociaux si elles procurent le plus grand bénéfice aux membres les plus désavantagés de la société (discrimination positive). Selon Amartya Sen, dans la continuité de cette théorie, la justice exige que les inégalités sur les biens premiers naturels soient aussi en partie corrigées pour atteindre l'égalité de certaines « capacités », c'est-à-dire une égalité dans le choix du mode de vie qu'un individu souhaite mener. En conséquence, selon la philosophie égalitariste libérale, une distribution de richesses au sein d'une société est injuste si les plus défavorisés étaient contraints dans leur développement de leur talent pour pouvoir mener la vie qu'ils souhaitent.

Cette approche par les « capacités » met en avant que les inégalités ne sont pas seulement monétaires, mais elles peuvent aussi être appréciées en termes d'accès à l'éducation, la formation professionnelle, la santé, l'accès au service public... Au sein de cette thèse, l'inégalité étudiée se mesure avant-tout en termes de rémunérations. Cependant, il faut garder à l'esprit que les inégalités sont multiples et interdépendantes. Par exemple, l'inégalité spatiale de revenus au sein d'une métropole peut entraîner des différences de taux de taxation local entre les différentes localités qui la composent (Fourrey (2018)), et par conséquent cela peut entraîner une disparité d'offre de services

publics avec des localités pauvres offrant moins de services que les autres, ce qui se traduirait par une remise en cause de l'égalité des chances. Il en résulterait, dans une logique égalitaire libérale, une distribution spatiale des revenus injuste.

Quel que soit le contrat social en vigueur dans la société, les travaux inclus dans cette thèse entendent apporter des éclaircissements et des avancées dans l'analyse de la distribution des richesses. Ils développent notamment une réflexion sur la mesure de l'inégalité et les différentes approches qui permettent d'en déterminer les causes. De fait, lorsque la question de la distribution des richesses est traitée, il est important pour le décideur public de connaître la contribution des différents facteurs aux inégalités pour qu'il puisse mettre en place des politiques efficaces en faveur d'une plus grande égalité, si tel est le souhait de la société. Ainsi, l'originalité première de cette thèse est de concevoir un outil qui permette de calculer l'importance des différents facteurs dans l'inégalité et de prendre en considération les interactions entre les différents facteurs. Grâce à la méthodologie développée, les résultats escomptés d'une politique publique égalisatrice peuvent être appréciés, que ce soit en termes de conséquences sur le niveau des inégalités, mais aussi sur les raisons qui expliquent ce résultat.

Au premier abord, les notions de justice et d'équité de la distribution des richesses dans une société sont abordées après la mesure de l'inégalité. Pour autant, le premier chapitre montre que l'outil principal qui donne une indication sur la concentration des richesses, c'est-à-dire les indices d'inégalité, sont emprunts de subjectivité. Dans ce chapitre, inspiré notamment par l'ouvrage de Villar (2017), les principes sur lesquels les indices d'inégalité s'appuient sont étudiés afin de mieux apprécier leurs résultats, et les principaux indices utilisés dans la littérature sont exposés. La discussion de ce chapitre permet aussi de faire le lien entre des concepts de justice, d'équité et d'inégalité.

Pour la mise en place d'une politique publique qui viserait à réduire les inégalités de revenus, il est très intéressant pour les décideurs politiques de connaître les groupes sociaux les plus touchés par les inégalités, ou encore les facteurs les plus créateurs ou réducteurs d'inégalités. C'est pourquoi le second chapitre aborde la décomposition des indices d'inégalité. Dans un premier temps, les décompositions par sous-populations sont présentées, puis celles qui permettent de répartir les inégalités entre sources de revenus (ex. : salaire, prime, allocations...). Ces deux approches peuvent parfois être combinées, ces décompositions sont alors dites multiples. Au sein de ce chapitre, une troisième section présente une méthode de décomposition par attribut qui repose sur un outil statistique et qui permet de répartir les inégalités de revenus entre les différentes caractéristiques personnelles. Une autre décomposition par attribut, et qui peut être aussi une décomposition multiple, est présentée dans la dernière section. Cette approche repose sur un concept de solution de la théorie des jeux, la célèbre valeur de Shapley (1953). La décomposition à la Shapley est l'objet principal de cette thèse, c'est une méthode qui combine de nombreux avantages : elle peut être multiple, elle

peut considérer plusieurs effets de groupes en même temps et elle n'est pas spécifique à un indice d'inégalité.

Le chapitre 3 applique la décomposition à la Shapley pour évaluer la contribution de différentes caractéristiques personnelles aux inégalités de salaire dans la fonction publique française en 2010. Cette étude a donné lieu à une publication dans le *Revue d'économie politique*, co-écrite avec Isabelle Lebon et Frédéric Chantreuil (Chantreuil *et al.* (2018)), et elle a bénéficié du financement de la Direction Générale de l'Administration et de la Fonction Publique et du Défenseur des Droits (convention n°2200730964), ainsi qu'une aide de l'Etat gérée par l'Agence nationale de la Recherche au titre du programme Investissements d'avenir portant la référence ANR-10-EQPX-17. Ce projet a permis de mettre en avant les différences dans la contribution du genre aux inégalités rémunérations dans la fonction publique selon les différentes régions françaises établie à l'époque, par catégorie (*i.e.* niveau hiérarchique) et par versant (*i.e.* secteur de la fonction publique). De ce travail de recherche, il ressort que les femmes sont moins payées que les hommes, dans toutes les régions françaises, quelle que soit la catégorie ou la composante de la fonction publique considérée. Et que, de plus, il existe des différences interrégionales notables.

Cependant, une question reste en suspens, quel serait le niveau d'inégalité observé s'il n'y avait plus de disparités de salaires entre les hommes et les femmes, à compétences et expériences données ? Les chapitres 4 et 5 s'attellent à cette question. Le chapitre 4, avec une approche théorique, entend démontrer que l'importance d'une caractéristique individuelle dans la création d'inégalités dépend d'une part de sa contribution marginale pure, c'est-à-dire la différence entre l'inégalité observée lorsque toutes les caractéristiques sont prises en compte et l'inégalité observée avec ces mêmes caractéristiques à l'exception du genre, et dépend d'autre part de ces interactions avec les autres caractéristiques individuelles. Dans ce chapitre, une nouvelle formulation de la valeur de Shapley est fournie, elle exprime l'importance d'un attribut comme la différence de sa contribution marginale pure avec ses interactions par paire, positives ou négatives, avec les autres attributs. Ce chapitre est issu d'une recherche commune avec Sébastien Courtin, Frédéric Chantreuil et Isabelle Lebon, et fait l'objet d'un article publié prochainement dans la revue *Social Choice and Welfare* (Chantreuil *et al.* (2019)). Par une étude empirique examinant la variation des importances des différentes caractéristiques individuelles entre deux périodes, le chapitre 5 confirme la décomposition possible de l'importance entre la contribution marginale pure et une somme pondérée d'interactions. Ce chapitre 5, issu d'un travail commun avec Isabelle Lebon et Frédéric Chantreuil, permet notamment de mettre en lumière que même si l'importance du genre aux inégalités de salaires en France en 2011, mesurée par l'indice de Gini, est de 6,98%, donner le même salaire aux hommes et aux femmes à âge et niveau d'éducation égaux diminuerait les inégalités de seulement 1,09%. La différence entre l'importance et l'impact du genre sur les inégalités étant due aux interactions du genre avec les autres caractéristiques individuelles (âge et niveau de diplôme notamment).

Chapitre 1

Les indices d'inégalité

Sommaire

1.1	Les propriétés des indices d'inégalité	12
1.2	La variance	15
1.3	La courbe de Lorenz et l'indice de Gini	16
1.3.1	La courbe de Lorenz	16
1.3.2	L'indice de Gini	17
1.4	Les indices d'inégalité de Theil	21
1.4.1	L'indice de Theil	21
1.4.2	L'indice second de Theil	22
1.4.3	Les indices généralisés de Theil	23
1.5	Les indices d'Atkinson	25
1.5.1	Une approche normative	25
1.5.2	Les indices d'Atkinson	26
1.6	Fonction d'Évaluation Sociale	29
1.7	Discussion et conclusion	31

1.1 Les propriétés des indices d'inégalité

Avant de décomposer les indices d'inégalité, il semble opportun d'étudier l'objet en question. Ainsi, cette partie s'attache à présenter les principaux indices d'inégalité utilisés dans la littérature. Cette synthèse s'inspire notamment de l'ouvrage de Villar (2017).

Un indice d'inégalité est une fonction I qui attribue à une distribution de revenus $\mathbf{y}$ pour n individus, tel que $\mathbf{y} \in \mathfrak{D} \equiv \cup_{n=1}^{\infty} \mathbb{R}^n$, un réel non négatif donnant une indication sur la concentration des revenus dans la population étudiée, ainsi $I : \mathbf{y} \rightarrow \mathbb{R}_+$. Les indices d'inégalité diffèrent dans leur construction et le poids attaché aux différentes parties de la population. Depuis plusieurs décennies, plusieurs indices synthétiques ont été proposés pour apprécier le niveau des inégalités observés dans les sociétés. D'après les travaux de Dalton (1920), Kolm (1976a,b), Shorrocks (1980) et Chakravarty (1999), quelques propriétés devraient être respectés par les indices d'inégalité :

— Normalisation :

L'indice d'inégalité est nul pour une distribution parfaitement égalitaire, sinon il est positif.

Ainsi : $I(k\mathbf{1}^n) = 0, \forall k > 0$, avec $\mathbf{1}^n$ le vecteur unité de dimension n .

— Symétrie (ou Anonymat) :

Seul le revenu des individus compte, une permutation des revenus entre individus ne doit pas faire varier la mesure des inégalités.

Si $\mathbf{y}_b$ est obtenu à partir de $\mathbf{y}_a$ par permutation de revenus entre individus alors l'inégalité mesurée ne varie pas : $I(\mathbf{y}_a) = I(\mathbf{y}_b)$.

— Invariance à la réplication¹ :

Si la distribution de revenus dans deux groupes distincts est exactement la même, alors l'indice d'inégalité mesuré dans chaque groupe doit être égal à l'indice d'inégalité mesuré sur l'ensemble de la population (*i.e.* les deux groupes fusionnés). Ainsi, l'indice d'inégalité ne doit dépendre que des fréquences relatives des distributions de revenus.

Pour une distribution $\mathbf{y} = \{y_1, y_2, \dots, y_n\}$ que l'on réplique k fois :

$$\mathbf{y} = \underbrace{\{y_1, \dots, y_1\}}_{k \text{ répliques}}, \underbrace{\{y_n, \dots, y_n\}}_{k \text{ répliques}}$$

l'indice d'inégalité reste inchangé :

$$I^k(\mathbf{y}) = I(\mathbf{y}), \forall k \geq 2$$

1. Ou principe de population de Dalton

— **Principe de transfert :**²

Suivant une proposition émise par Pigou (1912), Dalton (1920) propose une nouvelle condition aux indices d'inégalité : le principe de transfert. Un transfert de Pigou-Dalton est un transfert d'un « riche » à un « pauvre » de telle sorte que leur position relative dans la hiérarchie des revenus ne change pas. Un indice d'inégalité respecte le principe de transfert s'il diminue suite à un transfert de Pigou-Dalton.

Formellement, $\mathbf{y}_a$ et $\mathbf{y}_b$ sont deux distributions de revenus ordonnées dans l'ordre croissant, sachant que $\mathbf{y}_b$ est obtenue à partir de $\mathbf{y}_a$ après un transfert de Pigou-Dalton d'un individu j vers un individu i . Ainsi, si l'indice d'inégalité respecte le principe de transfert, $I(\mathbf{y}_b) \leq I(\mathbf{y}_a)$, tels que $y_j^b - y_i^b \leq y_j^a - y_i^a$ et $r_i^p \leq r_j^p \forall p \in \{a, b\}$, avec y_i^p et r_i^p le revenu de l'individu i et la position de l'individu i respectivement dans la distribution $\mathbf{y}_p$.

— **Continuité :**

La continuité exige qu'une petite variation des revenus individuels entraîne une petite variation de l'indice. Autrement dit, deux distributions quasi-similaires auront deux indices d'inégalité proches.

— **Invariance relative (Homogénéité de degré 0) :**

Les indices d'inégalité sont des fonctions homogènes de degré zéro. Ainsi, multiplié le revenu de chaque individu de la population par une constante (ex. un taux de change) ne modifie pas le niveau des inégalités.

Formellement, $I(\lambda \mathbf{y}) = I(\mathbf{y}), \forall \lambda > 0$.

Les indices d'inégalité qui respectent la normalisation, la symétrie, l'invariance à la réplication, le principe de transferts et la continuité sont appelé des indices réguliers. Si de plus, ils respectent le principe d'invariance relative alors ils font partis de la classe des indices réguliers relatifs. Les indices d'inégalité qui ne sont pas relatifs peuvent être linéairement homogènes (absolus) ou homogène de degré un :

— **Homogénéité linéaire :**

Un indice d'inégalité est linéairement homogène si l'ajout ou le retrait d'un même montant à tous les individus ne modifie pas le niveau des inégalités.

Formellement, un indice d'inégalité est absolu si pour toutes distributions de revenus $\mathbf{y}$, $\mathbf{y} \in \mathbb{D}_+$ tel que $\mathbf{x} = (\epsilon, \epsilon, \dots, \epsilon)$, avec $I(\mathbf{y} + \mathbf{x}) = I(\mathbf{y})$.

2. Le principe de transfert fait l'objet de nombreuses discussions (voir par exemple la thèse de Dubois (2016) pour une présentation exhaustive de ce sujet). Par ailleurs, le principe de transfert est ici considéré dans l'étude de l'inégalité de revenu. Cependant, quand l'objet étudié est autre, ce principe de transfert peut ne plus être adapté (voir par exemple : Fleurbaey et Michel (2001), Bazen et Moyes (2012), Moyes (2012) et Gravel *et al.* (2014)).

— Homogénéité de degré 1 :

Un indice est homogène de degré un si la multiplication de l'ensemble des revenus par une constante conduit à multiplier l'inégalité mesurée par cette même constante.

Formellement, un indice d'inégalité est homogène de degré un si pour toutes distributions $\mathbf{y} \in \mathcal{D}_+$ et $\lambda > 0$, $I(\lambda\mathbf{y}) = \lambda I(\mathbf{y})$.

Un indice d'homogénéité linéaire n'est pas adapté pour étudier les inégalités de revenus puisqu'il induit qu'une société à deux individus, où l'un a un revenu de 100€ et l'autre de 1000€, est aussi inégalitaire que s'ils avaient respectivement 100 100€ et 101 000€, alors que dans la différence de richesse est beaucoup plus faible dans le second cas. Un indice d'homogénéité de degré 1 peut avoir un intérêt lorsque le « volume des inégalités » est pris en considération, c'est-à-dire pour faire une différence entre une situation où un individu gagne 100€ et l'autre 1 000€ d'une situation où l'un gagne 1 000€ et l'autre 10 000€. Dans les deux cas de l'exemple précédent, le second individu a un revenu 10 fois plus important que le premier, l'indice relatif estimera donc que les deux situations de cette société à deux individus est similaire d'un point de vue inégalitaire. L'avantage considérable d'un indice relatif est qu'il permet directement de comparer différentes sociétés dans le temps ou à un temps donné, puisqu'il n'est pas sensible aux différents échelles monétaires entre ces sociétés. Un indice d'homogénéité de degré 1 peut aussi le faire, mais après avoir transformé les revenus de telle sorte qu'ils soient comparables dans le temps ou entre différentes sociétés à un temps donné. Cependant, dans ce dernier cas, l'indice sera sensible à l'année de base choisie pour comparer des revenus dans le temps.

Un autre principe a été mis en avant par Cowell et Flachaire (2018), celui de la monotonie à la distance.

— Monotonie à la distance :

Un indice d'inégalité respecte le principe de monotonie à la distance si l'inégalité mesurée augmente après que le revenu d'un individu s'écarte du point de référence (*i.e.* moyenne, médiane...), toutes choses égales par ailleurs.

Formellement, avec deux distributions $\mathbf{y}_a, \mathbf{y}_b \in \mathcal{D}_+$ identiques à l'exception d'une variation du revenu d'un individu i tel que $|y_{a_i} - \bar{y}_a| < |y_{b_i} - \bar{y}_b|$, où $\bar{y}_a$ et $\bar{y}_b$ représentent le point de référence respectif à chaque distribution, alors $I(\mathbf{y}_a) < I(\mathbf{y}_b)$.

Comme l'ont mis en avant Cowell et Flachaire (2018), la majorité des indices d'inégalité utilisés dans la littérature sont construits tel qu'au numérateur se trouve un indice de dispersion et au dénominateur la moyenne. Cette construction ne permet pas d'assurer le principe de monotonie à la distance car si un riche (par rapport à la moyenne) devient plus riche alors le numérateur et le dénominateur augmente, ce qui peut potentiellement conduire à une baisse de l'indice d'inégalité.

Ce principe est très intéressant dès lors que l'évolution des inégalités dans le temps est étudiée puisqu'un indice d'inégalité qui ne respecte pas ce principe pourrait donner une image faussée de l'évolution des inégalités³. Par exemple, dans le cas d'un indice basé sur la moyenne, si la classe moyenne riche a un revenu plus important alors le revenu moyen augmente et la classe moyenne pauvre devient plus pauvre et contient potentiellement plus d'individus, il paraît donc justifié de dire que les inégalités augmentent, cependant un indice qui ne respecte pas l'axiome de monotonie à la distance peut conduire à la conclusion inverse.

Avant de présenter les principaux indices d'inégalité réguliers utilisés dans la littérature, nous allons discuter de la variance, un indice de dispersion qui ne peut être assimilé à un indice d'inégalité. La différence entre les deux réside dans le fait qu'un indice d'inégalité fait la différence entre une dispersion de revenu qui a lieu dans le haut de la distribution par rapport à une dispersion qui a lieu dans le bas de la distribution.

1.2 La variance

Cette sous-section étudie la variance comme indice de dispersion d'une distribution, et il apparaît que même si cet indice respecte les propriétés des indices d'inégalité réguliers, il n'est pas considéré comme tel.

La variance est définie comme la moyenne du carré des écarts des revenus individuels à la moyenne. Cet indice est noté par :

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \mu)^2 \quad (1.1)$$

$$= \frac{1}{2} \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n (y_i - y_j)^2 \quad (1.2)$$

Avec y_i le revenu de l'individu i dans une population de n individus avec un revenu moyen μ . La seconde expression de Kendall et Stuart (1977) (équation 1.2) présente la variance comme la moitié de la moyenne quadratique des différences de revenus pris deux à deux, sans faire référence à une valeur centrale. La variance est une fonction dérivable, elle satisfait les propriétés de normalisation (compris entre 0 dans une situation parfaitement égalitaire et $(n-1)\mu^2$, lorsqu'un individu a tout), de symétrie, d'invariance à la réplique et respecte le principe de transferts de Dalton. Cependant la variance ne satisfait pas la propriété de l'invariance relative, c'est un indice absolu. Le principe d'invariance relative peut être obtenu en divisant la variance par le carré de sa moyenne. Cela donne

3. L'article de Cowell et Flachaire (2018) propose une comparaison de l'évolution des inégalités mesurées par l'indice de Gini (qui ne respecte pas le principe de monotonie à la distance) et l'indice de Theil 2 (qui le respecte) au Royaume-Uni et aux États-Unis.

le carré du coefficient de variation :

$$CV^2 = \frac{\sigma^2}{\mu^2} \quad (1.3)$$

L'écart-type, défini comme la racine carrée de la variance, est aussi souvent utilisé pour mesurer la dispersion d'une distribution. L'avantage de cet indicateur est d'être exprimé dans la même unité que les bornes de la distribution. De la même manière que la variance (équation 1.3), l'écart-type peut être divisé par la moyenne pour obtenir un indicateur qui respecte la propriété de l'invariance relative, le coefficient de variation (CV) est alors obtenu :

$$CV = \frac{\sigma}{\mu} \quad (1.4)$$

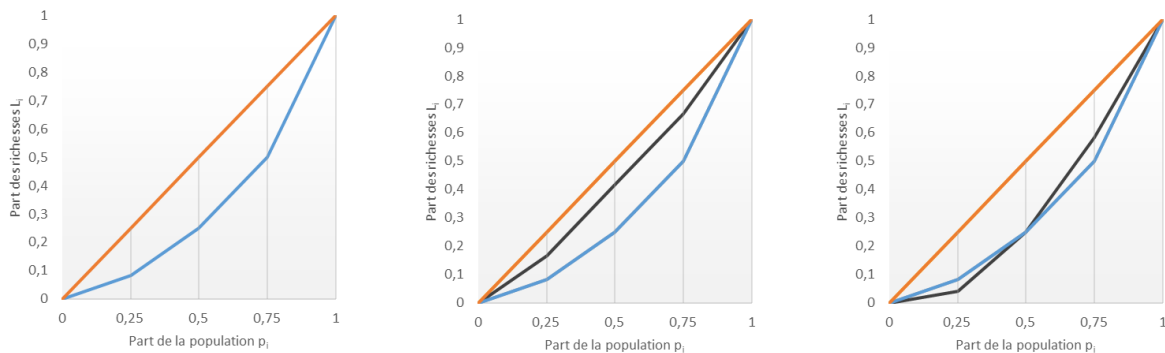
Le coefficient de variation respecte l'ensemble des propriétés des indices d'inégalité, cependant au même titre que la variance, il ne constitue pas pour autant un indice d'inégalité.

En effet, même si la variance et le coefficient de variation respectent les propriétés des indices d'inégalité réguliers, ils restent des indices de dispersion. La principale raison est qu'ils ne prennent pas correctement en compte la situation des moins bien lotis, dans le sens où ils ne font pas de différence entre un écart au-dessus de la valeur de référence (*i.e.* moyenne, médiane) ou en-dessous. La différence entre un indice de dispersion et d'inégalités repose avant-tout sur un jugement de valeur. De plus, au-delà des différents axiomes que devrait respecter un indice d'inégalité, la pondération des différents segments de la distribution fait aussi partie des choix du chercheur. Cette pondération reflète la sensibilité aux inégalités, c'est-à-dire l'importance donnée aux pauvres par rapport aux riches. Ce paramètre de sensibilité est rarement explicite dans les études des inégalités, ceci peut notamment être expliqué par le fait que l'indice le plus couramment utilisé, l'indice de Gini, ne fait pas apparaître ce paramètre. Pour autant l'indice de Gini et d'autres indices d'inégalité ont des versions généralisées où le paramètre de sensibilité apparaît explicitement, et dont la détermination de la valeur incombe au chercheur.

1.3 La courbe de Lorenz et l'indice de Gini

1.3.1 La courbe de Lorenz

Les indices d'inégalité sont des indices de concentration des richesses, c'est-à-dire qu'ils sont basés sur la comparaison entre la part des richesses détenues par les individus et la part qu'ils représentent dans la société. Cette idée est représentée graphiquement par la courbe de Lorenz (1905) (Graphiques 1.1). Pour construire ces graphiques, la distribution des revenus y des individus est classée par ordre croissant : $y_1 \leq y_2 \leq \dots \leq y_n$. La part cumulée de la population des plus pauvres



Graphique 1.1 – Courbes de Lorenz

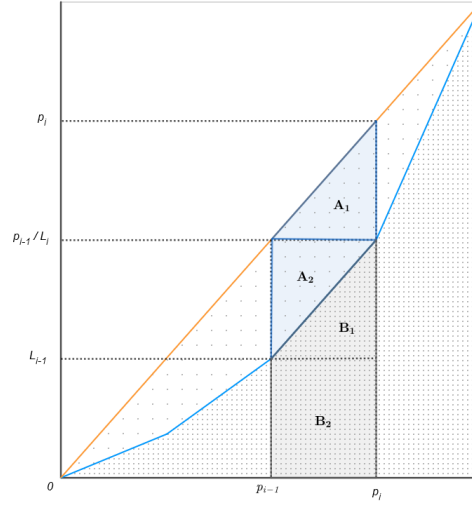
aux plus riches est représentée en abscisse et la part cumulée des richesses est représentée en ordonnée. Ainsi, la bissectrice de l'angle de l'origine de ces graphiques représente le cas parfaitement égalitaire, où la part des richesses détenues par des individus est égale à la part de ces individus dans la population totale. La différence entre la diagonale et chaque courbe représentant la concentration des richesses du cas étudié nous donne une mesure des inégalités.

Graphiquement, à l'aide du critère de la dominance de Lorenz, les deux distributions sont comparables si leur courbe ne se croisent pas (Graphique 1.1 central), celle du dessus représentant une distribution relativement plus égalitaire. De plus, il est intéressant de noter que quand une courbe domine une autre en tout points, tous les indices synthétiques donneront une valeur moins importante pour la première que pour la seconde. Ce qui n'est pas vérifié quand les courbes se croisent (Dasgupta *et al.* (1973), Rothschild et Stiglitz (1973)). Lorsqu'elles se croisent (Graphique 1.1 à droite), il est difficile de dire qu'elle est la situation la plus inégalitaire, cela conduit à utiliser des indices synthétiques basés sur cette courbe pour pouvoir comparer les distributions correspondantes. L'un des indices synthétiques les plus populaires, basé sur la courbe de Lorenz, est l'indice de Gini.

1.3.2 L'indice de Gini

L'indice de Gini a été proposé par Gini (1921). Cet indice correspond à l'aire comprise entre la diagonale et la courbe de Lorenz (zone en pointillés espacés dans le graphique 1.2, notée P_1) par rapport à l'aire totale comprise sous la diagonale (P_1 plus la zone en pointillés resserrés, notée P_2). Ainsi, plus l'aire comprise entre la diagonale et la courbe de Lorenz est importante, plus les inégalités sont importantes et plus l'indice de Gini est élevé. L'indice de Gini est noté G et il est égal à $\frac{P_1}{P_1 + P_2}$, sachant que $P_1 + P_2 = 0,5$, alors $G = 2 * P_1$.

Dans le graphique 1.2, l'aire représentée par les zones A_1, A_2, B_1 et B_2 est égale à la somme de l'aire comprise dans un rectangle (A_2, B_1 et B_2) et de celle comprise dans un triangle (A_1) et est



Graphique 1.2 – Courbe de Lorenz et Indice de Gini

égale à :

$$(p_i - p_{i-1})p_{i-1} + \frac{1}{2}(p_i - p_{i-1})^2 = \frac{1}{2}(p_i - p_{i-1})(p_i + p_{i-1})$$

Avec n individus classés par ordre croissant de revenus ($y_1 \leq y_2 \leq \dots \leq y_n$), p_i et L_i les parts cumulées respectives de la population et des richesses observées au rang de l'individu i . L'aire comprise dans les zones B_1 et B_2 est elle aussi la somme de l'aire d'un rectangle et d'un triangle :

$$(p_i - p_{i-1})L_{i-1} + \frac{1}{2}(p_i - p_{i-1})(L_i - L_{i-1}) = \frac{1}{2}(p_i - p_{i-1})(L_i + L_{i-1})$$

Ainsi, l'aire comprise dans la zone bleue du graphique 1.2 est la différence des deux équations précédentes :

$$\frac{1}{2}(p_i - p_{i-1})(p_i + p_{i-1}) - \frac{1}{2}(p_i - p_{i-1})(L_i + L_{i-1}) = \frac{1}{2}(p_i - p_{i-1})[(p_i + p_{i-1}) - (L_i + L_{i-1})]$$

De manière similaire, l'aire comprise entre la diagonale et la courbe de Lorenz entre deux points adjacents est sommée pour chaque point afin de calculer P_1 :

$$P_1 = \frac{1}{2} \sum_{i=1}^n (p_i - p_{i-1})[(p_i + p_{i-1}) - (L_i + L_{i-1})]$$

Sachant que $G = 2 * P_1$ et que $p_i - p_{i-1} = \frac{1}{n}$ pour tout i , alors :

$$G = \sum_{i=1}^n (p_i - p_{i-1})[(p_i + p_{i-1}) - (L_i + L_{i-1})] \quad (1.5)$$

$$G = \frac{1}{n} \sum_{i=1}^n [(p_i + p_{i-1}) - (L_i + L_{i-1})] \quad (1.6)$$

$$G = \frac{1}{n} \sum_{i=1}^n [(p_i - L_i) + (p_{i-1} - L_{i-1})] \quad (1.7)$$

Par ailleurs, $p_0 - L_0 = 0 - 0 = 0$ et que $p_n - L_n = 0 - 0 = 0$, ainsi $\sum_{i=1}^n p_{i-1} - L_{i-1} = \sum_{i=1}^n (p_i - L_i)$ et donc :

$$G = \frac{2}{n} \sum_{i=1}^n (p_i - L_i) \quad (1.8)$$

Cet indice respecte les propriétés de normalisation, de symétrie, d'invariance à la réplication, le principe des transferts et la continuité. L'indice de Gini peut être formulé de plusieurs façons dont voici un exemple proposé par Sen (1973) :

$$G = \frac{n+1}{n} - \frac{2}{n} \sum_{i=1}^n (n+1-i) \frac{y_i}{n\mu} \quad (1.9)$$

L'équation 1.8 nous dit que l'indice de Gini est égal à deux fois la moyenne des écart entre la part de la population et la part des richesses dans la distribution des revenus. Chaque écart entre ces deux mesures est traité symétriquement. Autrement dit, l'écart observé entre les 20% les plus pauvres et la richesse cumulée qu'ils détiennent a la même pondération dans l'indice de Gini que l'écart observé entre les 80% les plus pauvres et la part cumulée de leurs richesses. L'indice est compris entre 0 et $(n-1)/n$, et tend donc vers 1 quand $n \rightarrow \infty$.

L'équation 1.9, proposée par Sen (1973), montre que l'indice de Gini est une fonction linéaire d'une somme pondérée des parts de revenus, avec une pondération dépendant du rang de l'individu. Plus son rang est important plus le poids de la part de revenu détenu par l'individu i décroît, avec un poids n pour le plus pauvre et un poids égal à l'unité pour le plus riche. Ainsi, contrairement à l'indice de la variance, l'indice de Gini donne plus de poids aux pauvres (*i.e.* ayant un rang plus faible). Dans cette formule, l'impact d'un transfert de Dalton d'un montant δ d'un individu i vers un individu plus pauvre j est déterminé et il est égal à : $\delta \frac{2}{n^2} [j-i] < 0$. Ainsi, l'impact d'un transfert ne dépend pas du revenu des individus mais de leur position dans la hiérarchie des revenus.

Une version généralisée de l'indice de Gini permet d'estimer les inégalités selon un paramètre

de sensibilité aux inégalités ν , qui donne plus ou moins de poids en fonction du rang des individus :

$$G_\nu = 1 - \sum_{i=1}^n \left[\left(\frac{(n+1-i)^\nu - (n-i)^\nu}{n^\nu} \right) \cdot \frac{y_i}{\mu} \right] \quad (1.10)$$

L'équation 1.10 a été proposée par Donaldson et Weymark (1980) et dépend d'un paramètre $\nu > 1$. Quand $\nu = 2$, G_ν est égal à l'indice de Gini original (équation 1.8). Quand $1 < \nu < 2$, l'indice donne plus de poids aux riches qu'aux pauvres, et le contraire si $\nu > 2$. Ainsi le paramètre ν peut être interprété comme un paramètre d'aversion aux inégalités. Lorsque $\nu \rightarrow \infty$, $G_\nu \rightarrow 1 - \frac{y_1}{\mu}$, ainsi seul le rapport du revenu du plus pauvre à la moyenne est pris en compte. Lorsque $\nu \rightarrow 1$ alors $G_\nu \rightarrow 0$. L'indice de Gini généralisé est compris entre 0 et $1 - \frac{\nu}{n(\nu-1)}$.

Les trois distributions représentées dans le graphique 1.2 sont maintenant considérées, avec y_1 la distribution représentée sur les trois schémas, y_2 représentée uniquement sur le schéma du milieu et y_3 représentée uniquement sur le schéma de droite. La distribution y_1 est plus inégalitaire que y_2 car la courbe de concentration de richesses de y_1 est plus éloignée du cas parfaitement égalitaire que la courbe associée à y_2 . Cependant, dans le schéma de droite, il est difficile de dire qu'elle est entre y_1 et y_3 la distribution la plus inégalitaire car les deux courbes se croisent. Dans le tableau 1.1 les revenus sont ordonnés dans l'ordre croissant associé à chaque distribution sont donnés, ainsi que l'indice de Gini avec $\nu = 2$, $\nu = 3$ et sa limite quand $\nu \rightarrow \infty$.

Tableau 1.1 – Les indices de Gini

Rang individu	y_1	y_2	y_3
1	10	20	5
2	20	30	25
3	30	30	40
4	60	40	50
Gini ($\nu = 2$)	0,3333	0,1250	0,3125
Gini ($\nu = 3$)	0.4688	0.1875	0.4844
Gini ($\nu \rightarrow \infty$)	0.6667	0.3333	0.8333

Grâce aux indices synthétiques de Gini, les différentes distributions peuvent être ordonnées, même lorsque les courbes de Lorenz se croisent, cependant l'ordre dépend de la sensibilité aux moins bien lotis. Ainsi, l'indice de Gini ($\nu = 2$) estime que la distribution y_2 est moins inégalitaire que les deux autres, et que la distribution y_1 est plus inégalitaire que y_3 . Cependant, l'indice de Gini avec $\nu = 3$ estime que la troisième distribution est plus inégalitaire que la première. Ceci s'explique par le fait d'une situation plus défavorable du moins bien lotie dans y_3 et que la sensibilité

à cet individu est plus importante quand v augmente. L'importance du paramètre de sensibilité apparaît alors dans l'ordre des différentes distributions, le choix de l'évaluateur n'est donc pas sans conséquence.

1.4 Les indices d'inégalité de Theil

1.4.1 L'indice de Theil

L'indice de Theil est basé sur la théorie de l'information (Khinchin (1957), Kullback (1959)). Dans un premier temps, les auteurs définissent une fonction d'information qui évalue la pertinence de savoir qu'un événement soit survenu en fonction de la probabilité de cet événement. Cette fonction est décroissante par rapport à la probabilité que l'événement survienne, et s'il est certain que l'événement se produise alors la fonction d'information a une valeur nulle. Par ailleurs, la somme de la fonction d'information de deux événements indépendants est égal à l'information estimée des deux événements joints. S'il est supposé que la fonction d'information est dérivable, alors cette propriété additionnée aux précédentes conduisent à ce que la fonction d'information soit égal au log de l'inverse de la probabilité que l'événement survienne. Avec n événements de probabilité w_i , l'entropie $Q(\mathbf{w})$ est définie comme l'information espérée de cette situation :

$$Q(\mathbf{w}) = \sum_{i=1}^n w_i \log\left(\frac{1}{w_i}\right), \quad (1.11)$$

avec $\mathbf{w}$ le vecteur des probabilités associées à chaque événement.

L'entropie maximale est obtenue quand tous les événements ont la même probabilité $1/n$ et est ainsi égale à $\log n$. Theil (1967) utilise ce cadre pour étudier la concentration des revenus. Il introduit deux changements. Le premier consiste à remplacer la probabilité d'un événement w_i par la part de revenu détenu par un individu $s_i = y_i/n\mu$, avec μ le revenu moyen. Le second est que son indice est la différence entre d'une part le cas où tous les individus possèdent la même part de revenu (*i.e.* l'entropie maximale), et d'autre part l'entropie observée dans la distribution $Q(\mathbf{s})$. Ainsi, l'indice de Theil est défini par la différence entre une situation parfaitement égalitaire et la répartition des revenus observée :

$$T = \log n - Q(\mathbf{s}), \quad (1.12)$$

avec $\mathbf{s}$ le vecteur des parts de revenus observées. Sachant que $Q(\mathbf{s}) = \sum_{i=1}^n s_i \log \frac{1}{s_i} = -\sum_{i=1}^n s_i \log s_i$

et que $\log n = \sum_{i=1}^n s_i \log n$, l'équation 1.12 peut être réécrite telle que :

$$T = \sum_{i=1}^n s_i \log n + \sum_{i=1}^n s_i \log s_i \quad (1.13)$$

$$= \sum_{i=1}^n s_i \log n s_i \quad (1.14)$$

Cet indice peut être interprété comme une fonction de distance entre les parts des individus dans la population et leur part de revenu, pondérée par la part des revenus. Ainsi, la distance entre la part d'un individu dans la population et sa part de richesse aura d'autant plus de poids que sa part des richesses est importante. Dans une situation parfaitement égalitaire où tous les individus de la population ont la même part de revenu, $s_i = 1/n$ et $T = 0$. Si un individu possède tout et les autres rien, alors $T \rightarrow \log n$. L'indice peut être éventuellement normalisé en le divisant par $\log n$ pour que l'indicateur soit compris entre $[0, 1]$, cependant cela se fait au prix d'un renoncement au principe de l'invariance à la réplication puisque cet indice normalisé dépendrait de la taille de la population. Sachant que $s_i = y_i/\mu$, une autre version de l'indice de Theil (équation 1.14) peut être obtenue :

$$T = \frac{1}{n} \sum_{i=1}^n \frac{y_i}{\mu} \log \frac{y_i}{\mu} \quad (1.15)$$

L'indice de Theil satisfait l'ensemble des propriétés énoncées des indices d'inégalité. L'impact d'un transfert de Dalton δ d'un individu i à un individu plus pauvre j est égal à : $\delta \frac{1}{n\mu} \log \frac{y_j}{y_i} < 0$. Contrairement à l'indice de Gini cet impact ne dépend pas du rang des individus concernés par le transfert mais de leur revenus, plus la différence de revenu est importante plus l'impact est conséquent.

1.4.2 L'indice second de Theil

L'indice second de Theil (*Écart logarithmique moyen*), noté T^* , est similaire au premier à la différence que les pondérations ne reposent plus sur la part des revenus mais sur la part dans la population. L'équation obtenue est alors :

$$T^* = \frac{1}{n} \sum_{i=1}^n \log \frac{\mu}{y_i} \quad (1.16)$$

Dans une situation parfaitement égalitaire, c'est-à-dire où la part dans la population de chaque individu correspond à sa part de revenu, alors $T^* = 0$. Dans le cas contraire, où un possède tout, l'indice n'est pas borné. Cet indice peut être réécrit comme étant le logarithme du rapport entre la

moyenne arithmétique des revenus et la moyenne géométrique des revenus $\tilde{\mu}$:

$$\begin{aligned}
T^* &= \frac{1}{n} \sum_{i=1}^n (\log \mu - \log y_i) \\
&= \log \mu - \frac{1}{n} \sum_{i=1}^n \log y_i \\
&= \log \mu - \log \tilde{\mu} \\
&= \log \frac{\mu}{\tilde{\mu}}
\end{aligned} \tag{1.17}$$

Ce second indice de Theil respecte l'ensemble des propriétés énoncées. L'impact d'un transfert de Dalton est donné par : $\delta_n^1(\frac{1}{y_i} - \frac{1}{y_j})$. Un des critères de choix entre l'indice premier et l'indice second de Theil repose donc sur la méthode de pondération de l'écart entre la part d'un revenu d'un individu et son poids dans la population totale. Dans le premier cas, l'indice donne plus de poids aux écarts lorsque le revenu des individus concernés est élevé. Alors que dans le second cas, le poids de l'écart dans la mesure de l'inégalité augmente avec le poids de l'individu dans la population.

1.4.3 Les indices généralisés de Theil

Les indices de Theil présentés précédemment sont deux cas spécifiques des indices généralisés d'entropie. Shorrocks (1980) montre qu'il existe une unique famille de fonctions qui satisfait l'ensemble des propriétés énoncées précédemment (Section 1.1). Le théorème de Shorrocks énoncé par Villar (2017) est donné comme suit :

Théorème 1 (Shorrocks 1980). *Une fonction $I : \mathbf{y} \rightarrow \mathbb{R}$ satisfait les propriétés de normalisation, de symétrie, d'invariance à la réplique, le principe des transferts, de dérivabilité, d'invariance relative et de décomposition additive⁴, si et seulement si elle est de la forme :*

$$I_\theta(n, \mathbf{y}) = \frac{1}{n} \frac{1}{\theta(\theta - 1)} \sum_{i=1}^n \left[\left(\frac{y_i}{\mu} \right)^\theta - 1 \right] \tag{1.18}$$

(ou si c'est une transformation proportionnel de cette fonction, $\alpha I_\theta(n, \mathbf{y})$ avec $\alpha > 0$), avec $\theta \in \mathbb{R} \setminus 0$.

Pour $\theta = 1$, l'indice obtenu est l'indice de Theil. Pour $\theta = 0$ et $\alpha = 1/n$, l'indice obtenu est l'indice second de Theil. Pour $\theta = 2$, le demi de l'écart-type au carré est obtenu. Sachant que l'écart-type au carré correspond à la variance du vecteur de revenu centré sur sa moyenne.

Le paramètre θ est une mesure de la sensibilité de l'indice par rapport à la localisation où les transferts ont lieu. Plus θ est faible et tend vers $-\infty$, plus l'indice est sensible aux transferts vers le bas de la redistribution. Au contraire, plus θ augmente, plus l'indice devient régressif. Lorsque $\theta > 2$, l'indice d'inégalité diminue seulement si les revenus des plus riches sont égalisés⁵. L'effet

4. Cet axiome sera traité dans la section suivante.

5. Sachant que le principe des transferts est toujours respecté

d'un transfert de Dalton est donné par :

$$\begin{cases} \delta \frac{1}{n\mu^{\theta(\theta-1)}} (y_j^{\theta-1} - y_i^{\theta-1}) & \theta \neq 1 \\ \delta \frac{1}{n\mu} \log \frac{y_j}{y_i} & \theta = 1 \end{cases} \quad (1.19)$$

Dans la théorie informationnelle, une fonction d'information ($\phi(w)$) décroissante en la probabilité d'occurrence d'un événement w , dont la valeur est nulle quand la probabilité est de 1 et qui est additive, est de la forme $-\log(w)$. Si la propriété additive est relâchée, deux formes fonctionnelles apparaissent :

$$\phi(\beta, w) = \begin{cases} \frac{1-w^\beta}{\beta} & \beta \neq 0 \\ -\log w & \beta = 0 \end{cases} \quad (1.20)$$

L'utilisation de cette nouvelle fonction d'information dans la première formule de Theil conduit à une nouvelle définition de l'impact d'un transfert de Dalton, noté dT , sur les inégalités :

$$dT = ds[\phi(s_i) - \phi(s_j)], \quad (1.21)$$

avec $ds = \delta/n\mu$. Ainsi, plus la différence de revenu est importante entre les deux agents participant aux transferts, plus la réduction des inégalités est importante pour un montant de transfert donné.

En utilisant la fonction d'information (équation 1.20) et en posant $\theta = \beta + 1$, une nouvelle formule de l'indice de Theil généralisé est obtenue :

$$T_\theta = \frac{n^{1-\theta}}{\theta(\theta-1)} \frac{1}{n} \sum_{i=1}^n \left[\left(\frac{y_i}{\mu} \right)^\theta - 1 \right], \forall \theta \neq 0, 1 \quad (1.22)$$

L'équation 1.22 est une transformation linéaire de l'équation 1.18. Autrement dit, l'approche axiomatique de Shorrocks ou la généralisation de la fonction d'information⁶ conduisent toutes les deux à généraliser les indices d'inégalité basés sur le principe de l'entropie :

$$I_\theta(n, y) = \begin{cases} \frac{1}{n} \frac{1}{\theta(\theta-1)} \sum_{i=1}^n \left[\left(\frac{y_i}{\mu} \right)^\theta - 1 \right] & \theta \neq 0, 1 \\ \frac{1}{n} \sum_{i=1}^n \log \frac{\mu}{y_i} & \theta = 0 \\ \frac{1}{n} \sum_{i=1}^n \frac{y_i}{\mu} \log \frac{y_i}{\mu} & \theta = 1 \end{cases} \quad (1.23)$$

Dans les trois cas, le cas parfaitement égalitaire conduit à $I_\theta = 0$. La borne supérieure est quant à elle différente selon la valeur de θ . Si $\theta > 0$, la borne supérieure est $\frac{n^{\theta-1}-1}{\theta(\theta-1)}$ si $\theta \neq 1$ et $\log n$ si $\theta = 1$. Enfin, si $\theta \leq 0$, il n'y a pas de borne supérieure.

6. Généralisation par le relâchement de l'hypothèse d'additivité, les autres conditions sont maintenues.

Dans le tableau 1.2, les trois distributions du tableau 1.1 sont reprises. Pour chacune des distributions, l'indice de Theil, l'indice second de Theil et l'indice généralisé de Theil pour $\theta = 3$ sont donnés. Selon le paramètre fixé, l'ordre des distributions n'est pas le même. Ainsi, T et T^* classent les distributions de la même façon : $T_\theta(y_3) > T_\theta(y_1) > T_\theta(y_2)$, pour $\theta = 0, 1$, alors que T_3 classe les distributions tel que : $T_3(y_1) > T_3(y_3) > T_3(y_2)$. Ceci n'a rien de surprenant puisque un θ plus fort implique une plus faible prise en compte des moins bien lotis, et que la situation du plus pauvre dans y_3 est plus défavorable que dans y_2 . C'est pour cette même raison que l'indice second de Theil est bien plus fort que T_3 pour la troisième distribution. Lorsque ces résultats sont comparés avec ceux obtenus avec l'indice de Gini, deux remarques s'imposent. La première c'est que l'ordre des séries selon l'indice de Gini correspond ici à l'ordre de T_3 , c'est-à-dire l'indice le plus défavorable aux moins bien lotis parmi les indices de Theil. Deuxièmement, le rapport entre la première et la seconde série est plus marqué avec les indices de Theil, il est de l'ordre de 10 pour ces derniers contre 3 pour l'indice de Gini.

Tableau 1.2 – Les indices de Theil

Indice :	y_1	y_2	y_3
$T (\theta = 1)$	0,1874	0,0283	0,1961
$T^* (\theta = 0)$	0,2027	0,0294	0,2939
$T_\theta (\theta = 3)$	0,2222	0,0278	0,1493
Gini	0,3333	0,1250	0,3125

1.5 Les indices d'Atkinson

1.5.1 Une approche normative

Les indices d'Atkinson font partis des indices d'inégalité normatifs, c'est-à-dire qu'ils modélisent explicitement les inégalités comme une perte de bien-être social. Les indices d'inégalité étudiés précédemment reposent eux aussi sur des jugements, mais de manière implicite, notamment par le choix de la sensibilité d'un indicateur aux moins bien lotis. Deux approches normatives sont différenciées, d'une part celle qui évalue les inégalités par une fonction de bien-être social à la Arrow, c'est-à-dire en agrégeant les utilités individuelles en un indicateur de bien-être social. D'autre part, l'approche qui évalue une distribution de revenus sans passer par les utilités individuelles. Les indices d'Atkinson font partie des indices d'inégalité basés sur des fonctions de bien-être social à la Arrow.

Dans le contexte des inégalités, les alternatives sociales offertes aux individus sont limitées aux différentes redistributions de revenus. Communément, la fonction de bien-être social dépend de la

taille du revenu total, souvent représentée par la moyenne, et de la distribution du revenu, représentée par un indice d'inégalité. Pour que cette fonction soit décroissante par rapport au niveau des inégalités, il est supposé qu'elle est strictement quasi-concave. De cette manière, quand tous les individus ont la même fonction d'utilité qui dépend uniquement de leur revenu, et qui est strictement concave, un transfert de revenu d'un riche à un pauvre améliore le bien-être social puisque la baisse d'utilité du riche est plus que compensée par une hausse d'utilité du pauvre (*i.e.* le principe de transfert de Dalton est donc respecté). La concavité de la fonction d'utilité des agents est un principe essentiel qui conduit à ce que la situation d'égalité parfaite soit celle qui maximise le bien-être social. Cependant, une fonction d'utilité sociale en forme de U inversé est possible, où l'utilité maximale serait obtenue pour un certain niveau d'inégalité. Ceci peut s'expliquer par exemple par le fait que les individus préfèrent vivre dans une société méritocratique où le revenu est basé sur le mérite, et préféreraient donc avoir un certain niveau d'inégalité dans la société dans la mesure où les différences de revenu sont issues de différences dans le mérite pour l'obtenir.

Dans le cas où la fonction d'utilité est strictement concave, pour mesurer l'aversion aux inégalités, il peut être pertinent d'estimer l'élasticité de l'utilité marginale, noté $\epsilon(y)$, c'est-à-dire de mesurer le changement relatif de l'utilité marginale par rapport au changement relatif du revenu individuel :

$$\epsilon(y) = -\frac{\partial(\partial u/\partial y)}{\partial y} \frac{y}{\partial u/\partial y} = \frac{\partial^2 u/\partial y^2}{\partial u/\partial y} y \quad (1.24)$$

Dans un premier temps, il est plus simple de faire l'hypothèse que cet indice est constant. Ainsi, quel que soit le niveau de revenu d'un individu, son aversion aux inégalités est la même et correspond à la valeur de l'élasticité de l'utilité marginale ϵ . Cette restriction permet de définir la forme fonctionnelle des fonctions d'utilité :

$$u_e(y) = \begin{cases} \frac{y^{1-\epsilon}}{1-\epsilon} & \epsilon \neq 1 \\ \log y & \epsilon = 1 \end{cases} \quad (1.25)$$

avec $\epsilon > 0$.

Par ailleurs, si $\epsilon = 0$ alors $u_e(y) = y$, autrement dit si la fonction d'utilité n'est pas strictement concave, il n'y a pas d'aversion aux inégalités. Plus ϵ est important, plus le poids des revenus faibles augmente. Pour $\epsilon = 2$, l'aversion aux inégalités est déjà très prononcée car l'individu qui a un revenu égal à la moitié de la moyenne aura un poids 4 fois plus élevé que celui au revenu moyen.

1.5.2 Les indices d'Atkinson

L'indice de Atkinson (1970), noté A , estime le coût social des inégalités en comparant le revenu moyen dans une société (μ) au revenu équivalent égalitaire (ξ), c'est-à-dire le revenu distribué de

manière uniforme dans la société et qui permet d'obtenir le bien-être social similaire. Ainsi, l'indice suivant est donné :

$$A = 1 - \frac{\xi}{\mu} \quad (1.26)$$

avec $A = 0$ si $\xi = \mu$ et $A \rightarrow 1$ quand $\xi \rightarrow 0$.

Concrètement, si $A = 0,4$ alors ξ correspond à 60% du revenu moyen, le même niveau de bien-être pourrait être obtenu avec seulement 60% du revenu total distribué également dans la société.

Atkinson utilise une fonction de bien-être utilitariste, qui correspond donc à la somme des utilités individuelles. La fonction d'utilité est la même pour tous, et l'aversion aux inégalités est constante (équation 1.25). Ces deux principes conduisent à une fonction de bien-être social, $W(\cdot)$, de la forme suivante :

$$W(y) = \begin{cases} \sum_{i=1}^n \frac{y_i^{1-\epsilon}}{1-\epsilon} & \epsilon \neq 1 \\ \sum_{i=1}^n \log y_i & \epsilon = 1 \end{cases} \quad (1.27)$$

Par conséquent, si $\epsilon = 0$ l'aversion à l'inégalité est nulle et le bien-être social dépend uniquement du revenu total distribué dans l'économie, peu importe comment il est distribué, donc $A = 0$ pour toutes les distributions. Plus ϵ est élevé, plus le bas de la distribution importe pour estimer les inégalités. Dans le cas extrême où $\epsilon \rightarrow \infty$, seul le revenu du plus pauvre compte, la fonction de bien-être obtenue correspond au *Leximin social* (Hammond (1975), Rawls (1971)) : $W(y) = \min_{y_i}(y_1, \dots, y_n)$.

A partir de l'équation 1.27, une version généralisée de l'indice d'Atkinson est obtenue :

$$A_\epsilon = \begin{cases} 1 - \left[\frac{1}{n} \sum_{i=1}^n \left(\frac{y_i}{\mu} \right)^{1-\epsilon} \right]^{\frac{1}{1-\epsilon}} & \epsilon \neq 1 \\ 1 - \frac{\tilde{\mu}(y)}{\mu(y)} & \epsilon = 1 \end{cases} \quad (1.28)$$

avec $\tilde{\mu}(y)$ la moyenne géométrique des revenus distribués y .

Pour une distribution de revenus donnée, cet indice est croissant en ϵ , le paramètre d'aversion aux inégalités (Cowell et Jenkins (1995)). De plus, pour toutes valeurs de $\epsilon > 0$, $A_\epsilon = 0$ lorsqu'il y a une égalité parfaite. La borne supérieure, qui correspond à une situation où un individu possède tout et les autres rien, dépend de la valeur du paramètre ϵ :

$$A_\epsilon \in \begin{cases} [0, 1 - n^{\frac{-\epsilon}{1-\epsilon}}] & 0 < \epsilon < 1 \\ [0, 1] & \epsilon = 1 \\ [0, ?] & \epsilon > 1 \end{cases}$$

Dans le 1^{er} cas, la borne supérieure tend vers 1 quand $n \rightarrow \infty$. Et dans le dernier cas, la borne supérieure n'est pas définie, même si $A_\epsilon \rightarrow 1$ quand $y \rightarrow 0$.

L'indice d'Atkinson satisfait donc le principe de normalité, ainsi que les principes de symétrie, d'invariance à la réplication, le principe des transferts, de continuité, d'invariance relative. L'impact d'un transfert de Dalton δ d'un individu i à un individu plus pauvre j est donné par :

$$\delta \frac{(1 - A_\epsilon)^\epsilon}{n\mu^{1-\epsilon}} [y_i^{-\epsilon} - y_j^{-\epsilon}] < 0 \quad (1.29)$$

Empiriquement, pour $\epsilon = 1$, l'indice d'Atkinson donne des valeurs plus faibles que celui de Gini. Cependant, quand plusieurs sociétés sont comparées avec ces indices normalisés (rapportés à leur moyenne respective), Villar (2017) montre que l'indice d'Atkinson discrimine plus que l'indice de Gini en donnant des valeurs plus faibles aux sociétés avec un niveau d'inégalités plus faible que la moyenne, *vice-versa*. Enfin, pour $\epsilon > 0$ et $\theta = 1 - \epsilon$, l'indice d'entropie généralisé et l'indice d'Atkinson généralisé sont ordinalement équivalents, autrement dit ils classent toutes les distributions de revenus de la même manière, même si les indices utilisés sont différents. Ainsi dans le tableau 1.3, l'ordre des distributions de A_1 doit être le même que celui de *Theil2* ($\theta = 0$), ce qui est le cas. De plus, quel que soit la valeur du paramètre de sensibilité, l'ordre des distributions ne varie pas entre les indices d'Atkinson. Par ailleurs, seul l'indice de Gini estime que la première distribution est plus inégalitaire que la troisième.

Tableau 1.3 – Les indices d'Atkinson

Indice :	y ₁	y ₂	y ₃
A₁	0,1835	0,0290	0,2546
A_{0.5}	0,0937	0,0143	0,1132
A_{1.5}	0,2644	0,0439	0,4050
Gini	0,3333	0,1250	0,3125
Theil	0,1874	0,0283	0,1961
Theil 2	0,2027	0,0294	0,2939

Les indices d'Atkinson mesure les inégalités dans la distribution des revenus via une fonction d'utilité individuelle. Cette approche est sujette à plusieurs critiques. En premier lieu, il pourrait être préférable d'étudier directement les inégalités de bien-être plutôt que les inégalités de revenus. Par ailleurs, les indices d'Atkinson reposent sur des choix propres à l'auteur quant à la forme de la fonction d'utilité. Même si seuls les indices d'Atkinson sont étudiés, c'est-à-dire avec une aversion constante aux inégalités, il reste à décider de la valeur du paramètre d'aversion ϵ . De plus, la fonction d'utilité est supposée constante pour l'ensemble des individus, si cette hypothèse était rejetée, il faudrait alors concevoir un moyen de comparer différentes fonctions d'utilité sans tomber dans le théorème d'impossibilité de Arrow⁷. Sen (1973) a proposé une autre approche normative

7. Selon ce théorème, il n'existe pas de règle de vote parfaite qui permettrait de convertir les préférences indivi-

en utilisant une fonction de bien-être social qui mesure les inégalités de revenus sans utiliser de fonction d'utilité individuelle. Cette fonction est nommée Fonction d'Évaluation Sociale (FES).

1.6 Fonction d'Évaluation Sociale

Une fonction d'évaluation sociale $V(\cdot)$ donne à chaque distribution une valeur, avec de plus grandes valeurs pour de meilleures distributions. La forme de cette fonction peut être restreinte dans un premier temps par deux propriétés posées par Sen : le traitement symétrique des individus et la stricte quasi-concavité (*i.e.* la fonction est décroissante par rapport au niveau des inégalités).

Par ailleurs, Sen (1973) a défini le revenu équivalent égalitaire généralisé, y^e , tel que si ce revenu par tête est distribué à tous les individus alors le niveau de bien-être dans la société obtenu serait le même que celui obtenu avec la distribution effective. Ainsi, $V(\mathbf{y}^e) = V(\mathbf{y})$, avec $\mathbf{y}^e$ la distribution du revenu équivalent et $\mathbf{y}$ la distribution des revenus observée. Comme $V(\mathbf{y})$ est symétrique et quasi-concave, $y^e \leq \mu(\mathbf{y})$, avec $\mu(\mathbf{y})$ la moyenne des revenus observés. Ainsi, la famille des indices d'inégalité suivante est définie :

$$S_v = 1 - \frac{y^e}{\mu(\mathbf{y})} \quad (1.30)$$

Si la FES correspond à la somme des utilités individuelles alors l'indice S_v est égal à l'indice d'Atkinson. Comme dans le cas d'Atkinson, cette fonction peut être interprétée comme une mesure monétaire de la perte de bien-être causée par les inégalités. A ces trois principes posés pour obtenir des indices d'inégalité à partir de la FES, peut s'ajouter la continuité, l'homogénéité de degré 1 et le centrage sur la moyenne. Ce dernier stipule que si tous les revenus sont égaux alors le bien-être associé à cette distribution égalitaire a une valeur correspondant au revenu moyen.

Ces trois principes induisent qu'il existe une distribution de revenu égalitaire, notée $\gamma(\mathbf{y})$, qui donne le même niveau de bien-être dans la société que la distribution observée, et que cette solution est unique. Si les propriétés de symétrie et de stricte quasi-concavité sont ajoutées, alors $\gamma(\mathbf{y}) = y^e$. Ainsi pour chaque distribution, il existe un revenu équivalent tel que :

$$V(\mathbf{y}) = y^e \quad (1.31)$$

duelles ordonnées en un classement des préférences complet et transitif au niveau communautaire tout en respectant certains critères : universalité (*i.e.* toutes les préférences sont permises), pas de dictateur (*i.e.* cas où l'utilité d'un seul individu importerait), efficacité de Pareto (*i.e.* il n'existe pas une alternative où l'utilité agrégée serait plus importante) et l'indépendance des alternatives non-pertinentes (*i.e.* la préférence entre deux alternatives ne dépend pas d'une autre alternative).

D'après l'équation (1.30) :

$$\begin{aligned} V(\mathbf{y}) &= \mu(\mathbf{y})[1 - S_v(\mathbf{y})] \\ &= \mu(\mathbf{y}) - \mu(\mathbf{y})S_v(\mathbf{y}) \end{aligned} \quad (1.32)$$

Ainsi le bien-être associé à la distribution $\mathbf{y}$ est la différence entre le revenu moyen représentant une situation parfaitement égalitaire et le produit $\mu(\mathbf{y})S_v(\mathbf{y})$ qui décrit la perte de bien-être par individu due aux inégalités exprimées en termes monétaires.

Si l'indice d'inégalité d'Atkinson est mesuré avec $\epsilon = 1$ comme mesure des inégalités, alors $S_v = A_1(\mathbf{y}) = 1 - (\tilde{\mu}/\mu)$, avec $\tilde{\mu}$ la moyenne géométrique. En substituant S_v dans l'équation (1.32), alors :

$$V_{A_1}(\mathbf{y}) = \tilde{\mu}(\mathbf{y}) \quad (1.33)$$

Ceci correspond à la mesure de bien-être social utilisé par les Nations Unis pour introduire le concept d'équité dans l'indice de développement humain (IDH).

Si l'indice d'Atkinson prend des valeurs supérieures à 1, ce qui peut être le cas quand $\epsilon > 1$ (équation 1.5.2), alors l'indice de bien-être devient négatif, et la fonction de bien-être n'est pas monotone par rapport au revenu moyen. Par ailleurs, la fonction de bien-être dépend de l'unité de mesure des revenus puisque c'est une mesure monétaire de la perte de bien-être induite par les inégalités. Une solution à ce problème est de raisonner en termes relatifs, c'est-à-dire :

$$Z_v(\mathbf{y}) = \frac{\mu(\mathbf{y})S_v(\mathbf{y})}{V(\mathbf{y})} = \frac{S_v(\mathbf{y})}{1 - S_v(\mathbf{y})} \quad (1.34)$$

Cette expression nous donne la perte relative de bien-être, c'est-à-dire la part de revenu additionnel dont pourrait bénéficier la société en l'absence d'inégalité.

Les propriétés de dérivation et d'homogénéité de degré 1 permettent de réécrire la fonction d'évaluation sociale :

$$V(\mathbf{y}) = \sum_{i=1}^n \alpha_i(\mathbf{y})y_i, \quad (1.35)$$

avec $\alpha_i(\mathbf{y}) = \partial V / \partial y_i$. Ainsi, le bien-être peut être exprimé comme une somme pondérée de revenus individuels. Cette pondération reflète la courbe des fonctions de bien-être indifférenciées. Cette équation correspond à l'approche des biens premiers de Sen (1976). Ceci peut être résumé par le théorème suivant :

Théorème 2 (Sen, 1976). *Une fonction d'évaluation sociale $V : \mathbb{R}_{++}^n \rightarrow \mathbb{R}$ qui satisfait les propriétés de symétrie, stricte quasi-concavité, dérivation, homogénéité de degré 1 et centrée peut être exprimée comme : $V(\mathbf{y}) = \sum_{i=1}^n \alpha_i(\mathbf{y})y_i = \mu(\mathbf{y})[1 - I(\mathbf{y})]$, avec $\alpha_i(\mathbf{y})$ une fonction homogène continue de degré zéro et $I(\mathbf{y})$ un indice d'inégalité relatif.*

Sen (1976) montre que la fonction d'évaluation sociale peut être exprimée avec l'indice de Gini

avec une pondération $\alpha_i^G(\mathbf{y}) = (n - i + 1)/n$, ainsi la FES suivante est obtenue :

$$V_G(\mathbf{y}) \approx k\mu(\mathbf{y})[1 - G(\mathbf{y})], \quad (1.36)$$

avec k une constante positive ($k = n/2$) et G l'indice de Gini. Cette approximation tend vers l'égalité quand le nombre d'individus n devient grand. Ainsi le poids des individus est inversement proportionnel à leur rang.

Herrero et Villar (1989) ont quant à eux exprimé la fonction d'évaluation sociale en fonction de l'indice premier de Theil (1.15). La pondération associée à cette fonction est $\alpha_i^T(\mathbf{y}) = [1 - \log(y_i/\mu(\mathbf{y}))]/n$, avec la FES suivante :

$$V^T(\mathbf{y}) = \mu(\mathbf{y})[1 - T(\mathbf{y})], \quad (1.37)$$

avec T l'indice premier de Theil. Ainsi le poids de chaque individu dans l'indice premier de Theil dépend de leur revenu. Il est égal à $1/n$ pour les individus au revenu moyen, puis augmente progressivement pour les individus qui s'écartent par le bas de la moyenne et diminue progressivement pour ceux qui s'écartent par le haut de la moyenne. Par ailleurs, le poids devient négatif pour les revenus supérieur à $e\mu(\mathbf{y})$ (où e est la base du logarithme naturel), ainsi tous les individus qui ont un revenu moyen supérieur à trois fois la moyenne ont une pondération négative dans l'indice de Theil. Autrement dit, le revenu de ces individus diminue la FES et donc le bien-être social mesuré dans la société.

1.7 Discussion et conclusion

Au sein de ce premier chapitre, les principaux indices d'inégalité utilisés dans la littérature ont été présentés, notamment au regard de certains principes souhaitables pour apprécier au mieux les inégalités dans une société. Les différents exemples de ce chapitre ont illustré que tous les indicateurs ne classent pas dans le même ordre les distributions, et que cet ordre ne dépend pas uniquement de l'indice mais aussi du choix d'un paramètre de sensibilité aux inégalités. L'appréciation des inégalités dans une société repose donc sur des choix, celui de l'indice et de ses paramètres.

Dans ce chapitre et tout au long de cette thèse, les inégalités sont appréciées en termes de revenus. Le lien entre inégalités de revenus et inégalités de bien-être est possible si le revenu des individus est représentatif de leur niveau de bien-être de telle sorte qu'ils puissent être comparés entre eux. Pour pouvoir comparer une relation de bien-être social entre différents états de la société, il faut supposer qu'il existe une fonction de bien-être social définie sur la distribution des revenus. Les indices d'Atkinson supposent une fonction de bien-être social utilitariste où le bien-être total de la société est égal à la somme des utilités individuelles. Dans la littérature, l'indice d'Atkinson est

parfois dit « prioritariste » parce que les fonctions d'utilités sont supposées concaves, permettant ainsi de donner plus de poids aux individus les moins bien lotis (et donc permettre les transferts de Dalton). Cette concavité est supposée dans l'ensemble des indices présentés puisqu'ils respectent le principe de transfert. Cet axiome permet d'introduire un jugement de valeur dans la conception des indices d'inégalité, il impose une préférence pour une société égalitaire car le transfert de revenu d'un riche à un pauvre conduit à une situation préférable.

La relation entre l'évaluation du bien-être et son inégale répartition dans une société est ténue, car les deux évaluations reposent sur les valeurs partagées dans la société, mais cette relation n'est pas nécessairement monotone. Par exemple, il est possible qu'une société soit plus désirable en termes de bien-être qu'une autre, alors qu'elle est plus inégalitaire, *vice-versa*. Dans l'approche utilitariste classique, l'ajout d'un individu à une population aboutit toujours à une situation plus désirable, à partir du moment où l'utilité du nouvel individu est positive (*i.e.* dont la vie vaut la peine d'être vécue), puisque la somme des utilités est plus grande. Pour autant, la nouvelle distribution du bien-être dans la société peut être plus inégalitaire si, par exemple, le bien-être du dernier venu est relativement très faible ou très élevé. Au demeurant, le niveau d'utilité seuil à partir duquel l'ajout d'un individu a un effet réducteur sur l'inégalité estimée ne semble pas être défini pour les différents indices d'inégalité étudiés, et pourrait faire l'objet de recherches futures. Même si toutefois, l'ajout d'un nouvel individu a un impact quasi-nul sur l'inégalité dans les grandes populations.

De plus, dans l'approche utilitariste, une petite société à l'utilité globale forte pourrait être préférée à une grande société à l'utilité globale faible, à partir du moment où la somme des utilités dans le premier cas est plus importante que dans le second, peu importe les inégalités observées. Ce paradoxe est plus connu sous le nom de la « conclusion répugnante » (Parfit (1984)), et est notamment soulevé par les théoriciens de l'éthique de la population (voir Zuber (2017)). Pour répondre à ce problème, les utilitaristes ont développé le critère de l'utilitarisme à niveau critique, ainsi l'ajout d'un individu est jugé bénéfique s'il a une utilité supérieure ou égale au niveau critique, qui peut être un niveau constant (Blackorby *et al.* (2005)) ou encore l'utilité moyenne observée dans la société (Gravel (2011)). Le critère de bien-être de rang est une approche alternative dans laquelle s'inscrit l'indice de Gini, celle-ci consiste à pondérer le bien-être des individus selon son rang dans la distribution, avec des poids plus importants pour les revenus les plus faibles. Ce critère de bien-être de rang se différencie en deux points principaux des critères utilitaristes. Tout d'abord, ce sont des critères à valeurs variables, c'est-à-dire que l'augmentation de bien-être de la société par l'ajout d'un individu est de plus en plus faible, à utilité donnée. La pondération dépendant du nombre total d'individus dans la société, quel que soit le rang d'un individu, plus la population est grande plus son poids dans la société diminue. La deuxième différence notable est que les critères de bien-être de rang sont sensibles au contexte. Autrement dit, le bien-être apporté à la société par un individu dépend de l'utilité des autres individus vivant dans cette société (notamment selon son rang).

Par ailleurs, l'évaluation du bien-être dans une société pour une période donnée ne dépend pas des individus morts avant ou nés après si et seulement si l'utilité expérimentée tout au long de la vie par ces individus hors-période étudiée est la même pour tous les états possibles de la société (hypothèse d'indépendance de l'utilité des morts - Blackorby *et al.* (1993)). Par conséquent, un ordre de préférence entre plusieurs états de la société au regard de son bien-être est possible si l'utilité des individus qui ne sont pas présents pour la période étudiée (*i.e.* ils sont morts ou pas encore nés) est la même dans chaque état. Dans les faits, lorsque l'efficacité d'une politique publique est évaluée en termes de bien-être pour une période donnée, il est supposé que quel que soit l'état de la société demain, le bien-être des populations futures ne sera pas impacté par cette politique, ainsi leur bien-être serait le même quelque soit l'état de la société aujourd'hui. Cette hypothèse est souvent postulée implicitement, car l'accès aux données sur le bien-être des populations futures est dans la pratique difficile. Par conséquent, l'étude de l'inégalité de bien-être dans une société à un instant t , sans considérer les populations futures, peut représenter un intérêt limité s'il est supposé que la situation inégalitaire actuelle influe sur les inégalités futures et que l'utilité des individus d'aujourd'hui et de demain est liée aux inégalités dans la société.

Ce premier chapitre a évoqué différents éléments à prendre en considération lorsque le choix d'un indice d'inégalité s'impose pour étudier la répartition des richesses dans une société. D'autre part, le choix d'un indice d'inégalité pour une étude peut aussi dépendre de sa possibilité à être décomposé en différents éléments permettant de mieux comprendre l'origine et la nature des inégalités. Cela peut notamment permettre d'estimer les contributions relatives des différents groupes composant la société aux inégalités, ou encore celle des différentes sources de revenus. La décomposition d'un indice d'inégalité est donc d'intérêt pratique dans l'évaluation des politiques publiques. C'est pourquoi le chapitre suivant s'attache à présenter les différentes méthodes de décomposition des indices d'inégalité développées dans la littérature. Par la suite, un outil est développé pour réaliser des décompositions d'indices d'inégalité par facteurs. Cet outil peut être utilisé avec divers indices d'inégalité, ainsi il semble opportun d'avoir pris connaissance des principaux indices d'inégalité et leurs propriétés. De plus, grâce à l'état de l'art des méthodes de décomposition effectué dans le chapitre 2, les avantages de l'outil pourront être mis en lumière.

Chapitre 2

Les décompositions des indices d'inégalité

Sommaire

2.1	Introduction	36
2.2	La décomposition par sous-populations	37
2.2.1	Décomposition agrégative, additive et cohérente	37
2.2.2	Décomposition additive et cohérente des indices d'entropie	38
2.2.3	Décomposition additive de l'indice de Gini	40
2.2.4	La décomposition faible des α – Gini	43
2.3	La décomposition par facteurs	46
2.3.1	Décomposition d'un indice d'inégalité par facteurs	46
2.3.2	La décomposition par facteurs de l'indice de Gini	48
2.3.3	La décomposition tridimensionnelle de la variation de l'indice de Gini	48
2.4	La décomposition multiple basée sur des méthodes statistiques : les régressions RIF	49
2.4.1	Effet de composition et effet structurel	50
2.4.2	Les régressions RIF	57
2.5	La décomposition de Shapley des indices d'inégalité	61
2.5.1	Introduction au concept de Shapley	62
2.5.2	Décomposition d'un modèle hiérarchique	64
2.5.3	Normalisation des distributions de facteurs équivalents	69
2.5.4	Incidence de l'indice d'inégalité sur la décomposition	71
2.6	Discussion et conclusion	72

2.1 Introduction

Dans la littérature économique, il existe deux approches principales pour décomposer les indices d'inégalité. La première approche, la décomposition par sous-population consiste à décomposer l'inégalité observée entre chaque groupe inclus dans la société. Par exemple, il peut être utile de savoir si les inégalités de revenus sont plus ou moins forte pour le groupe des femmes par rapport à celui des hommes, et de connaître la part des inégalités qui est due aux différences de revenus entre les femmes et les hommes. C'est deux effets sont notés respectivement effet intragroupe et effet intergroupe. L'autre approche dominante consiste à décomposer les inégalités entre facteurs (ou sources de revenu). Dans une perspective d'évaluation de politiques publiques, il peut être fort intéressant d'estimer les contributions aux inégalités des différentes sources de revenu (ex. travail, capital, transferts sociaux...). Une contribution quasi-nulle des transferts sociaux serait signe de son inefficacité à corriger les inégalités. Ces deux premières approches principales ont ensuite été combinées pour donner lieu à des décompositions multiples. La décomposition multiple est une approche qui prend en compte à la fois les effets de groupes et de sources de revenu. Ce chapitre aborde également les décompositions multiples qui prennent en considération des dimensions supplémentaires, notamment une dimension temporelle. Une approche alternative, notamment adoptée dans l'article Fourrey *et al.* (2018), propose une décomposition par attribut. Les attributs font référence aux caractéristiques personnelles qui peuvent être source d'inégalité de revenus. Cette approche est analogue à celle des facteurs, puisqu'elle décompose les inégalités de revenus entre différents attributs dont la somme des vecteurs est égale à la distribution des revenus. Elle permet aussi de traiter des effets de groupes en considérant les groupes comme un facteur, ainsi la décomposition par attributs s'inscrit dans le prolongement des deux décompositions classiques, par sous-populations et par facteurs. Toutefois, la décomposition par attribut nécessite généralement une étape préalable, qui est la désagrégation du salaire entre les différents attributs.

Par ailleurs, toutes les décompositions ne se basent pas sur la même méthodologie. Il existe différents procédés pour décomposer un indice d'inégalité, certains reposent sur des arrangements mathématiques des indices, d'autres utilisent des outils statistiques et enfin certains trouvent leur origine dans la théorie des jeux. Dans un premier temps, les décompositions par sous-population, sources de revenu et multiples qui reposent sur des arrangement mathématiques sont présentées. Ensuite, une décomposition statistique par facteur de l'indice de Gini sera évoquée. Enfin, la décomposition à la Shapley des indices d'inégalité est exposée, c'est une décomposition par attribut tirée de la théorie des jeux coopératifs. Cette dernière est présentée en détail et est l'objet principal de cette thèse. Cette méthode a plusieurs avantages, elle permet de réaliser des décompositions multiples, elle peut considérer divers effets de groupe en même temps et elle s'applique à une vaste variété d'indice d'inégalité.

2.2 La décomposition par sous-populations

2.2.1 Décomposition agrégative, additive et cohérente

Un indice d'inégalité peut être décomposé par sous-populations s'il existe un arrangement mathématique de la fonction de l'indice d'inégalité qui permet de différencier les inégalités provenant de la dispersion des revenus au sein des groupes (*i.e.* effet intragroupe) de celles issues des différences entre les groupes (*i.e.* effet intergroupe).

La décomposition additive des mesures d'inégalité en sous-groupes a été dans un premier temps abordée par Theil (1967) et fut axiomatisée et développée par Shorrocks (1980, 1984, 1988). Le principe d'additivité ne doit pas être confondue avec le principe plus faible d'agrégativité. La décomposition agrégative ou additive ont toutes les deux pour objectif de séparer l'effet des inégalités intergroupe de l'effet des inégalités intragroupe. La différence entre ces deux approches réside dans le traitement de l'inégalité intragroupe. Lorsque la décomposition est agrégative, la part des inégalités associée à un groupe est marginale, elle correspond à l'inégalité intragroupe observée moins l'inégalité intragroupe observée lorsque une distribution égalitaire est imposée dans le groupe en question. Le problème majeur de cette approche est que la somme de l'inégalité intragroupe mesurée dans chaque groupe n'est pas nécessairement égale à l'inégalité intragroupe totale. L'axiome de décomposition additive quant à lui impose que la somme de l'inégalité intragroupe soit égale à l'inégalité intragroupe totale. L'approche consiste à associer une pondération à la mesure de l'inégalité de chaque groupe, cette pondération peut correspondre à la part dans la population totale du groupe en question ou encore sa part de richesse détenue dans l'économie considérée (Bourguignon (1979)).

Par la suite Bourguignon (1979) et Shorrocks (1980) ont proposé un axiome d'additivité. Formellement, d'après Shorrocks (1980), lorsque la décomposition d'un indice d'inégalité respecte le principe d'additivité, alors :

$$I(\mathbf{x}) = \sum_g w_g(\mu, \mathbf{n}) I(\mathbf{x}^g) + I(\mu_1 \mathbf{1}_{n_1}, \dots, \mu_g \mathbf{1}_{n_g}), \quad (2.1)$$

Le premier terme de l'équation représente l'effet intragroupe et correspond à la somme des inégalités mesurées sur les distributions de chaque groupe $\mathbf{x}^g$, pondérées par un poids propre à chaque groupe, w_g , qui dépend du vecteur des moyennes observées dans les différents groupes, noté μ , et du vecteur du nombre d'individus dans chaque groupe, noté $\mathbf{n}$. Le second terme correspond à l'effet intergroupe qui est l'inégalité mesurée si chaque individu avait pour revenu le revenu moyen de son groupe μ_g . Une nouvelle distribution de revenus dans chaque groupe g est ainsi obtenue, composée de n_g fois le revenu moyen observé dans le groupe μ_g , notée $\mu_g \mathbf{1}_{n_g}$ avec $\mathbf{1}_{n_g}$ le vecteur unitaire de taille n_g .

Cet axiome d'additivité est plus restrictif que l'axiome d'agrégativité, mais moins restrictif que l'axiome d'additivité de Bourguignon (1979) dans lequel l'élément intragroupe correspondait à la somme des inégalités de chaque groupe pondérées par la part de revenu, et où seul l'indice de Theil respecte les conditions nécessaires à son application. Néanmoins, Shorrocks (1984) montre que sa notion d'additivité n'est pas pleinement satisfaisante, car la somme des poids associés à chaque groupe n'est pas toujours égale à l'unité. Lorsque la somme des poids est égale à l'unité, alors l'indice intragroupe sera une moyenne pondérée des indices d'inégalité mesurées dans chaque groupe. Si la somme n'est pas égal à l'unité, alors la différence entre l'unité et la somme des poids sera proportionnelle à l'inégalité intergroupe. Ainsi, les poids utilisés pour la mesure de l'inégalité intragroupe ne sont pas indépendants de l'inégalité intergroupe. Ce qui peut poser problème lorsque l'évaluateur veut estimer le niveau des inégalités si les différences entre les moyennes des groupes étaient nulles, c'est-à-dire s'il n'y avait plus d'inégalité intergroupe. Pour ce faire, il faudrait modifier les moyennes associées à chaque groupe, ce qui modifierait donc les poids de chaque groupe utilisés pour mesurer l'inégalité intragroupe puisqu'ils sont basés sur la moyenne. Ce qui induit une dépendance entre l'inégalité intergroupe et l'inégalité intragroupe, il est donc préférable que les poids ne dépendent pas de la moyenne.

Par la suite, Shorrocks (1988) propose le principe de cohérence en sous-groupes pour suppléer aux principes d'agrégation et d'additivité. Pour vérifier le principe de cohérence, un indice d'inégalité I s'écrit :

$$I = f(I_1, \dots, I_g, \dots, I_k; p_1, \dots, p_g, \dots, p_k; s_1, \dots, s_g, \dots, s_k) \quad (2.2)$$

Ainsi l'indice d'inégalité est une fonction de l'inégalité mesurée dans chaque groupe g , de la part de population de chaque groupe et de la part de revenu détenue par chaque groupe, notés respectivement I_g , p_g et s_g . Sachant que f est une fonction croissante de ses k premiers arguments. L'idée essentielle du concept de cohérence est que si les inégalités augmentent dans un groupe alors les inégalités globales augmentent aussi, toutes les autres distributions restant inchangées.

2.2.2 Décomposition additive et cohérente des indices d'entropie

L'indice d'Atkinson, l'écart-type et *a fortiori* le coefficient de variation ne sont pas décomposables additivement. Il en va de même de l'indice de Gini qui de plus n'est pas cohérent. Cependant, la sous-section suivante montre que dans un cas très restreint où les distributions associées à chaque groupe ne se chevauchent pas alors l'indice de Gini est aussi une mesure additivement décomposable (Ebert (1988)). Seuls les indices d'entropie respectent cette propriété de cohérence. D'ailleurs Cowell (1988) postule que l'indice de Gini, la variance des logarithmes et l'écart moyen relatif sont trois mauvais indicateurs d'inégalité puisqu'ils ne respectent pas ce principe de cohérence. De fait, les inégalités peuvent augmenter au sein de chaque groupe alors que l'inégalité totale diminue,

ceteris paribus.

Comme il a été évoqué précédemment, les indices d'entropie sont décomposables additivement, mais seuls l'indice premier et second de Theil ont une somme des pondérations égale à l'unité. La décomposition additive de l'indice premier de Theil nous donne :

$$T = \sum_{g=1}^G \frac{n_g \mu_g}{n \mu} T^g + \sum_{g=1}^G \frac{n_g \mu_g}{n \mu} \log \frac{\mu_g}{\mu}, \quad (2.3)$$

avec dans le premier terme, l'inégalité intragroupe qui est la somme des indices de Theil mesuré dans chaque groupe (T^g) pondérés par la part de revenu détenue par ce groupe ; et dans le second terme, l'effet intergroupe qui est la somme du logarithme du rapport entre le revenu moyen du groupe μ_g et le revenu moyen de tous les individus, pondérée par la part de revenu détenue par ce groupe.

La décomposition additive obtenue entre un effet intragroupe et un effet intergroupe de l'indice second de Theil est :

$$T^* = \sum_{g=1}^G \frac{n_g}{n} T^{*g} + \sum_{g=1}^G \frac{n_g}{n} \log \frac{\mu}{\mu_g}, \quad (2.4)$$

avec la somme des indices second de Theil dans chaque groupe, noté T^{*g} , pondérés par la part du groupe g dans la population pour l'effet intragroupe ; et la somme des écarts logarithmiques entre le revenu moyen et le revenu moyen de chaque groupe, pondérée par la part de ce groupe dans la population.

La décomposition additive de l'indice de Theil permet à la fonction d'évaluation sociale (FES) basée sur cet indice d'être aussi décomposable, ainsi la formule suivante est obtenue pour une société composée de g groupes de n_g individus :

$$V^T(\mathbf{y}) = \sum_{g=1}^G \frac{n_g}{n} \mu_g (1 - T(\mathbf{y}^g)) - \sum_{g=1}^G \frac{n_g}{n} \mu_g \log \frac{\mu_g}{\mu} \quad (2.5)$$

$$= \sum_{g=1}^G \frac{n_g}{n} V^T(\mathbf{y}^g) - V_B^T(\mu \mathbf{1}_n) \quad (2.6)$$

Le premier terme correspond à la somme des FES de Theil estimées sur la distribution de revenus de chaque groupe $\mathbf{y}^g$, pondérées par les parts de population respectives à chaque groupe. Le second terme correspond à la perte sociale engendrée par les inégalités intergroupes, et il est égal à la somme du logarithme des écarts du revenu moyen de chaque groupe μ_g à la moyenne μ , pondérée

par la part de revenus respectivement détenus par les groupes.

2.2.3 Décomposition additive de l'indice de Gini

Cependant certains auteurs rejettent le principe jugé trop restrictif de cohérence défini par Shorrocks (Pyatt (1976), Silber (1989), Yitzhaki et Lerman (1991), Dagum (1997a,b)).

Dagum (1997b) propose une décomposition par groupes de l'indice de Gini entre une composante intragroupe et une composante intergroupe brute. Cette dernière composante est nécessairement positive, même si les groupes ont les mêmes moyennes. La composante intergroupe brute est décomposable entre une composante intergroupe nette et une composante de transvariation (concept introduit par Gini (1921)). L'inégalité intergroupe nette est une mesure des inégalités moyennes entre les groupes, puisqu'elle mesure les inégalités générées par les revenus des groupes les plus riches en moyenne avec ceux des groupes les moins riches. La transvariation évalue les inégalités provenant de la zone de chevauchement entre les distributions des différents groupes étudiés, c'est-à-dire la zone où les plus hauts revenus du groupe au revenu moyen le plus faible sont supérieurs aux revenus les plus faibles du groupe à la moyenne la plus forte.

Pyatt (1976) avait proposé une décomposition en deux termes de l'indice de Gini, le premier terme, l'effet intragroupe, comprenait l'effet de chevauchement¹. Ainsi la décomposition de Pyatt (1976) en deux termes et celle de Dagum (1997b) en trois termes sont similaires uniquement s'il n'y a pas de chevauchement entre les distributions de revenu des sous-populations.

Formellement, Dagum (1997b, 1987) pose 8 définitions pour obtenir une décomposition de l'indice de Gini en trois termes :

— **Définition 1 :** La moyenne absolue des différences de revenus pour chaque paire de revenu entre deux groupes j et h est notée :

$$\Delta_{jh} = \frac{\sum_{i=1}^{n_j} \sum_{r=1}^{n_h} |y_{ij} - y_{rh}|}{n_j * n_h} = E|Y_j - Y_h|$$

avec Y_i le vecteur de revenus des individus inclus dans le groupe i .

— **Définition 2 :** L'indice de Gini entre deux groupes est noté :

$$G_{jh} = \frac{\Delta_{jh}}{\mu_j + \mu_h}$$

Avec μ_i la moyenne des revenus dans le groupe i . De plus, $G_{jh} = G_{hj}$ et $\Delta_{jh} = \Delta_{hj}$.

1. Silber (1989) utilise la même approche et sépare l'effet du chevauchement des autres effets.

— **Définition 3 :** Un groupe j est plus influent qu'un groupe h si $\mu_j > \mu_h$. La relation « plus influent que » définit un ordre partiel strict sur l'ensemble des paires de groupes. Il est ainsi asymétrique et transitif.

— **Définition 4 :** L'influence économique brute d_{ij} est une moyenne pondérée des différences de revenu $y_{jh} - y_{hr}$ pour tous les revenus y_{jh} des membres du groupe j qui ont un revenu supérieur à un revenu y_{hr} des membres du groupe h , tel que j est plus influent qu'un groupe h :

$$d_{jh} = \int_0^{\inf} dF_j(y) \int_0^y (y-x) dF_h(x)$$

où $F(.)$ est la fonction de distribution cumulative.

— **Définition 5 :** Le moment de premier ordre de transvariation p_{jh} entre le groupe j et h , avec $\mu_j > \mu_h$, est une moyenne pondérée des différences de revenu $y_{hr} - y_{ji}$ pour toutes les paires de revenus, sachant que l'un des revenus est tiré du $h^{ème}$ groupe tel que $y_{hr} > y_{ji}$:

$$p_{jh} = \int_0^{\inf} dF_h(y) \int_0^y (y-x) dF_j(x)$$

La différence moyenne entre le groupe j et le groupe h est :

$$d_{jh} + p_{jh} = \Delta_{jh}$$

D'après les définitions 4 et 5 :

$$0 \leq p_{jh} \leq \frac{1}{2} \Delta_{jh} \leq d_{jh} \leq \Delta_{jh}$$

De plus, $p_{jh} = 0$ et $d_{jh} = \Delta_{jh}$ si les deux distributions ne se chevauchent pas. Et $p_{jh} = d_{jh} = \frac{1}{2} \Delta_{jh}$ si $\mu_j = \mu_h$.

— **Définition 6 :** L'influence économique nette entre deux sous-populations j et h , sachant que $\mu_j > \mu_h$ est $d_{jh} - p_{jh}$. L'influence économique nette est non-négative et son maximum possible est Δ_{jh} .

— **Définition 7 :** L'influence économique relative, D_{jh} , entre les groupes j et h , avec $\mu_j > \mu_h$, est le ratio entre l'influence économique net et son maximum possible :

$$D_{jh} = \frac{d_{jh} - p_{jh}}{\Delta_{jh}} = \frac{d_{jh} - p_{jh}}{d_{jh} + p_{jh}}$$

Cette mesure peut être nulle, puisque d_{jh} , p_{jh} et Δ_{jh} sont compris entre $[0, 1]$, et que $p_{jh} \leq d_{jh}$.

— **Définition 8 :** Le terme $G_{jh} * D_{jh}$ mesure la contribution nette à l'inégalité totale de l'effet in-

tergroupe. Le terme $G_{jh} * (1 - D_{jh})$ mesure la contribution de la transvariation de revenu entre le groupe j et h . La somme de ces deux termes mesure la contribution brute de l'inégalité intergroupe à l'inégalité totale.

À partir de l'équation de l'indice de Gini inter-individuel :

$$G = \frac{1}{2n^2\mu} \sum_{i=1}^n \sum_{r=1}^n |y_i - y_r| \quad (2.7)$$

et des huit définitions précédentes, l'indice de Gini peut être décomposé en trois termes :

$$G = \sum_{j=1}^k G_{jj} p_j s_j + \sum_{j=2}^k \sum_{h=1}^{j-1} G_{jh} (p_j s_h + p_h s_j) \quad (2.8)$$

$$= \sum_{j=1}^k G_{jj} p_j s_j + \sum_{j=2}^k \sum_{h=1}^{j-1} G_{jh} D_{jh} (p_j s_h + p_h s_j) + \sum_{j=2}^k \sum_{h=1}^{j-1} G_{jh} (1 - D_{jh}) (p_j s_h + p_h s_j) \quad (2.9)$$

$$= \text{Effet Intragroupe} + \text{Effet Intergroupe net} + \text{Effet Transvariation} \quad (2.10)$$

Ainsi, l'équation 2.8 est adaptée s'il n'y a pas de chevauchement dans les distributions, puisque le moment de premier ordre de transvariation p_{jh} entre le groupe j et h serait nul pour toutes paires de groupes et l'influence économique relative D_{jh} serait égale à 1. Dans le cas contraire, l'équation 2.9 est nécessaire pour prendre en compte l'effet de transvariation.

Prenons l'exemple suivant (Tableau 2.1) avec les revenus d'une population, notée $\mathbf{Y}$, répartie en trois groupes, tel que $\mathbf{Y} = (\mathbf{y}_1, \mathbf{y}_2, \mathbf{y}_3)$:

Tableau 2.1 – Distribution de revenus

Identifiant	$\mathbf{y}_1$	$\mathbf{y}_2$	$\mathbf{y}_3$
1	20	30	70
2	30	30	80
3	35	40	90
4	55	40	100
Revenu moyen	35	35	85

Les résultats de la décomposition de Gini en trois termes sont données dans le tableau 2.2. Ces résultats indiquent que l'effet des inégalités intragroupes est plus important dans les groupes 1 et 3. De plus, les inégalités intergroupe nettes du groupe 3 avec le groupe 2 et le groupe 1 sont les principales sources de l'inégalité totale (39.41% de l'inégalité totale pour chacun de ces effets). Ces deux effets sont égaux puisque les revenus moyens du groupe 1 et 2 sont les mêmes, c'est

Tableau 2.2 – Décomposition de Gini en trois termes

	Groupe 1	Groupe 2	Groupe 3
Effet Intragroupe	0,014785	0,005376	0,013441
	Groupe 1 / Groupe 2	Groupe 1 / Groupe 3	Groupe 2 / Groupe 3
Effet Intergroupe net	0	0,107527	0,107527
Effet Transvariation	0,024194	0	0

pour cette même raison que l'effet intergroupe entre les deux premiers groupes est nul. L'effet de transvariation du groupe 3 avec le groupe 1 et 2 est nul car il n'y a pas de chevauchement des distributions de l'un avec l'autre. Toutefois, l'effet de transvariation est positif entre le groupe 1 et 2, car les distributions se chevauchent.

Par ailleurs, Yitzhaki et Lerman (1991) proposent eux-aussi une décomposition en trois groupes qui se différencie sur deux points. Premièrement, l'effet intergroupe ne résulte pas des différences entre les revenus moyens de chaque groupes, mais du rang moyen entre chaque groupe. Ainsi, il peut y avoir un effet intergroupe nul alors que les moyennes sont différentes (et que les rangs moyens sont similaires). Le troisième terme lié au chevauchement (de « stratification » pour les auteurs) dépend lui aussi du rang, une forte stratification implique une faible variabilité de rang entre les groupes, et donc exerce un effet négatif sur les inégalités.

L'intérêt essentiel d'une décomposition en trois termes de l'indice de Gini est qu'elle ne se limite pas à la comparaison de moyenne, elle prend en considération les inégalités de revenus entre des individus qui appartiennent à des groupes différents (*i.e.* entre les riches des groupes pauvres et les pauvres des groupes riches). Ainsi l'hétérogénéité qui existe au sein de chaque groupe n'est pas masquée par la moyenne dans la comparaison intergroupe².

2.2.4 La décomposition faible des α – Gini

La décomposition entre trois composantes de Dagum a ensuite été généralisée par Chamani (2006a,b) et Ebert (2010) à l'ensemble des indices de la famille des α – Gini. D'après Ebert (2010), la décomposition additive est trop restrictive car le terme intergroupe repose seulement sur la moyenne de chaque groupe. Il propose alors une nouvelle décomposition en deux termes, le premier correspond à la composante intragroupe, identique à celle de Shorrocks (1980), et le second correspond à la composante intergroupe qui repose sur l'ensemble des comparaisons par paires de revenus individuels. Pour Ebert, une mesure d'inégalité faiblement décomposable s'écrit

2. Pour plus de détails sur la décomposition de Dagum et des autres méthodes de décompositions de l'indice de Gini, voir par exemple Mussard *et al.* (2006), Mussard et Terraza (2009) et Mornet *et al.* (2014).

pour tout $n = (n_1 + n_2)$ avec $n_1 \geq 1$ et $n_2 \geq 1$:

$$I(\mathbf{x}^1, \mathbf{x}^2, n_1 + n_2) = \gamma_1(\mathbf{n}) \frac{\mu_1^\alpha}{\mu^\alpha} I(\mathbf{x}^1, n_1) + \gamma_2(\mathbf{n}) \frac{\mu_2^\alpha}{\mu^\alpha} I(\mathbf{x}^2, n_2) + \beta(\mathbf{n}) \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} \frac{\mu_{1i,2j}^\alpha}{\mu^\alpha} I(x_i^1, x_j^2, 2) \quad (2.11)$$

avec $\gamma_1(\mathbf{n})$, $\gamma_2(\mathbf{n})$ et $\beta(\mathbf{n})$ des fonctions de pondérations strictement croissantes telles que $\mathbf{n} = (n_1, n_2)$ est le vecteur des tailles des sous-groupes de la population. Le vecteur de revenu associé au groupe j est noté $\mathbf{x}^j$ tel quel $\mathbf{x}^j = (x_1^j, \dots, x_{n_j}^j)$. Le paramètre α représente un paramètre de sensibilité assimilé au degré d'aversion pour l'inégalité, avec $\alpha \in \mathbb{R}_+$ ³.

La décomposition forte à la Bourguignon (1979) et Shorrocks (1980, 1984) et la décomposition faible de Ebert (2010) proposent toutes les deux un effet intragroupe comme le résultat d'une somme pondérée des inégalités au sein de chaque groupe. Cependant, dans la première approche, l'effet intergroupe est estimé à partir d'un revenu de référence pour chaque groupe (ex. : la moyenne), alors que dans la seconde approche, l'effet intergroupe dépend des comparaisons des revenus des individus de groupes différents.

De plus, les mesures faiblement décomposables respectent un principe de concentration. Le principe de concentration postule qu'une augmentation de la concentration des revenus (réduction de la distance entre les revenus et la moyenne) diminue le niveau d'inégalité. C'est un principe plus faible que celui des transferts de Pigou-Dalton puisqu'il n'impose pas qu'une plus grande concentration soit la conséquence d'un transfert d'un riche vers un pauvre, tout en préservant leur rang. Autrement dit, le principe de concentration prend en considération tous les individus de la société, alors que le principe de transfert est défini à partir d'un transfert entre deux individus.

À partir de cette propriété, une nouvelle classe de mesures est définie, appelée $\alpha - Gini$. Cette famille comprend l'indice de Gini ($\alpha = 1$) et également le coefficient de variation au carré ($\alpha = 2$) qui appartient à la famille des mesures dérivées de l'entropie généralisée introduite par Shorrocks (voir section 1.4.3). Pour $\alpha \geq 1$, cet indice respecte les différents principes usuels : la continuité, la normalisation ⁴, la symétrie, l'invariance relative, l'invariance à une réplique et le principe de transfert. La décomposition en sous-groupes de Dagum a ensuite été étendue aux mesures de type $\alpha - Gini$ par Chameni (2011). Chameni part de la formulation de l'indice de Gini interpersonnel suivante :

3. Ebert (2010) propose deux formulations pour une décomposition faible, la première présentée ici concerne les indices « compromis », c'est-à-dire ceux dont l'estimation dépend directement de la moyenne (ex. : Indice de Gini, Coefficient de variation au carré). Une seconde formulation concerne des indices absolus et relatifs. Par la suite, Mornet (2016) propose une formulation de la décomposition faible qui ne dépend pas de ces conditions d'invariance, avec une pondération qui dépend à la fois du nombre d'individus dans chaque groupe et de leur revenu moyen respectif.

4. L'indice a donc une borne inférieure égale à zéro. Cependant, contrairement à l'indice de Gini classique, sa borne supérieure peut excéder largement 1.

$$G^\alpha = \frac{1}{2n^2\mu^\alpha} \sum_{i=1}^n \sum_{r=1}^n |x_i - x_r|^\alpha \quad (2.12)$$

L'indice de Gini entre deux sous-populations est donné par :

$$G_{jh}^\alpha = \frac{\sum_{i=1}^{n_j} \sum_{r=1}^{n_h} |x_{ij} - x_{rh}|^\alpha}{n_j n_h (\mu_j^\alpha + \mu_h^\alpha)} \quad (2.13)$$

$\forall j, h \in 1, \dots, k$, G_{jj}^α correspondant au cas particulier où $j = h$.

À partir de cette formulation, une décomposition sur une population P entre k sous-groupes est proposée :

$$G^\alpha = \sum_{j=1}^k G_{jj}^\alpha p_j s_j + \sum_{j=2}^k \sum_{h=1}^{j-1} G_{jh}^\alpha D_{jh} (p_j s_h^\alpha + p_h s_j^\alpha) + \sum_{j=2}^k \sum_{h=1}^{j-1} G_{jh}^\alpha (1 - D_{jh}) (p_j s_h^\alpha + p_h s_j^\alpha) \quad (2.14)$$

avec $k \in 1, \dots, n$, p_j et s_j les parts respectives de la population j dans la population totale et la part des richesses détenue par j . Si $\alpha = 1$ alors cette décomposition revient à la décomposition de Gini en trois termes par Dagum. De plus, si $\alpha > 1$ alors les différents termes de la décomposition et le coefficient total $\alpha - Gini$ ne sont plus compris entre 0 et 1.

Parallèlement, dans la continuité des travaux de Dagum (1997a,b), Mussard *et al.* (2006) et Chameni (2011), Mornet *et al.* (2013) proposent une décomposition multi-niveaux (α, β) de l'indice α -Gini. Cette décomposition dépend de deux paramètres de sensibilité, α et β . L'aversion aux inégalités intragroupe dépend du premier paramètre, alors que la sensibilité aux inégalités intergroupe (*i.e.* effet intergroupe net et l'effet de transvariation) dépend des deux paramètres. Autrement dit, α représente l'aversion aux inégalités et β la sensibilité de l'effet intergroupe à la non transvariation (*i.e.* à la distance entre les distributions de revenus de chaque groupe). C'est une décomposition par sous-population multi-niveaux qui permet de décomposer l'indice α -Gini par sexe, puis par âge au sein de chaque groupe de sexe (Mornet *et al.* (2013)).

Certains auteurs privilégient les indices d'entropie généralisés puisqu'ils sont réguliers (*i.e.* ils respectent les axiomes de normalisation, le principe des transferts, la symétrie, l'invariance à une réplique), qu'ils sont invariants à une translation (homogène de degré 0) et additivement décomposables (donc cohérents). Or, seuls les indices premier et second de Theil ont une somme des poids intragroupe égale à l'unité, et donc permettent l'indépendance entre la mesure de l'inégalité intragroupe et de l'inégalité intergroupe. Enfin, seul l'indice second de Theil permet de respecter les principes énoncés précédemment ainsi que le principe de monotonie à la distance (Cowell et Flachaire (2018)). Toutefois, l'indice de Gini a aussi ses avantages. Il est lui aussi régulier, il est généralisable et permet ainsi d'inclure un paramètre d'aversion aux inégalités. Il est décomposable en trois termes, ce qui permet de prendre en considération l'effet de chevauchements entre les distribu-

tions des sous-populations. Néanmoins, il ne respecte pas le principe de monotonie à la distance et ne permet pas une décomposition cohérente au sens de Shorrocks. Ainsi, s'il est accepté qu'il peut y avoir des situations où l'inégalité diminue lorsque le revenu d'un individu s'écarte de la valeur de référence (ex. : la moyenne), alors l'indice de Gini peut être adapté à l'étude menée. D'autant plus, si un intérêt particulier est porté à l'effet de transvariation dans une étude de décomposition des inégalités par sous-populations. Cependant, si l'objectif est de mener une décomposition additive en deux termes (effets intragroupe et effets intergroupe) et cohérente avec un indice d'inégalité qui respecte le principe de monotonie à la distance, alors le second indice de Theil sera plus adapté.

2.3 La décomposition par facteurs

2.3.1 Décomposition d'un indice d'inégalité par facteurs

Cette section introduit le concept de la décomposition par facteurs, c'est-à-dire que l'objectif est ici d'établir la contribution de chaque facteur aux inégalités de revenu. Le revenu de chaque individu est composé de plusieurs sources de revenu :

$$\mathbf{x} = (y_1, \dots, y_n) = \left(\sum_{i=1}^m y_1^i, \dots, \sum_{i=1}^m y_n^i \right), \quad (2.15)$$

où chaque revenu individuel est composé de m sources de revenu.

Les termes facteurs et sources de revenu sont utilisés sans faire de distinction. Toutefois, le terme facteur sera plus approprié pour parler des caractéristiques individuelles (e.g. âge, genre...). Dans les deux cas, la somme des distributions associées à chaque facteur ou à chaque source de revenu doit être égale à la distribution du revenu total.

Shorrocks (1982) pose six axiomes que devraient respecter la décomposition d'un indice d'inégalité par facteurs :

Axiomes (Shorrocks 1982).

1. $I(\mathbf{x})$: Symétrique, continue et normalisé

L'indice d'inégalité $I(\mathbf{x})$ est une fonction continue et symétrique. De plus, la normalisation indique que si tous les revenus sont égaux alors l'indice d'inégalité est nul.

2. Continuité et traitement symétrique des facteurs

L'ensemble de facteurs, noté M , est composé de m facteurs disjoints et exhaustifs. La contribution du facteur i à $I(\mathbf{x})$ peut être représenté par $S_i(x^1, \dots, x^i, \dots, x^m; M)$, une fonction à deux arguments, le premier représente les distributions de revenus associées à chaque facteur i et le second l'ensemble des facteurs considérés. Ces fonctions sont supposées continues et

les facteurs traités symétriquement (peu importe l'ordre dans lequel sont considérés les facteurs).

3. Indépendance par rapport au niveau de désagrégation

La contribution d'un facteur ne doit pas dépendre de la manière dont les autres facteurs sont groupés.

4. Décomposition cohérente⁵

La somme des contributions, S_i , associées à chaque facteur i est égale l'inégalité totale. Ainsi, $\sum_{i=1}^m S_i(x^1, \dots, x^m; M) = \sum_{i=1}^m S_i(x^1, \dots, x^m; \mathbf{x}) = I(\mathbf{x})$.

5. Symétrie de la population et normalisation des distributions de facteurs équivalents

L'ordre par lequel les individus sont classés n'importe pas dans l'estimation de la contribution des facteurs. Et la contribution d'un facteur au niveau d'inégalité totale est nulle si son montant distribué est le même pour tous les individus.

6. Symétrie de deux facteurs

Ceci est une restriction au principe de traitement symétrique des facteurs énoncé précédemment. Les contributions de deux facteurs, dont l'un est une permutation de l'autre, sont égales lorsqu'il n'y a que deux facteurs. Cependant, elles seront différentes s'il y a plus de deux facteurs car la corrélation des deux facteurs en question avec les autres ne sera pas la même.

Si l'ensemble des axiomes de Shorrocks sont vérifiés, cela implique le théorème suivant :

Théorème 3 (Shorrocks 1982). Une règle de décomposition par sources de revenu qui respectent les six axiomes énoncés par Shorrocks (1982) implique la relation suivante :

$$s^i(I) = \frac{S^i(\mathbf{x}^i, \mathbf{x})}{I(\mathbf{x})} = \frac{\text{Cov}(\mathbf{x}^i, \mathbf{x})}{\sigma^2(\mathbf{x})}$$

Avec $s^i(I)$ la contribution relative d'un facteur i à l'indice $I(\mathbf{x})$, et $S^i(\mathbf{x}^i, \mathbf{x})$ la contribution absolue qui dépend de la distribution de la source i et de la distribution du revenu total $\mathbf{x}$. Comme l'a mis en avant Mussard (2007), ce théorème implique trois résultats. Premièrement, la règle de décomposition est unique. Deuxièmement, la contribution relative de chaque facteur à l'inégalité totale est indépendante de l'indice d'inégalité choisi. Ainsi, peu importe l'indicateur sélectionné, si celui-ci permet le respect des six axiomes de Shorrocks alors la contribution relative de chaque

5. La décomposition cohérente de Shorrocks (1982) n'est pas la même que celle de Shorrocks (1988). Les deux font référence à une notion d'additivité, mais la dernière inclut une relation croissante avec ses premiers arguments (Équation 2.2).

facteur sera la même que celle obtenue avec la covariance et la variance. Troisièmement, la variance et la covariance respectent les six conditions exposées par Shorrocks.

Ainsi, la décomposition naturelle de la variance par sources de revenu permet de respecter les différents axiomes définis par Shorrocks (1982). Cependant, comme le premier chapitre s'est attaché à le montrer, la variance n'est pas un indice d'inégalité mais de dispersion. Par ailleurs, certains axiomes sont discutables, comme par exemple l'hypothèse d'indépendance par rapport au niveau de désagrégation. Après tout, la contribution d'un facteur aux inégalités pourrait être appréciée relativement aux autres facteurs. En résumé, il est nécessaire de renoncer à quelques-uns de ces principes pour décomposer d'autres indicateurs d'inégalité par sources de revenu, ce qui apparaît dans la partie suivante avec la décomposition de l'indice de Gini.

2.3.2 La décomposition par facteurs de l'indice de Gini

Dans la continuité de Fei *et al.* (1978), Fields (1979) et Pyatt *et al.* (1980)⁶, Lerman et Yitzhaki (1985) ont proposé une décomposition de l'indice de Gini par sources de revenu. Ainsi l'indice de Gini mesuré sur la distribution du revenu total, G , peut être décomposé entre m sources de revenu :

$$G = \sum_{i=1}^m S_i G_i R_i \quad (2.16)$$

Avec S_i la part de la source i dans le revenu total, G_i l'indice de Gini mesuré sur la distribution de cette source et R_i la corrélation de la source i avec la distribution du revenu total. Par ailleurs, Stark *et al.* (1986) ont montré que la dérivée partielle du coefficient de Gini à une faible variation ϵ de la source i est égale à :

$$\frac{\partial G}{\partial \epsilon} = S_i (G_i R_i - G)$$

Le variation en pourcentage de l'indice suite à une petite variation en pourcentage de la distribution de la source i est égale à la contribution originale de la source i aux inégalités moins la part de la source i dans le revenu total :

$$\frac{\partial G / \partial \epsilon}{G} = \frac{S_i G_i R_i}{G} - S_i$$

La décomposition par facteurs de Lerman et Yitzhaki (1985) (Équation 2.16) ne respecte pas les axiomes de Shorrocks qui supposent une indépendance par rapport au rang puisque l'indice de Gini dépend du classement des ménages selon leur revenu, cependant ces auteurs estiment que leur décomposition du Gini reste « désirable ». Pour Lerman et Yitzhaki, l'indice de Gini respectent des concepts essentiels, tels que la dominance stochastique⁷ et le principe des transferts.

6. Voir Mussard et Terraza (2009) pour un historique plus complet.

7. C'est-à-dire que l'indice de Gini respecte les conditions nécessaires pour évaluer qu'une distribution $\mathbf{x}_1$ est plus

2.3.3 La décomposition tridimensionnelle de la variation de l'indice de Gini

La décomposition multiple a pour objectif d'estimer l'effet de plusieurs dimensions sur l'inégalité. Ainsi, cette décomposition peut proposer une décomposition par sources de revenu et/ou par sous-populations, et prendre en compte en même temps une dimension temporelle (ex. : Mussard et Savard (2012)). Mussini (2013) utilise une décomposition matricielle pour décomposer la variation de l'indice Gini selon trois dimensions : le temps, les sources de revenu et les sous-populations. L'effet temporel peut lui aussi être décomposé entre un effet de reclassement des individus (*i.e.* de changement de rang) dans la distribution de revenu et un effet de croissance des revenus disproportionnelle. Puisque l'indice de Gini est relatif, ce dernier effet serait nul si tous les individus expérimentaient la même croissance de leur revenu.

Pour m sources de revenu, k sous-populations, Mussini (2013) obtient donc l'équation suivante pour décomposer en trois dimensions le changement observé entre deux périodes de l'indice de Gini :

$$\Delta G = \sum_{i=1}^m \sum_{g=1}^k [R_g(Y^i) - P_g(Y^i)] + \sum_{i=1}^m \sum_{g=2}^k \sum_{h=1}^{k-1} [R_{gh}(Y^i) - P_{gh}(Y^i)] \quad (2.17)$$

Avec $R_g(Y^i)$ qui représente l'effet de reclassement dans la distribution de la source i dans le groupe g , et $R_{gh}(Y^i)$ qui estime l'effet intergroupe du reclassement par paire de groupes g, h . Cet effet est toujours non négatif, et si tous les individus ne changent pas de rang en passant d'une distribution à une autre, alors cet effet est nul. $P_g(Y^i)$ exprime l'effet de la croissance disproportionnelle pour la source i dans le groupe g , et $P_{gh}(Y^i)$ mesure cet effet intergroupe par paire. Cet effet peut être positif ou négatif. S'il est positif alors les disparités relatives de revenu diminuent entre les deux périodes, diminuant ainsi l'inégalité. Au contraire, s'il est négatif, alors la croissance des revenus disproportionnelle a induit plus de disparités et a donc participé à l'augmentation des inégalités.

2.4 La décomposition multiple basée sur des méthodes statistiques : les régressions RIF

Lorsque l'objectif de la décomposition d'un indice d'inégalité de revenu est d'estimer la contribution de chaque caractéristique personnelle (*i.e.* attribut), la première étape consiste à désagréger le revenu et à le répartir entre les différents attributs. Si l'évaluation étudie les différences observées dans le niveau d'inégalité entre deux groupes ou entre deux périodes, une étape concomitante ou préalable consiste à désagréger la distribution des revenus entre un effet composition et un effet

inégalitaire qu'une distribution $\mathbf{x}_2$ lorsque la première a une courbe de Lorenz qui domine celle de l'autre (Yitzhaki 1982).

structurel, et à estimer la contribution de chaque attribut à chaque effet.

L'effet structurel fait référence à la variation de la statistique étudiée (ex. : moyenne, variance...) qui est imputable à des variations dans les rendements des caractéristiques individuelles entre deux groupes ou deux périodes. Alors que l'effet de composition estime la part de la variation qui est due à un changement dans la distribution des caractéristiques individuelles (ex. : un groupe plus diplômé qu'un autre).

Dans un premier temps, cette section présente l'approche standard qui permet une distinction entre effet structurel et effet de composition d'après les travaux de Oaxaca (1973) et Blinder (1973), puis une extension de cette approche est proposée avec les méthodes de repondération (DiNardo *et al.* (1996)). Dans un second temps, l'approche par les régressions de fonctions d'influences recentrées (RIF) (Firpo *et al.* (2007, 2009, 2018)) est détaillée. Celle-ci permet notamment de décomposer par attributs l'indice de Gini, tout en différenciant les effets de composition et structurels. Une présentation plus exhaustive et détaillée de ces méthodes est disponible grâce à Fortin *et al.* (2011).

2.4.1 Effet de composition et effet structurel

La décomposition à la Oaxaca-Blinder

Les premiers auteurs à proposer cette approche ont été Oaxaca (1973) et Blinder (1973) avec la fameuse décomposition à la Oaxaca-Blinder (OB par la suite). Cette décomposition permet d'une part de décomposer la variation du salaire moyen entre un effet composition et un effet structurel, et d'autre part d'évaluer la contribution de chaque variable explicative du salaire pour chaque effet.

Dans un premier temps Y_{ti} est défini, le salaire de l'individu i dans le groupe g avec $g = 0, 1$. Le salaire d'un individu peut uniquement être observé dans un des deux groupes, donc Y_{0i} ou Y_{1i} . Ainsi, $Y_i = Y_{1i} \cdot G_i + Y_{0i} \cdot (1 - G_i)$ avec $G_i = 1$ si l'individu i appartient au groupe 1 et $G_i = 0$ s'il appartient au groupe 0. Il existe aussi un vecteur de variables explicatives $X \in \mathcal{X} \subset \mathbb{R}^K$ observables dans chaque groupe et une fonction de rendements $\pi(\cdot)$ qui associe pour chaque groupe les variables observées X et les variables inobservées ϵ à la distribution des salaires F_Y tel que $Y_{i,g} = \pi_g(X_i, \epsilon_i)$ pour $g = 0, 1$.

Dans l'approche standard d'OB, le revenu d'un individu i dans un groupe g est une fonction linéaire des variables explicatives individuelles et d'un terme d'erreur :

$$Y_{ti} = X_i' \beta_g + \epsilon_{ti}, \quad (2.18)$$

avec une espérance du terme d'erreur conditionnellement à X et G nul : $E[\epsilon_{ti}|X_i, G = g] = 0$ ⁸, et β_g

8. Cela suppose notamment que l'équation de salaire soit bien spécifiée et qu'il n'y ait pas de variables omises.

les rendements linéaires associés aux caractéristiques individuelles au sein du groupe g .

La différence totale de salaire moyen entre les deux groupes est notée $\Delta_O^\mu = E[Y|G = 1] - E[Y|G = 0]$ et peut être décomposée entre un effet structurel et un effet composition :

$$\begin{aligned}\Delta_O^\mu &= E[Y|G = 1] - E[Y|G = 0] \\ &= E[E(Y|X, G = 1)|G = 1] - E[E(Y|X, G = 0)|G = 0] \\ &= (E[X|G = 1]'\beta_1 + E[\epsilon_1|G = 1]) - (E[X|G = 0]'\beta_0 + E[\epsilon_0|G = 0]),\end{aligned}\tag{2.19}$$

avec $E[\epsilon_g|G = g] = 0$, en supposant que le terme d'erreur obtenu dans chaque groupe g a une espérance conditionnelle nulle $E[\epsilon_g|X, G = g]$, avec $g = 0, 1$. En ajoutant et en soustrayant le salaire moyen contrefactuel que le groupe 1 aurait obtenu avec les rendements du groupe 0, $E[X|G = 1]'\beta_0$, on a alors :

$$\begin{aligned}\Delta_O^\mu &= E[X|G = 1]'\beta_1 - E[X|G = 0]'\beta_0 \\ &= E[X|G = 1]'\beta_1 - E[X|G = 1]'\beta_0 + E[X|G = 1]'\beta_0 - E[X|G = 0]'\beta_0 \\ &= \underbrace{E[X|G = 1]'\beta_1 - E[X|G = 1]'\beta_0}_{\Delta_{S,OB}^\mu} + \underbrace{E[X|G = 1]'\beta_0 - E[X|G = 0]'\beta_0}_{\Delta_{X,OB}^\mu} \\ &= \Delta_{S,OB}^\mu + \Delta_{X,OB}^\mu,\end{aligned}\tag{2.20}$$

Où le premier terme de l'équation représente l'effet structurel et le second l'effet composition. Ces deux effets peuvent être exprimés comme la somme de la contribution de chaque variable à chaque effet :

$$\Delta_{S,OB}^\mu = \sum_{k=1}^K E[X^k|G = 1]'\beta_{1,k} - \beta_{0,k}\tag{2.21}$$

$$\Delta_{X,OB}^\mu = \sum_{k=1}^K [E[X^k|G = 1] - E[X^k|G = 0]]'\beta_{0,k}\tag{2.22}$$

L'effet structurel reflète uniquement une différence dans les rendements des caractéristiques individuelles, observées et inobservées, seulement s'il est possible d'établir la distribution des variables observées et inobservées comme celle du groupe 1. Ceci est possible sous deux hypothèses : l'ignorance et le support commun. La première stipule que la distribution des variables inobservées sachant X ne dépend pas de l'appartenance à un groupe. Formellement :

Hypothèses. [Ignorance] : Si (T, X, ϵ) ont une distribution jointe. Pour tout x dans $\mathcal{X}$: ϵ est indépendant de T sachant $X = x$.

La seconde hypothèse suppose que les différentes valeurs que peuvent prendre les caractéristiques individuelles sont observées dans les deux groupes. Formellement :

Hypothèses. [Support commun] : Pour tout x dans $\mathcal{X}$, $p(x) = P[G = 1|X = x] < 1$ et $P[G = 1] > 0$.

De fait, la décomposition à la OB est une approche simple qui permet d'évaluer la contribution de chaque caractéristique au revenu moyen pour chaque effet, structurel et de composition. Il suffit de remplacer les moyennes et les estimations par les moindres carrés ordinaires (MCO) des rendements $\hat{\beta}_g$ dans l'équation 2.23⁹, on a alors :

$$\begin{aligned}\hat{\Delta}_O^\mu &= \bar{X}_1\hat{\beta}_1 - \bar{X}_1\hat{\beta}_0 + \bar{X}_1\hat{\beta}_0 - \bar{X}_0\hat{\beta}_0 \\ &= \underbrace{\bar{X}_1(\hat{\beta}_1 - \hat{\beta}_0)}_{\Delta_{S,OB}^\mu} + \underbrace{(\bar{X}_1 - \bar{X}_0)\hat{\beta}_0}_{\Delta_{X,OB}^\mu} \\ &= \Delta_{S,OB}^\mu + \Delta_{X,OB}^\mu\end{aligned}\tag{2.23}$$

Cependant, quelques limites de cette approche sont à noter.

Premièrement, des difficultés dans l'interprétation des résultats peuvent émerger quand les variables explicatives inclues des variables catégorielles (ex. : le niveau de diplôme, genre, etc.). Pour intégrer une variable catégorielle dans une régression, un point de référence pour ces variables doit être choisi arbitrairement en omettant une des catégories. L'effet attribué à une catégorie est alors dépendant de la catégorie de référence, par exemple le surplus de salaire du fait d'avoir un diplôme du supérieur par rapport à un niveau BAC ne sera pas le même que par rapport à ne pas être diplômé. Par conséquent, lorsqu'une décomposition détaillée de l'effet structurel est menée, le choix de la catégorie de référence influence la contribution estimée de la variable explicative à cet effet. En conséquence, l'effet de composition et structurel total ne sont pas affectés par le choix de la catégorie de référence mais les effets détaillés le sont (Oaxaca et Ransom (1999), Gardeazabal et Ugidos (2004))¹⁰. Dans le cas précédent, le groupe de référence pour l'effet structurel était le groupe 0 mais cela aurait aussi bien pu être le groupe 1. Cette sensibilité des contributions des variables à l'effet structurel au groupe de référence doit être prise en compte lors de l'interprétation des résultats.

De plus, Firpo *et al.* (2007, 2018) privilégient l'utilisation du paramètre β_C à β_0 dans l'estimation de l'effet structurel (Équation 2.20) puisque l'estimation des paramètres dépend de la composition des distributions X . Ainsi la différence $(\beta_1 - \beta_C)$ reflète uniquement une différence de rendements des caractéristiques individuelles, alors que la différence $(\beta_1 - \beta_0)$ peut être aussi due à des différences de composition car les deux estimateurs ne sont pas évalués sur les mêmes

9. Dans une estimation par MCO, le vecteur des estimateurs $\hat{\beta}_g$ pour chaque groupe g est estimé en minimisant la somme des carrés des erreurs de prédiction : $\hat{\beta}_g = \arg \min_{\beta} \sum (Y_{ti} - X_{ti}\beta_g)^2$. Les estimations sont sans biais si $E[\hat{\beta}] = \beta$ avec $\hat{\beta} = \beta + (X^T X)^{-1} X^T \epsilon_g$

10. Des solutions ont été proposées, notamment par Gardeazabal et Ugidos (2004) et Yun (2005, 2008), qui consiste à modifier les variables catégorielles avant l'estimation ou d'imposer des restrictions sur les estimateurs. Cependant, aucune ne semble satisfaisante et aucune s'impose dans la littérature.

distributions.

Enfin, la décomposition d'OB décompose la variation du salaire moyen, elle ne permet pas de différencier la contribution des variables et l'importance des effets composition et structurel selon différents points de la distribution des salaires. De plus, cette décomposition repose sur l'hypothèse que le salaire dépend linéairement des caractéristiques individuelles (Barsky *et al.* (2002)). D'autres méthodes ont ainsi été développées par la suite pour palier ces différents problèmes. La section suivante présente les méthodes basées sur les régressions des fonctions d'influence recentrées (Firpo *et al.* (2007, 2009, 2018)). À présent, les méthodes de repondération développées par DiNardo *et al.* (1996) sont présentées, elles permettent de ne pas imposer une espérance conditionnelle linéaire et elles sont applicables à des indicateurs plus diverse que la moyenne ¹¹.

Les méthodes de repondération

Pour étudier la différence d'une statistique d'intérêt entre deux groupes, DiNardo *et al.* (1996) (DFL par la suite) ont proposé une méthode basée sur la pondération. À la différence de l'approche d'OB où la distribution contrefactuelle résultait de l'application des paramètres de rendements du groupe 0 à la distribution X du groupe 1, dans l'approche DFL la distribution contrefactuelle résulte de l'application des paramètres du groupe 0 à la distribution X du groupe 0 mais repondéré par une fonction de pondération $w_C(X)$ de telle sorte que les caractéristiques du groupe 0 soient similaires à celle du groupe 1. En conséquence, l'effet structurel ne résulte pas d'une comparaison entre β_1 et β_0 mais entre β_1 et β_C dans l'approche pondérée, ce qui permet de s'assurer que l'effet structurel ne soit pas dû à une différence de composition entre les deux groupes qui influencerait sur l'estimation des paramètres (Firpo *et al.* (2007)).

Formellement, nous avons deux groupes avec N_0 et N_1 individus respectivement dans le groupe 0 et dans le groupe 1. La probabilité qu'un individu i appartienne au groupe 1 est définie par p et la probabilité qu'un individu i appartienne au groupe 0 sachant X par $p(x) = Pr[G = 1|X = x]$. Définissons maintenant trois fonctions de repondération :

$$w_1 \equiv \frac{1}{p} \quad w_0 \equiv \frac{1}{1-p} \quad w_C(X) \equiv \frac{(1-p(x)).p}{p(x).(1-p)}.$$

Les paramètres estimés par les MCO dans chaque groupe $g = 0, 1$ sont :

$$\hat{\beta}_g = \left(\sum_{i=1}^{N_g} \hat{w}_g \cdot X_i \cdot X_i' \right)^{-1} \cdot \sum_{i=1}^{N_g} \hat{w}_g \cdot Y_{ti} \cdot X_i, \quad (2.24)$$

11. Il existe d'autres méthodes présentées par Fortin *et al.* (2011), notamment les méthodes quantiles (Machado et Mata (2005)) et les méthodes basées directement sur les fonctions de redistribution (Chernozhukov *et al.* (2013)).

Et les paramètres estimés par les MCO dans le contrefactuel sont :

$$\hat{\beta}_C = \left(\sum_{i=1}^{N_0} \hat{w}_C(X) \cdot X_i \cdot X_i' \right)^{-1} \cdot \sum_{i=1}^{N_0} \hat{w}_C(X) \cdot Y_{0i} \cdot X_i, \quad (2.25)$$

Avec :

$$\hat{X}_C = \sum_{i=1}^{N_0} \hat{w}_C(X) \cdot X_i, \quad (2.26)$$

où $\lim_{N_0 \rightarrow \infty} (\bar{X}_C) = \lim_{N_1 \rightarrow \infty} (\bar{X}_1) = E(X|G=1)$.

Si l'espérance conditionnelle de Y sachant X est linéaire, les deux régressions pondérée et non-pondérée donneraient le même vecteur des estimateurs β_0 . Cependant, si l'espérance conditionnelle n'est pas linéaire, alors les estimateurs de β_0 différeront, puisque la méthode par les MCO minimise un terme d'erreur sur deux échantillons différents ¹².

Ainsi, l'effet global estimé de différence de moyenne entre les deux groupes évalué avec les méthodes de repondération est :

$$\begin{aligned} \Delta_O^\mu &= \underbrace{(\bar{X}_1 \hat{\beta}_1 - \bar{X}_C \hat{\beta}_C)}_{\Delta_{S,R}^{\hat{\mu}}} + \underbrace{(\bar{X}_C \hat{\beta}_C - \bar{X}_0 \hat{\beta}_0)}_{\Delta_{X,R}^{\hat{\mu}}} \\ &= \Delta_{S,R}^{\hat{\mu}} + \Delta_{X,R}^{\hat{\mu}} \end{aligned} \quad (2.27)$$

L'effet structurel et l'effet composition peuvent être tous les deux décomposés entre un effet pure et un terme d'erreur qui doit tendre vers zéro :

$$\begin{aligned} \Delta_{S,R}^{\hat{\mu}} &= \underbrace{\bar{X}_1 (\hat{\beta}_1 - \hat{\beta}_C)}_{\Delta_{S,p}^{\hat{\mu}}} + \underbrace{(\bar{X}_1 - \bar{X}_C) \hat{\beta}_C}_{\Delta_{S,e}^{\hat{\mu}}} \\ &= \Delta_{S,p}^{\hat{\mu}} + \Delta_{S,e}^{\hat{\mu}} \end{aligned} \quad (2.28)$$

et

$$\begin{aligned} \Delta_{X,R}^{\hat{\mu}} &= \underbrace{(\bar{X}_C - \bar{X}_0) \hat{\beta}_0}_{\Delta_{X,p}^{\hat{\mu}}} + \underbrace{\bar{X}_C (\hat{\beta}_C - \hat{\beta}_0)}_{\Delta_{X,e}^{\hat{\mu}}} \\ &= \Delta_{X,p}^{\hat{\mu}} + \Delta_{X,e}^{\hat{\mu}} \end{aligned} \quad (2.29)$$

Le terme d'erreur de l'effet structurel, $\Delta_{S,e}^{\hat{\mu}}$, permet de vérifier que le facteur de pondération est bien estimé puisque la différence $(\bar{X}_1 - \bar{X}_C)$ devrait être égale à zéro, ou du moins tendre vers zéro avec de grands échantillons. L'effet de composition est aussi composé d'un terme d'erreur, $\bar{X}_C (\hat{\beta}_C - \hat{\beta}_0)$, qui doit être égal à zéro si le modèle est bien linéaire puisqu'il ne devrait pas y avoir de différence entre $\hat{\beta}_C$ et $\hat{\beta}_0$ ¹³. Ceci est donc un avantage par rapport à l'approche standard

12. Dit autrement, dans le cas où la relation entre X et Y est linéaire, peu importe le poids accordé à chaque point puisque les points seront reliés par une même droite, et donc les estimateurs par les MCO seront les mêmes. Cependant, dans le cas où une relation convexe est estimée avec un modèle linéaire par les MCO, si plus de poids est donné à gauche de la distribution alors la droite qui minimise les erreurs tendra se déplacer à l'horizontal, et donc les estimateurs seront différents.

13. Voir note de bas de page 12.

d'OB puisque la méthode de pondération permet d'évaluer si le modèle linéaire est bien adapté pour estimer l'espérance conditionnelle de Y selon X . Cependant, comme dans le cas standard, la méthode de repondération est réalisée sous les hypothèses d'ignorance et de support commun.

Les fonctions de repondération sont maintenant étudiées plus en détails. Pour pouvoir mettre en œuvre la méthode de décomposition par repondération, les fonctions de pondérations w_0 , w_1 et $w_C(X)$ doivent être définies. Les deux premières fonctions sont estimées à partir du nombre d'observations dans chaque groupe. Avec $p = Pr(G = 1) = N_1/N$, le poids associé à chaque individu est $w_1 = 1/p$ dans le groupe 1 et $w_0 = 1/(1 - p)$ dans le groupe 0. Par conséquent, le poids entre individus d'un même groupe est le même.

Pour estimer la fonction de pondération pour le contrefactuel, $w_C(X)$, il faut estimer la probabilité qu'un individu du groupe 0 appartienne au groupe 1 selon ses caractéristiques observées X . L'idée est de remplacer la distribution des X du groupe 0 ($F_{X_0}(\cdot)$) par la distribution des X du groupe 1 ($F_{X_1}(\cdot)$) au moyen d'un facteur de pondération pour obtenir la distribution contrefactuelle ($F_{X_C}(\cdot)$). Par conséquent, le facteur de pondération est égal au ratio de deux fonctions des distributions marginales des variables explicatives X :

$$w_C(X) = \frac{dF_{X_1}(X)}{dF_{X_0}(X)} \quad (2.30)$$

Cette expression peut être simplifiée par la règle de Bayes, qui stipule dans notre cas que la probabilité d'observer l'individu j dans le groupe 0 sachant qu'il appartient au groupe 1 est :

$$P(T_j = 1|G = 0) = \frac{P(G = 0|T_j = 1).P(T_j = 1)}{\sum_{i=1} N_1 P(G = 0|T_i = 1).P(T_i = 1)} \quad (2.31)$$

Vu que les deux groupes sont exclusifs et exhaustifs, alors $P(X) = P(X|G = 1) + P(X|G = 0) = 1$. Ainsi :

$$dF_{X_1}(X) = P(X|G = 1) = \frac{P(G = 1|X).dF(X)}{\int_x P(G = 1|X).dF(X)} = \frac{P(G = 1|X)}{P(G = 1)} \quad (2.32)$$

De manière similaire :

$$dF_{X_0}(X) = P(X|G = 0) = \frac{P(G = 0|X)}{P(G = 0)} \quad (2.33)$$

Le facteur de pondération est alors égal à :

$$\begin{aligned} w_C(X) &= \frac{P(X|G=1)}{P(X|G=0)} \\ &= \frac{P(G=1|X)/P(G=1)}{P(G=0|X)/P(G=0)} \\ &= \frac{(1-p(X)).p}{p(X).(1-p)} \end{aligned} \quad (2.34)$$

Ainsi, $P(G = 1|X)$ peut être estimé par un modèle de probabilité, et le facteur de pondération grâce aux probabilités estimées.

Pour estimer $P(G = 1|X)$, DFL proposent d'utiliser un modèle probit ou logit, alternativement Hirano *et al.* (2003) proposent un modèle logit non-paramétrique ¹⁴. Par ailleurs, Firpo *et al.* (2018) normalisent les pondérations de telle sorte que la somme des poids de chaque fonction de pondération soit égal à un. Ainsi, le poids de chaque individu pour chaque pondération est divisé par la somme respective des poids estimés.

Un des avantages de cette méthode est qu'une fois la distribution contrefactuelle estimée, la décomposition peut s'appliquer à une variété de statistiques : quantiles, variance, Gini... et pas uniquement la moyenne comme dans le cas standard d'OB. De plus, cette méthode intègre un terme d'erreur dans l'effet de composition qui permet de vérifier si l'espérance conditionnelle est bien linéaire comme l'approche d'OB le suppose. Ainsi, même dans le cas de la moyenne, il est préférable d'utiliser la méthode de repondération. De plus, la méthode est aisément applicable. Enfin, comme l'a montré Firpo *et al.* (2007), il est possible d'estimer les erreurs types des différents éléments obtenus lors de la décomposition, cependant d'après Fortin *et al.* (2011) il est plus simple d'évaluer l'inférence des résultats par bootstrap.

La principale limite de cette approche est la difficulté à estimer la contribution de chaque variable explicative à chaque effet. DFL propose une méthode pour mesurer la contribution d'une variable binaire à l'effet composition et l'effet structurel d'une variation dans la distribution des salaires entre deux périodes. Pour tous types de variables, Butcher et Dinardo (2002) et Altonji *et al.* (2008)) proposent une méthode pour estimer la contribution de chaque variable à l'effet de composition. Elle consiste à ajouter séquentiellement les variables explicatives dans le modèle de probabilité $P(G = 1|X)$ utilisé pour estimer le facteur de pondération. Ainsi, pour estimer l'effet de composition dû au premier facteur considéré, $P(G = 1|X_1)$ est estimé, puis $w_{X_1}(X_1)$ et la statistique d'intérêt sur ce nouveau contrefactuel pondéré. Puis pour estimer l'effet de composition de la seconde variable considérée, il faut refaire la même procédure mais cette fois-ci avec le facteur de pondération $w_{X_2|X_1}(X_1)$ estimé à partir de $P(G = 1|X_1, X_2)$. Le problème dans cette approche est que le résultat est dépendant de l'ordre dans lequel sont ajoutées les variables explicatives.

Comme l'a mis en avant Gelbach (2009), ce problème peut être assimilé au problème de la

14. Le modèle probit est un modèle de régression binomiale où la variable dépendante est une variable aléatoire binaire (i.e; prenant pour valeur 0 ou 1) tel que $P(Y = 1|X) = \theta(X'\beta)$ où $\theta(\cdot)$ est une fonction de répartition de la loi normale centrée réduite. Le modèle logit (ou logistique) est aussi un modèle de régression binomiale, où la loi de probabilité est modélisée à partir d'une loi logistique : $P(Y = 1|X) = \frac{e^{X'\beta}}{1+e^{X'\beta}}$. Dans l'approche standard ces deux modèles sont estimés par maximum de vraisemblance, ce qui suppose une certaine distribution pour le terme d'erreur, qui dans le cas d'un modèle logit est supposé suivre une distribution logistique. D'autres méthodes, comme la méthode du score maximum, sont dites non paramétriques lorsqu'aucune hypothèse n'est faite sur la distribution des résidus lors de la phase d'estimation des paramètres β .

variable omise, puisque considérer la première variable sans considérer les autres variables explicatives peut sous-estimer ou surestimer l'impact de cette variable. Seule la dernière variable explicative n'est pas soumise à ce problème, puisque son effet de composition est estimé en prenant en compte l'ensemble des autres variables. Ainsi le problème de dépendance au sentier peut être résolu en estimant l'effet de composition de chaque variable explicative quand celle-ci est considérée en dernière. Cependant, un nouveau problème s'ajoute, puisque la somme des ces contributions sera différente de l'effet de composition total.

Par conséquent, la décomposition détaillée par facteur des effets de composition et structurel dans le cas d'une décomposition par repondération n'est pas aisée, c'est pourquoi cette méthode a été combinée avec la méthode des régressions RIF par Firpo *et al.* (2009, 2018).

2.4.2 Les régressions RIF

Firpo *et al.* (2007, 2009, 2018) (FFL par la suite pour Firpo, Fortin et Lemieux) proposent une approche en deux temps pour réaliser une décomposition qui permet de connaître la contribution de chaque variable explicative à l'effet composition et à l'effet structurel afin d'expliquer la différence de distribution d'une variable revenu entre deux groupes. La première étape consiste à utiliser la méthode de repondération décrite précédemment (Section 2.4.1) pour établir la distribution contrefactuelle et ainsi décomposer la différence de distributions entre un effet structurel et un effet composition. Ensuite, grâce aux fonctions d'influence recentrées, ils estiment la contribution de chaque variable à chaque effet.

Les fonctions d'influence, notées $IF(\cdot)$, mesurent l'effet d'une perturbation marginale d'une distribution d'une variable sur un indicateur calculé sur cette même distribution. Ces fonctions permettent notamment d'estimer la sensibilité d'un indicateur à des observations aberrantes (Hampel (1974)). Cet indicateur (v) peut être la moyenne, un quantile, le coefficient de Gini, etc...

Ainsi, la fonction d'influence mesure la variation d'un indicateur calculé sur une distribution F lorsqu'elle tend vers une nouvelle distribution H . Dans un premier temps, $H = \delta_y$ est définie, avec δ_y une distribution qui donne du poids uniquement au point y . Ainsi, la fonction d'influence s'écrit :

$$IF(y; v, F) = \lim_{\epsilon \rightarrow 0^+} \frac{v(F_y) - v(F)}{\epsilon} \quad (2.35)$$

Avec $F_y = (1 - \epsilon)F + \epsilon\delta_y$, $0 \leq \epsilon \leq 1$.

Cette fonction décrit l'effet d'une variation infinitésimale (ϵ) au point y sur l'indicateur v de la distribution F . Par définition, $\int_{-\infty}^{\infty} IF(y; v, F) dF = 0$.

Dans le cas où l'indicateur de la distribution est la moyenne, c'est-à-dire où $v = \mu$, alors :

$$IF(y; \mu, F) = \lim_{\epsilon \rightarrow 0^+} \frac{(1 - \epsilon) \cdot \mu_F + \epsilon \cdot y - \mu_F}{\epsilon} = y - \mu_F \quad (2.36)$$

Ainsi la fonction d'influence de la moyenne correspond à la différence entre le point y et la moyenne, la somme des valeurs prises par cette fonction mesurée à chaque point de la distribution est bien égale à zéro.

La fonction d'influence recentrée consiste à ajouter à la fonction d'influence, l'indicateur estimé de la distribution initiale (Firpo *et al.* (2007)). La fonction d'influence recentrée obtenue est $RIF(y; v, F) = v(F) + IF(y; v, F)$, d'après l'intégrale ci-dessus :

$$\int RIF(y; v, F) dF(y) = \int v(F) + IF(y; v, F) \cdot dF(y) = v(F) \quad (2.37)$$

Dans le cas de la moyenne :

$$RIF(y; \mu; F) = y - \mu_F, \quad (2.38)$$

et donc

$$\int RIF(y; \mu; F) = \mu_F. \quad (2.39)$$

D'après la loi des espérances répétées, l'indicateur étudié peut être exprimé en termes d'espérance conditionnelle de la fonction d'influence recentrée :

$$v(F) = \int E[RIF(Y; v, F) | X = x] \cdot dF_X(x) \quad (2.40)$$

Les régressions des fonctions d'influence recentrées sont similaires à des régressions standards, à la différence que la variable dépendante est remplacée par la fonction d'influence recentrée. L'espérance de cette fonction, $E[RIF(Y; v, F) | X = x]$, peut être estimée par un modèle linéaire. Les coefficients mesurent l'effet moyen d'une variation dans les observations d'une variable X_i sur v , en gardant les autres variables explicatives constantes. Ainsi, si l'indicateur statistique utilisé est la moyenne, les coefficients de la régression RIF seront identiques à une régression standard par les moindres carrés ordinaires.

Pour rappel, la méthode de repondération permet d'obtenir la distribution contrefactuelle F_C , c'est-à-dire la distribution de revenus qui aurait prévalu si la structure de salaire était celle du groupe $G = 0$ avec la composition de l'autre groupe $G = 1$. Ainsi, cette méthode décompose l'effet structurel et l'effet pondération de la différence de distribution de revenus entre les deux groupes. La différence entre l'indicateur mesuré dans le groupe $G = 0$ et $G = 1$, notés respectivement v_0 et v_1 , peut être lui aussi décomposé grâce à l'indicateur mesuré dans la distribution contrefactuelle v_C . Dès lors :

$$\Delta_O^v = (v_1 - v_C) + (v_C + v_0) = \Delta_S^v + \Delta_X^v. \quad (2.41)$$

Le premier terme de l'équation Δ_S^v représente l'effet structurel alors que le second terme Δ_X^v reflète l'effet de composition, sous les hypothèses usuelles d'ignorance et de support commun.

Pour connaître la contribution de chaque variable à chaque effet, des régressions RIF sont réalisées en fonction des variables explicatives X et du groupe g . La notation de Firpo *et al.* (2018) est reprise, ainsi $m_g^v(x)$ représente la régression-RIF de l'indicateur statistique v pour le groupe g selon la distribution des variables explicatives x , telle que :

$$m_g^v(x) \equiv E[RIF(Y_g; v_g, F_g)|X, G = g], \quad (2.42)$$

pour $g = 0, 1$, et de la même manière pour la distribution contrefactuelle :

$$m_C^v(x) \equiv E[RIF(Y_0; v_C, F_C)|X, G = 1]. \quad (2.43)$$

Par conséquent, les égalités suivantes sont obtenues :

$$v_g = E[m_g^v(X)|G = g], \quad t = 0, 1 \quad (2.44)$$

$$v_C = E[m_C^v(X)|G = 1] \quad (2.45)$$

De ce fait :

$$\Delta_S^v = E[m_1^v(X)|G = 1] - E[m_C^v(X)|G = 1] \quad (2.46)$$

$$\Delta_X^v = E[m_C^v(X)|G = 1] - E[m_0^v(X)|G = 0]. \quad (2.47)$$

Dans le cas de la moyenne, la décomposition proposée par Firpo *et al.* (2018) sera équivalente à la décomposition d'Oaxaca-Blinder uniquement si l'espérance conditionnelle des fonctions d'influence recentrées sont linéaires en x . Par ailleurs, si l'espérance conditionnelle des fonctions d'influence recentrées n'est pas linéaire en x alors un terme d'erreur sera introduit dans l'effet de composition, encore faut-il pouvoir apprécier son importance. Plus formellement, l'effet composition peut être réécrit tel que :

$$\Delta_X^v = (E[X|G = 1] - E[X|G = 0])'\beta_0^v + R^v, \quad (2.48)$$

avec le terme d'erreur $R^v = E[X|G = 1]'\beta_C^v - \beta_0^v$, et avec β_0^v et β_C^v les estimateurs respectifs de la régression-RIF dans le groupe $g = 0$ et dans le contrefactuel. Si les deux estimateurs sont égaux, ce qui doit être le cas si la modélisation est bien adaptée pour apprécier l'espérance conditionnelle,

alors le terme d'erreur est nul.

De plus, à partir du premier terme de l'équation précédente, la contribution de chaque variable explicative k à l'effet composition est déterminée, telle que :

$$\Delta_X^v - R^v = \sum_{k=1}^K (E[X^k|G=1] - E[X^k|G=0])' \beta_{0,k}^v. \quad (2.49)$$

Les RIF donnent une approximation de 1^{er} ordre du vrai effet de composition Δ_X^v par $\hat{\Delta}_X^v - \hat{R}^v$, Firpo *et al.* (2018) parlent alors « d'effets politiques ». Une estimation de l'erreur peut être donnée en comparant d'une part la différence entre l'indicateur statistique mesurée sur la distribution contrefactuelle et la distribution du groupe $G=0$, $v_C - v_0$, et d'autre part l'estimation de $(E[X|G=1] - E[X|G=0])' \beta_0^v$ obtenue grâce aux régressions-RIF. Si cette différence est trop importante alors les régressions-RIF, du moins celles avec une spécification linéaire, ne sont sûrement pas adaptées à l'étude réalisée.

L'effet structurel est donné par $\Delta_S^v = E[m_1^v(X)|G=1] - E[m_C^v(X)|G=1]$, qui peut aussi s'exprimer par :

$$\Delta_S^v = E[X|G=1]'(\beta_1^v - \beta_C^v), \quad (2.50)$$

avec β_1^v et β_C^v les paramètres estimés respectivement à partir de la distribution de revenu du groupe $G=1$ et la distribution contrefactuelle.

L'avantage de l'approche de l'effet structurel de FFL est d'utiliser β_C^v au lieu β_0^v , puisque la différence $(\beta_1^v - \beta_C^v)$ reflète uniquement un effet structurel, alors que la différence $(\beta_1^v - \beta_0^v)$ peut être due à des différences de composition entre les deux groupes. Par ailleurs, comme dans le cas d'OB, la contribution de chaque variable à l'effet structurel est dépendante au choix du groupe de référence.

De fait, les régressions RIF sont des approximations linéaires d'indicateurs non-linéaires (exceptée la moyenne avec une espérance conditionnelle linéaire). Cette approximation peut ne pas être adaptée (Rothe (2015)), il est ainsi important d'estimer la robustesse des résultats, notamment par des méthodes bootstrap (Firpo *et al.* (2018)).

La décomposition de l'indice de Gini avec les régressions RIF

Une décomposition de l'indice de Gini est possible avec une régression RIF. Tout d'abord, l'indice de Gini mesuré sur la distribution de revenu y est défini par :

$$G(F_Y) = 1 - 2\mu^{-1}R(F_Y) \quad (2.51)$$

avec F_Y la fonction de distribution cumulée de y , $R(F_Y) = \int_0^1 GL(p; F_Y) dp$ avec $p(y) = F_Y(y)$ et où $GL(p; F_Y)$ est la courbe de Lorenz généralisée de F_Y , donnée par $GL(p; F_Y) = \int_{-\infty}^{F_Y^{-1}(p)} z dF_Y(z)$. Alors que la courbe de Lorenz indique les proportions de revenus cumulée, avec $L(1; F_Y) = 1$, la courbe de Lorenz généralisée indique les proportions du revenu moyen cumulée, avec $GL(1; F_Y) = \mu_y$. En somme, la première nous indique la part de richesse détenue par $p\%$ de la population, alors que la seconde donne une indication sur la richesse totale détenue ces $p\%$ de la population. Autrement dit, la courbe de Lorenz généralisée prend en compte un effet volume de richesses distribuées dans l'économie, ce qui peut être pertinent pour comparer des économies différentes¹⁵.

De cette équation, Monti (1991) dérive la fonction d'influence recentrée de l'indice de Gini :

$$RIF(y; G, F_Y) = 2 \frac{y}{\mu} \left[F_Y(y) - \frac{(1+G)}{2} \right] + 2 \left[\frac{(1-G)}{2} - GL(p; F_Y) \right] + G \quad (2.52)$$

où $(1-G)/2$ et $(1+G)/2$ correspondent respectivement à l'aire au-dessus et en-dessous de la courbe de Lorenz. Cette fonction est continue et convexe. Sa première dérivée est égale à $2/\mu [F_Y(y) - (1+G)/2]$, et atteint un minimum lorsque $F_Y(y) = (1+G)/2$.

L'avantage de cette méthode est qu'elle peut être réalisée localement (et non globalement comme dans la procédure de Chernozhukov *et al.* (2013)), il n'est donc pas nécessaire de réaliser une régression à chaque points et de prendre en considération le problème de la monotonie. Par ailleurs, même si la décomposition détaillée de l'effet structurel dépend du groupe de référence, un autre avantage de cette approche est que la décomposition détaillée est indépendante au sentier (*i.e.* elle ne dépend de l'ordre par lequel les variables explicatives sont considérées).

La section suivante aborde la décomposition des indices d'inégalité grâce à la valeur de Shapley. Cette méthode permet aussi une décomposition par attributs, autrement dit elle est aussi capable de prendre en compte plusieurs effets de groupes en les traitant comme des attributs (ex. : âge, diplôme...). Cette méthode peut aussi donner lieu à une décomposition multi-niveaux par exemple elle peut traiter les effets du genre puis les contributions des attributs au sein de chaque groupe. L'avantage fondamental par rapport à toutes les méthodes de décomposition est que cette approche est applicable à une large variété d'indices d'inégalité.

2.5 La décomposition de Shapley des indices d'inégalité

La décomposition par sous-populations des indices d'inégalité a pour objectif d'estimer la contribution de chaque groupe à l'inégalité, ainsi que l'effet des différences entre groupes sur les

15. Voir Atkinson (1970). Ce dernier stipule que si la courbe de Lorenz généralisée d'une distribution est supérieure en tous points à une autre, et que le décideur public a une aversion pour l'inégalité, alors le bien-être est supérieur dans le premier cas.

inégalités. Comme la section 2.2.2 l’a évoqué, les indices de Theil, notamment l’indice de l’écart logarithmique moyen, sont des indices bien adaptés pour cette décomposition. Toutefois la décomposition en trois termes de l’indice de Gini peut avoir son intérêt si l’évaluateur estime que l’effet de chevauchement des distributions est une information pertinente à prendre en considération.

Puis, les décompositions par sources ont été étudiées, elles consistent à déterminer la contribution des facteurs (ex. sources de revenu) à l’inégalité. D’après Shorrocks (1982), la décomposition de la variance est l’approche la plus adaptée pour ce type de décomposition. Même si certains auteurs se sont écartés des principes de Shorrocks pour proposer d’autres décompositions, notamment celle de l’indice de Gini proposée par Lerman et Yitzhaki (1985).

Les deux approches, décompositions par sous-populations et décompositions par sources, sont complémentaires dans l’évaluation des politiques publiques. Une méthode de décomposition qui permettrait de décomposer à la fois par sous-populations et par sources pourrait par exemple déterminer l’impact d’une réforme des cotisations sociales selon les différentes catégories socioprofessionnelles. Des décompositions multidimensionnelles ont par la suite été proposées pour prendre en compte à la fois des effets de groupes et des effets de sources, tout en considérant une dimension temporelle (Mussard et Savard (2012)). Cet effet temporel a par ailleurs été divisé entre deux effets : un effet de reclassement des individus et un effet de croissance disproportionnelle des revenus (Jenkins et Van Kerm (2006), Mussini (2013)).

La décomposition à la Shapley permet, elle aussi, de prendre en compte à la fois des effets de sous-populations et des effets de sources, soit au moyen de modèles hiérarchiques (ex. : une décomposition par sources au sein de chaque groupe), soit en les traitant de manière similaire (*i.e.* au même niveau). Plus généralement, cette dernière approche permet de décomposer l’inégalité entre toutes les caractéristiques individuelles (ou attributs par la suite), sachant que ceux-ci peuvent faire référence à des effets de groupes (ex. : âge, diplôme, genre...). Par conséquent, la première étape consiste à désagréger le revenu entre les différentes caractéristiques individuelles, ce qui peut se faire par des méthodes statistiques et se combiner avec les méthodes de décomposition citées par Fortin *et al.* (2011). D’autre part, la décomposition à la Shapley s’applique à tous les indices d’inégalité, ce qui est un avantage non-négligeable de cette approche.

2.5.1 Introduction au concept de Shapley

La décomposition de Shapley est un concept de théorie des jeux coopératifs qui permet de distribuer entre différents joueurs le bénéfice ou le coût qu’ils produisent conjointement. L’importance d’un joueur dans un jeu étant la moyenne pondérée de ses contributions marginales pour l’ensemble des coalitions possibles de joueurs. Par la suite, ce concept a été appliqué à différents sujets : partage de surplus, modèles de taxation, allocation de parts de marché ou encore dans l’allocation de

pouvoir politique. Dans la continuité des premiers travaux de Auvray et Trannoy (1992), Chantreuil et Trannoy (1999, 2011) et Shorrocks (1999, 2013) ont proposé une approche axiomatique de la décomposition des indices d'inégalité basée sur la méthode de décomposition de Shapley. L'objectif de cette méthode est d'identifier et de hiérarchiser les facteurs les plus influents, autrement dit d'estimer l'importance des facteurs à l'inégalité.

Soit un ensemble d'individus N , $N := \{1, \dots, a, \dots, n\}$, avec $n \geq 2$, et un ensemble de sources de revenu $M := \{1, \dots, i, \dots, m\}$, avec $m \geq 2$. Une coalition de sources de revenu quelconque est notée S , et $\mathcal{M}$ représente l'ensemble de ces coalitions possibles. Une situation $x := [y_a^i]$, où $y_a^i \geq 0$ est le montant de chaque source de revenu i reçu par l'individu a . x peut être vu comme une matrice $n \times m$, où la ligne a indique le montant de chaque source de revenu reçu par l'individu a , et la colonne i indique la distribution de la source i entre les n individus. L'ensemble des x est noté par $\mathcal{X} := \{x | y_a^i \geq 0, \forall a, i\}$. La distribution de la source i est donnée par $x^i := (y_1^i, \dots, y_n^i)^T$, avec T pour indiquer la transposée du vecteur. La distribution agrégée est $X := \sum_{i \in M} x^i$ et le revenu moyen est $\mu(X) := \frac{1}{n} \sum_{i \in M} x^i$. De manière similaire, la distribution agrégée des sources incluses dans S est $X^S := \sum_{i \in S} x^i$ et $\mu(X^S) := \frac{1}{n} \sum_{i \in S} x^i$ le revenu moyen correspondant.

L'indice d'inégalité est une fonction $I : \mathbb{R}^n \rightarrow \mathbb{R}_+$, tel que $I(\cdot)$ est un indice régulier relatif. La distribution de revenus est définie selon les sources de revenu prises en compte, si aucune n'est prise en considération alors $Y(\emptyset) = 0^n$, sinon pour tout $S \in \mathcal{M}$:

$$Y(S, \lambda) = (\sum_{i \in S} y_1^i + \lambda, \dots, \sum_{i \in S} y_n^i + \lambda) \quad (2.53)$$

avec $\lambda \in \mathbb{R}^+$. Selon la valeur de λ , Chantreuil et Trannoy (2011) définissent différents jeux d'inégalité. Ils introduisent en premier lieu le jeu d'inégalité à revenus nuls (jeu nul par la suite) avec $\lambda = 0$, dans lequel les distributions des revenus non compris dans S sont remplacées par un vecteur nul. Puis ils introduisent le jeu d'inégalité égalisateur (jeu égalitaire par la suite) avec $\lambda = \mu(X^{\setminus S})$. Dans ce jeu, les distributions non présentes dans S sont remplacées par des vecteurs égalitaires dont le montant fixé pour tous correspond au revenu moyen des sources en question.

Le concept de solution proposé par Shapley (1953) est noté $Sh_i(M; V_I(X))$. Il dénote l'importance de la source i dans un ensemble de sources M , avec une fonction caractéristique $V_I(\cdot)$. La fonction caractéristique est une fonction de coût ou de bénéfice qui permet d'estimer la valeur du jeu pour n'importe quelle coalition. Dans notre cas, la fonction caractéristique est notée $V_I(\cdot)$, elle évalue l'inégalité d'une distribution de revenus $Y(S, \lambda)$ selon un indice d'inégalité $I(\cdot)$, ainsi $V_I(S) = I(Y(S, \lambda))$.

Dans la version probabiliste de la valeur de Shapley, l'importance d'une source de revenu i aux inégalités est égale à :

$$Sh_i = \sum_{S \subseteq \mathcal{M} \setminus \{i\}} \frac{s!(m-s-1)!}{m!} [V_I(X \cup i) - V_I(X)] \quad (2.54)$$

Autrement dit, l'importance d'un facteur i aux inégalités est la moyenne de l'ensemble de ses contributions marginales à toutes les coalitions possibles de facteurs.

Chantreuil et Trannoy (1999) montrent que la décomposition de Shapley des indices d'inégalité est efficace car la somme des importances est égale à la valeur du jeu total. Cet axiome d'efficacité permet la hiérarchisation des différents facteurs grâce aux importances relatives (sh_i) :

$$\sum_{i=1}^m Sh_i = I(X) \quad (2.55)$$

$$\sum_{i=1}^m sh_i = 1 \quad (2.56)$$

Avec $sh_i = Sh_i / I(X)$.

Cette propriété permet aux indices décomposables en sous-groupes comme les indices cohérents d'être multi-décomposables. En effet, une décomposition de Shapley par facteurs de chaque composante (inter et intragroupe) est cohérente.

Cependant, comme Chantreuil et Trannoy (1999) l'ont spécifié, la méthode de décomposition de Shapley ne respecte pas toutes les propriétés que devraient respecter une décomposition par facteurs selon Shorrocks (1982). Cette décomposition est dépendante du niveau de désagrégation et elle ne normalise pas nécessairement les distributions de facteurs équivalents. Ceux deux points sont étudiés plus précisément dans les deux sections suivantes.

2.5.2 Décomposition d'un modèle hiérarchique

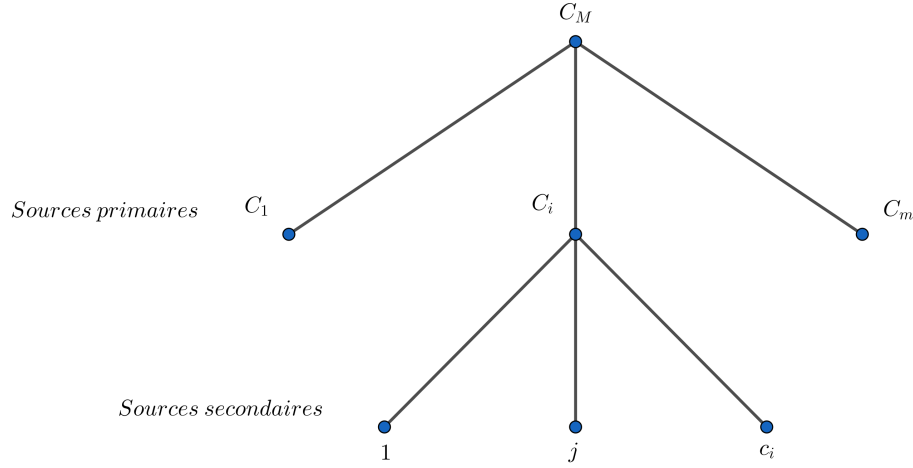
La valeur d'Owen

Chantreuil et Trannoy (2013) considèrent le cas où chaque source de revenu est composé de plusieurs sources. Par exemple le revenu total peut être composé des revenus du marché et des revenus sociaux. Le premier pouvant être composé des salaires et des revenus du capital, et le second des allocations familiales¹⁶ et d'allocations sous conditions de ressources. Dans la décomposition à la Shapley, l'importance des sources dépend du nombre de sources considérées. Une solution à ce problème de dépendance est de considérer la valeur de Owen (1977), qui est une généralisation de la valeur de Shapley dans le sens où un groupe n'est plus composé que d'une seule source mais de plusieurs. La valeur d'Owen est une décomposition à la Shapley en deux étapes : la première

16. Ici, les allocations familiales sont supposées sans condition de ressource.

consiste à estimer l'importance de chaque source primaire, puis la seconde définit l'importance de chaque source secondaire à leur source primaire respective. Dans ce cas, la somme des importances des différentes sources d'une source primaire est égale à l'importance de la source primaire à l'indice d'inégalité. Autrement dit, la décomposition est robuste à l'agrégation (Shorrocks (2013)).

Graphique 2.1 – Arbre des sources de revenu



En présence d'un ensemble de m sources de revenu primaires $C_M := \{C_1, \dots, C_i, \dots, C_m\}$, avec $m \geq 2$. Chacune des sources C_i est constituée de c_i sources secondaires telles que $C_i := \{1, \dots, j, \dots, c_i\}$ (voir graphique 2.1). Une coalition de sources de revenu primaires quelconque est notée S , composée de s sources, et une coalition de sources de revenu secondaires quelconque est notée H , composée de h sources. L'ensemble des sources secondaires est noté K , il comprend k sources telles que $k = \sum_{i=1}^m c_i$. Une partition de l'ensemble K est notée P et une coalition de revenus (primaires et/ou secondaires) de cette partition est notée C , comprenant c sources.

Dans sa version probabiliste, la valeur d'Owen pour une source j est égale à :

$$Ow_j = \sum_{C \in 2^P, C_i \notin C} \sum_{H \subseteq C_i, j \notin H} \frac{c!(m-c-1)!h!(c_i-h-1)!}{m!h!} [V_I(C \cup H \cup j) - V_I(C \cup H)] \quad (2.57)$$

Autrement dit, seules les coalitions qui permettent de respecter l'ordre des groupes (*i.e.* la partition) sont considérées dans l'estimation de la valeur d'Owen. D'après le théorème de Khmel'nitskaya et Yanovskaya (2007), la valeur d'Owen permet de respecter la monotonie (*i.e.* l'augmentation de l'inégalité d'une source entraîne l'augmentation de l'importance de cette source, toutes autres sources restant inchangées.), tout en respectant les axiomes d'efficacité, de symétrie intragroupe (traitement égalitaire des facteurs égaux au sein d'un groupe) et de symétrie intergroupe (traitement égalitaire de groupes égaux).

Nested Shapley

Une autre méthode pour prendre en compte la structure hiérarchique des sources de revenu est l'utilisation du Nested-Shapley (Chantreuil et Trannoy (1999)). Au même titre que la méthode d'Owen, l'importance des facteurs secondaires dépend du facteur primaire mais ne dépend pas de la manière dont les autres facteurs primaires sont désagrégés. Même si les approches sont similaires, l'interprétation des résultats est différente (Sastre et Trannoy (2000)). Pour estimer l'importance d'un facteur, la valeur d'Owen prend en compte des coalitions de facteurs secondaires appartenant à des facteurs primaires auxquels le facteur en question n'appartient pas, ce que ne fait pas la méthode du Nested Shapley.

La méthode du Nested Shapley (ou du Shapley emboîtés) s'opère en deux temps. Dans un premier temps, l'importance associée à chaque source secondaire est estimée au sein de leur groupe respectif. Cependant, contrairement à une décomposition à la Shapley classique, l'importance estimée ici prend en considération les autres sources primaires dont les vecteurs de revenus respectifs sont égalisés au niveau de leur moyenne pour neutraliser leur effet inégalitaire.

Ensuite, dans un second temps, la différence entre la somme de l'importance des sources secondaires estimées précédemment avec l'importance de leur source primaire respective est prise en compte pour assurer l'efficacité du Nested Shapley. Ainsi, l'importance des sources secondaires est révisée, concrètement une part équitable de cette différence (*i.e.* entre la valeur de Shapley d'une source primaire et la somme de l'importance des sources secondaires estimée dans la première étape) est ajoutée à l'importance des sources secondaires.

En reprenant les notations précédentes, l'inégalité d'une source primaire peut être donnée par :

$$\begin{aligned} V_I(C_i) &= \sum_{j \in C_i} \sum_{H \subset C_i, j \notin H} \frac{h!(c_i - h - 1)!}{c_i!} [V_I(H \cup \{j\} | \mu(X^{C_m \setminus C_i})) - V_I(H | \mu(X^{C_m \setminus C_i}))] \\ &= \sum_{j \in C_i} S h_j^* \end{aligned} \quad (2.58)$$

Avec $S h_j^*$, l'importance mesurée par la valeur de Shapley d'une source secondaire à sa source primaire, sachant que les autres sources primaires sont prises en compte par leur vecteur de revenu moyen ($\mu(X^{C_m \setminus C_i})$). La valeur du Nested Shapley alors obtenue est :

$$NS h_j = S h_j^* + \frac{1}{c_i} [S h_i - V_I(C_i)] \quad (2.59)$$

Soit :

$$\begin{aligned}
 NSh_j = & \sum_{H \subset C_i, j \notin H} \frac{h!(c_i - h - 1)!}{c_i!} [V_I(H \cup \{j\} | \mu(X^{C_m \setminus C_i})) - V_I(H | \mu(X^{C_m \setminus C_i}))] \\
 & + \frac{1}{c_i!} \left[\sum_{S \subset C_M, C_i \notin S} \frac{s!(m - s - 1)!}{m!} [V_I(S \cup C_i) - V_I(S)] - V_I(C_i) \right]
 \end{aligned} \tag{2.60}$$

Ainsi, selon Sastre et Trannoy (2002), l'importance estimée d'un facteur secondaire dans le cas du Nested Shapley ne dépend pas des autres facteurs secondaires, contrairement à la valeur d'Owen. L'autre différence notable est que le Nested Shapley compare uniquement des coalitions de sources de même niveau. Ces deux différences les incitent à privilégier la méthode du Nested Shapley pour sa cohérence dans l'interprétation économique.

Néanmoins, le chapitre 4 démontre que $\sum_{S \subset C_M, C_i \notin S} \frac{s!(m-s-1)!}{m!} [V_I(S \cup C_i) - V_I(S) - V_I(C_i)]$ peut être interprété comme la somme des interactions de la source i avec les autres sources primaires. Puisque si les inégalités mesurées n'étaient pas plus fortes ($V_I(S \cup C_i) - V_I(S) - V_I(C_i) > 0$) ou au contraire plus faibles ($V_I(S \cup C_i) - V_I(S) - V_I(C_i) < 0$) quand les sources sont considérées ensemble plutôt que séparément, alors il n'y aurait pas d'interactions (*i.e.* $Sh_i - V_I(C_i) = 0$). Autrement dit, la somme de l'importance des sources secondaires à l'inégalité de sa source primaire serait égale à l'importance de sa source primaire à l'inégalité totale (*i.e.* $\sum_{j \in C_i} Sh_j^* = V_I(C_i) = Sh_i$). Ainsi, le Nested Shapley consiste à répartir équitablement l'interaction d'une source primaire avec les autres sources primaires à l'importance de ses sources secondaires. L'importance d'une source secondaire avec le Nested Shapley dépend donc de la composition des sources primaires et de leurs interactions. De plus, la répartition équitable des interactions pourrait ne pas être justifiée s'il s'avérait qu'une source secondaire interagit plus qu'une autre avec une source primaire.

Lorsqu'une structure hiérarchique exogène des sources de revenu est définie, il faut utiliser une méthode de décomposition qui permette que l'importance d'une source secondaire ne dépende pas de la manière dont sont agrégées les autres sources de revenu. Selon Sastre et Trannoy (2002), la méthode du Nested Shapley est à privilégier pour son interprétation économique. Cependant, il ne faudrait pas négliger le fait que la valeur d'Owen respecte des propriétés axiomatiques que ne respecte pas le Nested Shapley (Winter (2002), Khmelnitskaya et Yanovskaya (2007)), notamment la monotonie. Ainsi, si l'évaluateur souhaite que l'importance d'une source de revenu aux inégalités soit plus importante lorsque sa distribution est plus inégalitaire, sachant que les distributions des autres sources sont inchangées, alors il devrait adopter une décomposition hiérarchique à la Owen plutôt qu'à la Nested Shapley. Par ailleurs, la valeur d'Owen ne suppose pas une répartition équitable de l'interaction des sources secondaires d'une source primaire avec les autres sources primaires. La notion d'interaction sera développée dans les chapitres 4 et 5.

Exemple comparatif : Owen VS Nested Shapley

Une distribution de revenus hiérarchisée est définie (Tableau 2.3) afin d'illustrer la différence entre l'approche d'Owen et celle du Nested Shapley. Trois sources primaires sont considérées : revenus du travail, transferts sociaux et revenus du capital. Chacune de ces sources est subdivisée entre deux sources secondaires, respectivement : salaire et prime, allocations familiales et allocations logements, revenus du capital financier et revenus du capital immobilier.

Tableau 2.3 – Distribution de revenus hiérarchisée

Indiv.	Revenus du travail			Transferts Sociaux			Revenus du capital			Total
	Salaire	Prime	Total	Alloc. Famille	Alloc. Logement	Total	Financier	Immobilier	Total	
1	20	10	30	10	20	30	20	30	50	110
2	35	15	50	10	10	20	50	50	100	170
3	45	25	70	10	0	10	100	50	150	230
Total	100	50	150	30	30	60	170	130	300	510
$V_I(.)$	0,0327	0,0196	0,0523	0	0,0261	0,0261	0,1046	0,0261	0,1307	0,1569

Note : $V_I(.)$ est calculée avec l'indice de Gini selon un jeu égalitaire.

Dans un premier temps, le procédé du Nested Shapley est présenté dans le Tableau 2.4. La décomposition est ici réalisée dans le cadre d'un jeu égalitaire en prenant en compte l'indice de Gini comme indicateur d'inégalité. L'indice de Gini n'étant pas additif, la somme de l'inégalité mesurée dans chaque source primaire n'est pas égale à l'inégalité totale, néanmoins la somme de l'importance de chaque source mesurée par la valeur de Shapley est bien égale à l'inégalité mesurée sur la distribution du revenu total. Dans la seconde partie du tableau, les valeurs de Shapley associées à chaque source secondaire au sein de leur source primaire sont dénotées par NSh_j^* . Comme expliqué dans la section 2.5.2, ces dernières sont calculées sans prendre en compte les revenus des sources primaires dont elles ne font pas parties, ainsi la somme des NSh_j^* d'une source primaire est égale à l'inégalité de cette source primaire ($V_I(C_i)$). Ensuite, dans un second temps, ces valeurs de Shapley sont corrigées (NSh_j) pour que leur somme soit égale à l'inégalité totale. Concrètement, la correction au sein de chaque source primaire consiste à attribuer à chaque source secondaire une part équitable (ici la moitié) de la différence entre l'inégalité de cette source ($V_I(C_i)$) et la valeur de Shapley de cette source (NSh_{C_i}).

Maintenant, le tableau 2.5 compare l'importance des sources de revenu calculées avec la méthode du Nested Shapley avec celles obtenues avec la décomposition à la Owen et de Shapley (celle-ci ne prenant pas en compte la structure hiérarchique des revenus). Les contributions relatives des sources secondaires par rapport à la valeur de Shapley de leur source primaire sont très proches dans les trois cas, excepté pour la source primaire des transferts sociaux. Dans le cas d'Owen et

Tableau 2.4 – Décomposition à la Nested Shapley

Sources primaires	Revenus du travail		Transferts sociaux		Revenus du capital		Somme
$V_I(C_i)$	0,0523		0,0261		0,1307		0,2092
NSh_{Ci}	0,0436		-0,0087		0,1220		0,1569
Sources secondaires	Salaire	Primes	Famille	Logement	Financier	Immobilier	Somme
NSh_j^*	0,0327	0,0196	0,0000	0,0261	0,1046	0,0261	0,2092
Nsh_j^*	62,50%	37,50%	0,00%	100,00%	80,00%	20,00%	
NSh_j	0,0283	0,0153	-0,0174	0,0087	0,1002	0,0218	0,1569
Nsh_j	65,00%	35,00%	200,00%	-100,00%	82,14%	17,86%	

Note : $Nsh_j^* = NSh_j^*/V_I(C_i)$ et $Nsh_j = NSh_j/NSh(C_i)$.

de Shapley, un revenu distribué égalitairement (dans un jeu égalitaire) a une valeur nulle, alors que dans le cas du Nested Shapley la valeur peut être négative (selon les autres sources secondaires considérées). Ainsi, la valeur d'Owen indique clairement le résultat attendu d'une source de revenu égalitaire, contrairement à la méthode du Nested Shapley. De plus, la valeur associée à l'allocation logement est positive dans le cas du Nested Shapley car c'est la seule source secondaire créatrice d'inégalité au sein de la source primaire des transferts sociaux, alors que cette valeur est négative avec la méthode d'Owen et de celle de Shapley, car cette allocation diminue l'inégalité totale (*i.e.* relativement à toutes les autres sources). Enfin, la valeur relative des sources primaires à l'inégalité totale est similaire entre les deux méthodes prenant en compte la structure hiérarchique (Nested Shapley et Owen) et diffère de celle estimée avec la valeur de Shapley.

Tableau 2.5 – Décomposition à la Nested Shapley, Owen et Shapley

	Revenus du travail		Transferts sociaux		Revenus du capital		
Nested Shapley	Salaire	Primes	Famille	Logement	Financier	Immobilier	Gini
NSh_j	0,0283	0,0153	-0,0174	0,0087	0,1002	0,0218	0,1569
NSh_j/NSh_{Ci}	65,00%	35,00%	200,00%	-100,00%	82,14%	17,86%	
$NSh_j/I(C_M)$	18,06%	9,72%	-11,11%	5,56%	63,89%	13,89%	100%
$NSh_{Ci}/I(C_M)$	27,78%		-5,56%		77,78%		100%
Owen	Salaire	Primes	Famille	Logement	Financier	Immobilier	Gini
Ow_j	0,0272	0,0163	-	- 0,0087	0,0991	0,0229	0,1569
Ow_j/Ow_{Ci}	62,50%	37,50%	0,00%	100,00%	81,25%	18,75%	
$Ow_j/I(C_M)$	17,36%	10,42%	0,00%	-5,56%	63,19%	14,58%	100%
$Ow_{Ci}/I(C_M)$	27,78%		-5,56%		77,78%		100%
Shapley	Salaire	Primes	Famille	Logement	Financier	Immobilier	Gini
Sh_j	0,0292	0,0172	-	- 0,0144	0,1011	0,0237	0,1569
Sh_j/Sh_{Ci}	62,91%	37,09%	0,00%	100,00%	80,98%	19,02%	
$Sh_j/I(C_M)$	18,61%	10,97%	0,00%	-9,17%	64,44%	15,14%	100%
$Sh_{Ci}/I(C_M)$	29,58%		-9,17%		79,58%		100%

Note de lecture : Pour chaque méthode utilisée, les lignes du tableau donnent dans l'ordre : l'importance de chaque source secondaire, l'importance relative de ces sources par rapport à la contribution de la source primaire à laquelle elles appartiennent, l'importance relative des sources secondaires par rapport à l'inégalité totale (*i.e.* quand l'ensemble C_M de sources primaires est pris en compte), et enfin l'importance relative de la source primaire à l'inégalité totale.

2.5.3 Normalisation des distributions de facteurs équivalents

Le seconde divergence évoquée entre la décomposition de Shapley des indices d'inégalité et les propriétés de Shorrocks, c'est que cette décomposition ne respecte pas la normalisation des distributions de facteurs équivalents. Autrement dit, l'importance d'un facteur dont la distribution est égalitaire peut ne pas être nulle.

Pour mieux expliciter ce propos, le cas d'une distribution de trois sources de revenu (i, j, h) répartie entre trois individus (a, b, c) est étudiée. L'importance est ici calculée selon deux approches : le jeu égalitaire ($\lambda = \mu(X^S)$) et le jeu nul ($\lambda = 0$). Dans la première approche, la distribution d'une source de revenu non présente dans S est remplacée par sa moyenne. Dans la seconde approche, un vecteur nul est imposé pour les sources non présentes dans S . Dans les deux cas, la distribution est remplacée par un vecteur égalitaire d'un montant fixé, c'est pourquoi les indices absolus, invariants à une translation, ne sont pas étudiés. Pour comparer la pertinence de ces deux approches, un exemple de décomposition est présenté dans le tableau 2.6, où les inégalités sont mesurées par l'indice de Gini.

La valeur de Shapley associée à la source égalitaire h est nulle dans le cas du jeu égalitaire et est négative dans le jeu nul. L'approche par le jeu égalitaire pourrait être privilégiée puisque celle-ci permet de donner une importance nulle à un facteur égalitaire, ce qui sera toujours vérifié

Tableau 2.6 – Jeu égalitaire VS Jeu nul avec indice de Gini

		Source i	Source j	Source h	Total
		$Y(x^i)$	$Y(x^j)$	$Y(x^h)$	$Y(X)$
	Individu a	3	7	10	20
	Individu b	5	5	10	20
	Individu c	7	3	10	20
		Sh_i	Sh_j	Sh_h	$I(x)$
Importance	$\lambda = 0$	0,02	0,02	-0,04	0
	$\lambda = \mu(X^S)$	0	0	0	0

puisque le vecteur de revenu associé à une source distribuée de manière égalitaire sera toujours le même, qu'il soit dans une coalition S ou pas. Tandis que dans le jeu nul, il est modifié : il passe d'un vecteur égalitaire à un vecteur nul. De même, l'importance d'un vecteur égalitaire dans un jeu nul sera toujours négative, puisque l'indice d'inégalité mesuré sur la distribution des revenus comprenant la source en question sera toujours plus faible que la mesure de l'inégalité faite sur une distribution où le revenu de tous les individus est retranché du même montant. Ceci est vrai à deux conditions : (1) tant que l'indice n'est pas absolu et (2) si le montant retranché est positif. Si le montant retranché est négatif, par exemple dans l'étude de l'importance d'un impôt forfaitaire (ex. timbres fiscaux, redevances...) aux inégalités de revenus nets, le raisonnement contraire s'applique et l'importance est toujours positive. Le signe de l'importance d'une source de revenu égalitaire selon la somme distribuée et le type de jeu considéré est résumé dans le tableau 2.7. En somme, si le chercheur estime que l'importance d'une distribution égalitaire doit être nulle alors il doit privilégier l'approche par jeu égalitaire, au contraire s'il estime que cette importance peut être négative ou positive car cela fait sens économiquement parlant alors il doit privilégier les jeux nuls¹⁷.

Tableau 2.7 – Signe de l'importance d'une source égalitaire

		Constante	
		$c > 0$	$c < 0$
Jeux nuls	$\lambda = 0$	$Sh_i < 0$	$Sh_i > 0$
Jeux égaux	$\lambda = \mu(X^S)$	$Sh_i = 0$	$Sh_i = 0$

Par ailleurs, la distribution de la source i correspond à l'opposé de la source j , leurs importances respectives sont les mêmes dans les deux configurations, positif dans le cas du jeu nul et nul dans le cas du jeu égalitaire. Pour autant, dans les faits deux distributions similaires, c'est-à-dire qui ont le même niveau d'inégalité prise une à une isolément, n'auront probablement pas la même importance

17. Ceci ne serait pas vérifié pour les indices invariants à une translation. C'est pourquoi l'étude est restreinte aux indices relatifs.

puisque l'importance d'une source dépend des autres sources. Par exemple, la source h du tableau 2.6 pourrait être inégalitaire et distribuée en faveur des individus b et c , comme la distribution de la source i . Il en résulterait une inégalité en faveur de b et c provoquée par les sources i et h (*i.e.* importances positives) que la source j compenserait en partie (*i.e.* importance négative).

Dans le tableau 2.8, les importances relatives des facteurs d'une distributions sont données selon plusieurs indices d'inégalité. L'intérêt de cette sous-section est que la distribution de la source j , qui compense en partie les inégalités créées par les autres sources, a des importances positives dans le cas du jeu nul et des importances négatives dans le jeu égalitaire. Ce résultat, en plus de la normalisation des distributions des facteurs équivalents, nous incite à préférer les jeux égalitaires aux jeux nuls.

De plus, d'après les travaux de Sastre et Trannoy (2002), l'importance serait plus sensible au niveau de désagrégation des différentes sources de revenu dans le cas d'un jeu nul que dans un jeu égalitaire. Ainsi, si à cet argument est ajouté le fait qu'une distribution égalitaire a une importance nulle dans un jeu égalitaire, ce qui est dans l'esprit des indices d'inégalité, alors le jeu égalitaire semble plus adapté pour décomposer les indices d'inégalité. Dans la prochaine sous-section, un dernier argument est présenté qui tend à faire préférer le jeu égalitaire au jeu nul.

2.5.4 Incidence de l'indice d'inégalité sur la décomposition

Shorrocks (1982) démontre que si les six propriétés qu'il a défini pour réaliser la décomposition d'un indice d'inégalité par sous-groupes sont respectées, alors l'importance relative de chaque groupe devrait être mesurée par une décomposition de la variance. Cependant, dans notre cas, les propriétés de normalisation des distributions des facteurs équivalents et la dépendance au niveau de désagrégation ne sont pas respectées. Une des conséquences possibles du non-respect de ces axiomes est que l'importance relative des facteurs dans une décomposition à la Shapley des indices d'inégalité est sensible au choix de l'indice.

Dans le tableau 2.8, la décomposition de plusieurs indices d'inégalité est réalisée à partir de la même distribution de revenu. Quel que soit le type de jeu considéré, les indices d'Atkinson et de Theil sont proches alors que l'indice de Gini donne des estimations des importances qui s'en écartent substantiellement, même si l'écart est moins important dans le cas du jeu égalitaire. Toutefois, quand bien même l'ordre des sources en termes d'importances relatives reste inchangé selon les différents indicateurs, il faut noter que différentes mesures d'inégalité peuvent conduire à différentes conclusions.

Par ailleurs, dans le cas d'un jeu nul, seule une source égalitaire peut avoir une valeur de Shapley négative et une source qui tend à réduire l'inégalité totale (ex. transferts sociaux) aura une valeur positive, proche de zéro dans le cas des indices de Theil et d'Atkinson (plus sensibles aux bas de la

Tableau 2.8 – Indices d'inégalité et importances relatives

		Source i	Source j	Source h	Total
		$Y(x^i)$	$Y(x^j)$	$Y(x^h)$	$Y(X)$
	Individu a	30	30	50	110
	Individu b	50	20	100	170
	Individu c	70	10	150	230
		Sh_i	Sh_j	Sh_h	$I(x)$
$\lambda = 0$	Gini	21,19%	16,47%	62,34%	100%
	Theil	10,59%	4,25%	85,16%	100%
	Theil 2	8,78%	3,33%	87,89%	100 %
	Atkinson	9,70%	3,79%	86,51%	100 %
$\lambda = \mu(\mathbf{X}^S)$	Gini	27,78 %	-5,56 %	77,78 %	100%
	Theil	33,71%	-17,05%	83,34%	100%
	Theil 2	34,50%	-17,86%	83,36%	100 %
	Atkinson	34,10%	-17,45%	83,35%	100 %

distribution). Dans le cas d'un jeu égalitaire, une source égalitaire aura une valeur associée nulle et une source qui tend à réduire l'inégalité aura une contribution négative, ce qui est plus appréciable d'un point de vue de l'interprétation économique.

2.6 Discussion et conclusion

Dans ce chapitre, plusieurs méthodes de décompositions des indices d'inégalité ont été traitées. Ces méthodes peuvent reposer sur des arrangements mathématiques des fonctions d'inégalité, ou sur des outils statistiques, ou encore sur un concept de la théorie des jeux coopératifs, en l'occurrence la valeur de Shapley. Elles peuvent s'appliquer à une décomposition par sous-populations et/ou par sources de revenu, ou encore plus généralement par attributs. Elles sont aussi applicables en prenant en considération la structure hiérarchique des revenus.

Par la suite, la décomposition à la Shapley est la principale méthode étudiée. Celle-ci peut être applicable en une seule étape si les différents attributs considérés pour la décomposition ont des montants de revenus qui sont pré-déterminés (ex. : décomposition par sources de revenu), ou elle peut être appliquée après une étape de désagrégation du revenu entre différentes caractéristiques individuelles. De plus, c'est un outil qui peut être utilisé dans une variété de domaines de l'économie appliquée, notamment en indiquant l'importance d'un facteur aux inégalités observées dans une société. Cette décomposition a plusieurs avantages. En premier lieu, ces résultats sont facilement interprétables, elle évite notamment l'introduction d'un terme résiduel dans l'équation de la décomposition (Shorrocks (2013)). De plus cette décomposition peut s'appliquer à tous les indices

d'inégalité, tout en considérant que l'interprétation des résultats obtenus dépend en partie de l'indice d'inégalité choisi. Enfin, cette décomposition permet de prendre en considération des effets de groupes et des effets de facteurs en même temps. Par exemple, des modèles hiérarchiques peuvent être utilisés pour estimer l'importance de chaque sous-population aux inégalités, puis appliquer une décomposition par facteurs au sein de chaque groupe. Une méthode alternative, adoptée dans le chapitre suivant, consiste à apprécier les effets de groupes comme des attributs et permet ainsi d'estimer l'importance de plusieurs effets de groupes en même temps.

Chapitre 3

L'importance du genre à l'inégalité salariale dans la fonction publique : une analyse régionale

Sommaire

3.1	Introduction	76
3.2	Mesure de l'intensité des inégalités liées au genre dans la fonction publique, des disparités importantes entre les régions	79
3.2.1	Méthodologie et données	79
3.2.2	Comparaison interrégionale de l'intensité des inégalités liées au genre sur l'ensemble de la fonction publique	81
3.3	Quel lien avec les spécificités régionales ?	83
3.3.1	Spécificités dans la structure de la fonction publique	83
3.3.2	Un parallèle avec des caractéristiques socio-démographiques et le secteur privé	87
3.4	Intensité régionale des inégalités liées au genre par catégorie et par versant .	89
3.5	Discussion et conclusion	97
3.6	Annexes	99

3.1 Introduction

L'arrivée sur le marché du travail des générations d'après-guerre s'est accompagnée d'une forte augmentation de la participation féminine. Cependant, l'accroissement de la place des femmes dans la vie professionnelle a rapidement soulevé la question du niveau des rémunérations qui leurs sont versées relativement à celles reçues par leurs homologues masculins. Ainsi, dès 1972, une loi (du 22 décembre) vise à assurer l'égalité des rémunérations des hommes et des femmes ayant « un travail de valeur égale », une notion précisée par la loi du 13 juillet 1983 qui l'associe à l'exigence d'un ensemble comparable de connaissances consacrées par un titre, un diplôme ou une pratique professionnelle. En dépit de ces premiers textes, des différences de rémunération que ni la durée, ni la « valeur » du travail ne peuvent expliquer ont continué d'exister. Aussi, une nouvelle loi sur « l'égalité salariale entre les femmes et les hommes », votée le 23 mars 2006, a affiché l'objectif de supprimer totalement les écarts salariaux entre les unes et les autres dans un délai de 5 ans.

Cet arsenal législatif prouve l'implication des gouvernements français successifs dans la lutte contre une moindre rémunération des femmes qui ne serait pas fondée sur des éléments objectifs. Plusieurs études récentes ont pourtant montré qu'un certain manque à gagner des femmes par rapport aux hommes s'observe aussi dans les emplois placés directement sous leur autorité, ceux de la fonction publique (Guégot (2011), DGAFP (2012), Duvivier *et al.* (2015), Fremigacci *et al.* (2016) et Lebon *et al.* (2016)). Si tous ces travaux ne démontrent pas que ces différences de rémunération entre hommes et femmes sont infondées, ils établissent néanmoins qu'elles existent quel que soit le segment considéré de la fonction publique. Notre analyse qui s'intéresse également au secteur public, entend aborder une dimension rarement traitée dans les travaux portant sur les écarts de rémunération entre les hommes et les femmes ; il s'agit de rechercher les éventuelles spécificités régionales en la matière. En effet, il existe de nombreuses comparaisons internationales, avec par exemple Insee et Statistique publique (2017) (p.169) pour une comparaison des écarts de salaires moyens, brut annuel et brut horaire, entre femmes et hommes dans les pays de l'Union Européenne en 2014. Ou encore Christofides *et al.* (2013) pour une étude basée sur une méthode paramétrique avec une décomposition à la Oaxaca-Ransom de l'écart de rémunération entre femmes et hommes entre effets expliqués et inexpliqués dans les pays de l'Union Européenne en 2007, tout en prenant en compte des différences d'effets entre le bas et le haut de la distributions des salaires. Enfin, citons également Lucifora et Meurs (2006) pour une comparaison inter-européenne spécifique au secteur public. Cependant, la littérature n'a que très rarement abordé la dimension territoriale des inégalités liées au genre, du moins dans le cas de la France. Ainsi, seules sont renvoyées les statistiques d'écarts régionaux de salaire moyen entre les femmes et les hommes établis par l'Insee pour le secteur privé et semi-public pour l'année 2010 (Insee (2018)) et à une ancienne comparaison interdépartementale, Charraud et Saada (1974) qui s'appuie sur des données salariales de 1970.

Aucune comparaison dans le temps ne sera néanmoins possible, car ces auteurs se contentent de dresser une carte à partir d'écarts salariaux relatifs considérés par intervalle sans que les valeurs exactes soient précisées, et qui ne sont pas corrigées des différences de temps de travail. A notre connaissance, aucune comparaison de ce type n'a jamais été établie concernant le secteur public, et c'est cette comparaison qui va être l'objet de notre analyse.

A cet effet, les données pour l'année 2010 du Système d'Information sur les Agents des Services Publics (SIASP par la suite, Insee (2010)) sont utilisées. C'est une base exhaustive qui retrace l'ensemble des rémunérations reçues annuellement par les agents travaillant pour l'un des trois versants de la fonction publique : la fonction publique d'Etat (FPE), la fonction publique hospitalière (FPH), la fonction publique territoriale (FPT). C'est donc à partir des véritables distributions de rémunérations qu'il est possible d'appréhender, pour chaque région, le niveau de ce qui pourrait apparaître comme des différences liées au genre, et d'établir une comparaison interrégionale.

Se pose alors la question de l'outil adéquat pour mener à bien l'identification d'éventuels différentiels de rémunération entre les hommes et les femmes. Deux types de réponse sont envisageables, la première repose sur l'évaluation de la rentabilité des caractéristiques individuelles, la seconde sur la décomposition des inégalités. Dans la lignée des articles de Blinder (1973) et Oaxaca (1973), les travaux appartenant au premier de ces deux pans de la littérature (voir par exemple Meurs et Ponthieux (2000)) cherchent à déterminer la part de l'écart de salaire moyen des hommes et des femmes qui relève de différences objectivement fondées. A ce titre, le premier élément qui doit être avancé, est celui du temps de travail en moyenne plus réduit chez les femmes et à l'origine d'une importante partie des différences de rémunération. Cependant, le passage à l'équivalent plein temps permet de s'en abstraire. Viennent ensuite les facteurs dont se déduit la productivité des travailleurs, à savoir leur niveau de diplôme, leur expérience professionnelle et leur ancienneté dans l'entreprise. La rentabilité de ces facteurs de capital humain est estimée pour les hommes d'une part pour les femmes d'autre part. La partie du décalage salarial qui provient d'une moindre rentabilité chez les femmes que chez les hommes de facteurs similaires, peut être considérée comme liée à des problématiques de genre, sans que l'on puisse exclure définitivement qu'elle soit liée à d'autres caractéristiques inobservées.

Si elles sont riches de leur exhaustivité sur les rémunérations versées, les données du SIASP sont extrêmement pauvres concernant les caractéristiques individuelles des agents dont ni le niveau de diplôme, ni le concours d'entrée dans la fonction publique, pas plus que la date de ce concours ne sont connus. Faute des éléments nécessaires pour réaliser une analyse de la rentabilité des caractéristiques individuelles, notre étude s'inscrit uniquement dans la littérature qui s'intéresse à la décomposition des indices d'inégalités. Usuellement, les travaux en question permettent d'expliquer les inégalités observées soit en termes de sources de revenu, en déterminant quelle proportion des inégalités totales est attribuable à la dispersion de chacune des composantes de la rémunéra-

tion (par exemple les salaires et les primes), soit en termes de sous-populations, en répartissant les inégalités totales au prorata des inégalités internes mesurées sur chacune des sous-populations qui composent l'échantillon. C'est l'optique qui est notamment adoptée dans l'article de Terraza *et al.* (2005) qui étudient la distribution des salaires à travers la décomposition spécifique de l'indice de Gini. Il s'agit alors de subdiviser les inégalités observées entre une part expliquée par les différences salariales des hommes entre eux, respectivement des femmes entre elles, une part expliquée par la différence de salaire moyen entre les hommes et les femmes et une part qui dépend du recouvrement plus ou moins important entre les distributions de salaires masculines et féminines. Si cette décomposition permet d'étudier de façon détaillée l'évolution dans le temps de la structure des salaires, elle n'a pas vocation à estimer directement les écarts de rémunération entre les hommes et les femmes. Se démarquant des approches précédentes, la présente étude s'appuie sur une nouvelle solution pour appréhender l'existence et l'importance de tels écarts.

Il s'agit d'une décomposition à la Shapley (1953) des inégalités. Cette méthodologie, inspirée des travaux de Chantreuil et Trannoy (2011, 2013), doit être adaptée pour permettre de déterminer l'importance (*i.e.* la contribution totale) à l'inégalité d'une caractéristique individuelle, ou attribut, des agents composant l'échantillon (Chantreuil et Lebon (2015)). Le principal attribut étudié ici est le genre qui va capter le niveau des inégalités expliquées par les différences de rémunération entre les hommes et les femmes. Un montant nul des inégalités liées au genre signifierait une absence de différentiel entre les rémunérations des hommes et des femmes, alors qu'un montant strictement positif attesterait d'écarts d'autant plus importants que ce montant serait élevé sans pour autant préjuger de l'origine de ces écarts.

L'utilisation de cette méthode sur les données de chacune des régions françaises, 26 en 2010 DOM compris, permet d'établir une cartographie de la valeur des inégalités déterminées par les décalages existant entre les rémunérations des femmes et des hommes. Cet indice reflétant l'intensité des inégalités liées au genre est calculé pour chaque région sur l'ensemble des salariés de la fonction publique, mais également séparément sur chaque versant et sur chaque catégorie hiérarchique des postes. Il montre qu'aucune région, quel que soit le type de salariés de la fonction publique considéré, n'est totalement exempte d'un certain manque à gagner des femmes par rapport aux hommes. Cependant, l'intensité du phénomène est très variable d'une région à l'autre. Les différences ainsi observées peuvent ensuite être rapprochées des spécificités régionales de la fonction publique ou du marché du travail en général pour tenter d'en comprendre l'origine.

La deuxième section présente la méthodologie et les données utilisées, ainsi que les résultats régionaux établis au niveau de la fonction publique dans son ensemble. La section 3 cherche à rapprocher les disparités interrégionales observées de la structure de l'emploi dans la fonction publique de chaque région et d'autres caractéristiques régionales. L'absence de lien direct entre la structure de l'emploi public et les inégalités liées au genre dans ce secteur conduit dans la sec-

tion 4 à détailler la décomposition par catégorie et par versant et permet de mettre en lumière les informations supplémentaires qui s'en déduisent. La section 5 conclut.

3.2 Mesure de l'intensité des inégalités liées au genre dans la fonction publique, des disparités importantes entre les régions

3.2.1 Méthodologie et données

Les analyses présentées dans cet article ont été menées à partir des données, pour l'année 2010, du Système d'Information sur les Agents des Services Publics (SIASP). Cette base retrace les rémunérations et les durées annuelles du travail pour l'ensemble des salariés de la fonction publique et donne des précisions sur la localisation de leur établissement d'exercice, ouvrant ainsi la possibilité de réaliser des études territorialisées. Une fois éliminés les emplois aidés, occasionnels ou payés à l'acte, ainsi que les lignes comportant des informations manquantes ou aberrantes¹, les données comprennent 4 637 701 agents des services publics (voir Lebon *et al.* (2016) pour plus de précisions concernant le traitement de la base de données). La rémunération globale de chaque agent correspond à la somme de son salaire et de ses primes qui sont malheureusement indissociables des autres suppléments de rémunération. Cette rémunération est ramenée à son équivalent temps plein annualisé afin d'obtenir des résultats indépendants des différences de durée du travail dues au temps partiel comme aux heures complémentaires ou supplémentaires. Les discontinuités qui peuvent exister concernant le taux de rémunération horaire, notamment en ce qui concerne les salariés à 80 ou 90% payés un peu plus que proportionnellement à ces quotités, les primes des travailleurs à temps partiel ou bien le paiement des heures complémentaires ou supplémentaires, créent à l'évidence quelques biais dans les rémunérations équivalent temps plein obtenues. Dans la plupart des cas, il est impossible de savoir comment ces biais potentiels influencent la rémunération relative des hommes et des femmes. Mais, en ce qui concerne la correction des temps partiels des agents travaillant à 80 et 90%, il est clair qu'elle déforme légèrement la distribution des rémunérations équivalent temps plein en faveur des femmes qui sont de très loin les plus nombreuses à être concernées par ces durées.

La méthodologie mise en œuvre pour mesurer les inégalités liées au genre au niveau national comme à l'échelle de chaque région est celle proposée par Chantreuil et Lebon (2015). Elle suppose de réécrire la rémunération de chaque agent comme la somme de trois éléments : un élément lié à son âge qui sert de proxy à la durée de la carrière qui n'est pas connue, cet élément est identique

1. Ces restrictions ont conduit à éliminer 1 128 259 observations.

pour tous les agents du même âge dans la sous-population concernée et égal à leur rémunération moyenne ; un élément spécifique au genre prenant la même valeur pour tous les agents du même âge et du même sexe et égal à la différence entre cette rémunération moyenne et celles des agents du même âge et du même sexe ; et un élément spécifique à chaque agent.

La rémunération R d'un agent de la fonction publique i d'âge a et de genre g (avec $g = f, m$) s'écrit donc :

$$R(a, g, i) = \bar{R}(a) + [\bar{R}(a, g) - \bar{R}(a)] + [R(a, g, i) - \bar{R}(a, g)] \quad (3.1)$$

Le premier élément de cette somme prend une valeur strictement positive. Le deuxième élément est positif ou négatif selon que les agents de genre g et d'âge a gagnent en moyenne plus ou moins que l'ensemble des agents d'âge a dans la sous-population considérée. Le dernier élément peut être positif ou négatif selon que cet agent gagne plus ou moins que la moyenne des agents de son sexe et de son âge appartenant à la même sous-population.

Sont ainsi construites trois distributions, dites respectivement « âge », « genre » et « reste », dont le cumul permet de reconstituer la distribution observée des rémunérations. Sur cette base, il est possible d'appliquer à l'indice servant à mesurer les inégalités globales de rémunération, il s'agit ici de l'indice de Gini, une décomposition à la Shapley (Chantreuil et Trannoy (2011, 2013)) capable de déterminer la part des inégalités attribuable à chacun des trois éléments précédemment décrits. Analytiquement, ces trois attributs sont ainsi traités comme autant de sources de revenu.

La décomposition mise en œuvre pour étudier les inégalités de genre dans la fonction publique, permet donc d'appréhender et de décrire les disparités de rémunération qui peuvent exister entre les hommes et les femmes, et ce en prenant en compte l'ensemble de la distribution salariale (et pas seulement des différences qui peuvent exister au niveau de la moyenne ou de la médiane). Les résultats pourraient être affinés par l'introduction dans la décomposition de facteurs supplémentaires tels que le diplôme ou le concours d'entrée dans la fonction publique, mais, ainsi qu'il a été relevé précédemment, ces informations sont totalement absentes de notre base de données. Notons que le manque d'éléments tangibles sur ces questions, comme sur la durée réelle de la carrière, pourrait induire un biais dans nos résultats. Ainsi, si, en général, les femmes de la fonction publique étaient moins diplômées que les hommes, si elles étaient entrées par des concours hiérarchiquement moins élevés, ou si elles avaient connu plus d'interruptions de carrière, ces différences objectivement fondées se trouveraient néanmoins captées par la variable « genre ». Cette observation doit inciter à la prudence dans l'interprétation des résultats.

L'indice qui mesure les inégalités totales, ici l'indice de Gini, est décomposé à la Shapley, ce qui permet d'obtenir la relation suivante :

$$Gini = Sh_\alpha + Sh_\gamma + Sh_\rho \quad (3.2)$$

Où Sh_α représente l'importance de l'âge, Sh_γ celle du genre, et Sh_ρ celle du reste.

L'indice de Gini est notée $G(.)$ pour la distribution correspondante, pour chaque individu, à la somme des termes indiqués dans la parenthèse. Sachant que les distributions des éléments de la rémunération associées à l'âge, au genre et au reste sont notés respectivement α , γ et ρ^2 , et que μ_j (avec $j = \alpha, \gamma, \rho$) représente la valeur moyenne de la variable indicée sur l'ensemble de la distribution, la valeur de Shapley de chaque attribut se calcule de la façon suivante :

$$\begin{aligned}
 Sh_\alpha &= \frac{1}{3}G(\alpha, \mu_\gamma, \mu_\rho) + \frac{1}{6}[G(\alpha, \gamma, \mu_\rho) - G(\mu_\alpha, \gamma, \mu_\rho)] + \frac{1}{6}[G(\alpha, \mu_\gamma, \rho) - G(\mu_\alpha, \mu_\gamma, \rho)] \\
 &\quad + \frac{1}{3}[G(\alpha, \gamma, \rho) - G(\mu_\alpha, \gamma, \rho)] \\
 Sh_\gamma &= \frac{1}{3}G(\mu_\alpha, \gamma, \mu_\rho) + \frac{1}{6}[G(\alpha, \gamma, \mu_\rho) - G(\alpha, \mu_\gamma, \mu_\rho)] + \frac{1}{6}[G(\mu_\alpha, \gamma, \rho) - G(\mu_\alpha, \mu_\gamma, \rho)] \\
 &\quad + \frac{1}{3}[G(\alpha, \gamma, \rho) - G(\alpha, \mu_\gamma, \rho)] \\
 Sh_\rho &= \frac{1}{3}G(\mu_\alpha, \mu_\gamma, \rho) + \frac{1}{6}[G(\alpha, \mu_\gamma, \rho) - G(\alpha, \mu_\gamma, \mu_\rho)] + \frac{1}{6}[G(\mu_\alpha, \gamma, \rho) - G(\mu_\alpha, \gamma, \mu_\rho)] \\
 &\quad + \frac{1}{3}[G(\alpha, \gamma, \rho) - G(\alpha, \gamma, \mu_\rho)]
 \end{aligned} \tag{3.3}$$

L'importance du genre, Sh_γ , obtenue à l'aide de cette décomposition est ainsi égale à la somme pondérée :

- du Gini d'une distribution de rémunération dans laquelle l'âge et le reste prendraient leur valeur moyenne sur l'ensemble de la distribution et le genre ses valeurs réelles $[G(\mu_\alpha, \gamma, \mu_\rho)]$;
- de la différence entre le Gini d'une distribution dans laquelle les attributs âge et genre prendraient leurs valeurs réelles et le reste sa valeur moyenne et le Gini d'une distribution où le genre serait aussi ramené à sa valeur moyenne $[G(\alpha, \gamma, \mu_\rho) - G(\alpha, \mu_\gamma, \mu_\rho)]$;
- de la différence entre le Gini d'une distribution dans laquelle le genre et le reste prendraient leurs valeurs réelles et l'âge sa valeur moyenne et le Gini d'une distribution où le genre serait aussi ramené à sa valeur moyenne $[(\mu_\alpha, \gamma, \rho) - G(\mu_\alpha, \mu_\gamma, \rho)]$;
- et de la différence entre le Gini de la distribution effective des revenus et le Gini de la distribution lorsque le genre est ramené à sa valeur moyenne $[G(\alpha, \gamma, \rho) - G(\alpha, \mu_\gamma, \rho)]$.

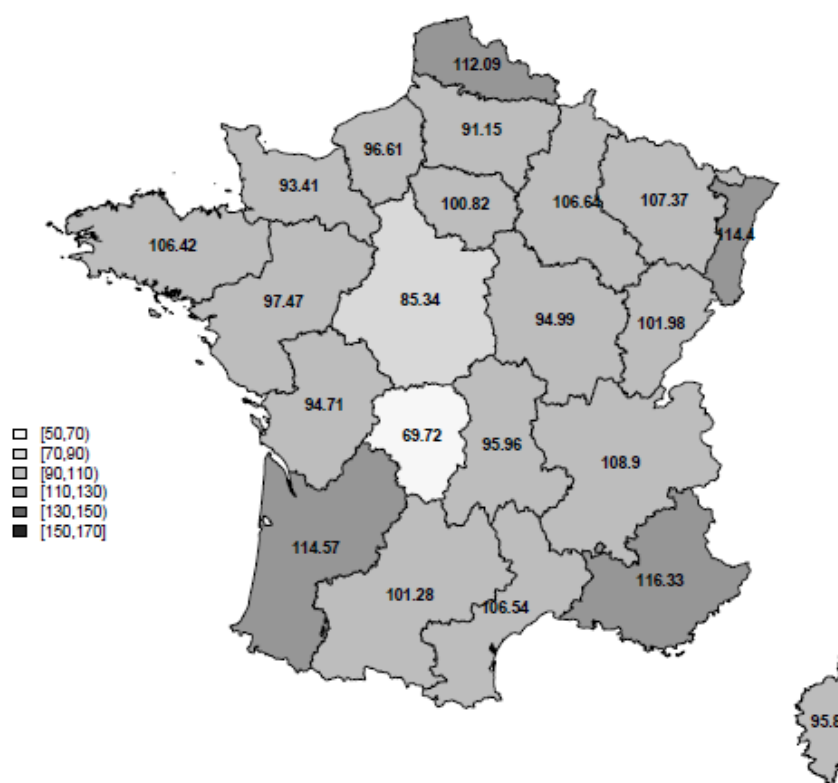
En l'absence de tout décalage entre les distributions des rémunérations respectives des hommes et des femmes, l'importance de l'attribut genre est nulle, soit $Sh_\gamma = 0$. Lorsqu'un tel décalage existe Sh_γ prend une valeur strictement positive, une valeur d'autant plus élevée que le décalage est important. De fait, quelle que soit la sous-population étudiée, les décalages observés existent et traduisent systématiquement un manque à gagner des femmes par rapport aux hommes. Il faut garder en tête cette réalité pour lire les résultats à venir concernant les valeurs de Sh_γ qui sont interprétées comme un indicateur de l'intensité des inégalités expliquées par le genre.

2. Soit, pour un individu i d'âge a et de genre g : $\alpha_i = \bar{R}(a)$, $\gamma_i = \bar{R}(a, g) - \bar{R}(a)$ et $\rho_i = R(a, g, i) - \bar{R}(a, g)$.

3.2.2 Comparaison interrégionale de l'intensité des inégalités liées au genre sur l'ensemble de la fonction publique

Carte 3.1 – Intensités régionales des inégalités liées au genre dans la fonction publique

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

Conformément à la méthodologie précédemment présentée, l'intensité des inégalités liées au genre a été évaluée pour chaque région française au niveau de la fonction publique dans son ensemble³. Pour des raisons de lisibilité et de meilleure comparabilité des résultats, tous les indicateurs d'intensité de toutes les sous-populations de l'ensemble des régions ont été ramenés à une base 100 unique. La base choisie est la valeur de l'indice d'intensité mesurée au niveau national de la fonction publique toutes catégories et tous versants confondus. Cette intensité qui va servir d'étalon, se calcule à partir des résultats obtenus pour la fonction publique dans son ensemble : l'importance du genre prend alors une valeur égale à 0,01975⁴. Lorsque l'indicateur associé à une

3. L'intégralité des résultats des décompositions est proposée en annexe dans le tableau 3.5

4. Sachant que l'indice de Gini sur la distribution de l'ensemble des rémunérations de la fonction publique est de 0,2209 (voir Annexe, tableau 3.6).

région prend une valeur supérieure (resp. inférieure) à 100 (sur la carte 3.1 pour les régions de France métropolitaine et dans le tableau 3.1 pour les DOM), cela signifie que dans la région en question l'intensité des inégalités de genre est supérieure (resp. inférieure) à ce qui est observé au niveau national.

La carte 3.1 et le tableau 3.1 font apparaître de très gros contrastes entre les régions françaises. Si l'Île-de-France affiche une intensité des inégalités de genre presque identique à celle de la fonction publique nationale, ce n'est pas le cas pour l'ensemble des régions. En effet, les scores des différentes régions métropolitaines s'échelonnent de moins de 70 dans le Limousin à plus de 115 en Provence-Alpes-Côte d'Azur (PACA). C'est cependant l'outremer qui semble le moins affecté par les inégalités de rémunérations entre hommes et femmes avec notamment un score de moins de 50 pour la Réunion. Plus généralement, ce sont les régions frontalières du sud, de l'est et du nord de la France métropolitaine qui semblent les plus concernées par les inégalités liées au genre. Alors que les régions centrales et de l'ouest, ainsi que la Corse et l'ensemble des DOM, sont moins impactés.

Tableau 3.1 – Intensités des inégalités liées au genre dans les DOM au niveau de la fonction publique, par catégorie et par versant (Fonction publique nationale, base 100)

Régions	Fonction publique	Catég. A	Catég. B	Catég. C	FPE	FPH	FPT
Guadeloupe	78,31	99,32	95,48	103,36	105,74	48,8	110,8
Martinique	62,36	97,29	82,61	76,34	109,97	42,3	93,94
Guyane	76,88	106,79	85,44	73,8	109,7	100,78	55,96
La Réunion	48,47	93,77	63,81	46,9	116,49	49,96	66,97

Champ : France entière, salariés de la fonction publique. Source : SIASP 2010, Insee, Traitement des auteurs.

L'origine de la faiblesse ou de l'importance des inégalités de genre ainsi obtenue est-elle à rechercher dans la composition de l'emploi public ? La réponse à cette question suppose de répertorier les caractéristiques régionales dans la structure de la fonction publique qui pourrait ensuite être confrontées avec les intensités d'inégalité qui ont été établies. Mais d'autres spécificités régionales pourraient également éclairer les résultats obtenus comme il sera constaté ultérieurement.

3.3 Quel lien avec les spécificités régionales ?

3.3.1 Spécificités dans la structure de la fonction publique

Parce qu'il est ici question d'un secteur bien particulier de l'économie, la première idée pour expliquer que l'intensité des inégalités de genre soit différente selon les régions est que celles-ci soient liées aux caractéristiques locales de la fonction publique. L'origine des disparités interrégio-

nales pourrait ainsi être à rechercher dans la structure par catégorie et par versant de l'emploi dans la fonction publique dans chaque région (voir les tableaux 3.2 et 3.3). Cette double subdivision de l'emploi public doit préalablement être rapidement présentée. Sont considérées comme des salariés de la fonction publique, les personnes qui, de façon générale, travaillent pour les administrations publiques. Dans les administrations publiques, les établissements employeurs sont de nature différente et cette typologie détermine l'appartenance des postes à l'un des trois versants de la fonction publique. Tous ceux qui relèvent de l'action de l'Etat, au niveau central (ministères,...) comme en région (préfectures, rectorats,...) ainsi que les établissements d'enseignement public (collèges, lycées, universités,...) appartiennent à la fonction publique d'Etat (FPE). Les collectivités territoriales (communes, départements,...) emploient des personnels qui relèvent de la fonction publique territoriale (FPT). Du fait de la décentralisation, la FPT a vu sa taille augmenter dans la période récente. La fonction publique hospitalière (FPH) regroupe principalement les agents travaillant dans les hôpitaux et les maisons de retraite publiques. Les trois catégories correspondent quant à elles à une classification hiérarchique des emplois. En référence au classement par PCS de l'Insee, les postes de la catégorie A peuvent être assimilés à des emplois de cadre, ceux de la catégorie B à des professions intermédiaires, et ceux de la catégorie C à des statuts d'ouvriers/employés.

Si une grande partie des écarts de rémunération captés par notre indicateur était déterminée par des difficultés d'accès des femmes aux emplois les plus valorisés de la fonction publique conformément aux résultats obtenus par d'autres études (voir par exemple Fremigacci *et al.* (2016)), la répartition régionale des emplois entre catégories pourrait en atténuer ou au contraire en exacerber la valeur. Une plus grande concentration des emplois dans les catégories B et C pourrait ainsi aboutir à une valeur plus faible de l'indice d'intensité, et inversement quand la région comprend une plus forte proportion de postes de la catégorie A, puisque c'est à cette catégorie qu'appartiennent ces emplois les mieux rémunérés. Au-delà de cette répartition des postes, un taux de féminisation plus élevé de la catégorie A pourrait révéler à l'inverse un accès plus aisé des femmes à l'ensemble des types d'emplois y compris les mieux rémunérés et donner une plus faible importance du genre aux inégalités régionales de rémunération, alors qu'une forte féminisation des catégories B et C conduirait au résultat opposé. Même s'il n'existe pas d'a priori quant au fonctionnement relatif de chacun des versants qui puisse faire attendre tel ou tel type de résultat d'une concentration plus ou moins forte de l'emploi public régional, respectivement dans la FPE, la FPH ou la FPT, c'est néanmoins l'occasion de s'interroger également sur la possibilité d'un éventuel impact de la répartition des postes entre les différents types d'administration.

Dressons donc un rapide bilan régional de l'emploi dans la fonction publique. La distribution des postes entre les régions suit approximativement celle de la population française (voir tableau 3.2 et Annexe, tableau 3.6) à ceci près que l'Île-de-France, avec 20,41% des emplois publics, en regroupe une proportion notablement supérieure à celle de la part de la région dans la population

Tableau 3.2 – Part de chaque région dans l’emploi public, répartition des postes entre les catégories A,B, C et taux de féminisation de chaque sous-population (en %)

Régions	Fonction publique		Catégorie A		Catégorie B		Catégorie C	
	Part	Taux	Répartition	Taux	Répartition	Taux	Répartition	Taux
	nationale	féminisation	régionale	féminisation	régionale	féminisation	régionale	féminisation
Île-de-France	20,41	63,28	34,49	62,52	22,92	62,59	42,6	64,27
Champagne-Ardenne	2,02	65,62	31,71	62,5	18,83	70,98	49,45	65,59
Picardie	2,72	66,56	32,29	64,09	18,7	72,7	49,01	65,85
Haute-Normandie	2,79	66,77	31,44	64,5	18,74	69,03	49,83	67,36
Centre	3,67	66,5	30,28	64,99	18,44	71,42	51,28	65,63
Basse-Normandie	2,25	65,03	31	62,4	18,93	71,02	50,07	64,39
Bourgogne	2,5	66,43	30,05	63,09	20,68	72,29	49,27	66,01
Nord-Pas-de-Calais	6,1	62,5	33,39	61,84	19,24	63,05	47,37	62,74
Lorraine	3,42	64,42	32,85	62,18	21,24	68,2	45,91	64,27
Alsace	2,51	64,56	36,61	61,56	21,39	69,68	42	64,58
Franche-Comté	1,73	66,54	33,32	62,07	19,75	73,68	46,93	66,7
Pays de la Loire	6,03	66,87	31,88	63,37	19,96	70,15	48,16	67,83
Bretagne	3,91	64,61	31,19	62,18	19,81	66,86	49	65,24
Poitou-Charentes	2,7	65,62	28,93	62,37	18,32	70,98	52,76	65,55
Aquitaine	4,93	63,6	29,99	62,19	19,28	65,96	50,73	63,54
Midi-Pyrénées	4,58	63,49	31,01	62,58	19,63	66,89	49,37	62,72
Limousin	1,33	64,4	26,73	61,61	21,12	70,71	52,15	63,27
Rhône-Alpes	9,05	66,15	34,68	63,96	19,75	70,45	45,58	65,95
Auvergne	2,21	64,17	29,36	62,28	18,66	69,07	51,98	63,47
Languedoc-Roussillon	3,99	61,75	31,33	61,16	19,04	64,23	49,62	61,18
Provence-Alpes-Côte d’Azur	7,48	61,71	31,05	62,19	19,85	62,94	49,1	60,91
Corse	0,51	53,98	27,92	57,91	19,09	58,42	52,99	50,32
Guadeloupe	0,74	61,03	31,43	62,13	16,89	62,85	51,69	59,76
Martinique	0,75	62,33	31,33	63,61	18,99	66,18	49,68	60,05
Guyane	0,36	55,54	35,52	57,02	18,08	55,43	46,4	54,44
La Réunion	1,31	54	35,17	57,2	13,94	58,28	50,89	50,63
France entière	100	64,08	32,39	62,58	20,15	66,65	47,47	64,02

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

Lecture : L’Île-de-France regroupe 20,41% des salariés de la fonction publique parmi lesquelles 63,28% de femmes. Ces salariés se répartissent entre les catégories A, B et C dans des proportions respectivement égales à 34,49%, 22,92% et 42,60% qui sont respectivement féminisées à hauteur de 62,52%, 62,59% et 64,27%.

nationale du fait de la concentration de l’administration centrale dans la capitale. La Bretagne apparaît à l’inverse quelque peu sous dotée en emplois du secteur public relativement au nombre de ses habitants.

Le secteur tout entier est extrêmement féminisé, à hauteur de 64% au niveau national, et la très grande majorité des régions affichent un taux de féminisation supérieur à 61% sachant que ce taux reste partout inférieur à 67%. Rompant avec cette homogénéité, trois régions affichent une proportion de salariées nettement plus faible, aux alentours de 54%, à savoir la Corse, la Guyane et la Réunion. Le moindre taux de féminisation de l’emploi public dans ces trois régions qui se vérifie pour chacune des trois catégories d’agents, est probablement à rapprocher d’un taux d’activité des femmes relativement plus faible que sur le reste du territoire (Insee (2016b)). Dans la plupart des régions, les taux de féminisation restent assez similaires dans les catégories A et C (presque toujours compris entre 62 et 65%), en revanche la catégorie B compte souvent une proportion de femmes plus élevée et plus disparate d’une région à l’autre, puisqu’elle dépasse 70% dans un grand nombre d’entre elles, mais elle descend en dessous de 60% dans plusieurs.

A contrario, la répartition intra-régionale de l'emploi public par catégorie fait apparaître une proportion assez stable des postes de la catégorie B d'une région à l'autre (comprise sur un intervalle de 18 à 22%, à l'exception de la Réunion et dans une moindre mesure de la Guadeloupe) et beaucoup plus d'hétérogénéité sur la part des A et des C avec dans les deux cas un écart de 10 points entre les régions dans lesquelles la concentration est la plus faible et celles dans lesquelles elle est la plus forte. Ainsi le schéma du Limousin (avec 26,73% de A et 52,15% de C) s'oppose-t-il radicalement à celui de l'Alsace (avec 36,61% de A et 42% de C).

La répartition par versant des emplois de la fonction publique permet d'expliquer certains des résultats mis en lumière au paragraphe précédent. Ainsi, la faible proportion d'emplois de la catégorie C en Alsace peut-elle s'expliquer par une proportion particulièrement réduite de salariés de la FPT (28,64%) qui appartiennent pour les trois-quarts d'entre eux à la catégorie C. De même, une très faible proportion d'agents de la FPH à la Réunion (13,32%) conduit logiquement à une faible proportion d'agents de la catégorie B, catégorie dont dépendent les infirmières en 2010. Il faut également remarquer en Île-de-France une part très élevée de salariés de la FPE, ce qui est logique du fait de l'implantation des ministères à Paris, mais sans que cela déséquilibre réellement la répartition par catégorie. Plus généralement, le tableau 3.3 révèle la diversité des régions en termes de répartition de l'emploi public par versant. Dans les trois cas, qu'il s'agisse de la FPE, de la FPH ou de la FPT, l'écart entre les régions qui comportent respectivement la plus faible proportion et la plus forte proportion de postes qui dépendent d'un versant particulier, est de 14 points.

L'analyse précédente a mis en lumière les importantes différences qui existent d'une région à l'autre sur des éléments tels que le taux de féminisation et les répartitions de postes par catégorie et par versant. Reste à savoir si un lien peut effectivement être envisagé entre ces éléments et l'intensité régionale des inégalités liées au genre.

Le faible taux de féminisation de l'emploi public observé dans les DOM où l'intensité des inégalités attribuable au genre est parmi les plus limitées, ne se retrouve pas dans les régions métropolitaines d'intensité comparable que sont le Limousin et le Centre. En outre, plusieurs des régions où cette intensité est à l'inverse forte, Nord-Pas-de-Calais, Aquitaine, PACA, comptent moins de salariées femmes que la moyenne nationale. Sans surprise, les distributions régionales d'intensité et de taux de féminisation n'apparaissent que peu corrélées.

De même, l'indicateur d'intensité ne semble pas pouvoir être associé aux proportions de postes appartenant aux différentes catégories hiérarchiques. Cette conclusion était également attendue étant données les disparités dans la situation des régions affichant les valeurs les plus faibles (resp. élevées) pour les inégalités associées à l'attribut genre. Ainsi, le Limousin fait apparaître des proportions particulièrement faible de catégorie A et forte de catégorie B, alors que la Réunion est dans un cas exactement inverse. Le Limousin possède aussi une part importante de postes C, alors qu'ils sont particulièrement peu nombreux en Guyane. Les répartitions de poste par catégories ne

Tableau 3.3 – Répartition de l’emploi public régional entre les trois versants de la fonction publique et taux de féminisation de chaque sous-population (en %)

Régions	FPE		FPH		FPT	
	Répartition régionale	Taux féminisation	Répartition régionale	Taux féminisation	Répartition régionale	Taux féminisation
Île-de-France	46,28	58,29	18,27	74,41	35,45	64,07
Champagne-Ardenne	39,74	59,16	28,02	79,92	32,25	61,17
Picardie	36,61	63,57	29,45	77,31	33,94	60,47
Haute-Normandie	36,57	62,43	24,93	79,85	38,5	62,44
Centre	36,76	62,02	26,68	80,63	36,56	60,7
Basse-Normandie	36,57	59,34	28,18	78,89	35,25	59,84
Bourgogne	36,24	61,62	29,46	79,6	34,3	60,19
Nord-Pas-de-Calais	39,61	58,23	23,56	75,02	36,83	59,09
Lorraine	42,5	59,04	26,3	79,39	31,19	59,12
Alsace	42,01	59,73	29,34	80,51	28,64	55,32
Franche-Comté	38,36	61,35	28,38	80,94	33,26	60,23
Pays de la Loire	37,32	61,47	26,86	81,08	35,82	61,84
Bretagne	39,77	57,84	24,72	79,02	35,51	62,15
Poitou-Charentes	33,86	60,87	26,02	79,11	40,12	60,89
Aquitaine	37,23	59,18	22,72	77,45	40,05	59,86
Midi-Pyrénées	38,95	59,31	21,09	78,07	39,96	59,88
Limousin	32,91	58,8	32,06	78,04	35,03	57,17
Rhône-Alpes	39,16	62,14	24,4	79,48	36,44	61,53
Auvergne	36,51	58,92	27,53	79,16	35,97	58,02
Languedoc-Roussillon	37,27	58,44	20,61	75	42,12	58,2
Provence-Alpes-Côte d’Azur	40,1	57,69	20,12	75,24	39,78	58,93
Corse	38,85	55,85	18,33	69,28	42,82	45,74
Guadeloupe	41	59,26	17,96	69,16	41,04	59,23
Martinique	38,54	61,3	24	69,94	37,46	58,51
Guyane	47,27	53,53	14,23	67,51	38,5	53,58
La Réunion	44,38	56,41	13,32	67,07	42,29	47,37
France entière	40,09	59,47	23,28	77,6	36,63	60,54

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

sont pas plus homogènes dans les départements à haut niveau d’intensité.

Contrairement aux inégalités liées au genre, la répartition par catégorie est très directement associée à l’importance de chaque versant dans chaque région. En effet, la forte concentration de A dans la FPE et de C dans la FPT détermine un lien direct entre les distributions d’emplois par catégorie et par versant, mais pas plus la seconde que la première ne semble pouvoir être, dans les faits, reliée à notre indicateur d’intensité.

Puisque les différences dans la structure de l’emploi public n’expliquent globalement pas les disparités interrégionales de l’intensité des inégalités liées au genre dans ce secteur, se pose alors la question si à l’inverse les spécificités régionales se reproduisent à l’intérieur de chaque catégorie d’agent et de chaque versant. Il s’agit de vérifier si les tendances régionales concernant des écarts de rémunération plus ou moins élevés entre les hommes et les femmes se vérifient ou non sur les différents segments de la fonction publique. Auparavant, l’analyse des résultats au niveau agrégé est menée par une mise en parallèle avec des spécificités régionales, autres que celles qui concernent la fonction publique.

3.3.2 Un parallèle avec des caractéristiques socio-démographiques et le secteur privé

Si les particularités locales de la fonction publique ne semblent pas impacter les inégalités liées au genre qui sont mesurées, il se pourrait que ces inégalités aient un lien avec d'autres caractéristiques régionales, et en particulier avec celle ayant trait à la disparité des situations des hommes et des femmes sur le marché du travail. C'est pourquoi les écarts de salaires moyens calculés par l'Insee pour chacune des régions pour l'année 2010 dans les secteurs privé et semi-public sont utilisés pour capter cette particularité régionale⁵.

Bien que cette étude de l'Insee (Insee (2018)) utilise une approche différente de la notre pour appréhender les disparités de traitement entre hommes et femmes, de très intéressantes similitudes apparaissent. Les écarts salariaux moyens entre les hommes et les femmes sont en effet particulièrement faibles dans les DOM, et c'est à la Réunion qu'ils sont le plus limités, ce qui coïncide très précisément avec les résultats obtenus avec la décomposition des inégalités dans le secteur public. De la même façon, les deux évaluations identifient le Limousin comme la région métropolitaine la moins concernée par les inégalités associées au genre. Le parallèle se poursuit lorsque l'on considère les régions affectées par les plus fortes inégalités. En effet, les trois régions dans lesquelles les décalages salariaux hommes-femmes sont les plus importants dans le secteur privé, à savoir l'Alsace, Rhône-Alpes et PACA sont également les trois régions pour lesquelles notre indicateur d'intensité prend la valeur la plus élevée. Plus généralement, les taux de corrélation entre les écarts salariaux moyens du secteur privé proposés au tableau 3.6 et les intensités des inégalités liées au genre données dans la carte 3.1 et au tableau 3.1 atteint un niveau très élevé : -0,86. Le signe négatif du taux de corrélation est lié au signe, lui-même négatif, des écarts de rémunération relatifs dans le secteur privé, car l'Insee mesure la différence entre le salaire moyen des femmes et le salaire moyen des hommes qui est inférieure à zéro dans toutes les régions, et la rapporte au salaire moyen des hommes.

Le rapprochement de deux ensembles de valeurs numériques issues de méthodologies très différentes risque de conduire à observer et à interpréter une corrélation fallacieuse. Pour éviter un tel écueil, les résultats des deux études sont ramenés à un simple classement des 26 régions dans l'ordre croissant (ou décroissant) des disparités de revenu entre hommes et femmes, afin de calculer la corrélation entre ces deux classements. Le taux de corrélation ainsi déterminé n'est que légèrement plus faible en valeur absolue que celui obtenu sur les valeurs numériques, puisqu'il est égal à 0,77. Cela confirme notre premier résultat. En dépit du caractère particulier de l'emploi public, il semble que les inégalités liées au genre en termes de rémunération obéissent largement à une logique territoriale dont les déterminants restent à identifier.

5. Les chiffres de l'Insee (Insee (2018)) sont proposés en Annexe, tableau 3.6

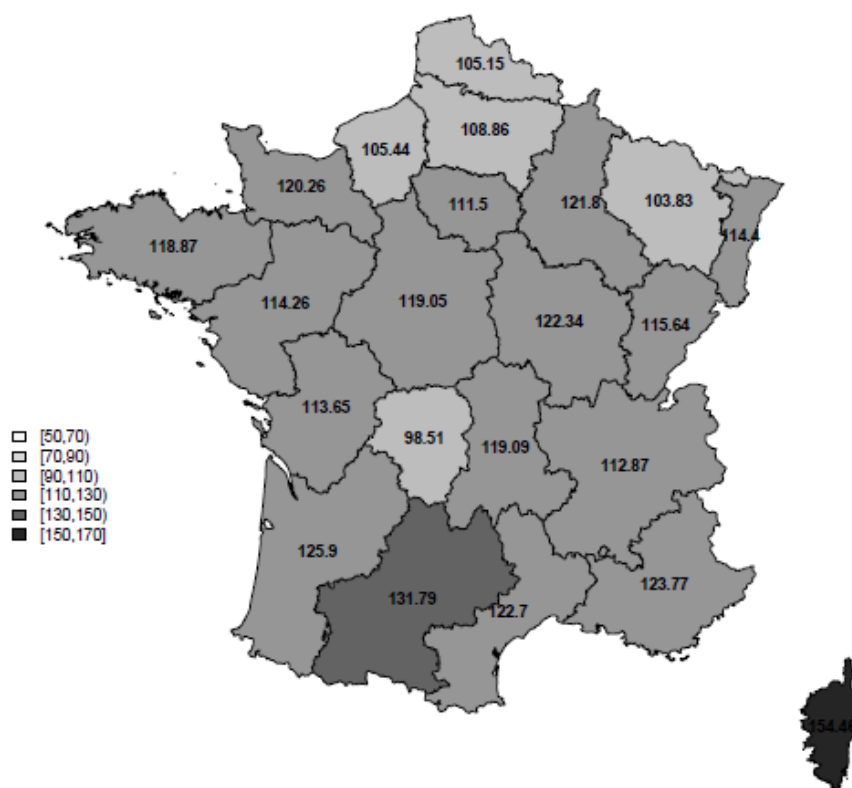
La présence de grandes métropoles régionales dans plusieurs des zones les plus impactées par les inégalités attribuables au genre pourrait inciter à rechercher un lien entre un habitat plus ou moins urbain et l'intensité de ces inégalités, mais aucune corrélation n'a pu être établie⁶. En revanche, une corrélation négative (d'un niveau assez moyen, -0,64) a été obtenue en rapprochant l'indice d'intensité régional et une variable a priori plus inattendue, à savoir le nombre d'enfants par famille⁷. Un nombre d'enfants élevé s'accompagnerait donc d'inégalités liées au genre dans la fonction publique plus faibles que dans les régions où il y en a moins. Il faut se méfier de toute sur-interprétation d'une corrélation qui pourrait n'être que le fruit du hasard, et les pistes d'explications qui existent, doivent être évoquées avec d'autant plus de prudence que l'absence du nombre d'enfants des agents dans la base SIASP ne permet pas de vérifications. Avant d'exposer ces pistes, deux éléments supplémentaires méritent d'être indiqués. D'une part, l'indicateur d'intensité est nettement plus corrélé (négativement) au nombre d'enfants qu'il ne l'est (négativement aussi) au taux d'activité des femmes dans la région (Insee (2016b)). D'autre part, les écarts salariaux du secteur privé sont beaucoup plus difficiles à relier au nombre d'enfants par famille que ceux du secteur public. La juxtaposition de ces résultats peut suggérer l'interprétation suivante. Dans les territoires où il y a en moyenne plus d'enfants, les femmes diplômées pourraient choisir plus fréquemment d'intégrer la fonction publique où la pénalité salariale liée à la maternité est plus réduite que dans le secteur privé (Meurs *et al.* (2010), Duvivier et Narcy (2015)), et les régions concernées pourraient ainsi concentrer dans ce secteur des salariées plutôt plus qualifiées qu'ailleurs, d'où, sur l'ensemble de la distribution des rémunérations, des écarts plus faibles avec les rémunérations des hommes que dans les autres régions. Cet effet de structure, s'il existe, n'aurait pas été correctement capté par le taux de féminisation des différentes catégories. Notons cependant que si un tel phénomène a été identifié par des travaux étrangers, notamment par Nielsen *et al.* (2004) pour la Suède, il n'est pas obtenu par Duvivier et Narcy (2015) dans le cas de la France. Une explication plus simple pourrait être liée au supplément familial de traitement qui est malheureusement indissociable dans les données des autres compléments de revenu (primes,...). S'il était touché plus fréquemment par les femmes, lorsqu'elles sont salariées de la fonction publique, que par les hommes, il pourrait en résulter un resserrement global des rémunérations des hommes et des femmes dans les régions où il y a le plus d'enfants. Faute de pouvoir étayer ou repousser l'une ou l'autre de ces hypothèses, il est impossible d'aller plus loin dans l'analyse.

6. Les données utilisées pour faire ce calcul (Insee et Statistique publique (2017)) portent sur la population en 2010.

7. Les données sont issues du recensement de la population (Insee (2013)).

Carte 3.2 – Intensités régionales des inégalités liées au genre dans la catégorie A

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

3.4 Intensité régionale des inégalités liées au genre par catégorie et par versant

Faute que la structure de la fonction publique ait pu expliquer l'intensité des inégalités liées au genre, il convient de se demander si les tendances régionales observées sur l'ensemble de la fonction publique se retrouvent lorsque l'on considère séparément chacun de ses segments. C'est à cette question que s'attache notamment la présente section. Les cartes 3.2 à 3.4 retracent ainsi les indices d'intensité pour les différentes régions métropolitaines respectivement pour les catégories A, B et C, celles mesurées pour les DOM étant proposées dans le tableau 3.1⁸.

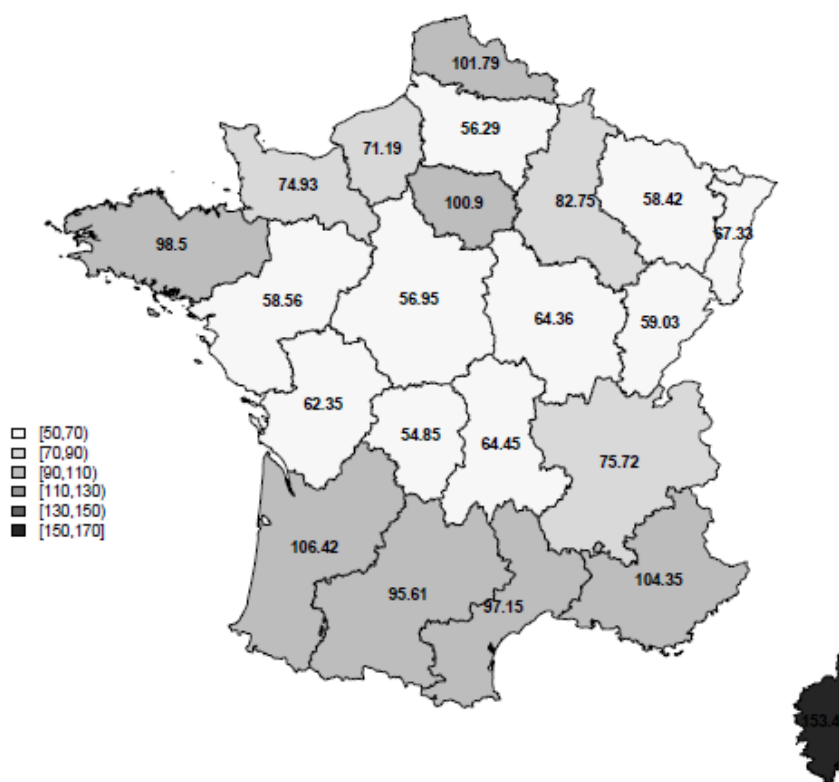
Un aperçu global de chacune des trois cartes montre une catégorie A dans laquelle les inégalités de genre sont élevées dans la totalité des régions métropolitaines puisqu'une seule d'entre elles, le Limousin, fait apparaître un score légèrement inférieur à 100, alors que c'est aussi le cas de trois

8. Voir, en Annexe, le tableau 3.7 pour une présentation des décompositions complètes.

des quatre DOM : la Guadeloupe, la Martinique et la Réunion. A l’opposé, la catégorie B affiche de faibles niveaux d’inégalité de genre, puisque seules quatre régions dépassent la valeur retenue comme étalon, dont une seule notablement, la Corse avec un indice d’intensité égal à 154. La catégorie C se caractérise par une variabilité considérable des intensités d’inégalité avec un rapport de 1 à 3 de l’éventail des scores qui vont de 46,9 à la Réunion à 154,64 en PACA.

Carte 3.3 – Intensités régionales des inégalités liées au genre dans la catégorie B

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

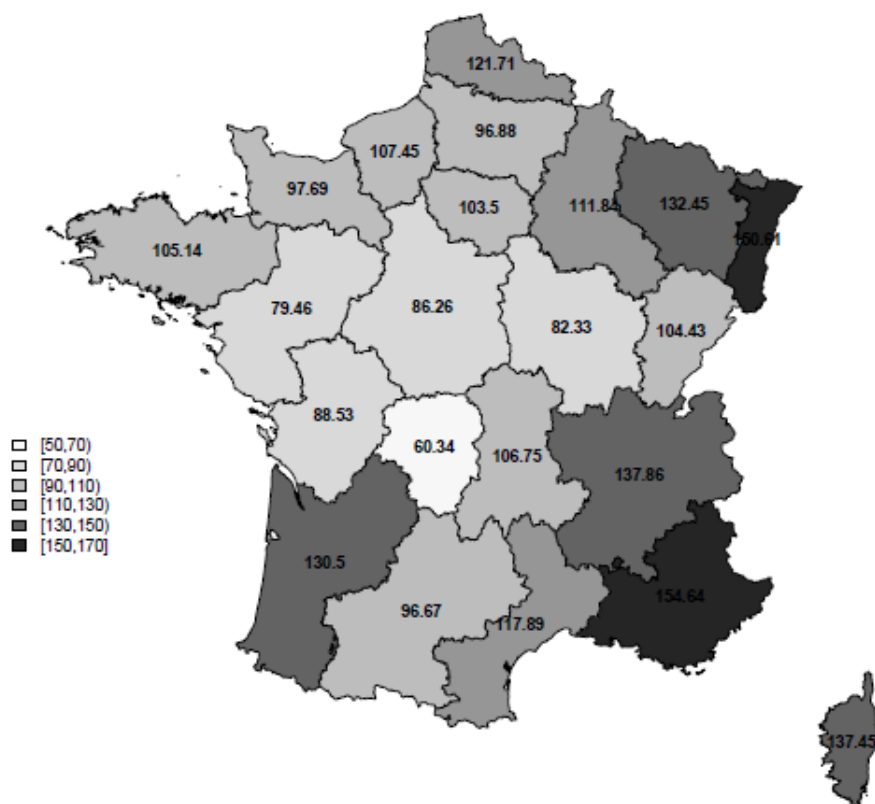
Source : SIASP 2010, Insee, Traitement des auteurs.

Une analyse plus détaillée permet de dégager un certain nombre de tendances intéressantes. La carte 3.2 fait ressortir, pour la catégorie A de la fonction publique, une intensité particulièrement élevée des inégalités expliquées par le genre dans les régions du sud de la France au premier rang desquelles la Corse et Midi-Pyrénées, suivies d’Aquitaine, de PACA et de Languedoc-Roussillon. Il faut cependant noter que Champagne-Ardenne et Bourgogne atteignent des scores presque équivalents.

Mais ce qui caractérise les cinq régions du sud précitées, c’est qu’elles comptent aussi parmi les plus concernées par les décalages de rémunération hommes/femmes dans la catégorie B. Dans

Carte 3.4 – Intensités régionales des inégalités liées au genre dans la catégorie C

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

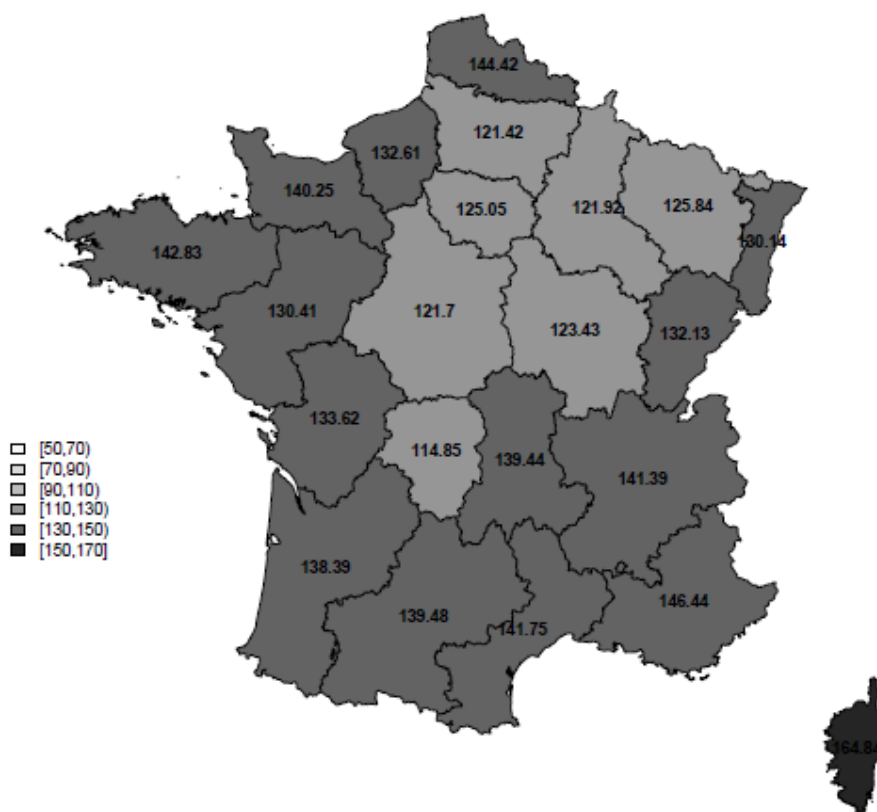
cette dernière catégorie, le Nord-Pas de Calais, l'Île-de-France et la Bretagne ont des scores proches de ceux des régions du sud. La catégorie C, représentée sur la carte 3.4, fait à nouveau apparaître quatre des cinq régions du sud, toutes sauf Midi-Pyrénées, comme particulièrement affectées par les inégalités liées au genre. C'est aussi le cas des régions Alsace, Lorraine et Rhône-Alpes pour cette catégorie. Dès lors, ce qui distingue les régions du sud de la France des autres régions qui ont été évoquées, c'est la régularité d'un niveau élevé de l'indicateur d'intensité des inégalités liées au genre quelle que soit la catégorie concernée.

A l'inverse, le Limousin et les DOM, en particulier la Réunion, affichent les plus faibles intensités d'inégalité liée au genre dans toutes les catégories. Quelques autres régions du centre de la France, comme les Pays de Loire, le Poitou-Charentes, le Centre et la Bourgogne, ont également de faibles scores d'inégalité, si la catégorie A est omise.

Les cartes 3.5 à 3.7 et le tableau 3.1 illustrent l'évaluation par versant des intensités régionales

Carte 3.5 – Intensités régionales des inégalités liées au genre dans la FPE

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

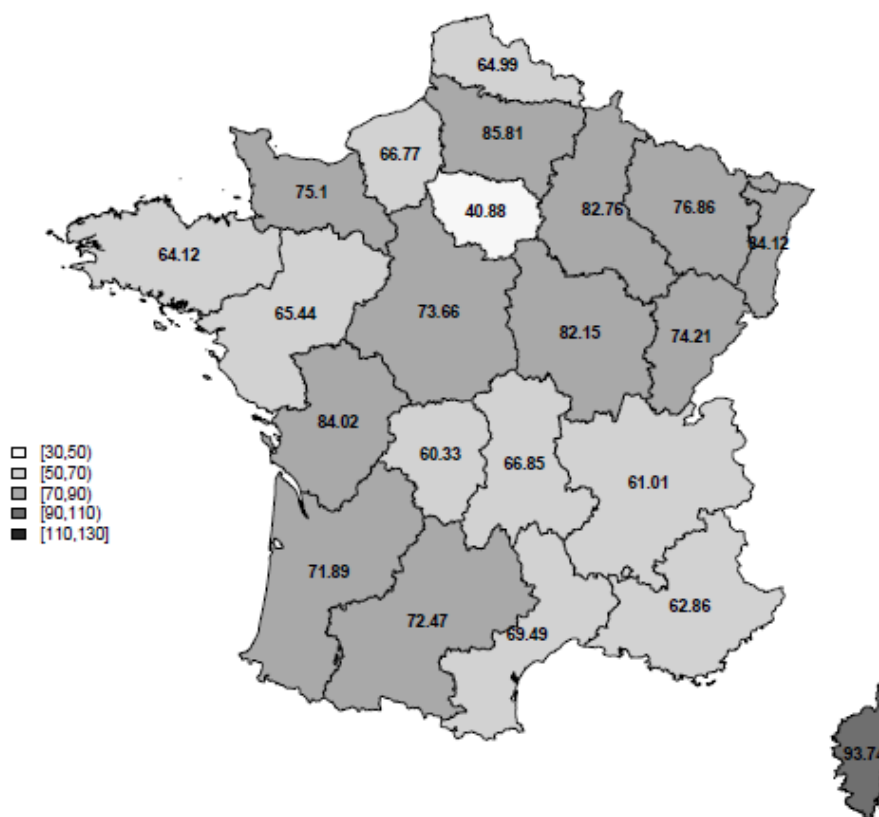
d'inégalité attribuable au genre⁹. La FPE apparaît particulièrement affectée avec un score supérieur à 100 quelle que soit la région considérée, et donc une intensité d'inégalité bien supérieure à la fonction publique nationale. La carte 3.5 souligne en outre la forte homogénéité des inégalités de genre sur ce versant, puisque le score de la région la plus touchée, en l'occurrence la Corse (avec 164,8), n'est égal qu'à une fois et demi celui de la région métropolitaine la moins concernée, le Limousin (avec 114,85). Il faut cependant noter que trois des DOM, la Guadeloupe, la Martinique et la Guyane, ont des scores un peu inférieurs à celui du Limousin.

Dans le cas de la FPH (carte 3.6), l'homogénéité est à rechercher dans la faiblesse des inégalités de genre, puisque le score d'aucune région, même la plus affectée qui est à nouveau la Corse, n'atteint un niveau égal à 100. En revanche, la très grande hétérogénéité des scores précédemment observés dans la catégorie C se reflète naturellement dans la situation de la FPT dont les trois-quarts

9. Voir, en Annexe, le tableau 3.8 pour une présentation des décompositions complètes

Carte 3.6 – Intensités régionales des inégalités liées au genre dans la FPH

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

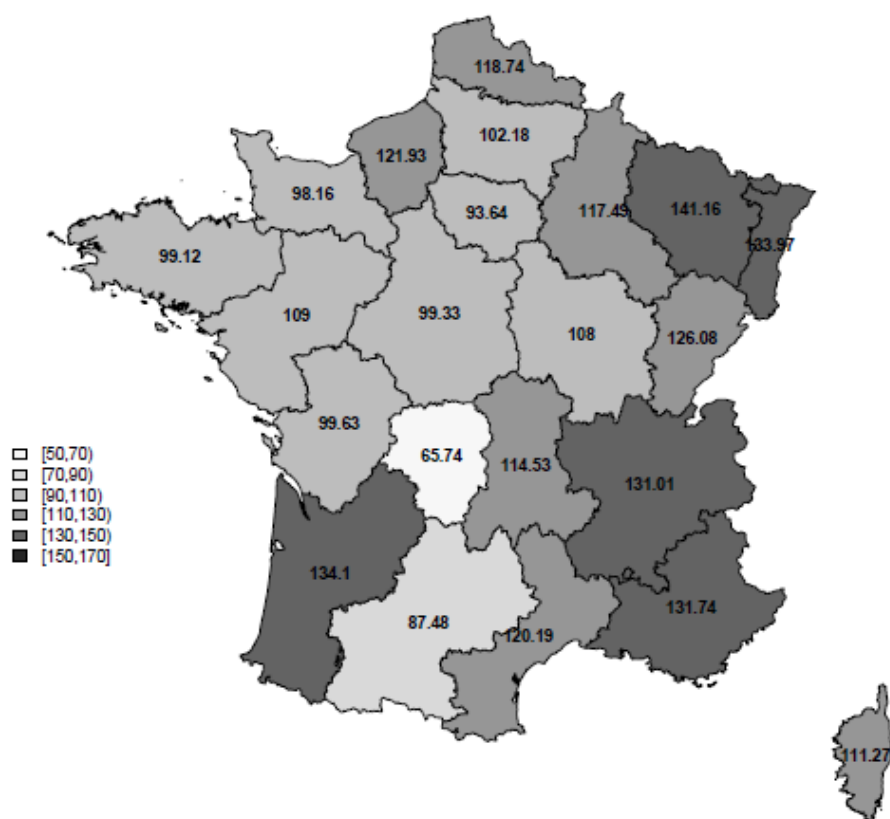
des salariés sont classés dans cette catégorie (carte 3.7).

Peut-on identifier des régularités régionales observables pour tous les versants ? Seules quatre régions appartiennent dans les trois cas à l'ensemble des régions ayant des intensités d'inégalités liées au genre relativement restreintes, il s'agit de deux régions d'outremer, à savoir la Martinique et la Réunion, et de deux régions métropolitaines, à savoir le Limousin et l'Île-de-France. Il est plus difficile de trouver l'équivalent du côté des scores d'inégalité élevés, car aucun groupe clairement identifié de régions ne ressort comme fortement inégalitaire quel que soit le versant considéré. Quelques régions comptent parmi les plus affectées dans deux versants simultanément, mais pas dans les trois. On peut ainsi citer la Corse pour la FPE et la FPH, PACA et Rhône-Alpes pour la FPE et la FPT ou l'Alsace pour la FPH et la FPT.

Afin de synthétiser les résultats précédemment présentés, les distributions des indicateurs d'intensité obtenus dans les différentes régions pour chaque catégorie d'une part, pour chaque versant d'autre part, doivent être agrégées. A cet effet, les positions ordonnées de la plus faibles à la plus

Carte 3.7 – Intensités régionales des inégalités liées au genre dans la FPT

Fonction publique nationale, base 100



Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

élevés des intensités de chacune des régions sont sommées respectivement pour les différentes catégories et pour les différents versants. Cela permet d'obtenir les deux classements des régions présentés au tableau 3.4 qui doivent permettre d'abstraire des biais dus aux différences de structure de la fonction publique.

Il s'agit, par construction, de deux classements différents, mais qui présentent néanmoins d'importantes similitudes. Ainsi trois des quatre régions classées comme les moins affectées par les inégalités de genre au regard de ces indicateurs sont identiques, il s'agit de la Réunion, du Limousin et de la Martinique. Il est intéressant de remarquer que les trois régions précitées sont aussi celles qui ont la plus faible intensité d'inégalité au niveau global de la fonction publique (voir carte 3.1 et tableau 3.1), ce qui tendrait à montrer que les différences dans la structure par emploi des régions ne déforment pas réellement la mesure des inégalités de genre.

A l'autre extrémité de chacune des listes, quatre des six dernières régions apparaissent aussi

Tableau 3.4 – Classements synthétiques des régions selon leurs intensités d'inégalité liée au genre par catégorie et par versants

Classement par catégories		Classement par versants	
Région	Rang (détails par catégorie)	Région	Rang (détails par versant)
Limousin	1e (A(3e); B(1e); C(2e))	Limousin	1e (FPE(4e); FPH(5e); FPT(2e))
La Réunion	2e (A(1e); B(8e); C(1e))	Martinique	1e (FPE(3e); FPH(2e); FPT(6e))
Martinique	3e (A(2e); B(15e); C(4e))	La Réunion	3e (FPE(5e); FPH(4e); FPT(3e))
Picardie	3e (A(9e); B(2e); C(10e))	Île-de-France	4e (FPE(10e); FPH(1e); FPT(5e))
Pays de la Loire	5e (A(13e); B(5e); C(5e))	Guadeloupe	5e (FPE(1e); FPH(3e); FPT(14e))
Centre	6e (A(17e); B(3e); C(7e))	Guyane	6e (FPE(2e); FPH(26e); FPT(1e))
Poitou-Charentes	6e (A(12e); B(7e); C(8e))	Centre	7e (FPE(7e); FPH(16e); FPT(9e))
Guyane	8e (A(8e); B(17e); C(3e))	Pays de la Loire	8e (FPE(13e); FPH(10e); FPT(13e))
Lorraine	9e (A(5e); B(4e); C(22e))	Midi-Pyrénées	9e (FPE(19e); FPH(15e); FPT(4e))
Guadeloupe	10e (A(4e); B(18e); C(12e))	Bretagne	10e (FPE(23e); FPH(8e); FPT(8e))
Franche-Comté	11e (A(15e); B(6e); C(14e))	Bourgogne	11e (FPE(9e); FPH(20e); FPT(12e))
Bourgogne	12e (A(21e); B(9e); C(6e))	Picardie	11e (FPE(6e); FPH(24e); FPT(11e))
Haute-Normandie	12e (A(7e); B(12e); C(17e))	Basse-Normandie	13e (FPE(20e); FPH(18e); FPT(7e))
Basse-Normandie	14e (A(19e); B(13e); C(11e))	Auvergne	14e (FPE(18e); FPH(12e); FPT(16e))
Auvergne	15e (A(18e); B(10e); C(16e))	Champagne-Ardenne	14e (FPE(8e); FPH(21e); FPT(17e))
Île-de-France	16e (A(10e); B(22e); C(13e))	Haute-Normandie	14e (FPE(15e); FPH(11e); FPT(20e))
Nord-Pas-de-Calais	17e (A(6e); B(23e); C(20e))	Poitou-Charentes	17e (FPE(16e); FPH(22e); FPT(10e))
Rhône-Alpes	17e (A(11e); B(14e); C(24e))	Rhône-Alpes	18e (FPE(21e); FPH(6e); FPT(22e))
Alsace	19e (A(14e); B(11e); C(25e))	Nord-Pas-de-Calais	19e (FPE(24e); FPH(9e); FPT(18e))
Bretagne	20e (A(16e); B(21e); C(15e))	Franche-Comté	20e (FPE(14e); FPH(17e); FPT(21e))
Midi-Pyrénées	21e (A(25e); B(19e); C(9e))	Languedoc-Roussillon	21e (FPE(22e); FPH(13e); FPT(19e))
Champagne-Ardenne	22e (A(20e); B(16e); C(18e))	Provence-Alpes-Côte d'Azur	22e (FPE(25e); FPH(7e); FPT(23e))
Languedoc-Roussillon	23e (A(22e); B(20e); C(19e))	Aquitaine	23e (FPE(17e); FPH(14e); FPT(25e))
Aquitaine	24e (A(24e); B(25e); C(21e))	Lorraine	23e (FPE(11e); FPH(19e); FPT(26e))
Provence-Alpes-Côte d'Azur	25e (A(23e); B(24e); C(26e))	Alsace	25e (FPE(12e); FPH(23e); FPT(24e))
Corse	26e (A(26e); B(26e); C(23e))	Corse	26e (FPE(26e); FPH(25e); FPT(15e))

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

dans les deux classements, à savoir, la Corse, PACA, l'Aquitaine et Languedoc-Roussillon. Deux d'entre elles, PACA et Aquitaine, ressortaient également lors de l'analyse globale de la fonction publique comme les plus affectées par les inégalités liées au genre. Le Languedoc-Roussillon bien que classé un peu plus loin faisait également partie des régions les plus concernées. En revanche, la quatrième, la Corse, semble nettement la plus inégalitaire lorsque l'on distingue les catégories ou les versants, alors qu'au niveau global de la fonction publique, elle n'affichait qu'une intensité d'inégalité très moyenne (voir carte 3.1). Ces résultats ne sont cependant pas contradictoires. En effet, si la fonction publique corse est nettement moins féminisée qu'elle ne l'est en moyenne dans les autres régions françaises, le décalage n'est pas le même pour toutes les catégories, puisque l'écart à la moyenne est inférieur à 5 points pour la catégorie A alors qu'il atteint presque 14 points pour la catégorie C. Dans ce cas, les fortes intensités d'inégalité liées au genre qui existent dans chaque catégorie, sont à l'évidence compensées par un biais dans la nature des postes occupés par les femmes qui appartiennent relativement plus fréquemment qu'ailleurs à la catégorie la mieux rémunérée de la fonction publique. Cette observation peut probablement se généraliser à l'ensemble des secteurs. Une publication de l'Insee de 2016 (Insee (2016a)) met notamment en évidence que, quel que soit le secteur, les femmes qui participent au marché du travail dans cette région, sont

relativement plus diplômées que sur l'ensemble du territoire national et cela expliquerait donc aussi bien nos résultats que la faiblesse de l'écart de rémunération entre les femmes et les hommes observé en moyenne dans le secteur privé (voir Annexe, tableau 3.6). Tout explicable qu'il soit, le résultat concernant la Corse oblige clairement à reconsidérer l'idée d'une absence généralisée d'impact de la structure de l'emploi public sur la mesure régionale des inégalités liées au genre.

Dans la même veine, la mesure des intensités par versant fait apparaître des régions dans lesquelles les inégalités liées au genre observées dans chacune des catégories semblent se compenser partiellement, ainsi l'Île-de-France ou Midi-Pyrénées, et des régions, telles que la Lorraine, où elles sont au contraire exacerbées. Ces résultats qui doivent inciter à ne pas sur-interpréter un chiffre en particulier, permettent néanmoins de confirmer quelques tendances générales quant à la géographie des inégalités attribuables au genre dans la fonction publique. Les DOM sont moins concernés par ce problème que la majorité des régions métropolitaines ; en métropole, les régions centrales, et en particulier le Limousin, le Centre et la Picardie, sont moins concernées que le reste du territoire. A l'opposé, en croisant les résultats sur la population globale et les différentes sous-populations étudiées, les écarts de rémunération les plus importants entre les hommes et les femmes salariés de la fonction publique s'observent dans plusieurs régions du sud de la France.

3.5 Discussion et conclusion

Le premier point qui ressort de notre étude, est que des décalages de rémunération entre les femmes et les hommes agents de la fonction publique s'observent en 2010 dans l'ensemble des 26 régions françaises en dépit de la volonté des gouvernements successifs de lutter contre les inégalités liées au genre dans le domaine professionnel. Partout, les femmes sont, en équivalent temps plein, moins bien payées que leurs homologues masculins, et ce résultat reste vrai à chaque niveau hiérarchique et dans chaque composante de la fonction publique. Quelques précautions s'imposent cependant, car si la base SIASP est riche de son exhaustivité, elle manque cruellement d'autres informations sur les salariés qui pourraient fonder objectivement ces différences. Aucun élément ne permet en effet de connaître la véritable ancienneté des agents dans la fonction publique ; de même les variables, telles que les éventuelles périodes de temps partiel, le niveau de diplôme et la nature du concours d'entrée dans la fonction publique sont absentes. Faute de ces données, il est difficile d'établir définitivement si et dans quelle mesure les femmes sont réellement désavantagées par rapport aux hommes dans leur déroulement de carrière.

Une étude complémentaire mériterait néanmoins d'être réalisée ; il s'agirait de dresser une typologie des fonctions liées aux postes occupés respectivement par les hommes et les femmes dans le secteur public. Différents travaux (voir par exemple Meurs et Ponthieux (2006)) ont en effet montré que les femmes, volontairement ou non, s'engagent plus fréquemment que les hommes dans des

métiers qui ouvrent des perspectives moins favorables en termes de rémunération et de carrière. Cette ségrégation professionnelle, qu'elle résulte ou non de leur choix, empêcherait que les rémunérations des hommes et des femmes puissent réellement converger à caractéristiques identiques en termes de niveau de diplôme et d'expérience. Une typologie comparable à celle de Meurs et Ponthieux (2006) pourrait être reconstituée à partir des codes des postes qui sont indiqués dans les données du SIASP. En dépit de l'intérêt fondamental que pourrait avoir une telle analyse pour la compréhension de l'origine des différences de rémunérations entre hommes et femmes dans la fonction publique au niveau national, il n'est pas certain que l'explication des disparités interrégionales en la matière soit à rechercher dans cette seule direction.

Les niveaux pris par l'indicateur d'intensité sont très différents d'une région à l'autre, coïncident très largement avec les écarts de salaire moyen femmes-hommes observées dans le secteur privé. Ce constat renvoie à l'idée d'une spécificité régionale du fonctionnement du marché du travail sur le plan des choix faits ou des opportunités offertes respectivement aux hommes et aux femmes quel que soit le secteur considéré. La présente analyse n'est évidemment pas suffisante pour poser une telle conclusion de façon définitive. Faute de disposer des données du SIASP sur plusieurs années, il est impossible de vérifier que la corrélation reste forte au cours du temps.

A ce stade, une seule piste d'explication aux différences régionales en termes d'inégalités attribuables au genre semble se dessiner, le nombre d'enfants par famille sans doute du fait des comportements qu'il induit chez les femmes sur le marché du travail, semble aller dans le sens d'une réduction de ces inégalités. Il ne s'agit à l'évidence que d'un élément d'explication parmi d'autres qui n'ont pas pu être identifiés. Plus généralement, il semble que quel que soit le secteur étudié, les travaux concernant la situation relative des hommes et des femmes sur le marché du travail, voire les politiques qui envisagent de remédier à ces différences gagneraient à prendre en considération l'existence de spécificités régionales ce qui n'est que très rarement le cas actuellement.

3.6 Annexes

Tableau 3.5 – Décomposition de l'inégalité de rémunération mesurée par l'indice de Gini selon les attributs âge, genre et reste

	Fonction publique		
Régions :	Age	Genre	Reste
Île-de-France	0,0419	0,0207	0,1672
Champagne-Ardenne	0,0327	0,0202	0,1596
Picardie	0,033	0,0167	0,1551
Haute-Normandie	0,0354	0,0178	0,1526
Centre	0,0338	0,0155	0,1538
Basse-Normandie	0,0357	0,0172	0,1535
Bourgogne	0,0347	0,0178	0,1569
Nord-Pas-de-Calais	0,0356	0,0201	0,1453
Lorraine	0,037	0,0198	0,1492
Alsace	0,0404	0,0214	0,1471
Franche-Comté	0,0353	0,0188	0,1521
Pays de la Loire	0,0427	0,0179	0,1444
Bretagne	0,0399	0,0197	0,1477
Poitou-Charentes	0,0362	0,018	0,1584
Aquitaine	0,0366	0,0224	0,1601
Midi-Pyrénées	0,0383	0,02	0,1623
Limousin	0,0415	0,0131	0,156
Rhône-Alpes	0,0395	0,0202	0,1481
Auvergne	0,0341	0,018	0,1574
Languedoc-Roussillon	0,0368	0,0204	0,1568
Provence-Alpes-Côte d'Azur	0,038	0,0223	0,1546
Corse	0,0344	0,0182	0,1597
Guadeloupe	0,0345	0,0163	0,1823
Martinique	0,0374	0,0136	0,1927
Guyane	0,0354	0,0169	0,1938
La Réunion	0,0333	0,0126	0,2455
France entière	0,038	0,0197	0,1632

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

Tableau 3.6 – Répartition de la population française et Écart de salaire net moyen EQTP entre les femmes et les hommes dans les secteurs privé et semi-public, 2010, par région (en %)

Régions	Répartition population nationale (en %)	Ecart de salaire moyen entre femmes et hommes (en %) $\frac{F-H}{H}$
Île-de-France	18,24	-20,83
Champagne-Ardenne	2,07	-17,95
Picardie	2,96	-17,95
Haute-Normandie	2,84	-20,42
Centre	3,94	-18,78
Basse-Normandie	2,28	-17,84
Bourgogne	2,54	-19,33
Nord-Pas-de-Calais	6,25	-19,09
Lorraine	3,64	-20,97
Alsace	2,86	-21,94
Franche-Comté	1,81	-20,43
Pays de la Loire	5,53	-19,56
Bretagne	4,95	-19,52
Poitou-Charentes	2,74	-15,18
Aquitaine	5	-18,45
Midi-Pyrénées	4,46	-20,51
Limousin	1,15	-14,88
Rhône-Alpes	9,64	-21,76
Auvergne	2,09	-17,19
Languedoc-Roussillon	4,08	-18,21
Provence-Alpes-Côte d'Azur	7,58	-21,1
Corse	0,48	-14,79
Guadeloupe	0,62	-11,63
Martinique	0,61	-11,5
Guyane	0,35	-12,02
La Réunion	1,27	-7,28
France entière	100	-19,6

Champ : France entière

Source : Insee

Tableau 3.7 – Décomposition de l'inégalité de rémunération mesuré par l'indice de Gini selon les attributs âge, genre et reste, France et régions, catégories de la fonction publique

Régions :	Catégorie A			Catégorie B			Catégorie C		
	Age	Genre	Reste	Age	Genre	Reste	Age	Genre	Reste
Île-de-France	0,0703	0,025	0,1559	0,0429	0,0148	0,1063	0,0207	0,0134	0,111
Champagne-Ardenne	0,0586	0,0241	0,1385	0,049	0,0107	0,0853	0,02	0,014	0,1064
Picardie	0,066	0,02	0,1194	0,0531	0,0076	0,0908	0,0162	0,0123	0,1139
Haute-Normandie	0,0642	0,0192	0,1203	0,0466	0,0094	0,091	0,0197	0,014	0,1119
Centre	0,064	0,022	0,1206	0,0462	0,0072	0,0885	0,0188	0,0103	0,1045
Basse-Normandie	0,0645	0,0221	0,1193	0,0573	0,0104	0,0882	0,0193	0,0124	0,1103
Bourgogne	0,0653	0,0231	0,1226	0,0525	0,0087	0,0906	0,0171	0,0098	0,1067
Nord-Pas-de-Calais	0,0662	0,0189	0,1158	0,0483	0,0131	0,0824	0,0214	0,0157	0,1071
Lorraine	0,0699	0,0192	0,1179	0,0571	0,0081	0,0893	0,0174	0,0171	0,1101
Alsace	0,0679	0,022	0,1253	0,0442	0,0086	0,0897	0,02	0,0209	0,1146
Franche-Comté	0,0623	0,021	0,1203	0,0462	0,0074	0,0859	0,0208	0,0129	0,1041
Pays de la Loire	0,0656	0,0213	0,1221	0,0673	0,0085	0,086	0,0223	0,0097	0,1051
Bretagne	0,0592	0,0226	0,1308	0,055	0,0134	0,0838	0,0241	0,0129	0,1006
Poitou-Charentes	0,0663	0,0218	0,1262	0,056	0,0084	0,0861	0,0187	0,0112	0,1115
Aquitaine	0,0615	0,0254	0,1385	0,055	0,0152	0,0901	0,0182	0,0165	0,1065
Midi-Pyrénées	0,0622	0,0261	0,1334	0,0605	0,0143	0,0925	0,0193	0,0125	0,1133
Limousin	0,0689	0,0191	0,1286	0,0638	0,0077	0,0852	0,0231	0,0073	0,1047
Rhône-Alpes	0,0666	0,0213	0,1231	0,0494	0,0103	0,0925	0,0202	0,0178	0,1063
Auvergne	0,0616	0,0225	0,1272	0,0524	0,0084	0,0846	0,0194	0,0135	0,1088
Languedoc-Roussillon	0,0632	0,0231	0,124	0,0505	0,0136	0,0925	0,0184	0,0147	0,1066
Provence-Alpes-Côte d'Azur	0,0628	0,025	0,1382	0,0478	0,0147	0,0949	0,0196	0,0199	0,1046
Corse	0,0542	0,0316	0,1432	0,0357	0,0194	0,0867	0,0204	0,0171	0,1017
Guadeloupe	0,0629	0,0156	0,0972	0,0841	0,0166	0,0941	0,0206	0,0144	0,1211
Martinique	0,0637	0,0159	0,1033	0,0909	0,0149	0,0959	0,0278	0,0123	0,1408
Guyane	0,0709	0,0194	0,1132	0,063	0,0152	0,1205	0,0297	0,0116	0,1338
La Réunion	0,0663	0,0141	0,0882	0,0736	0,0102	0,0955	0,0183	0,0102	0,2152
France entière	0,0447	0,0227	0,1607	0,047	0,0171	0,0978	0,0273	0,0174	0,1043

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

Tableau 3.8 – Décomposition de l'inégalité de rémunération mesuré par l'indice de Gini selon les attributs âge, genre et reste, France et régions, versants de la fonction publique

Régions	PFE			FPH			FPT		
	Age	Genre	Reste	Age	Genre	Reste	Age	Genre	Reste
Île-de-France	0,0518	0,026	0,1548	0,0541	0,0076	0,1466	0,0284	0,0165	0,1523
Champagne-Ardenne	0,0405	0,0207	0,129	0,0488	0,0159	0,1508	0,0181	0,0208	0,1596
Picardie	0,0456	0,0186	0,1068	0,0497	0,0167	0,1515	0,0157	0,0181	0,1648
Haute-Normandie	0,0485	0,021	0,1074	0,0517	0,0127	0,1484	0,02	0,0213	0,1542
Centre	0,044	0,0191	0,1121	0,0505	0,0139	0,1461	0,0175	0,0162	0,1489
Basse-Normandie	0,0462	0,0219	0,1064	0,0533	0,0143	0,146	0,0191	0,0162	0,1498
Bourgogne	0,0468	0,0199	0,114	0,0473	0,0158	0,1523	0,0181	0,0182	0,1522
Nord-Pas-de-Calais	0,0515	0,023	0,1037	0,0527	0,0123	0,1467	0,0184	0,0205	0,1547
Lorraine	0,049	0,0207	0,1143	0,0508	0,0143	0,1438	0,0173	0,025	0,1561
Alsace	0,0452	0,0221	0,1227	0,046	0,0158	0,1486	0,0256	0,0259	0,165
Franche-Comté	0,048	0,0211	0,1098	0,0469	0,0141	0,152	0,0197	0,022	0,1533
Pays de la Loire	0,0532	0,0216	0,1104	0,0553	0,0121	0,1401	0,0235	0,0188	0,151
Bretagne	0,042	0,0241	0,1229	0,059	0,012	0,1388	0,0232	0,016	0,1411
Poitou-Charentes	0,0457	0,0215	0,1129	0,0539	0,0166	0,1505	0,0188	0,0169	0,154
Aquitaine	0,0444	0,0249	0,1322	0,0511	0,0136	0,1469	0,0197	0,0231	0,1499
Midi-Pyrénées	0,0509	0,0258	0,1305	0,0492	0,0136	0,1475	0,0221	0,0152	0,1567
Limousin	0,0544	0,0199	0,1192	0,0506	0,011	0,1431	0,0236	0,0107	0,1472
Rhône-Alpes	0,0496	0,0238	0,115	0,0534	0,0117	0,1498	0,0229	0,0234	0,1533
Auvergne	0,0427	0,0224	0,1149	0,0492	0,0127	0,1502	0,0171	0,0193	0,1524
Languedoc-Roussillon	0,0504	0,0246	0,1194	0,0522	0,0134	0,1496	0,021	0,02	0,1455
Provence-Alpes-Côte d'Azur	0,0478	0,0268	0,1299	0,0487	0,0124	0,1604	0,0262	0,023	0,1461
Corse	0,036	0,0297	0,1359	0,0601	0,0187	0,1445	0,023	0,0189	0,1478
Guadeloupe	0,0578	0,0178	0,1129	0,0633	0,0096	0,1461	0,0138	0,0179	0,1486
Martinique	0,0547	0,0184	0,1144	0,0611	0,0085	0,1557	0,0162	0,0173	0,1722
Guyane	0,0533	0,0196	0,1267	0,0592	0,022	0,1627	0,0289	0,0093	0,1471
La Réunion	0,0628	0,0201	0,1097	0,0846	0,0102	0,134	0,0123	0,0171	0,2569
France entière	0,0411	0,0219	0,1454	0,0485	0,0151	0,1542	0,0311	0,0234	0,1572

Champ : France entière, salariés de la fonction publique.

Source : SIASP 2010, Insee, Traitement des auteurs.

Chapitre 4

Contribution marginale pure, interactions et l'importance d'un attribut

Sommaire

4.1	Introduction	104
4.2	Jeu inégalitaire	105
4.3	L'importance d'un attribut	107
4.3.1	Contribution marginale	107
4.3.2	Les interactions entre les attributs	108
4.4	Discussion et conclusion	112
4.5	Annexe	113

4.1 Introduction

L'impact d'une politique publique qui vise à réduire les inégalités peut être défini comme étant la différence entre le niveau d'inégalité avant la mise en place de la politique et après la production de ces effets. Par exemple, si le décideur public décide de supprimer les différences de salaires entre les hommes et les femmes, toutes choses égales par ailleurs, alors l'impact de la politique publique sera la différence entre d'une part l'inégalité mesurée lorsque toutes les caractéristiques personnelles contribuent à la formation du salaire, et d'autre part l'inégalité mesurée avec ces mêmes caractéristiques excepté le genre.

L'impact d'un attribut sur les inégalités correspond à sa contribution marginale pure (CMP par la suite), c'est-à-dire la différence entre l'inégalité totale et l'inégalité mesurée lorsque l'attribut n'est plus créateur de différences entre les individus. L'approche marginaliste qui consiste à agréger la somme des CMP des attributs peut ne pas être efficace, c'est-à-dire que la somme des CMP peut ne pas être égale à l'inégalité totale. Dans le contexte des jeux d'inégalité, Chantreuil et Trannoy (2011, 2013) et Shorrocks (2013) ont montré que seule la valeur de Shapley permet de réconcilier l'approche marginaliste avec une désagrégation efficace de l'inégalité entre ces différents attributs. Dans un jeu inégalitaire donné, la contribution d'un attribut à l'inégalité totale (*i.e.* son importance) est donc mesurée par la valeur de Shapley.

L'inefficacité de l'approche marginaliste pour désagréger l'inégalité entre une somme de poids associé à chaque attribut réside dans son ignorance de leurs interactions. La notion d'interaction dans les jeux non-additifs¹ a été introduite par Owen (1972)². Il est effectivement possible que les attributs ne soient pas indépendants les uns des autres dans la production d'inégalités. Par exemple, si le revenu est décomposé entre revenu du capital et revenu du travail, il se peut que la différence entre l'inégalité totale d'une part, moins l'inégalité mesurée sur les revenus du capital et l'inégalité mesurée sur les revenus du travail d'autre part, soit différente de zéro. Si la différence est positive, alors les deux catégories de revenus interagissent négativement dans le sens où leur interaction augmente l'inégalité. Au contraire, si la différence est négative alors ils interagissent positivement car leur interaction fait diminuer les inégalités (*i.e.* l'un compense en partie les inégalités créées par l'autre). Si la différence est nulle, alors les deux classes de revenus n'interagissent pas, le jeu est additif.

L'objectif de ce chapitre est de décomposer les indices d'inégalités grâce à la valeur de Shapley, de telle sorte que l'importance et l'impact de chaque attribut puissent être estimés. A priori,

1. On parle de jeux non-additifs quand la valeur de la fonction caractéristique associée à chaque joueur pris séparément ne somme pas à la valeur du jeu total. Les jeux reposant sur des indices d'inégalité ne sont pas additifs car la somme de l'inégalité mesurée pour source de revenu n'est pas égale à l'inégalité totale (*i.e.* quand toutes les sources sont prises en compte).

2. D'autres auteurs ont étudié les interactions dans ce cadre. Parmi eux, Murofushi et Soneda (1993), Grabisch (1996, 1997a,b), Grabisch *et al.* (2000), Kojadinovic (2003, 2004, 2005)

l'importance d'un attribut coïnciderait avec son impact, cependant ce constat ne se vérifie que dans certains cas spécifiques. C'est pourquoi ce chapitre démontre que l'importance d'un attribut aux inégalités selon la valeur de Shapley peut être décomposée entre d'une part sa contribution marginale pure (*i.e.* son impact), et d'autre part ses interactions par paires avec les autres attributs. Ainsi, un indice d'inégalité peut être réécrit comme la somme des contributions marginales pures des attributs moins la somme des interactions par paires d'attributs.

4.2 Jeu inégalitaire

Les notations utilisées ici sont similaires à celles de la section 2.5.1. L'indice d'inégalité est une fonction $I : \mathbb{R}^n \rightarrow \mathbb{R}_+$, tel que $I(\cdot)$ est un indice régulier relatif. La distribution de revenus est définie selon les sources de revenu prises en compte, si aucune n'est prise en considération alors $Y(\emptyset) = 0^n$, sinon pour tout $S \in \mathcal{M}$:

$$Y(S, \lambda) = (\sum_{i \in S} y_1^i + \lambda, \dots, \sum_{i \in S} y_n^i + \lambda) \quad (4.1)$$

avec $\lambda \in \mathbb{R}^+$.

Pour rappel, une fonction caractéristique $V_I(\cdot)$ est définie pour évaluer l'inégalité d'une distribution de revenus $Y(S, \lambda)$ selon un indice d'inégalités $I(\cdot)$, ainsi $V_I(S) = I(Y(S, \lambda))$. Parmi l'ensemble $\mathcal{V}$ des jeux λ -inégalitaires, deux types de jeux sont considérés, les jeux nuls avec $\lambda = 0$, et les jeux égalitaires avec $\lambda = \mu(X) - \mu(X^S)$. Dans les deux variantes du jeu, la distribution des revenus associée à un ou plusieurs attributs est remplacée par une distribution égalitaire (nulle ou égale à la moyenne), c'est pourquoi les indices absolus, invariant à une translation, ne sont pas traités. La section 2.5 a comparé les deux approches, notamment en mettant en avant qu'un attribut avec une distribution égalitaire (ex. : un revenu universel) aura une contribution nulle dans le jeu égalitaire, et une contribution négative dans un jeu nul (ou positive si l'attribut correspond à un revenu négatif, comme une taxe forfaitaire). Cette différence conduit à ce qu'un attribut qui tend à compenser les inégalités créées par les autres attributs ait une contribution négative dans le jeu égalitaire, et faiblement positive dans le jeu nul. De plus, d'après Sastre et Trannoy (2018), l'importance d'un attribut serait plus sensible au niveau de désagrégation des différentes sources du revenu dans le cas du jeu nul que dans le jeu égalitaire. Enfin, la sensibilité des contributions des différents attributs est moindre avec le jeu égalitaire. En pratique, le choix entre jeu nul et jeu égalitaire dépend de la nature du problème traité. Ainsi, si la politique publique évaluée consiste à égaliser les bonus au sein d'une entreprise, alors le jeu égalitaire sera plus adaptée pour estimer l'impact de cette réforme. Pour autant, si la politique mise en œuvre consiste à supprimer une allocation familiale par exemple, alors le jeu nul semble plus adapté.

D'autre part, l'impact de telles politiques publiques égalisatrices pourrait être estimée en calculant le niveau d'inégalité une fois les modifications apportées à la distribution de l'attribut considéré, c'est-à-dire en comparant le niveau d'inégalité *ex ante* et *ex post* la variation. Cependant, cette approche, certes rapide et efficace, ne permet pas de comprendre le résultat. En effet, comme ce chapitre le démontre, le résultat de telles politiques dépend des interactions entre les différents attributs, avec des interactions qui tendent à augmenter les inégalités et d'autres qui les réduisent.

Exemple 1 : Le tableau 4.1 décrit une situation x , avec trois attributs (i, j, h) et trois individus (a, b, c) . L'inégalité est mesurée avec l'indice de Gini. De cette situation x , le jeu λ -inégalitaire est défini, pour $\lambda = \mu(X) - \mu(X^S)$ (Tableau 4.2) et $\lambda = 0$ (Tableau 4.3).

Tableau 4.1 – Distributions des revenus

	Attribut i	Attribut j	Attribut h
Individu a	3	8	9
Individu b	5	10	15
Individu c	7	12	21
Moyenne	5	10	15

Tableau 4.2 – Jeu égalitaire

	i	j	h	$i + j$	$i + h$	$j + h$	$i + j + h$
Individu a	28	28	24	26	22	22	20
Individu b	30	30	30	30	30	30	30
Individu c	32	32	36	34	38	38	40
	$V_I(i)$	$V_I(j)$	$V_I(h)$	$V_I(i, j)$	$V_I(i, h)$	$V_I(j, h)$	$V_I(i, j, h)$
Gini	0,0296	0,0296	0,0889	0,0592	0,1185	0,1185	0,1481

Lecture : Dans le cas du jeu égalitaire, le revenu associé à l'attribut i de l'individu a est égale au montant qu'il perçoit pour cette source (3) plus la moyenne des deux autres sources j et h (respectivement 10 et 15), soit 28.

Tableau 4.3 – Jeu nul

	i	j	h	$i + j$	$i + h$	$j + h$	$i + j + h$
Individu a	3	8	9	11	12	17	20
Individu b	5	10	15	15	20	25	30
Individu c	7	12	21	19	28	33	40
	$V_I(i)$	$V_I(j)$	$V_I(h)$	$V_I(i, j)$	$V_I(i, h)$	$V_I(j, h)$	$V_I(i, j, h)$
Gini	0,1778	0,0889	0,1778	0,1185	0,1778	0,1422	0,1481

Lecture : Dans le cas du jeu nul, le revenu associé à la paire d'attributs $i + j$ de l'individu a est égale à la somme du montant qu'il perçoit pour chacune de ces sources (respectivement 3 et 8), soit 11.

4.3 L'importance d'un attribut

4.3.1 Contribution marginale

Comme l'introduction l'a mis en avant, une première approche pour mesurer l'importance d'un attribut consiste à calculer sa contribution marginale pure.

Définition 1. La contribution marginale pure. Pour un jeu λ -inégalitaire donné (M, V_I) , la contribution marginale pure (CMP) d'un attribut $i \in M$ est définie par :

$$CMP_i(M, V_I) = V_I(M) - V_I(M \setminus \{i\}) \quad (4.2)$$

La définition 1 stipule que la contribution marginale pure d'un attribut i est la différence entre la valeur de la fonction V_I , estimée sur la somme des distributions de tous les attributs inclus dans M , et la valeur de cette fonction estimée sur la somme des distributions de ces mêmes attributs mais dont la distribution de l'attribut concerné a été remplacée par un vecteur nul ($\lambda = 0$) ou un vecteur égalitaire égal à la moyenne ($\lambda = \mu(X) - \mu(X^S)$). Cette CMP peut être positive ou négative. Dans le premier cas, cela signifie que sans les disparités de revenu associées à cet attribut, l'inégalité serait plus faible. Au contraire, dans le second cas, sans l'inégalité de revenu intrinsèque à ce facteur alors l'inégalité totale serait plus forte.

Suite Exemple 1. Le jeu égalitaire précédent (Tableau 4.2) montre que l'approche marginaliste peut être efficace pour expliquer l'inégalité totale. En effet, comme le Tableau 4.4 le montre, la somme des contributions marginales pures égale à l'inégalité totale dans la situation x mesurée par l'indice de Gini.

Tableau 4.4 – Contribution Marginales Pures- Jeu égalitaire

$CMP_i(M, V_I)$	$CMP_j(M, V_I)$	$CMP_h(M, V_I)$	Somme
0,0296	0,0296	0,0889	0,1481 = $V_I(i, j, h)$

Cependant, cette solution est problématique car elle n'est pas toujours efficace. Le Tableau 4.5 illustre ce point avec le jeu nul (défini dans le Tableau 4.3), avec $\sum_{i \in M} CMP_i(M, V_I) \neq V_I(M) = I(x)$. En effet, la somme des contributions marginales pures n'est pas égale à l'inégalité totale.

Tableau 4.5 – Contributions Marginales Pures - Jeu nul

$CMP_i(M, V_I)$	$CMP_j(M, V_I)$	$CMP_h(M, V_I)$	Somme
0,0059	-0,0296	0,0296	0,0059 $\neq V_I(i, j, h)$

Par conséquent, la décomposition d'un indice d'inégalité avec l'approche marginaliste peut ne pas être satisfaisante. Chantreuil et Trannoy (2011, 2013) et Shorrocks (2013) montre que la valeur de Shapley (Shapley (1953)) est la seule fonction qui résout ce problème de non-efficacité avec l'approche marginale. La fonction de Shapley est ici rappelée.

Définition 2. La fonction de Shapley. Pour un jeu λ -inégalitaire (M, V_I) , avec $s = |S|$, la fonction de Shapley (Sh) d'un attribut $i \in M$ est donnée par :

$$Sh_i(M, V_I) = \sum_{\substack{S \subseteq M \\ i \notin S}} \frac{(m-s-1)!s!}{m!} [V_I(S \cup \{i\}) - V_I(S)] \quad (4.3)$$

Suite Exemple 1. Les importances des attributs obtenues avec la fonction de Shapley sont données dans le Tableau 4.6.

Tableau 4.6 – Importances des attributs

	$Sh_i(M, V_I)$	$Sh_j(M, V_I)$	$Sh_h(M, V_I)$	Sum
Jeu égalitaire	0,0296	0,0296	0,0889	0,1481 = $V_I(i, j, h)$
Jeu nul	0,0662	0,0039	0,0780	0,1481 = $V_I(i, j, h)$

Dans les deux jeux, la somme de l'importance des attributs i, j et h , mesurées avec la fonction de Shapley, est égale à l'inégalité totale.

Dans le jeu égalitaire, l'importance mesurée par la CMP (Tableau 4.4) et par la fonction de Shapley sont identiques.

Dans le jeu nul, la CMP de l'attribut j (Tableau 4.5) est négative, ce qui veut dire que supprimer les revenus issus de cet attribut augmenterait les inégalités. Cependant, le signe de l'importance de l'attribut j estimé avec la fonction de Shapley est positif, il indique donc que cet attribut a une contribution positive à l'inégalité.

Se pose alors la question de définir la meilleure approche pour expliquer l'importance de l'attribut j à l'inégalité. Par la suite, il sera démontré que la fonction de Shapley est plus adaptée car elle permet de distinguer l'importance d'un facteur de son impact (i.e sa CMP) sur les inégalités. La différence entre les deux réside dans la prise en compte des interactions entre paires d'attributs.

4.3.2 Les interactions entre les attributs

L'interaction entre deux attributs peut être appréciée par la différence entre le niveau d'inégalité estimé lorsque les deux attributs sont considérés ensemble et la somme de l'inégalité mesurée sur la distribution associée à chacun de ces attributs. Formellement, pour un jeu λ -inégalitaire (M, V_I)

et pour toutes paires d'attributs $i, j \in M$, l'interaction entre i et j dépend de la différence $V_I(\{i, j\}) - V_I(\{i\}) - V_I(\{j\})$. Cependant, les autres attributs inclus dans M peuvent aussi modifier l'interaction entre i et j . C'est pourquoi Murofushi et Soneda (1993) proposent d'estimer l'interaction entre deux attributs en prenant en compte toutes les coalitions possibles T pouvant être présentes lorsque i et j interagissent. Alors, la définition suivante de l'interaction entre deux attributs dans un jeu à M attributs en présence d'une coalition $T \setminus \{i, j\}$ est obtenue.

Définition 3. Interaction par paire. Pour un jeu λ -inégalitaire (M, V_I) et pour toutes paires d'attributs $i, j \in M$, $i \neq j$, l'interaction entre i et j en présence d'une coalition $T \subseteq M \setminus \{i, j\}$, est définie par :

$$Int(i, j, T) = V_I(\{i, j\} \cup T) - V_I(\{i\} \cup T) - V_I(\{j\} \cup T) + V_I(T) \quad (4.4)$$

Suite Exemple 1. L'interaction par paire pour chaque attribut est calculée, comme défini par l'équation (4.4). Les résultats sont donnés dans le Tableau 4.7.

Tableau 4.7 – Interactions

Attribut i	$Int(i, j, \{\emptyset\})$	$Int(i, j, \{h\})$	$Int(i, h, \{\emptyset\})$	$Int(i, h, \{j\})$
Jeu égalitaire	0,0000	0,0000	0,0000	0,0000
Jeu nul	-0,1481	0,0059	-0,1778	-0,0237
Attribut j	$Int(j, i, \{\emptyset\})$	$Int(j, i, \{h\})$	$Int(j, h, \{\emptyset\})$	$Int(j, h, \{i\})$
Jeu égalitaire	0,0000	0,0000	0,0000	0,0000
Jeu nul	-0,1481	0,0059	-0,1244	0,0296
Attribut h	$Int(h, i, \{\emptyset\})$	$Int(h, i, \{j\})$	$Int(h, j, \{\emptyset\})$	$Int(h, j, \{i\})$
Jeu égalitaire	0,0000	0,0000	0,0000	0,0000
Jeu nul	-0,1778	0,0059	-0,1244	-0,0237

Dans le cas particulier du jeu égalitaire, toutes les interactions par paires sont nulles. Ceci explique pourquoi l'approche marginale par la CMP et celle par l'équation de Shapley donne la même importance aux attributs. Cependant, ce résultat particulier peut ne pas se vérifier dans d'autres exemples.

Dans le jeu nul, la paire d'attribut i et j est caractérisée par une interaction négative lorsqu'elle est considérée seule ($T = \emptyset$). Plus formellement, $Int(i, j, \{\emptyset\}) = 0,1185 - (0,1778 + 0,0889) = -0,1481$. Par conséquent, lorsque i et j interagissent uniquement entre eux, le niveau d'inégalité est plus faible que la somme de l'inégalité lorsqu'ils sont considérés seuls. Cependant, ajouter l'attribut h à la paire d'attributs (i, j) modifie l'interaction, celle-ci devient positive : $Int(i, j, h) = 0,1481 - (0,1778 + 0,1422) + 0,1778 = 0,0059$. Autrement dit, dans une situation où la source h est créatrice d'inégalité, l'interaction entre i et j augmente l'inégalité.

A partir de cette définition des interactions, Murofushi et Soneda (1993) proposent une fonction d'interaction par paire d'attributs qui est une somme pondérée de l'ensemble des interactions par

paire possibles :

$$I_{ij} = \sum_{t=0}^{n-2} \frac{(n-t-2)!t!}{(n-1)!} \sum_{T \subseteq M \setminus \{i,j\}} Int(i, j, T) \quad (4.5)$$

Avec $t = |T|$ le nombre d'attributs dans la coalition T .

Ainsi, l'importance d'un attribut i par la fonction de Shapley peut être réécrite comme la somme de la valeur du jeu pour l'attribut i seul plus la moitié des interactions entre i et les autres attributs j (Grabisch (1997b)). L'interaction entre i et j étant répartie équitablement entre les deux.

$$Sh_i = V_I(i) + \frac{1}{2} \sum_{j \in M \setminus i} I_{ij} \quad (4.6)$$

Contrairement à l'équation précédente, et ceci dans un souci d'évaluation de politiques publiques, l'équation de la fonction de Shapley est réécrite non plus pour mettre en exergue la différence entre la contribution de i lorsqu'il est considéré seul ($V_I(i)$) de ses interactions par paires, mais en différenciant sa CMP ($V_I(M) - V_I(M \setminus i)$) de ses interactions.

Définition 4. La fonction de Shapley revisitée. Pour un jeu λ -inégalitaire (M, V_I) , la fonction de Shapley revisitée ($\overline{Sh}_i(M, V_I)$) d'un attribut $i \in M$ est définie par :

$$\overline{Sh}_i(M, V_I) = CMP_i - \sum_{\substack{j \in M \\ i \neq j}} INT(i, j, T) \quad (4.7)$$

Avec, $INT(i, j, T) = \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} (Int(i, j, T))$, $m = |M|$ le nombre d'attributs et $t = |T|$ le nombre d'attributs dans la coalition T .

Pour un jeu λ -inégalitaire (M, V_I) , la définition 4 stipule que l'importance d'un attribut est égale à la somme de sa contribution marginale pure, moins une somme pondérée de l'ensemble de ses interactions par paires. Autrement dit, pour une importance donnée, l'impact d'une politique publique qui supprime les inégalités créées par un attribut i sera d'autant plus fort que ses interactions sont faibles. Par exemple, prenons le cas de l'effet du genre sur les inégalités de salaires. Si l'importance du genre à l'inégalité est uniquement due à ses interactions avec les autres facteurs (ex. : les femmes sont moins présentes dans les secteurs les plus profitables, elles sont plus diplômées...), alors donner le même salaire aux hommes et aux femmes à toutes caractéristiques égales ne réduirait pas les inégalités car il n'y aurait pas d'effet du genre *per se*.

Le Tableau 4.8 présente les 8 cas possibles qui peuvent définir le signe de l'importance d'un attribut.

Quand le signe de la CMP et des interactions sont opposés (cas 1 et 2), alors l'importance

Tableau 4.8 – CMP, Interactions et Importance

	CMP_i	$\sum_{\substack{j \in M \\ i \neq j}} INT(i, j, T)$	$\overline{Sh}_i(M, V_I)$
Cas 1	+	-	+
Cas 2	-	+	-
Cas 3	+	+	+
Cas 4	+	+	-
Cas 5	-	-	-
Cas 6	-	-	+

d'un attribut sur l'inégalité est renforcée, positivement ou négativement. Quand les signes sont identiques, alors le signe de l'importance dépend la valeur de l'un par rapport à l'autre. Dans le cas 3 (4), les interactions sont plus faibles (fortes) que la CMP, par conséquent l'importance de i est positive (négative). Un raisonnement symétrique s'applique pour les cas 5 et 6.

Suite Exemple 1. $\overline{Sh}_i(M, V_I)$ est maintenant calculée pour chaque attribut.

Tableau 4.9 – Importances

Attribut i	CMP_i	$INT(i, j, T)$	$INT(i, h, T)$	$\overline{Sh}_i(M, V_I)$
Jeu égalitaire	0,0296	0,0000	0,0000	0,0296
Jeu nul	0,0059	-0,0227	-0,0376	0,0662
Attribut j	CMP_j	$INT(j, i, T)$	$INT(j, h, T)$	$\overline{Sh}_j(M, V_I)$
Jeu égalitaire	0,0296	0,0000	0,0000	0,0296
Jeu nul	-0,0296	-0,0227	-0,0108	0,0039
Attribut h	CMP_h	$INT(h, i, T)$	$INT(h, j, T)$	$\overline{Sh}_h(M, V_I)$
Jeu égalitaire	0,0889	0,0000	0,0000	0,0889
Jeu nul	0,0296	-0,0375	-0,0109	0,0780

Pour chaque type de jeu inégalitaire, la fonction de Shapley revisitée conduit au même résultat que la fonction de Shapley. Comme cela était le cas précédemment, puisque les interactions sont toutes nulles dans le cas du jeu égalitaire, la CMP est égale à l'importance estimée par la fonction de Shapley revisitée.

En conséquence, le résultat principal est déduit.

Proposition 1. Shapley revisitée. Pour un jeu λ -inégalitaire (M, V_I) , et pour chaque attribut :

$$\overline{Sh}_i(M, V_I) = Sh_i(M, V_I) \quad (4.8)$$

La proposition 1 indique que l'importance d'un attribut, mesurée par la fonction de Shapley, peut être exprimée comme la somme de sa contribution marginale pure avec la somme des interactions par paires de cet attribut avec les autres attributs.

Preuve 1 La preuve de cette proposition est donnée en annexe.

À partir de la proposition 1, le résultat suivant est obtenu.

Proposition 2. La décomposition des indices d'inégalité. Pour un jeu λ -inégalitaire (M, V_I) , l'inégalité $V_I(M) = I(x)$ est décomposable comme suit :

$$I(x) = \sum_{i \in M} CMP_i - \sum_{i \in M} \sum_{\substack{j \in M \\ i \neq j}} INT(i, j, T) \quad (4.9)$$

Preuve 2 L'axiome d'efficacité de la fonction de Shapley conduit à ce résultat.

La proposition 2 met en lumière que tous les indices d'inégalité peuvent être décomposés entre deux termes, le premier terme correspondant à la somme des contributions marginales pures associées aux attributs, et le second étant la somme pondérée de toutes les interactions par paires possibles.

4.4 Discussion et conclusion

Dans le chapitre 2, des décompositions par sous-populations et des décompositions par sources de revenu ont été présentées. Les décompositions multiples, qui à la fois estiment un effet de groupe et un effet de source, ont aussi été évoquées. La décomposition à la Shapley permet également de faire des décompositions par groupe et par source, et elle permettrait de faire des décompositions entre plusieurs effets de groupes en même temps en considérant les groupes comme des attributs. Ainsi, la décomposition à la Shapley des indices d'inégalité a plusieurs avantages : elle peut prendre plusieurs dimensions, elle est efficace (la somme des importances est l'inégalité totale), elle s'applique à une grande variété d'indices d'inégalité, et enfin, comme ce chapitre le démontre, cette décomposition peut distinguer la contribution marginale pure d'un attribut de son importance totale, la différence entre les deux résidant dans la prise en compte des interactions avec les autres attributs.

En effet, l'importance d'un attribut étant égal à sa contribution marginale pure, moins la somme des ses interactions par paires, qui peuvent être positives ou négatives. La contribution marginale pure d'un attribut correspond à l'impact qu'aurait la suppression des inégalités inhérentes à cet attribut sur l'inégalité totale. Il apparaît dès lors crucial, en termes d'anticipation des effets des politiques publiques égalisatrices, de prendre en compte les interactions pour comprendre pourquoi l'impact d'un attribut peut différer de sa contribution totale à l'inégalité.

Le chapitre suivant montre comment la prise en considération des interactions entre les différents attributs permet d'interpréter le résultat sur les inégalités de salaire d'une politique publique qui vise à supprimer les écarts de rémunération entre les hommes et les femmes.

4.5 Annexe

Preuve de la proposition 1

Pour un jeu λ -inégalitaire (M, V_I) et un attribut $i \in M$. Sa fonction de Shapley correspondante et sa fonction de Shapley revisitée sont données par $Sh_i(M, V_I)$ et $\overline{Sh}_i(M, V_I)$, respectivement.

Premièrement :

$$Sh_i(M, V_I) = \left(\sum_{\substack{S \subseteq M \\ i \notin S}} \frac{s!}{m.(m-1) \dots (m-s)} V_I(S \cup \{i\}) \right) - \left(\sum_{\substack{S \subseteq M \\ i \in S}} \frac{s!}{m.(m-1) \dots (m-s)} V_I(S) \right) = \alpha - \beta$$

De la même façon :

$$\begin{aligned} \overline{Sh}_i(M, V_I) &= \left(V_I(M) - \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} [V_I(T \cup \{ij\}) - V_I(T \cup \{i\})] \right) \\ &\quad - \left(V_I(M \setminus \{i\}) + \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} [V_I(T) - V_I(T \cup \{j\})] \right) \\ &= \left(V_I(M) - \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T \cup \{ij\}) + \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T \cup \{i\}) \right) \\ &\quad - \left(V_I(M \setminus \{i\}) + \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T) - \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T \cup \{j\}) \right) \\ &= \alpha' - \beta' \end{aligned}$$

Il faut donc prouver que $\alpha = \alpha'$ et $\beta = \beta'$ pour que $Sh_i(M, V_I) = \overline{Sh}_i(M, V_I)$.

1^{ère} étape : $\alpha = \alpha'$

Posons $\alpha' = V_I(M) - A + B$, avec

$$A = \sum_{j \neq i} \sum_{\substack{T \subseteq M \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T \cup \{ij\})$$

et

$$B = \sum_{j \neq i} \sum_{\substack{T \subseteq \mathcal{M} \\ i, j \notin T}} \frac{(m-t-2)!(t+1)!}{m!} V_I(T \cup \{i\})$$

i) Dans un premier temps A est considéré.

Avec $K = T \cup \{j\}$ et $|K| = k = t + 1$. Puis $K \subseteq M$ et $i \notin K$. Ainsi,

$$A = \sum_{\substack{K \subseteq \mathcal{M} \\ i \notin K}} \sum_{j \in K} \frac{(m-(k-1)-2)!((k-1)+1)!}{m!} V_I(K \cup \{i\})$$

De plus, puisque $j \in K$:

$$A = \sum_{\substack{K \subseteq \mathcal{M} \\ i \notin K}} k \frac{(m-(k-1)-2)!((k-1)+1)!}{m!} V_I(K \cup \{i\}),$$

Ainsi,

$$\begin{aligned} A &= \sum_{\substack{K \subseteq \mathcal{M} \\ i \notin K}} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \\ &= \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \end{aligned}$$

ii) Maintenant $V_I(M) - A$ est étudié :

$$\begin{aligned} V_I(M) - A &= V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \\ &= V_I(M) - \frac{(m-1).(m-1)!}{m.(m-1) \dots (m-(m-1))} V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \\ &= V_I(M) - \frac{(m-1).(m-1)!}{m!} V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \\ &= V_I(M) - \frac{(m-1)}{m} V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \\ &= \frac{1}{m} V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m-1) \dots (m-k)} V_I(K \cup \{i\}) \end{aligned}$$

iii) B est maintenant observé.

Avec $K' = T$ et $|K'| = k' = t$. Puisque $j \notin K'$:

$$\begin{aligned}
B &= \sum_{\substack{K' \subseteq \mathcal{M} \\ i, j \notin K'}} (m - k' - 1) \frac{(m - k' - 2)!(k' + 1)!}{m!} V_I(K' \cup i) \\
&= \sum_{\substack{K' \subseteq \mathcal{M} \\ i, j \notin K'}} \frac{(k' + 1)!}{(k' + 1)!} m.(m - 1) \dots (m - k') V_I(K' \cup i) \\
&= \sum_{\substack{K' \subseteq \mathcal{M} \\ i, j \notin K'}} \frac{(k' + 1)!}{m.(m - 1) \dots (m - k')} V_I(K' \cup i) \\
&= \sum_{K \subseteq \mathcal{M} \setminus i} \frac{(k + 1)!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\})
\end{aligned}$$

iv) Finalement,

$$V_I(M) - A + B = \frac{1}{m} V_I(M) - \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k.k!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\}) + \sum_{K \subseteq \mathcal{M} \setminus i} \frac{(k + 1)!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\})$$

Par conséquent,

$$\begin{aligned}
V_I(M) - A + B &= \frac{1}{m} V_I(M) + \sum_{K \subseteq \mathcal{M} \setminus i} \frac{(k + 1)! - k.k!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\}) \\
&= \frac{1}{m} V_I(M) + \sum_{K \subseteq \mathcal{M} \setminus i} \frac{(k + 1)! - k.k!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\}) \\
&= \frac{1}{m} V_I(M) + \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\}) \\
&= \sum_{K \subseteq \mathcal{M} \setminus i} \frac{k!}{m.(m - 1) \dots (m - k)} V_I(K \cup \{i\}) \\
&= \alpha
\end{aligned}$$

2^{ème} étape : $\beta = \beta'$

Sachant $\beta' = V_I(M \setminus \{i\}) + C - D$, avec

$$C = \sum_{j \neq i} \sum_{\substack{T \subseteq \mathcal{M} \\ i, j \notin T}} \frac{(m - t - 2)!(t + 1)!}{m!} V_I(T)$$

et

$$D = \sum_{j \neq i} \sum_{\substack{T \subseteq \mathcal{M} \\ i, j \notin T}} \frac{(m - t - 2)!(t + 1)!}{m!} V_I(T \cup \{j\})$$

i) Premièrement, D est étudié.

Avec $L = T \cup \{j\}$ et $|L| = l = t + 1$, avec $L \subseteq \mathcal{M}$, $i \notin L$. Ainsi,

$$D = \sum_{\substack{L \subseteq \mathcal{M} \\ i \notin L}} \sum_{j \neq i} \frac{(m - (l - 1) - 2)!((l - 1) + 1)!}{m!} V_I(L)$$

De plus, puisque $j \in L$:

$$D = \sum_{\substack{L \subseteq \mathcal{M} \\ i \notin L}} l \frac{(m - (l - 1) - 2)!((l - 1) + 1)!}{m!} V_I(L)$$

Finalement,

$$D = \sum_{\substack{L \subseteq \mathcal{M} \\ i \notin L}} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L)$$

ii) Deuxièmement, $V_I(M \setminus \{i\}) - D$ est étudié.

$$\begin{aligned} V_I(M \setminus \{i\}) - D &= V_I(M \setminus \{i\}) - \sum_{\substack{L \subseteq \mathcal{M} \\ i \notin L}} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\ &= V_I(M \setminus \{i\}) - \frac{(m-1).(m-1)!}{m.(m-1) \dots (m-(m-1))} V_I(M \setminus \{i\}) - \sum_{\substack{L \subset \mathcal{M} \\ i \notin L}} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\ &= V_I(M \setminus \{i\}) - \frac{(m-1)}{m} V_I(M \setminus \{i\}) - \sum_{\substack{L \subset \mathcal{M} \\ i \notin L}} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\ &= \frac{1}{m} V_I(M \setminus \{i\}) - \sum_{L \subset \mathcal{M} \setminus i} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \end{aligned}$$

iii) Troisièmement, C est étudié.

Avec $L' = T$ et $|L'| = l' = t$, avec $L' \subseteq \mathcal{M}$, $i, j \notin L'$.

Puisque $j \notin L'$,

$$\begin{aligned} C &= \sum_{\substack{L' \subseteq \mathcal{M} \\ i, j \notin L'}} (m - l' - 1) \frac{(m - l' - 2)!((l' + 1)!}{m!} V_I(L') \\ &= \sum_{\substack{L' \subseteq \mathcal{M} \\ i, j \notin L'}} \frac{(l' + 1)!}{m.(m-1) \dots (m-l')} V_I(L') \\ &= \sum_{L \subset \mathcal{M} \setminus i} \frac{(l+1)!}{m.(m-1) \dots (m-l)} V_I(L) \end{aligned}$$

iv) Pour conclure, il est obtenu :

$$\begin{aligned}
V_I(M \setminus \{i\}) + C - D &= \frac{1}{m} V_I(M) + \sum_{L \subset \mathcal{M} \setminus i} \frac{(l+1)!}{m.(m-1) \dots (m-l)} V_I(L) - \sum_{L \subset \mathcal{M} \setminus i} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\
&= \frac{1}{m} V_I(M \setminus \{i\}) + \sum_{L \subset \mathcal{M} \setminus i} \frac{(l+1)! - l!}{m.(m-1) \dots (m-l)} V_I(L) \\
&= \frac{1}{m} V_I(M \setminus \{i\}) + \sum_{L \subset \mathcal{M} \setminus i} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\
&= \sum_{\substack{L \subseteq \mathcal{M} \\ i \notin L}} \frac{l!}{m.(m-1) \dots (m-l)} V_I(L) \\
&= \beta
\end{aligned}$$

Chapitre 5

Mesurer l'impact d'une politique égalisatrice sur l'inégalité : une application à l'écart de rémunération entre les femmes et les hommes en France

Sommaire

5.1	Introduction	120
5.2	Les attributs, leur importance et leur impact	121
5.2.1	Les distributions associées aux attributs	121
5.2.2	L'importance de l'attribut genre	122
5.2.3	L'impact d'une politique égalisatrice efficace	124
5.3	L'impact de la suppression de l'écart de salaire entre les hommes et les femmes, une application sur données françaises	128
5.3.1	Les données	128
5.3.2	L'importance de l'attribut genre	129
5.3.3	L'impact du genre	130
5.3.4	L'interprétation des interactions	130
5.3.5	Un effet de composition sur l'importance du genre	132
5.4	Discussion et conclusion	134
5.5	Annexe : Distribution des observations	135

5.1 Introduction

En 2006, en France, une loi a été votée pour supprimer l'écart de salaire entre les femmes et les hommes sous 5 ans ¹. Au 1^{er} janvier 2006, le salaire moyen des femmes à plein temps représentait 81.5% de celui des hommes. Quelques années plus tard, ce chiffre n'a que légèrement augmenté pour atteindre 82,4% en 2011 (Observatoire des inégalités (2017)). Il semble donc que cette loi a échoué à réduire l'écart de rémunération entre les travailleurs des deux sexes, pour autant une question subsiste : quel aurait-été le niveau d'inégalité obtenu si l'application de cette loi avait été un succès ? Autrement dit, dans quelle mesure une telle politique publique égalisatrice réduirait-elle l'inégalité globale en France ?

Pour répondre à cette question, une décomposition à la Shapley semble l'approche la plus adéquate (Chantreuil et Trannoy (2011, 2013)). En effet, cette méthode permet de décomposer l'inégalité de salaire entre les différentes caractéristiques individuelles (ou attributs). Cela permet de définir la part des inégalités attribuable à l'écart de rémunération entre sexe, appelée importance du genre dans l'inégalité (Chantreuil et Lebon (2015)).

L'analyse menée dans ce chapitre montre que la suppression des différences de salaire dues à un attribut, le genre dans notre cas, conduit à une évolution des inégalités qui ne correspond pas à la simple soustraction de l'importance du genre dans l'inégalité. Ce résultat a priori contre-intuitif est expliqué par les interactions du genre avec les autres caractéristiques individuelles. Comme le chapitre 4 l'a démontré, l'impact sur l'inégalité totale de l'élimination des différences de salaire observées sur un attribut peut être affaibli ou accentué par les interactions de cet attribut avec les autres. Par exemple, il est possible que l'effet du genre sur l'inégalité soit plus élevé pour les femmes les plus âgées que pour les femmes les plus jeunes. Si tel est le cas, l'âge accentue l'inégalité attribuable au genre (et vice versa).

Ainsi, la présente analyse de l'écart de rémunération entre les hommes et les femmes est différente de celles communément réalisées dans la littérature (voir par exemple Blau et Kahn (2017) pour une revue de la littérature sur le sujet). L'approche classique des travaux sur ce sujet est celle proposée initialement par Oaxaca-Blinder (Oaxaca (1973), Blinder (1973)), dans laquelle la formation du salaire pour chaque genre est étudiée séparément, puis comparée afin d'expliquer l'écart de salaire. Par conséquent, cet écart peut être expliqué par un effet structurel et un effet de composition (voir section 2.4.1). Puis dans un second temps, l'importance des facteurs peut être calculée dans chaque effet. La démarche adoptée ici est différente, les deux groupes (hommes et femmes) ne sont plus étudiés séparément mais sont considérés en même temps. Ainsi, à la différence d'Oaxaca-Blinder, l'objectif n'est pas de comparer deux groupes avec une statistique calculée à partir de leur distribution de salaire respective, mais d'évaluer l'effet du groupe en lui-même sur l'inégalité (*i.e.*

1. Loi relative à l'égalité de salaire entre les femmes et les hommes N°2006-340 du 23/03/2006, publié au Journal Officiel n°71 du 24/03/2006.

les effets redistributifs d'un attribut sur l'inégalité, tout en considérant les interactions entre chaque attributs). Pour autant, l'outil de décomposition à la Shapley et la méthode d'Oaxaca-Blinder sont plus complémentaires que contradictoires, puisqu'il est possible de calculer l'effet de composition et l'effet structurel sur l'importance des différents attributs.

D'après la loi de 2006, tous les employeurs doivent assurer, pour un même travail ou un travail de valeur égale, une égalité de salaire entre les hommes et les femmes. L'équivalence entre deux emplois est appréciée selon les connaissances professionnelles, l'expérience, les responsabilités, et le stress physique et mental. Dans la présente analyse, l'écart de rémunération entre hommes et femmes est calculé pour un âge donné, comme proxy de l'expérience professionnelle, et pour un niveau d'éducation donné, pour un proxy des connaissances professionnelles. Malheureusement, par faute de données suffisantes, le niveau de responsabilité de l'employé, ainsi que son stress physique et mental ne peuvent pas être pris en compte.

Grâce aux données collectées par l'Enquête Emploi INSEE (2006, 2011, 2016), l'importance des caractéristiques individuelles (âge, niveau d'éducation, genre) dans l'inégalité sont calculées et leur évolution dans le temps peut être étudiée. L'effet de la loi de 2006 sur l'inégalité si celle-ci avait été un succès, est ensuite calculée. L'étude menée ici révèle que malgré l'importance non négligeable du genre pour expliquer l'inégalité, l'effet d'une politique qui donnerait le même salaire aux hommes et femmes à caractéristiques identiques serait faible. Une analyse contrefactuelle est également réalisée pour mettre en avant un effet de composition sur l'importance des attributs considérés dans l'inégalité. Le résultat global permet de mettre en lumière l'intérêt d'une décomposition à la Shapley pour anticiper les effets d'une politique égalisatrice. Dans la lignée du chapitre précédent, mais cette fois-ci dans une approche appliquée, ce travail met en évidence l'intérêt et la nécessité de prendre en compte les interactions entre les différents attributs pour pouvoir évaluer et comprendre le résultat d'une politique publique égalisatrice.

La deuxième section de ce chapitre présente la méthodologie adoptée, la distinction entre l'importance d'un attribut et son impact pourra ainsi être faite. La troisième section applique cette méthode sur des données françaises pour illustrer les différents mécanismes. La dernière section conclut.

5.2 Les attributs, leur importance et leur impact

5.2.1 Les distributions associées aux attributs

Selon la loi, l'écart de rémunération jugé illégal est apprécié à expérience et compétences égales, ainsi l'écart de salaire observé entre les hommes et les femmes est calculé à âge donné, comme proxy de l'expérience professionnelle, et pour un niveau d'éducation donné, comme proxy

des connaissances professionnelles. L'attribut du genre est donc associé avec un surplus ou une perte de salaire par rapport au salaire moyen des individus ayant le même âge et le même niveau de diplôme.

En conséquence, les distributions de salaire associées avec l'âge, l'éducation et le genre sont déterminées respectivement comme suit :

- Le salaire moyen des individus qui ont le même âge ;
- Pour un âge donné, la différence de salaire entre le salaire moyen des individus ayant le même niveau d'éducation et le salaire moyen de tous les individus ;
- Pour un âge donné et un niveau d'éducation donné, la différence entre le salaire moyen des individus du même sexe avec le salaire moyen de l'ensemble des individus.

L'attribut « caractéristiques inobservées » prend une valeur propre à chaque individu. Elle est égale à la différence entre le salaire de l'individu avec le salaire moyen des personnes de même âge, même niveau d'éducation et même genre. Par conséquent, le salaire d'un individu i pour un âge a_i , un niveau d'éducation e_i , un genre g_i et dont les caractéristiques inobservées sont notées u_i , peut être considéré comme la somme de ces quatre attributs :

$$y_i(a_i, e_i, g_i, u_i) = \bar{y}_{a_i} + (\bar{y}_{a_i, e_i} - \bar{y}_{a_i}) + (\bar{y}_{a_i, e_i, g_i} - \bar{y}_{a_i, e_i}) + (y_i(a_i, e_i, g_i, u_i) - \bar{y}_{a_i, e_i, g_i}) \quad (5.1)$$

Il y a donc quatre distributions, chacune associée à un attribut, et leur somme équivaut à la distribution du salaire initial. L'inégalité de salaire de cette distribution initiale peut être décomposée par la méthode de Shapley. Cette méthode permet de déterminer la part des inégalités attribuable aux disparités de salaire en termes d'âge, de niveau d'éducation, de genre ou des caractéristiques inobservées. La section suivante présente la méthodologie qui permet une telle décomposition.

5.2.2 L'importance de l'attribut genre

Comme mentionné précédemment, la décomposition mise en place permet d'associer chaque attribut avec une part de l'inégalité salariale, mesurée par l'indice de Gini, noté $\Gamma(Y)$ (Gini (1921)). L'indice de Gini calculé sur la distribution des salaires est noté $\Gamma(Y_A, Y_E, Y_G, Y_U)$, puisque la somme des distributions associées à chaque attribut est égale à la distribution des salaires : $Y_A + Y_E + Y_G + Y_U = Y$. Avec les vecteurs de rémunérations associés : $Y_A = |\bar{y}_a|$, $Y_E = |\bar{y}_{a,e}| - |\bar{y}_a|$, $Y_G = |\bar{y}_{a,e,g}| - |\bar{y}_{a,e}|$ et $Y_U = |Y| - |\bar{y}_{a,e,g}|$. Enfin, $\mu_A, \mu_E, \mu_G, \mu_U$ sont respectivement les valeurs moyennes associées à l'âge, au niveau d'éducation, au genre et aux caractéristiques inobservées. Au regard de ces notations, il convient de noter que :

$$\Gamma(Y_A, Y_E, Y_G, Y_U) = \Gamma(Y) \quad (5.2)$$

et

$$\Gamma(\mu_A, \mu_E, \mu_G, \mu_U) = 0 \quad (5.3)$$

Ainsi, $\Gamma(\mu_A, \mu_E, Y_G, \mu_U)$ mesure l'inégalité dans la distribution des valeurs associées au genre uniquement.

Pour reprendre les notations du chapitre 4, la fonction caractéristique $V_\Gamma(\cdot)$ évalue l'inégalité d'une coalition d'attributs S , notée $Y(S, \lambda)$, mesurée par l'indice de Gini $\Gamma(\cdot)$, ainsi $V_\Gamma(S) = \Gamma(Y(S, \lambda))$. L'approche retenue est un jeu égalitaire, avec $\lambda = \mu(X) - \mu(X^\delta)$, autrement dit pour neutraliser l'effet des disparités observées d'un attribut, sa distribution est remplacée par une distribution égalitaire où chaque individu reçoit la valeur moyenne correspondante.

Pour une décomposition efficace des inégalités, celles-ci sont décomposées par la fonction de Shapley (Section 4.3 du chapitre 4), ainsi :

$$Sh_A + Sh_E + Sh_G + Sh_U = \Gamma(Y) \quad (5.4)$$

Avec Sh_j , la valeur de Shapley calculée sur la distribution des rémunérations associé à l'attribut j (Y_j). À partir des notations précédentes, l'importance de chaque attribut est déterminée comme suit :

$$\begin{aligned} Sh_A = & \frac{1}{4}[\Gamma(Y_A, \mu_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U)] \\ & + \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, Y_U)] \end{aligned} \quad (5.5)$$

$$\begin{aligned} Sh_E = & \frac{1}{4}[\Gamma(\mu_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U)] \\ & + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U)] \\ & + \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U)] \\ & + \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, Y_U)] \end{aligned} \quad (5.6)$$

$$\begin{aligned}
Sh_G = & \frac{1}{4}[\Gamma(\mu_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U)] \\
& + \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, Y_U)]
\end{aligned} \tag{5.7}$$

$$\begin{aligned}
Sh_U = & \frac{1}{4}[\Gamma(\mu_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U)] \\
& + \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, Y_G, \mu_U)]
\end{aligned} \tag{5.8}$$

Ainsi, l'importance de chaque attribut apparait bien comme la somme pondérée de l'ensemble de ses contributions marginales.

5.2.3 L'impact d'une politique égalisatrice efficace

Sans effet intrinsèque du genre, le salaire moyen des hommes et des femmes serait identique pour un âge et un niveau d'éducation donnés. De ce fait, c'est le résultat qui devrait être obtenu si la politique mise en place parvient à remplir ses objectifs. Cela entraînerait une nouvelle distribution de salaire, où les valeurs associées à l'âge et au genre ne seraient pas modifiées, mais où les différences de rémunération générées par le genre seraient nulles : $(\bar{y}_{a_i, e_i, g_i} - \bar{y}_{a_i, e_i}) = 0, \forall i$. Il en résulterait une importance nulle du genre dans la nouvelle distribution, notée $Sh_{\tilde{G}}$:

$$Sh_{\tilde{G}} = 0 \tag{5.9}$$

L'inégalité restante est alors répartie entre l'importance des attributs d'âge, d'éducation et de caractéristiques inobservées, notées respectivement $Sh_{\tilde{A}}$, $Sh_{\tilde{E}}$ et $Sh_{\tilde{U}}$. Ces nouvelles importances prennent en considération le fait que les valeurs associées au genre sont maintenant toutes égales à la moyenne μ_G (qui est nulle par construction). Ainsi, les nouvelles importances obtenues sont égales à :

$$\begin{aligned}
Sh_{A_{\tilde{G}}} = & \frac{1}{4}[\Gamma(Y_A, \mu_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U)]
\end{aligned} \tag{5.10}$$

$$\begin{aligned}
Sh_{E_{\tilde{G}}} = & \frac{1}{4}[\Gamma(\mu_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& + \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U)]
\end{aligned} \tag{5.11}$$

$$\begin{aligned}
Sh_{U_{\tilde{G}}} = & \frac{1}{4}[\Gamma(\mu_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& + \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& + \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U)]
\end{aligned} \tag{5.12}$$

La somme de l'importance des attributs dans la nouvelle situation, décrites par les équations 5.9 à 5.12, est égale à l'inégalité de salaire calculée ex post :

$$Sh_{A_{\tilde{G}}} + Sh_{E_{\tilde{G}}} + Sh_{\tilde{G}} + Sh_{U_{\tilde{G}}} = \Gamma(Y_A, Y_E, \mu_G, Y_U) \tag{5.13}$$

Par conséquent, il semble naturel de postuler qu'une politique qui supprime l'inégalité due à l'attribut genre fait varier l'inégalité d'un montant égal à $\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, Y_U)$. Autrement dit, la variation de l'inégalité induite par une telle politique est l'impact de cette politique sur l'inégalité globale des salaires. Cependant, comme le chapitre précédent l'a démontré, l'impact sur l'inégalité d'un attribut lorsque sa distribution est remplacée par sa moyenne est différent de son importance. Une autre façon de montrer cela est de réécrire l'impact comme étant la différence entre les importances ex ante et ex post de la politique mise en œuvre :

$$\begin{aligned}
\Delta\Gamma &= (Sh_{A_{\bar{G}}} + Sh_{E_{\bar{G}}} + Sh_{\bar{G}} + Sh_{U_{\bar{G}}}) - (Sh_A + Sh_G + Sh_E + Sh_U) \\
&= -Sh_G + (Sh_{A_{\bar{G}}} - Sh_A) + (Sh_{E_{\bar{G}}} - Sh_E) + (Sh_{U_{\bar{G}}} - Sh_U)
\end{aligned} \tag{5.14}$$

Avec :

$$\begin{aligned}
Sh_{A_{\bar{G}}} - Sh_A &= \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
&+ \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U)] \\
&- \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, Y_U)]
\end{aligned} \tag{5.15}$$

et

$$\begin{aligned}
Sh_{E_{\bar{G}}} - Sh_E &= \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(Y_A, Y_E, \mu_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
&+ \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U)] \\
&- \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U)] \\
&- \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, Y_U)]
\end{aligned} \tag{5.16}$$

et

$$\begin{aligned}
Sh_{U_{\bar{G}}} - Sh_U &= \frac{1}{12}[\Gamma(\mu_A, \mu_E, \mu_G, Y_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(Y_A, \mu_E, \mu_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
&+ \frac{1}{12}[\Gamma(\mu_A, Y_E, \mu_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
&+ \frac{1}{4}[\Gamma(Y_A, Y_E, \mu_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U)] \\
&- \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, Y_G, \mu_U)]
\end{aligned} \tag{5.17}$$

D'après les équations 5.14 à 5.17, il est possible de déterminer quelle aurait été l'évolution de l'inégalité de salaire si le biais du genre avait été supprimé.

Dans l'intention de montrer l'interaction du genre avec les autres attributs, les équations 5.14 à 5.17 sont réécrites comme suit :

$$\begin{aligned}
Sh_{A_{\bar{G}}} - Sh_A &= -\frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U) + \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
&- \frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U) + \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
&- \frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, Y_U) + \Gamma(\mu_A, Y_E, \mu_G, Y_U)]
\end{aligned} \tag{5.18}$$

$$\begin{aligned}
Sh_{E_{\bar{G}}} - Sh_E = & -\frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& -\frac{1}{12}[\Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(Y_A, Y_E, \mu_G, \mu_U) + \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& -\frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U) + \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& -\frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, \mu_G, Y_U) + \Gamma(Y_A, \mu_E, \mu_G, Y_U)]
\end{aligned} \tag{5.19}$$

$$\begin{aligned}
Sh_{U_{\bar{G}}} - Sh_U = & -\frac{1}{12}[\Gamma(\mu_A, \mu_E, Y_G, Y_U) - \Gamma(\mu_A, \mu_E, Y_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, Y_U)] \\
& -\frac{1}{12}[\Gamma(Y_A, \mu_E, Y_G, Y_U) - \Gamma(Y_A, \mu_E, Y_G, \mu_U) - \Gamma(Y_A, \mu_E, \mu_G, Y_U) + \Gamma(Y_A, \mu_E, \mu_G, \mu_U)] \\
& -\frac{1}{12}[\Gamma(\mu_A, Y_E, Y_G, Y_U) - \Gamma(\mu_A, Y_E, Y_G, \mu_U) - \Gamma(\mu_A, Y_E, \mu_G, Y_U) + \Gamma(\mu_A, Y_E, \mu_G, \mu_U)] \\
& -\frac{1}{4}[\Gamma(Y_A, Y_E, Y_G, Y_U) - \Gamma(Y_A, Y_E, Y_G, \mu_U) - \Gamma(Y_A, Y_E, \mu_G, Y_U) + \Gamma(Y_A, Y_E, \mu_G, \mu_U)]
\end{aligned} \tag{5.20}$$

Les équations 5.18 à 5.20, avec l'équation 5.14, montrent que l'impact de la suppression du biais du genre, à âge et diplôme donnés, sur l'inégalité de salaire est égal à l'importance du genre, plus une somme pondérée des interactions entre le genre et les autres attributs. À l'instar du chapitre précédent, le terme d'interaction se réfère à la notion développée notamment par Murofushi et Soneda (1993) (équation 4.4). C'est-à-dire que l'évaluation de l'interaction entre deux attributs doit considérer le cas où ils sont les seuls créateurs d'inégalité, mais aussi les cas où d'autres attributs sont créateurs d'inégalité puisque ces derniers peuvent influencer la relation entre les deux attributs considérés.

Les données fournies par l'Enquête Emploi INSEE (2006, 2011, 2016) permettent dans un premier temps d'évaluer l'efficacité de la loi de 2006 par rapport à la situation relative des hommes et femmes sur le marché du travail en termes de salaire. Ensuite, l'enquête de 2011 permet de simuler l'impact sur l'inégalité de salaire mesurée par l'indice de Gini de la loi votée en 2006 si celle-ci avait été un succès. Dans l'étude présentée ici, l'écart de rémunération entre hommes et femmes est supprimé, ce qui conduit à une augmentation de salaire pour les femmes et à une diminution pour les hommes. D'autres solutions étaient envisageables. Par exemple, dans la littérature de la déprivation, il est parfois préféré de supprimer l'écart de rémunération lié au genre uniquement dans le cas où celui-ci désavantage l'individu, ce qui concrètement revient à aligner le salaire des femmes sur celui des hommes à caractéristiques données (del Río *et al.* (2011)). Cependant, la généralisation d'une telle politique à l'économie globale conduirait à une augmentation du coût du travail. Par conséquent, cela affecterait le fonctionnement du marché du travail, et rendrait plus hypothétique l'existence *ceteris paribus* d'une nouvelle distribution des salaires où seule une variation positive de rémunération est envisagée pour une partie de la population. C'est la raison d'une approche qui consiste à supprimer l'écart de rémunération de genre pour tous.

5.3 L'impact de la suppression de l'écart de salaire entre les hommes et les femmes, une application sur données françaises

5.3.1 Les données

L'Enquête Emploi réalisée par l'INSEE est utilisée pour évaluer l'inégalité de salaires en France. Au sein de cette enquête, toutes les personnes de 15 ans ou plus habitant dans des logements ordinaires à titre de résidences principales sont interrogées².

Tout d'abord, les observations pour lesquelles au moins une information est manquante concernant le salaire³, l'âge, le plus haut niveau de diplôme obtenu ou le genre ont été supprimées. Ensuite, seules les personnes salariées (selon la définition du BIT) âgées entre 20 et 60 ans ont été prises en compte. La première restriction du champ des salariés est nécessaire pour que l'étude soit en accord avec la loi de 2006 qui porte sur les inégalités salariales, et la restriction sur l'âge est adoptée à cause du manque d'observations en-deçà et au-delà de ces limites⁴.

Les individus qui déclarent travailler moins de 30h ou plus de 45h en moyenne hebdomadaire sur l'année sont supprimées. La borne inférieure de cette dernière restriction évite les effets propres au temps partiel, notamment au regard des carrières des femmes, et facilite donc l'interprétation des résultats. Et la borne supérieure est basée sur les limites légales de nombre d'heures de travail pouvant être réalisé par semaine⁵.

Par ailleurs, pour comparer les individus, il est nécessaire d'étudier l'inégalité en termes de salaires horaires, puisque deux individus peuvent avoir un même salaire mensuel mais ne pas avoir effectué le même nombre d'heures de travail pour l'obtenir. Après traitement des données, il reste 26033 observations en 2006, 38445 observations en 2011 et 23795 en 2016⁶.

Enfin, l'âge est regroupé 8 niveaux (par tranche de 5 années), l'éducation est divisée en 11 ni-

2. Un logement ordinaire est un logement séparé et indépendant (du point de vue de son accès). Sont donc exclus les foyers de travailleurs, les centres d'hébergement ou d'accueil, les hôpitaux, les maisons de retraite, les cités universitaires, etc... Depuis 2003, cette enquête est réalisée chaque semaine sur une base trimestrielle. Chaque logement est questionné six fois, avec un intervalle trimestriel entre chaque session.

3. Pour les salariées : rémunération mensuelle nette en euros retirée de la profession principale, redressée des non-réponses, des réponses fournies en tranches ainsi que des primes mensualisées. À ma connaissance, aucun détail n'est fourni par l'INSEE sur les méthodes de redressement.

4. Cette dernière restriction n'est pas inhabituelle, par exemple, Blau et Kahn (2017) se focalisent sur les travailleurs américains qui ont entre 25 et 64 ans.

5. En effet, en France, les employés ne peuvent pas travailler plus de 48h par semaine (voir Article L3121-35 du Code du Travail) et le nombre d'heures de travail réalisé en moyenne par semaine sur 12 semaines consécutives ne peut pas excéder 44 heures (voir Article L3121-36 du Code du Travail).

6. Pour éviter l'effet des valeurs extrêmes, notamment dues à des erreurs de saisies, les individus qui ont un salaire horaire inférieur à 5,5, 6 et 6,5 respectivement en 2006, 2011 et 2016 ne sont pas pris en compte. Ces restrictions sont basées sur le salaire minimum légal et la distribution observée des salaires.

veaux différents : Sans Diplôme (ou seulement le certificat d'études primaires), Brevet des collèges, CAP-BEP (ou autre diplôme de ce niveau), Baccalauréat général, Baccalauréat professionnel (ou autre diplôme de ce niveau), Baccalauréat + 2 ans, Licence ou Master 1 (Bac +3 ou +4), Master (Bac +5), Diplôme d'école d'ingénieur, Diplôme d'école du commerce, Doctorat en étude de la santé et Doctorat hors étude en santé. Les distributions des individus par âge, niveau d'éducation et genre sont données en Annexe 5.5.

5.3.2 L'importance de l'attribut genre

Le tableau 5.1 donne l'importance de chaque attribut en 2006, 2011 et 2016. Ces résultats sont plutôt similaires sur les trois années étudiées. En 2016, 10,74% de l'inégalité de salaire mesurée par l'indice de Gini est expliquée par l'attribut âge, 25,56% par le niveau d'éducation, 6,87% par le genre contre 56,83% par les caractéristiques inobservées. L'évolution de l'importance du genre d'une période à une autre est très faible, avec 6,98% en 2011 et 6,87% en 2016, ce qui remet sérieusement en cause l'efficacité de la loi votée en 2006. La prochaine section permet d'estimer plus précisément l'effet d'une politique publique efficace en quantifiant ce qu'aurait été la variation de l'inégalité globale si l'application de la loi avait effectivement permis la suppression des écarts de salaire dus au genre.

Tableau 5.1 – Décomposition par attribut de l'inégalité de salaire en 2006, 2011 et 2016

Année	Attribut	Importance	Importance Relative
2006	Age	0,02074	10,62%
	Éducation	0,04914	25,17%
	Genre	0,01292	6,62%
	Inobservé	0,11242	57,59%
	Gini	0,19522	100%
2011	Age	0,02049	10,87%
	Éducation	0,04431	23,50%
	Genre	0,01316	6,98%
	Inobservé	0,11058	58,65%
	Gini	0,18854	100%
2016	Age	0,02063	10,74%
	Éducation	0,04909	25,56%
	Genre	0,01320	6,87%
	Inobservé	0,10912	56,83%
	Gini	0,19203	100%

Source : Enquête Emploi (2006, 2011 et 2016) et auteurs.

5.3.3 L'impact du genre

L'impact du genre est évalué à partir du niveau d'inégalité qui pourrait être obtenu si les différences de salaire entre hommes et femmes, à âge et diplôme donnés, n'existaient plus. Le tableau 5.3 présente les détails de l'estimation de l'importance du genre en 2011. Les contributions marginales calculées à partir de la distribution observée sont données dans la 2^{ème} colonne, et celles calculées à partir de la distribution contrefactuelle, où W_G est remplacé par μ_G , sont données dans la 3^{ème} colonne. La 4^{ème} colonne indique la différence entre les deux précédentes. Une différence positive correspond à une interaction défavorable entre le genre et l'attribut considéré, puisque cette interaction augmente le niveau d'inégalité. Au contraire, si la différence est négative, alors l'interaction entre les attributs diminue l'inégalité.

Le tableau 5.2 résume les résultats en mettant en avant l'importance du genre et son impact en 2011. Comme le chapitre précédent l'a montré, l'impact d'un attribut fait référence à sa contribution marginale pure, c'est-à-dire sa contribution marginale à l'inégalité qui résulte de la différence entre le niveau d'inégalité de salaire et le niveau d'inégalité lorsque la distribution associée au genre est remplacée par sa distribution égalitaire moyenne (ex. $\Gamma(W_A, W_E, W_G, W_U) - \Gamma(W_A, W_E, \mu_G, W_U)$ pour l'impact du genre).

Compte tenu des interactions entre le genre et les autres attributs, si le salaire des hommes et des femmes était identique à âge et niveau de diplôme donnés, alors l'inégalité ne diminuerait non pas de 6,98% mais de 1,09%. Ceci s'explique par le fait qu'une large proportion de l'importance du genre est liée à ses interactions avec les autres attributs.

Tableau 5.2 – Importance du genre et son impact sur l'inégalité

Attribut	Importance	Importance Relative	Impact	Impact Relatif
Gender	0,01316	6,98%	0,00206	1,09%

Source : Enquête Emploi (2011) et auteurs.

Dans la sous-section suivante, une proposition d'interprétation de ces interactions est donnée selon des faits stylisés et la littérature existante sur ce sujet.

5.3.4 L'interprétation des interactions

Du fait des interactions entre les attributs, si le salaire des hommes et des femmes était égal à âge et niveau d'éducation donnés, alors l'inégalité, mesurée par l'indice de Gini, ne diminuerait pas de 6,98%, mais seulement de 1,09%.

Tableau 5.3 – Décomposition de l'inégalité de salaire en 2011

L'importance des attributs et leurs interactions

Importance des attributs en 2011	Distribution observée	Biais du genre supprimé	Différence
Sh_A = $\frac{1}{4}[\Gamma(W_A, \mu_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, \mu_G, \mu_U) - \Gamma(\mu_A, W_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, W_G, \mu_U) - \Gamma(\mu_A, \mu_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, \mu_G, W_U) - \Gamma(\mu_A, \mu_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, W_G, \mu_U) - \Gamma(\mu_A, W_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, \mu_G, W_U) - \Gamma(\mu_A, W_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, W_G, W_U) - \Gamma(\mu_A, \mu_E, W_G, W_U)]$ $+ \frac{1}{4}[\Gamma(W_A, W_E, W_G, W_U) - \Gamma(\mu_A, W_E, W_G, W_U)]$	0,02049 0,012741 0,001171 0,002425 0,000740 0,001083 0,000431 0,000693 0,001202	0,02254 0,012741 0,001171 0,004247 0,000740 0,001171 0,000431 0,000740 0,001294	+0,00205
Sh_E = $\frac{1}{4}[\Gamma(\mu_A, W_E, \mu_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, \mu_G, \mu_U) - \Gamma(W_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, W_G, \mu_U) - \Gamma(\mu_A, \mu_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, \mu_G, W_U) - \Gamma(\mu_A, \mu_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, W_G, \mu_U) - \Gamma(W_A, \mu_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, \mu_G, W_U) - \Gamma(W_A, \mu_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, W_G, W_U) - \Gamma(\mu_A, \mu_E, W_G, W_U)]$ $+ \frac{1}{4}[\Gamma(W_A, W_E, W_G, W_U) - \Gamma(W_A, \mu_E, W_G, W_U)]$	0,04431 0,022047 0,004273 0,004895 0,001830 0,003554 0,001521 0,001768 0,004426	0,04769 0,022047 0,004273 0,007349 0,001830 0,004273 0,001521 0,001830 0,004564	+0,00337
Sh_G = $\frac{1}{4}[\Gamma(\mu_A, \mu_E, W_G, \mu_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, W_G, \mu_U) - \Gamma(W_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, W_G, \mu_U) - \Gamma(\mu_A, W_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, \mu_E, W_G, W_U) - \Gamma(\mu_A, \mu_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, W_G, \mu_U) - \Gamma(W_A, W_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, W_G, W_U) - \Gamma(W_A, \mu_E, \mu_G, W_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, W_G, W_U) - \Gamma(\mu_A, W_E, \mu_G, W_U)]$ $+ \frac{1}{4}[\Gamma(W_A, W_E, W_G, W_U) - \Gamma(W_A, W_E, \mu_G, W_U)]$	0,01316 0,0009389 0,001308 0,000676 0,000264 0,000588 0,000218 0,000203 0,000515	0 	-0,01316
Sh_R = $\frac{1}{4}[\Gamma(\mu_A, \mu_E, \mu_G, W_U) - \Gamma(\mu_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, \mu_G, W_U) - \Gamma(W_A, \mu_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, \mu_G, W_U) - \Gamma(\mu_A, W_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, \mu_E, W_G, W_U) - \Gamma(\mu_A, \mu_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, W_E, \mu_G, W_U) - \Gamma(W_A, W_E, \mu_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(W_A, \mu_E, W_G, W_U) - \Gamma(W_A, \mu_E, W_G, \mu_U)]$ $+ \frac{1}{12}[\Gamma(\mu_A, W_E, W_G, W_U) - \Gamma(\mu_A, W_E, W_G, \mu_U)]$ $+ \frac{1}{4}[\Gamma(W_A, W_E, W_G, W_U) - \Gamma(W_A, W_E, W_G, \mu_U)]$	0,11058 0,039836 0,009772 0,007759 0,010413 0,007020 0,008682 0,007286 0,019810	0,11626 0,039836 0,009772 0,007759 0,013279 0,007020 0,009772 0,007759 0,021059	+0,00568
Gini	0,18854	0,18648	-0,00206

Source : Enquête Emploi (2011) et auteurs.

L'interaction entre le genre et le niveau d'éducation peut être expliquée en partie par des effets de composition, qui ne peuvent malheureusement pas être pris en compte avec les données disponibles. Un même niveau d'éducation n'offre pas les mêmes opportunités de salaire selon la

spécialité du diplôme et il apparaît que les femmes sont sous-représentées dans des types de diplômes très profitables en termes de perspectives de rémunération. Par exemple, en 2009, la part des femmes dans les classes préparatoires scientifiques était seulement de 30,5%, 26% dans les classes d'ingénieurs et 24% dans les IUT (Caille *et al.* (2011)).

Il y a aussi les effets de rupture de carrière. En 2006, 28% des femmes avait connu une interruption de carrière pour s'occuper de leurs enfants, contre 2% des hommes. La durée moyenne de ces ruptures de carrières est d'environ 4 ans et 3 mois (DARES (2010)). Ce phénomène est une combinaison de multiples effets, effet d'âge (*i.e.* l'âge des femmes à la naissance de leurs enfants), effet générationnel (*i.e.* le changement de pratiques sociales entre générations) et un effet d'éducation⁷. Par conséquent, l'effet des interruptions de carrières pour s'occuper des enfants peut être capté par les interactions du genre avec l'âge et le niveau d'éducation.

De plus, les effets de l'âge et de l'éducation se combinent avec l'effet du genre dans le sens où l'interaction des deux premiers avec le second augmente l'inégalité. En effet, l'échantillon utilisé comprend des individus âgés de 20 à 60 ans, il est donc probable que les femmes qui sont entrées sur le marché du travail dans les années 2000 n'ont pas fait face aux mêmes difficultés que celles qui sont entrées sur le marché du travail dans les années 70. Le manque de considération envers les femmes était probablement plus marqué quelques décennies auparavant qu'aujourd'hui. En conséquence, les opportunités de carrière étaient plus réduites auparavant, ce qui entraîne une situation contemporaine dans laquelle les postes à responsabilité, avec des salaires plus élevés, sont plus souvent occupés par des hommes que par des femmes. Par exemple, la part des femmes au sein de la catégorie socio-professionnelle des cadres et des professions intellectuelles était de 37,3% en 2007 et de 40,2% en 2012 (Observatoire des inégalités (2014)).

En résumé, une loi qui imposerait un salaire égal entre hommes et femmes à expérience et compétence égales n'aurait qu'un effet marginal sur l'inégalité global. Les résultats de ce travail mettent en évidence l'intérêt de prendre en considération les effets de long terme et les choix individuels pour analyser l'efficacité potentielle des politiques publiques égalisatrices. L'objectif de la méthode développée est d'aider à la mise en place des politiques efficaces, notamment pour réduire l'inégalité entre hommes et femmes sur le marché du travail en permettant de réelles possibilités d'anticipation de leurs effets.

5.3.5 Un effet de composition sur l'importance du genre

Pour aller plus loin dans l'interprétation des interactions entre le genre et les autres attributs, l'effet de composition en termes de genre est mis en avant. Cet effet indique la conséquence d'une

7. La part des femmes qui ont interrompu leur carrière pour s'occuper de leurs enfants et la durée de moyenne de l'interruption diminuent toutes les deux avec le niveau d'éducation. Elles sont respectivement de 42% et de 6 ans pour les moins diplômées, et 25% et 31 mois pour les plus diplômées (DARES (2010)).

sous-représentation ou d'une sur-représentation des hommes et des femmes par catégorie âge-diplôme, relativement à la proportion observée de ces deux groupes dans la société sur l'importance des attributs.

Une méthode similaire au score de propension, utilisée communément dans la littérature de l'évaluation des politiques publiques (DiNardo *et al.* (1996)), est appliquée pour réaliser cette séparation entre effet de composition et effet inexpliqué. Ainsi, l'échantillon des hommes est pondéré par un facteur $\Psi(X)$, qui est égal à :

$$\Psi(X) = \frac{P(\text{Genre}=Femme|Age,Education)/P(\text{Genre}=Femme)}{P(\text{Genre}=Homme|Age,Education)/P(\text{Genre}=Homme)}$$

Avec X la distribution relative à chaque attribut. Ici, la probabilité ($P(.)$) est calculée non paramétriquement, elle est basée sur les parts observées d'hommes et de femmes dans chaque groupe (*i.e.* chaque paire âge-niveau d'éducation correspondant à un groupe). Grâce à ce facteur de repondération, les poids des hommes sont réajustés alors que celui des femmes restent inchangés. Une nouvelle décomposition de l'inégalité est réalisée selon ces nouvelles pondérations. Ainsi, la différence entre les résultats issus de la distribution initiale et ceux obtenus par le contrefactuel indique à quel point la différence entre les poids par groupe comparés à ceux de la population totale impacte l'estimation de l'importance de chaque caractéristique personnelle. Les résultats sont donnés dans le tableau 5.4.

Tableau 5.4 – Importance initiale et contrefactuelle en 2011

	Distribution initiale (1)		Contrefactuelle (2)		Effets de composition (1) - (2)	
	Importance	Imp. Relative	Importance	Imp. Relative	Brute	Relative
Age	0,02049	10,87%	0,02008	10,49 %	0,00040	1,96 %
Éducation	0,04431	23,50%	0,04289	22,40%	0,00413	3,22%
Genre	0,01316	6,98%	0,01546	8,08%	- 0,00230	-17,50%
Inobservé	0,11058	58,65%	0,11303	59,03%	- 0,00245	-2,22%
Gini	0,18854	100%	0,19147	100%	-0,00293	-1,55%
Impact du Genre	0,00206	1,09%	0,00460	2,40%		

Source : Enquête Emploi (2011) et traitement des auteurs.

Dans chaque cas, les importances relatives sont approximativement les mêmes, ce qui indique un faible effet de composition. En effet, l'inégalité est plus élevée de 0,00293 points (+1,55%) dans la distribution contrefactuelle pour l'année 2011. L'effet de composition totale négatif indique que la composition dans la distribution initiale diminue l'inégalité par rapport à une distribution où la part des femmes et des hommes serait la même pour chaque groupe considéré que celle observée dans la population totale. L'effet de composition du genre et des attributs inobservés sont les principales causes de cet effet de composition total. S'il n'y avait pas eu d'effets de composition, l'importance du genre dans l'inégalité serait 17,5% plus forte. Ce dernier résultat suggère que les

femmes sont sous-représentées (sur-représentées) dans les catégories âge-diplôme souffrant le plus (le moins) de l'écart de rémunération lié au genre.

A la lumière de ces résultats, les interactions entre les attributs semblent peu liées à un effet de composition. Si les hommes et les femmes étaient répartis équitablement, supprimer l'écart de rémunération entre ces deux groupes réduirait l'inégalité de 2,40%, au lieu de 1,09% dans la distribution observée (Tableau 5.4).

5.4 Discussion et conclusion

La décomposition à la Shapley des indices d'inégalités permet de diviser l'inégalité totale entre les différents attributs. A priori, il est judicieux de supposer que l'importance d'un attribut correspond au montant par lequel diminuerait l'inégalité si les disparités de salaire associées à cet attribut étaient supprimées. Cependant, le résultat principal du chapitre précédent sous un angle théorique, et celui de ce chapitre dans une optique empirique, est que l'évolution de l'inégalité ne correspond pas nécessairement à l'importance d'un attribut. Ceci s'explique par l'effet des interactions entre les différents attributs.

Ces différents points ont été illustrés avec des données françaises pour étudier l'efficacité d'une loi votée en 2006, dont l'objectif était de supprimer l'inégalité de salaires entre les hommes et les femmes sous 5 ans. Les résultats montrent que, malheureusement, cet objectif n'a pas été atteint, puisque l'importance du genre dans l'inégalité totale en 2011 est similaire à celle de 2006. La méthode développée permet aussi de déterminer quel aurait été le niveau d'inégalité salariale si la loi avait été un succès. Il apparaît que l'inégalité, mesurée par l'indice de Gini, aurait diminué uniquement de 1,09% alors que l'importance du genre dans l'inégalité est de 6,98%. Cette diminution de l'inégalité, plus faible que l'importance du genre, est expliquée par les interactions entre le genre et les autres attributs : âge, niveau d'éducation et caractéristiques inobservées. L'existence de ces interactions met en lumière la complexité à déterminer les sources des inégalités de salaire entre hommes et femmes à laquelle une simple politique égalisatrice ne pourrait répondre que partiellement.

L'outil développé permet d'anticiper les effets sur la distribution des salaires d'une politique égalisatrice supprimant les disparités de salaires provoquées par un attribut. Cet outil pourrait aussi être utilisé dans le cadre d'une décomposition par sous-populations ou par sources de salaire. Ainsi, il est possible pour les décideurs politiques d'estimer l'efficacité de telles politiques au regard des leurs objectifs.

Malgré ces avantages, cette méthode souffre de quelques limites. Premièrement, la distribution associée à chaque attribut est dépendante de l'ordre par lequel les facteurs sont considérés dans la première étape (*i.e.* la désagrégation du salaire), et la multiplication des facteurs pris en consi-

dération peut mener à un manque d'observations pour certaines occurrences. Dans de futures recherches, d'autres méthodes de désagrégation pourront être utilisées pour dépasser ces problèmes. Par exemple, l'utilisation de modèles économétriques de type régression quantile ou régression spline pourront être utilisés pour désagréger le salaire. Cependant, même avec l'outil économétrique, il reste quelques limites, notamment s'il y a des variables catégorielles puisque l'importance de ce type de variable dépendra de la catégorie de référence choisie pour estimer le modèle. Par ailleurs, la décomposition à la Shapley des indices d'inégalités ne dispense pas le chercheur de faire des choix normatifs, notamment concernant le choix de l'indice d'inégalité utilisé ou encore la manière dont l'effet des disparités de salaires est éliminé. Dans ce chapitre, l'effet du genre a été supprimé pour tous les individus. Ce choix peut être remis en question puisqu'il entraîne une augmentation de salaire pour les femmes et une diminution de salaire pour les hommes. D'autres solutions sont possibles. Entre autres, dans la lignée de la littérature sur la déprivation, seul le désavantage en termes de rémunération expérimenté par les femmes pourrait être supprimé (voir par exemple del Río *et al.* (2011)). Cependant, une généralisation d'une telle politique au niveau de toute l'économie entraînerait une augmentation du coût du travail (*i.e.* une augmentation de la masse salariale). Par conséquent, cela affecterait le fonctionnement du marché du travail, il est alors difficile de prévoir les effets de cette hausse sur la distribution des salaires. C'est pourquoi le choix de cette étude s'est porté sur un jeu égalitaire, où le déficit de salaire des femmes est compensé par le surplus de celui des hommes, de telle sorte que la masse salariale ne varie pas.

5.5 Annexe : Distribution des observations

Tableau 5.5 – Distribution des employés à plein-temps - 2006

Distribution des employés à plein-temps par âge, diplôme et genre - 2006												
	Peu/pas diplômé	Brevet des collèges	BAC Pro	BAC Général	BAC +2	BAC +3-4	Master Bac +5	Ecole de commerce	Ecole d'ingénieur	Doctorat en santé	Autre doctorat	Total
Homme												
20	260	400	259	58	160	40	11	0	21	1	0	1210
25	325	425	355	114	366	122	86	8	81	4	8	1894
30	450	526	299	146	348	171	75	6	67	3	8	2099
35	553	837	188	96	268	123	54	6	55	7	15	2202
40	691	900	136	113	257	97	39	7	40	2	16	2298
45	739	861	123	137	188	83	26	7	25	5	9	2203
50	830	782	99	118	131	74	15	2	21	4	10	2086
55	487	370	70	76	89	74	9	1	16	9	17	1218
60	17	12	2	2	5	2	1	0	3	0	0	44
Total	4352	5113	1531	860	1812	786	316	37	329	35	83	15254
Femme												
20	85	120	178	76	209	62	17	2	8	0	0	757
25	109	158	207	139	436	248	105	17	26	6	3	1454
30	141	222	168	119	355	275	83	14	20	7	13	1417
35	220	381	164	126	262	149	56	7	10	7	7	1389
40	433	450	140	168	251	112	32	5	12	5	10	1618
45	532	447	135	199	194	99	29	3	6	12	9	1665
50	578	371	85	158	177	101	14	1	5	7	3	1500
55	389	227	45	67	97	81	8	1	0	8	4	927
60	26	7	3	5	7	2	1	0	1	0	0	52
Total	2513	2383	1125	1057	1988	1129	345	50	88	52	49	10779

Source : Auteurs, basés sur l'enquête emploi (INSEE) 2006.

Tableau 5.6 – Distribution des employés à plein-temps - 2011

Distribution des employés à plein-temps par âge, diplôme et genre - 2011												
	Peu/pas diplômé	Brevet des collèges	BAC Pro	BAC Général	BAC +2	BAC +3-4	Master Bac +5	Ecole de commerce	Ecole d'ingénieur	Doctorat en santé	Autre doctorat	Total
Homme												
20	270	445	415	82	210	89	39	2	29	3	0	1584
25	392	579	479	148	390	256	154	12	130	4	9	2553
30	418	651	475	211	471	254	142	21	115	9	29	2796
35	547	785	419	216	497	267	96	16	102	2	28	2975
40	793	1129	320	159	420	216	85	8	65	12	33	3240
45	923	1333	234	160	298	144	56	5	67	5	23	3248
50	1023	1192	195	167	244	108	33	7	43	10	23	3045
55	779	798	123	122	165	100	32	6	51	9	15	2200
60	43	31	11	12	18	18	3	0	6	2	1	145
Total	5188	6943	2671	1277	2713	1452	640	77	608	56	161	21786
Femme												
20	92	197	236	103	234	113	39	4	11	2	0	1031
25	146	230	329	158	529	317	260	17	48	16	5	2055
30	140	237	328	148	572	371	180	15	37	13	23	2064
35	221	335	288	197	492	399	115	24	24	19	18	2132
40	410	620	274	200	455	296	92	21	15	12	23	2418
45	620	762	259	276	398	191	49	7	14	14	17	2607
50	724	680	197	243	346	158	34	2	6	16	3	2409
55	620	474	111	193	227	119	30	3	4	8	5	1794
60	47	31	7	22	20	15	3	1	1	1	1	149
Total	3020	3566	2029	1540	3273	1979	802	94	160	101	95	16659

Source : Auteurs, basés sur l'enquête emploi (INSEE) 2011.

Tableau 5.7 – Distribution des employés à plein-temps - 2016

Distribution des employés à plein-temps par âge, diplôme et genre - 2016												
	Peu/pas diplômé	Brevet des collèges	BAC Pro	BAC Général	BAC +2	BAC +3-4	Master Bac +5	Ecole de commerce	Ecole d'ingénieur	Doctorat en santé	Autre doctorat	Total
Homme												
20	91	189	257	29	126	74	22	1	13	0	0	802
25	145	333	264	45	213	144	116	5	74	3	4	1346
30	194	344	307	56	231	147	130	6	88	5	19	1527
35	236	441	339	56	301	164	98	11	69	6	16	1737
40	265	563	333	86	299	185	83	7	85	3	20	1929
45	417	783	229	81	248	165	51	6	58	5	14	2057
50	472	849	169	69	231	110	48	2	45	0	21	2016
55	469	735	136	76	148	84	31	4	46	5	12	1746
60	51	40	19	5	23	11	2	0	6	1	2	160
Total	2340	4277	2053	503	1820	1084	581	42	484	28	108	13320
Femme												
20	31	87	141	36	103	84	47	0	5	0	0	534
25	45	143	171	34	196	208	191	6	23	11	3	1031
30	60	172	190	48	272	207	177	12	34	3	7	1182
35	95	179	233	66	318	255	120	6	11	12	12	1307
40	146	262	250	96	383	267	91	10	19	8	17	1549
45	208	455	250	91	315	247	74	9	14	6	20	1689
50	338	534	218	124	262	133	32	0	9	10	8	1668
55	313	432	156	101	201	107	31	2	7	8	2	1360
60	46	38	12	14	17	19	7	0	1	0	1	155
Total	1282	2302	1621	610	2067	1527	770	45	123	58	70	10475

Source : Auteurs, basés sur l'enquête emploi (INSEE) 2016.

Conclusion générale

Les différents chapitres de cette thèse ont eu pour objectif de présenter en différentes étapes la méthodologie nécessaire pour évaluer l'impact d'une politique publique égalisatrice. Le premier chapitre s'est intéressé à l'outil principal utilisé pour apprécier la distribution des revenus, c'est-à-dire aux indices d'inégalité. Le second chapitre a détaillé les différentes approches possibles pour décomposer un indice d'inégalité, que ce soit pour mettre en avant des effets de groupes, de sources et/ou encore d'attributs. Cet état de l'art a ensuite fait place à trois études originales. Ainsi, le troisième chapitre étudie les inégalités de salaire dans la fonction publique française, en mettant en lumière le fait que l'importance du genre dans l'inégalité est considérable mais disparate entre les régions administratives françaises. Le quatrième chapitre démontre théoriquement que l'importance d'un attribut peut se décomposer comme la somme de sa contribution marginale pure moins une somme d'interactions par paire avec les autres attributs. Le cinquième et dernier chapitre expose empiriquement l'effet de ces interactions sur l'impact d'une politique publique qui supprimerait l'écart de rémunération entre les hommes et les femmes.

Choisir un indice d'inégalité Le chapitre 1 a permis de mettre en lumière le fait que la mesure de l'inégalité est subjective, et qu'il n'y a donc pas une définition de l'inégalité évidente qui s'impose à tous. Tout d'abord, les principes sur lesquels reposent ces indices peuvent être discutés. Il est communément admis qu'il est préférable que l'indice soit relatif, ainsi si tous les revenus augmentent du même montant alors l'inégalité mesurée diminue. Pour autant, d'aucun pourrait préférer qu'il soit homogène linéairement et par conséquent que l'inégalité ne varie pas suite à une telle variation du revenu. Quoiqu'il en soit, même s'il y avait un accord sur les principes à satisfaire, d'autres paramètres peuvent rentrer en considération dans la mesure de l'inégalité, notamment celui de l'aversion à l'inégalité. En conséquence, le décideur public doit choisir parmi les différents indices aux paramètres divers celui qui reflète les valeurs de la société au regard de ce que devrait être la distribution des revenus. Cette notion de valeur repose sur le concept plus global de justice sociale, c'est d'ailleurs pourquoi il en a été question dans l'introduction générale de cette thèse et dans la discussion de ce premier chapitre. De futures études pourraient être utiles pour mieux comprendre la philosophie propre à chaque indice, notamment en évaluant le revenu seuil à partir

duquel l'ajout d'un individu ayant ce revenu augmente les inégalités dans la société. Ces études pourraient aussi porter sur les préférences des citoyens afin d'apprécier plus rigoureusement leurs souhaits quant à la distribution des richesses, cela permettrait donc de définir plus clairement le contrat social partagé par la société.

Décomposer les inégalités Le chapitre 2 présente un état de l'art sur les méthodes permettant de décomposer les indices d'inégalité pour évaluer l'importance de chaque groupe, de chaque source de revenus et/ou de chaque caractéristique individuelle dans l'inégalité totale. Ce chapitre a mis en lumière l'intérêt des décompositions à la Shapley. Cette approche efficace peut être utilisée pour une décomposition multi-niveaux, par exemple elle permet de calculer l'importance des sources de revenu dans l'inégalité de rémunération pour plusieurs groupes différents. Elle peut aussi donner lieu à une décomposition par attributs, ce qui permet la prise en considération de plusieurs effets de groupe en même temps et permet de les apprécier les uns par rapport aux autres. Enfin et surtout, la décomposition à la Shapley est applicable à une grande variété d'indices d'inégalité. L'intérêt de cette méthode est d'ailleurs démontré dans le chapitre 3.

L'importance du genre dans l'inégalité de rémunération dans la fonction publique française

Le chapitre 3 applique une décomposition à la Shapley pour calculer l'importance des écarts de rémunérations entre hommes et femmes dans l'inégalité salariale dans la fonction publique française en 2010. Cette étude a notamment mis en avant les disparités régionales de l'importance du genre dans l'inégalité, ceci reflète notamment des particularités régionales quant au fonctionnement du marché du travail. Avec un accès aux données plus important, cette enquête pourrait être approfondie pour prendre en compte d'autres caractéristiques individuelles, et ainsi permettre une analyse plus précise des écarts de rémunérations entre les hommes et les femmes, et éventuellement différencier ce qui est dû à des effets de composition et ce qui relève de la discrimination. Les deux chapitres suivants s'interrogent sur les effets qui auraient été obtenus avec la suppression de ces écarts sur les inégalités.

La prise en considération des interactions pour évaluer les politiques publiques égalisatrices

Le chapitre 4 prouve théoriquement que l'impact de la suppression de l'écart de rémunération entre hommes et femmes sur l'inégalité totale ne correspond pas à son importance. Il est démontré que l'importance d'un attribut est égale à son impact, moins une somme pondérée d'interactions par paires entre cet attribut avec les autres attributs. Sachant que l'impact d'un attribut correspond à sa contribution marginale pure, c'est-à-dire la différence entre le niveau d'inégalités quand tous les attributs sont pris en considération avec le niveau d'inégalité avec ces mêmes attributs excepté l'attribut en question. Autrement dit, l'importance d'un attribut correspond à son impact sur les

inégalités s'ils n'interagissait pas avec les autres attributs considérés. L'interaction entre deux attributs s'explique par le fait que la distribution associée à un attribut peut compenser ou au contraire aggraver les inégalités d'un autre. Une démonstration empirique de ce résultat est proposée dans le dernier chapitre, ce qui permet de mieux comprendre ce que peuvent être les interactions entre attributs.

L'importance et l'impact du genre dans l'inégalité de salaire en France Le dernier chapitre montre l'intérêt d'évaluer une politique publique égalisatrice avec une décomposition à la Shapley. Grâce à cette approche, l'étude montre que l'écart de rémunération entre hommes et femmes, à âge et niveaux de diplôme donnés, contribue à 6,98% à l'inégalité de salaire en France en 2011. Pour autant, supprimer cet écart ne réduirait les inégalités que de 1,09%. Cette différence s'explique par l'interaction entre l'attribut genre avec les autres attributs. Par exemple, si cet écart de rémunération augmente avec l'âge alors l'attribut genre aggrave les inégalités dues à l'âge, donc leur interaction est négative. Ce chapitre a aussi été l'occasion de montrer que cette approche peut être complétée par une décomposition entre effets de composition et effets structurels, présentée au chapitre 2, afin de mieux interpréter les résultats. Le travail exposé dans ce chapitre montre tout l'intérêt de la décomposition à la Shapley revisitée pour évaluer une politique publique égalisatrice.

Par la suite L'outil et la méthodologie développés au sein de cette thèse pour décomposer les indices d'inégalité pourront être améliorés. Un axe de recherche futur pourra notamment enrichir la compréhension des interactions. En effet, plusieurs causes influent sur l'interaction entre deux attributs : un effet de composition intragroupe (ex. : la répartition d'hommes et femmes au sein de chaque groupe d'âge ne correspond pas à la répartition observée dans la population totale), un effet de composition inter-groupe (*i.e.* issu d'une différence entre l'effectif d'hommes-femmes entre chaque groupe d'âge), un effet de différences des rendements intragroupe (*i.e.* le rendement en terme de revenu associé à un attribut n'est pas le même entre les individus d'un groupe d'âge partageant des caractéristiques similaires) ou encore des différences de rendements intergroupe (les rendements des caractéristiques individuelles entre les groupes d'âge sont différents). Plusieurs questions se posent alors : Comment apprécier ces différents effets de composition et effets structurels sur l'importance d'un attribut ? Un ordre doit-il être défini pour la prise en considération de ces effets ? Si oui, dans quel ordre prendre en compte ces effets ? etc...

Par ailleurs, dans cette thèse, deux méthodes sont considérées pour neutraliser l'inégalité : les jeux nuls et les jeux égalitaires. Cependant, y-aurait-il des cas où un mixte des deux seraient plus appropriés ? Par exemple, en neutralisant l'effet du genre en rendant nul le vecteur de revenu associé au genre et en neutralisant les autres effets par l'égalisation de leur vecteur. Si cela était souhaitable, que deviendrait-il de la décomposition à la Shapley revisitée ? Cette recherche reste à mener.

Enfin, l'outil développé est ici appliqué à l'étude de l'inégalité de revenu. Il semble opportun d'utiliser cet outil pour répondre à d'autres interrogations sociétales. Il peut notamment permettre d'apporter des éclaircissement dans l'étude la pauvreté : Quelles en sont ses causes ? Quelle est l'importance de chaque cause ? Comment les sources de la pauvreté interagissent entre-elles ? avec quelle intensité ? Voici donc là un bref exposé de futurs travaux de recherches qui s'inscrivent dans la continuité de cette thèse.

Bibliographie

- J. G. ALTONJI, P. BHARADWAJ et F. LANGE : Changes in the characteristics of American youth. *Working Paper, Yale University*, 2008.
- C. ARNSPERGER et P. VAN PARIJS : *Ethique économique et sociale*. La découverte édition, 2000.
- A. B. ATKINSON : On the measurement of inequality. *Journal of Economic Theory*, 2(3):244–263, 1970.
- C. AUVRAY et A. TRANNOY : Décomposition par source de l'égalié des revenus à l'aide de la valeur de Shapley. Sfax, 1992.
- R. BARSKY, J. BOUND, K. K. CHARLES et J. P. LUPTON : Accounting for the Black-White wealth gap : a nonparametric approach. *Journal of the American Statistical Association*, 97(459):663–673, 2002.
- S. BAZEN et P. MOYES : Elitism and stochastic dominance. *Social Choice and Welfare*, 39(1):207–251, 2012.
- J. BENTHAM : *An introduction to the principles of morals and legislation*. 1789.
- C. BLACKORBY, W. BOSSERT et D. DONALDSON : *Intertemporal population ethics : a welfarist approach*. Numéro 93-13 in Discussion paper. Vancouver : Dept. of Economics, University of British Columbia, 1993.
- C. BLACKORBY, W. BOSSERT et D. DONALDSON : *Population Issues in Social Choice Theory, Welfare Economics, and Ethics*. Cambridge University Press, Cambridge, 2005.
- F. BLAU et L. KAHN : The gender wage gap : extent, trends, and explanations. *Journal of Economic Literature*, 55(3):789–865, 2017.
- A. S. BLINDER : Wage discrimination : reduced form and structural estimates. *Journal of Human Resources*, 8(4):436–455, 1973.

- F. BOURGUIGNON : Decomposable income inequality measures. *Econometrica*, 47(4):901, 1979.
- K. F. BUTCHER et J. DINARDO : The immigrant and native-born wage distributions : evidence from United States censuses. *Industrial and labor relations review*, 56(1):97–121, 2002.
- J-P. CAILLE, F. DEFRESNE, S. KESKPAIK, C. LAMBERT, S. LEMAIRE, B. LE RHUM, A. REBEYROL et F. THOMAS : Filles et garçons sur le chemin de l'égalité de l'école à l'enseignement supérieur. Rapport technique, Direction de l'évaluation, de la prospective et de la performance. Ministère de l'éducation nationale., 2011.
- S.R. CHAKRAVARTY : Measuring inequality : the axiomatic approach. *In Handbook of Income Inequality Measurement*, pages 163–186. Silber J., kluwer academic publishers édition, 1999.
- N. CHAMENI : Linking Gini to Entropy : measuring inequality by an interpersonal class of indices. *Economics Bulletin*, 4(5):1–9, 2006a. URL <http://www.accessecon.com/pubs/EB/2006/Volume4/EB-06D00002A.pdf>.
- N. CHAMENI : A note on the decomposition of the coefficient of variation squared : comparing entropy and Dagum's methods. *Economics Bulletin*, 4(8):1–8, 2006b. URL <http://www.accessecon.com/pubs/EB/2006/Volume4/EB-05D60003A.pdf>.
- N. CHAMENI : A generalisation of the Gini coefficient : measuring economic inequality. *Mimeo*, 2011.
- F. CHANTREUIL, S. COURTIN, K. FOURREY et I. LEBON : A note on the decomposability of inequality measures. *Social Choice and Welfare*, 2019.
- F. CHANTREUIL, K. FOURREY et I. LEBON : Panorama régional de la contribution du genre aux inégalités dans la fonction publique ., *Revue d'économie politique*, 128(1):137, 2018.
- F. CHANTREUIL et I. LEBON : Gender contribution to income inequality. *Economics Letters*, 133:27–30, 2015.
- F. CHANTREUIL et A. TRANNOY : Inequality decomposition values : the trade-off between marginality and consistency. *THEMA Discussion Paper, Université Cergy Pontoise*, 1999.
- F. CHANTREUIL et A. TRANNOY : Inequality decomposition values. *Annals of Economics and Statistics*, (101/102):13–36, 2011.
- F. CHANTREUIL et A. TRANNOY : Inequality decomposition values : the trade-off between marginality and efficiency. *The Journal of Economic Inequality*, 11(1):83–98, 2013.

- A. CHARRAUD et K. SAADA : Les écarts de salaires entre hommes et femmes. *Economie et statistique*, 59(1):3–17, 1974.
- V. CHERNOZHUKOV, I. FERNANDEZ-VAL et B. MELLY : Inference on counterfactual distributions. *Econometrica*, 81(6):2205–2268, 2013.
- L. N. CHRISTOFIDES, A. POLYCARPOU et K. VRACHIMIS : Gender wage gaps, ‘sticky floors’ and ‘glass ceilings’ in Europe. *Labour Economics*, (21):86–102, 2013.
- F. A. COWELL : Inequality decomposition : three bad measures. *Bulletin of Economic Research*, 40 (4):309–312, 1988.
- F. A. COWELL et E. FLACHAIRE : Inequality measurement and the rich : why inequality increased more than we thought ? page 29, 2018.
- F. A. COWELL et S. P. JENKINS : How much inequality can we explain? A methodology and an application to the United-States. *The Economic Journal*, 105(429):421, 1995.
- C. DAGUM : Decomposition and interpretation of Gini and the generalized entropy inequality measures. volume 157th Meeting, pages 200–205. Business and Economic Statistics Section, 1997a.
- C. DAGUM : A new approach to the decomposition of the Gini income inequality ratio. *Empirical Economics*, 22:515–531, 1997b.
- Camilo DAGUM : Measuring the economic affluence between populations of income receivers. *Journal of Business and Economics Statistics*, 5(1):5–12, 1987.
- H. DALTON : The measurement of the inequality of incomes. *The Economic Journal*, 30(119):348, 1920.
- DARES : Interruptions de carrière professionnelle et salaires des hommes et des femmes en 2006. Rapport technique 11, Ministère de l’économie, de l’industrie et de l’emploi. Ministère du travail, des relations sociales, de la famille, de la solidarité et de la ville., 2010.
- P. DASGUPTA, A. SEN et D. STARRETT : Notes on the measurement of inequality. *Journal of Economic Theory*, 6(2):180–187, 1973.
- C. del Río, C. GRADÍN et O. CANTÓ : The measurement of gender wage discrimination : the distributional approach revisited. *Journal of Economic Inequality*, 9:57–86, 2011.
- DGAFF : Rapport annuel sur l’état de la fonction publique [2011-2012]. Politiques et pratiques de ressources humaines : Faits et chiffres. Rapport technique, Ministère de la réforme de l’Etat, de la décentralisation et de la fonction publique, La documentation Française, Paris, 2012.

- J. DiNARDO, N. FORTIN et T. LEMIEUX : Labor market institutions and the distribution of wages, 1973-1992 : a semiparametric approach. *Econometrica*, 64(5):1001–1044, 1996.
- D. DONALDSON et J. A. WEYMARK : A single-parameter generalization of the Gini indices of inequality. *Journal of Economic Theory*, 22:67–86, 1980.
- M. DUBOIS : *Essais sur les principes de transferts dans un cadre welfariste-parétien avec séparabilité forte*. Thèse de doctorat, Université de Montpellier, 2016.
- C. DUVIVIER, J. LANFRANCHI et M. NARCY : Les sources de l'écart de rémunération entre hommes et femmes au sein des trois versants de la fonction publique. Rapport technique, TEPP - Rapport de recherche, 2015.
- C. DUVIVIER et M. NARCY : The motherhood wage penalty and its determinants : a public-private comparison. *Labour*, 29(4):415–443, 2015.
- U. EBERT : On the decomposition of inequality : partitions into nonoverlapping sub-groups. *In Measurement in Economics*. Physica, Heidelberg, eichhorn w. édition, 1988.
- U. EBERT : The decomposition of inequality reconsidered : weakly decomposable measures. *Mathematical Social Sciences*, 60(2):94–103, 2010.
- John C. H. FEI, G. RANIS et S. W. Y. KUO : Growth and the family distribution of income by factor components. *The Quarterly Journal of Economics*, 92(1):17, 1978.
- G. S. FIELDS : Inequality in urban colombia : a decomposition analysis. *Review of Income and Wealth*, 25(3):327–341, 1979.
- S. FIRPO, N. FORTIN et T. LEMIEUX : Decomposing wage distributions using recentered influence function regressions. *Working Paper*, 2007.
- S. FIRPO, N. FORTIN et T. LEMIEUX : Unconditional quantile regressions. *Econometrica*, 77(3):953–973, 2009.
- S. FIRPO, N. FORTIN et T. LEMIEUX : Decomposing wage distributions using recentered influence function regressions. *Econometrics*, 6(28), 2018.
- M. FLEURBAEY et P. MICHEL : Transfer principles and inequality aversion, with an application to optimal growth. *Mathematical Social Sciences*, 42(1):1–11, 2001.
- N. FORTIN, T. LEMIEUX et S. FIRPO : Decomposition methods in economics. *In Handbook of Labor Economics*, volume IV.A, pages 1–102. Orley Ashenfelter and David Card, Amsterdam : North-Holland, 2011.

- K. FOURREY : Local taxation and income segregation in France. *In Divides cities. Understanding intra-urban inequalities*, pages 30–31. OECD Publishing, Paris, 2018.
- K. FOURREY, I. LEBON et F. CHANTREUIL : Measuring the impact of equalizing public policies on inequality : a decomposition approach. *Mimeo, Université Caen Normandie*, 2018.
- F. FREMIGACCI, L. GOBILLON, D. MEURS et S. ROUX : Égalité professionnelle entre les hommes et les femmes : des plafonds de verre dans la fonction publique ? *Economie et statistique*, (488-489):97–121, 2016.
- J. GARDEAZABAL et A. UGIDOS : More on identification in detailed wage decompositions. *Review of Economics and Statistics*, 86(4):1034–1036, 2004.
- Jonah B GELBACH : When do covariates matter ? And which ones, and how much ? 2009.
- C. GINI : Measurement of inequality of income. *Economic Journal*, 31:22–43, 1921.
- M. GRABISCH : The representation of importance and interaction of features by fuzzy measures. *Pattern Recognition Letters*, 17:567–575, 1996.
- M. GRABISCH : Alternative representations of discrete fuzzy measures for decision making. *Int. J. of Uncertainty, Fuzziness, and Knowledge Based Systems*, 5:587–607, 1997a.
- M. GRABISCH : k-order additive discrete fuzzy measures and their representation. *Fuzzy Sets and Systems*, (92):167–189, 1997b.
- M. GRABISCH, M. ROUBENS et J-L. MARICHAL : Equivalent representations of set functions. *Math. Oper. Res.*, 25(2):167–189, 2000.
- N. GRAVEL : Comparing societies with different numbers of individuals on the basis of their advantage. *In Social Ethics and Normative Economics. Essays in Honour of Serge-Christophe Kolm*. Fleurbaey, M., Salles, M., and Weymark J. A., Springer-Verlag Berlin Heidelberg édition, 2011.
- N. GRAVEL, B. MAGDALOU et P. MOYES : Ranking distributions of an ordinal attribute. *Document de travail AMSE*, 50:1–25, 2014.
- F. GUÉGOT : *L'égalité professionnelle hommes-femmes dans la fonction publique : rapport au Président de la République*. Collection des rapports officiels. La Documentation française, Paris, 2011.
- P. J. HAMMOND : A note on extreme inequality aversion. *Journal of Economic Theory*, 11(3):465–467, 1975.

- F. R. HAMPEL : The influence curve and its role in robust estimation. *Journal of the American Statistical Association*, 69(346):383–393, 1974.
- F. HAYEK : *La Constitution de la liberté*. 1960.
- C. HERRERO et A. VILLAR : Comparaciones de renta real y evaluacion del bienestar. *Herri ekonomiaz. Revista de economia publica*, 2:79–102, 1989.
- K. HIRANO, G. W. IMBENS et G. RIDDER : Efficient estimation of average treatment effects using the estimated propensity score. *Econometrica*, 71(4):1161–1189, 2003.
- INSEE : Enquête emploi en continu (version fpr). *ADISP-CMH*, 2006.
- INSEE : Système d’information sur les agents des services publics. Guide d’utilisation du fichier poste., 2010.
- INSEE : Enquête emploi en continu (version fpr). *ADISP-CMH*, 2011.
- INSEE : Couples - Familles - Ménages en 2010. Recensement de la population - Base infracommunale (IRIS), 2013. URL <https://www.insee.fr/fr/statistiques/2028571>.
- INSEE : Enquête emploi en continu (version fpr). *ADISP-CMH*, 2016.
- INSEE : Les Femmes sur le marché du travail insulaire. Moins souvent actives et davantage au chômage. Rapport technique 9, 2016a.
- INSEE : Population active, emploi et chômage en 2013 - France métropolitaine, 2016b. URL <https://www.insee.fr/fr/statistiques/2020404/?geo=METRO-1>.
- INSEE : Indicateurs régionaux sur les inégalités entre les femmes et les hommes, 2018. URL <https://www.insee.fr/fr/statistiques/2513786>.
- INSEE et STATISTIQUE PUBLIQUE : Femmes et hommes, l’égalité en question. Rapport technique, 2017.
- S. P. JENKINS et P. VAN KERM : Trends in income inequality, pro-poor income growth, and income mobility. *Oxford Economic Papers*, 58(3):531–548, 2006.
- M. KENDALL et A. STUART : *The advanced theory of statistics*, volume Vol. 1 : Distribution Theory. Charles Griffin & Company Limited, London, 4 édition, 1977.
- A. KHINCHIN : *Mathematical formulations of information theory*. Dover Publications, New York, 1957.

- A. B. KHMELNITSKAYA et E. B. YANOVSKAYA : Owen coalitional value without additivity axiom. *Mathematical Methods of Operations Research*, 66(2):255–261, 2007.
- I. KOJADINOVIC : Modeling interaction phenomena using fuzzy measures : on the notions of interaction and independence. *Fuzzy Sets and Systems*, 135:317–340, 2003.
- I. KOJADINOVIC : Estimation of the weights of interacting criteria from the set of profiles by means of information-theoretic functionals. *European Journal of Operational Research*, 155:741–751, 2004.
- I. KOJADINOVIC : An axiomatic approach to the measurement of the amount of interaction among criteria or players. *Fuzzy Sets and Systems*, 152:417–435, 2005.
- S-C. KOLM : Unequal inequalities I. *Journal of Economic Theory*, 12:416–442, 1976a.
- S-C. KOLM : Unequal inequalities II. *Journal of Economic Theory*, 13:82–111, 1976b.
- S. KULLBACK : *Information theory and statistics*. Wiley, New York, 1959.
- I. LEBON, J-P. GUIRONNET, F. GAVREL et F. CHANTREUIL : La contribution des écarts de rémunération entre les femmes et les hommes à l'inégalité des rémunérations dans la fonction publique : une approche par la décomposition des inégalités. *Economie et statistique*, (488-489):151–168, 2016.
- R. I. LERMAN et S. YITZHAKI : Income inequality effects by income source : a new approach and applications to the United-States. *The Review of Economics and Statistics*, 67(1):151, 1985.
- J. LOCKE : *Les deux Traités du gouvernement civil*. 1690.
- M-O LORENZ : Methods for measuring concentration of wealth. *Journal of the American Statistical Association*, 9(70):209–219, 1905.
- C. LUCIFORA et D. MEURS : The public sector pay gap in France, Great Britain and Italy. *Review of Income and Wealth*, 52(1):43–59, 2006.
- J. A. F. MACHADO et J. MATA : Counterfactual decomposition of changes in wage distributions using quantile regression. *Journal of Applied Econometrics*, 20(4):445–465, 2005.
- K. MARX : *Le Capital. Critique de l'économie politique*. 1867.
- D. MEURS, A. PAILHÉ et S. PONTHEUX : Child-related Career Interruptions and the Gender Wage Gap in France. *Annals of Economics and Statistics*, (99/100):15–46, 2010.

- D. MEURS et S. PONTHEUX : Une mesure de la discrimination dans l'écart de salaire entre hommes et femmes. *Economie et statistique*, (337-338):135–158, 2000.
- D. MEURS et S. PONTHEUX : L'écart des salaires entre les femmes et les hommes peut-il encore baisser? *Economie et statistique*, 398-399:99–129, 2006.
- Y. MICHAUD : *Egalité et inégalités*. UTLS. Odile Jacob, 2006.
- A. C. MONTI : The Study of the Gini Concentration Ratio by Means of the Influence Function. 1991.
- P. MORNET : On the axiomatization of the weakly decomposable inequality indices. *Mathematical Social Sciences*, 83:71–78, 2016.
- P. MORNET, S. MUSSARD, F. SEYTE et M. TERRAZA : La décomposition de l'indicateur de Gini en sous-groupes. *Revue française d'économie*, 29(2):179, 2014.
- P. MORNET, C. ZOLI, S. MUSSARD, J. SADEFO-KAMDEM, F. SEYTE et M. TERRAZA : The (α, β) -multi-level α -Gini decomposition with an illustration to income inequality in France in 2005. *Economic Modelling*, 35:944–963, 2013.
- P. MOYES : Comparisons of heterogeneous distributions and dominance criteria. *Journal of Economic Theory*, 147(4):1351–1383, 2012.
- T. MUROFUSHI et S. SONEDA : Techniques for reading fuzzy measures (III) : interaction index. Sapporo, Japan, 1993.
- S. MUSSARD : La décomposition des mesures d'inégalité en sources de revenu : méthodes et applications. *L'Actualité économique*, 83(3):415, 2007.
- S. MUSSARD, M. N. PI ALPERIN, F. SEYTE et M. TERRAZA : Extensions of Dagum's Gini decomposition. *Statistica & Applicazioni IV*, pages 5–30, 2006.
- S. MUSSARD et L. SAVARD : The Gini multi-decomposition and the role of Gini's transvariation : application to partial trade liberalization in the Philippines. *Applied Economics*, 44(10):1235–1249, 2012.
- S. MUSSARD et M. TERRAZA : Décompositions des mesures d'inégalité : le cas des coefficients de Gini et d'entropie. *Recherches économiques de Louvain*, 75(2):151, 2009.
- M. MUSSINI : On decomposing inequality and poverty changes over time : A multi-dimensional decomposition. *Economic Modelling*, 33:8–18, 2013.

- H. S. NIELSEN, M. SIMONSEN et M. VERNER : Does the Gap in Family-friendly Policies Drive the Family Gap ? *Scandinavian Journal of Economics*, 106(4):721–744, 2004.
- R. OAXACA : Male-Female wage differentials in urban labor markets. *International Economic Review*, 14(3):693, 1973.
- R. OAXACA et M. R. RANSOM : Identification in detailed wage decompositions. *Review of Economics and Statistics*, 81(1):154–157, 1999.
- OBSERVATOIRE DES INÉGALITÉS : Une répartition déséquilibrée des professions entre les hommes et les femmes, 2014. URL http://www.inegalites.fr/spip.php?page=article&id_article=1048&id_groupe=15&id_mot=103&id_rubrique=114.
- OBSERVATOIRE DES INÉGALITÉS : L'évolution des inégalités de salaires entre hommes et femmes, 2017. URL http://www.inegalites.fr/spip.php?page=article&id_article=1482&id_groupe=15&id_mot=104&id_rubrique=114.
- G. OWEN : Multilinear extension of games. *Management Science*, (18):64–79, 1972.
- G. OWEN : Values of games with a priori unions. In *Mathematical Economics and Game Theory*, volume 141 de *Lecture Notes in Economics and Mathematical Systems*,. Henn R., Moeschlin O., Berlin, Heidelberg, springer édition, 1977.
- D. PARFIT : *Reasons and Persons*. Clarendon Press, Oxford, 1984.
- A. C. PIGOU : *Wealth and Welfare*. MacMillan and Co., Limited, London, 1912.
- G. PYATT : On the interpretation and disaggregation of Gini coefficients. *The Economic Journal*, 86 (342):243, 1976.
- G. PYATT, C-N. CHEN et J. FEI : The distribution of income by factor components. *The Quarterly Journal of Economics*, 95(3):451, 1980.
- J. RAWLS : *A theory of justice*. 1971.
- J. ROEMER : *Free to Lose : An introduction to Marxist Economic Philosophy*. Havard University Press, 1988.
- C. ROTHE : Decomposing the composition effect. *Journal of Business Economics and Statistics*, 33 (323-37), 2015.
- M. ROTHSCHILD et J. E. STIGLITZ : Some further results on the measurement of inequality. *Journal of Economic Theory*, 6(2):188–204, 1973.

- M. SASTRE et A. TRANNOY : A marginalist approach to Inequality Decomposition by factor components : an application to OECD countries using the LIS database. *Mimeo THEMA*, 2000.
- M. SASTRE et A. TRANNOY : Shapley inequality decomposition by factor components : some methodological issues. *Journal of Economics*, 9:51–89, 2002.
- M. SASTRE et A. TRANNOY : Shapley inequality decomposition by factor components : some methodological issues. *THEMA, Université de Cergy-Pontoise*, page 14, 2018.
- A. SEN : *On Economic Inequality*. Clarendon Press, Oxford, 1973.
- A. SEN : Real National Income. *Review of economic studies*, 43(1):19–39, 1976.
- A. SEN : *Commodities and Capabilities*. North-Holland Sole distributors for the U.S.A. and Canada, Elsevier Science Publishing Co, 1985.
- A. SEN : *Inequality Reexamined*. Russell Sage Foundation Clarendon Press Oxford University Press, 1992.
- L. S. SHAPLEY : A value for n-person game. *In Contributions to the Theory of Games*, volume 2, pages 307–317. Kuhn, H.W., Tucker, A.W., princeton university press édition, 1953.
- A. F. SHORROCKS : The class of additively decomposable inequality measures. *Econometrica*, 48 (3):613, 1980.
- A. F. SHORROCKS : Inequality decomposition by factor components. *Econometrica*, 50(1):193, 1982.
- A. F. SHORROCKS : Inequality decomposition by population subgroups. *Econometrica*, 52(6):1369, 1984.
- A. F. SHORROCKS : Aggregation issues in inequality measurement. *In Measurement in Economics*, pages 429–452. Physica-Verlag, New York, eichhorn w. édition, 1988.
- A. F. SHORROCKS : Decomposition procedures for distributional analysis : a unified framework based on the shapley value. *Mimeo, University of Essex*, 1999.
- A. F. SHORROCKS : Decomposition procedures for distributional analysis : a unified framework based on the shapley value. *The Journal of Economic Inequality*, 11(1):99–126, 2013.
- J. SILBER : Factor components, population subgroups and the computation of the Gini Index of inequality. *The Review of Economics and Statistics*, 71(1):107, 1989.

- O. STARK, J. E. TAYLOR et S. YITZHAKI : Remittances and inequality. *The Economic Journal*, 96 (383):722, 1986.
- M. TERRAZA, F. SEYTE, S. MUSSARD et M. KOUBI : Évolution des inégalités salariales en France entre 1976 et 2000 : une étude par la décomposition de l'indicateur de Gini. *Économie & prévision*, (169-171):139–169, 2005.
- H. THEIL : *Economics and information theory*. North-Holland, Amsterdam, 1967.
- A. VILLAR : *Lectures on inequality, poverty and welfare*. Numéro 685 in *Lecture Notes in Economics and Mathematical Systems*. 2017.
- E. WINTER : The Shapley value. In *Handbook of game theory with economic applications*, volume 3, pages 2025–2054. North-Holland, 2002.
- S. YITZHAKI et R. I. LERMAN : Income stratification and income inequality. *Review of Income and Wealth*, 37(3):313–329, 1991.
- M-S. YUN : A simple solution to the identification problem in detailed wage decompositions. *Economic Inquiry*, 43(4):766–772, 2005.
- M-S. YUN : Identification problem and detailed Oaxaca decomposition : A general solution and inference. *Journal of Economic and Social Measurement*, (33):27–38, 2008.
- S. ZUBER : Éthique de la population : l'apport des critères de bien-être dépendant du rang. *Revue économique*, 68(1):73, 2017.

Résumé

Les différents chapitres de cette thèse ont eu pour objectif de présenter en différentes étapes la méthodologie nécessaire pour évaluer l'impact d'une politique publique égalisatrice. Le premier chapitre s'est intéressé à l'outil principal utilisé pour apprécier la distribution des revenus, c'est-à-dire aux indices d'inégalité. Le second chapitre a détaillé les différentes approches possibles pour décomposer un indice d'inégalité, que ce soit pour mettre en avant des effets de groupes, de sources et/ou encore d'attributs. Cet état de l'art a ensuite fait place à trois études originales. Ainsi, le troisième chapitre étudie les inégalités de salaire dans la fonction publique française, en mettant en lumière le fait que l'importance du genre dans l'inégalité est considérable mais disparate entre les régions administratives françaises. Le quatrième chapitre démontre théoriquement que l'importance d'un attribut peut se décomposer comme la somme de sa contribution marginale pure moins une somme d'interactions par paire avec les autres attributs. Le cinquième et dernier chapitre expose empiriquement l'effet de ces interactions sur l'impact d'une politique publique qui supprimerait l'écart de rémunération entre les hommes et les femmes.

Mots clés : Indice d'inégalité, Décomposition d'indices par attributs, Contribution du Genre, Interactions, Impact des politiques égalisatrices

Abstract

The different chapters in this thesis aim to present the methodology necessary to evaluate the impact of any equalizing public policies. The first chapter deals with the main tool used to appreciate an income distribution, that are the inequality indices. The second chapter details the different approaches possible to decompose an inequality index, whether to put forward group effects, income source effects and/or attribute effects. This state of the art is followed by three original studies. In that respect, the third chapter considers the income inequality in the French public administration, and it highlights the fact that the importance of gender to inequality is considerable but disparate across the different French administrative regions. The fourth chapter demonstrates theoretically that the importance of an attribute can be decomposed as a sum of its pure marginal contribution minus a sum of pairwise interactions with the others attributes considered. Finally, the fifth chapter exposes empirically the interactions effects on a public policy impact that aims to eliminate the gender wage gap.

Key words : Inequality index, Decomposition by attributes, Gender contribution, Interactions, Equalizing policies impact