

ANNEXE :

MÉTHODE D'ANALYSE : RÉGRESSION GAUSSIENNE ET BAYÉSIENNE

Le présent *Commentaire* a pour but d'estimer la probabilité d'automatisation d'une profession donnée, en fonction des compétences qu'elle recouvre. La base de données O*Net dont je me suis servie contient des renseignements détaillés sur 954 professions, y compris le niveau de chaque compétence prescrite et l'importance de cette dernière dans l'exercice des fonctions, ainsi que le niveau de diverses activités requises.

Le modèle repose sur la base conceptuelle suivante : les êtres humains restent plus compétents que les robots en ce qui concerne certaines compétences ou tâches. Les emplois nécessitant une importante maîtrise de ces compétences ne peuvent pas être entièrement automatisés, car l'être humain s'avère toujours plus performant qu'un robot pour exécuter les tâches connexes. À l'inverse, les emplois ne nécessitant pas ces compétences (ou à un niveau minime) ont davantage de chances d'être automatisables.

Les variables liées aux compétences décrites dans le tableau 1 détaillent un ensemble de variables explicatives, ou *vecteur d'attributs*, pour chaque profession. La variable dépendante (possibilité ou non d'automatiser un emploi) est incomplète et, en grande partie, inconnue. Je sais uniquement quels emplois sont entièrement automatisables et, avec une certitude relative, quels emplois sont actuellement impossibles à automatiser. Le degré d'automatisation possible des emplois partiellement automatisables est inconnu.

Pour gérer cette difficulté, j'ai employé un processus de régression gaussienne et bayésienne, qui s'avère une méthode puissante de classification à noyaux. Elle exploite des « données d'apprentissage » (les renseignements confirmés) pour déterminer par voie de probabilité les valeurs restantes inconnues. De plus, cette méthode ne restreint pas la relation entre les compétences et le potentiel d'automatisation aux relations constantes (l'impact d'une compétence donnée sur les chances

d'automatisation n'est pas constant à tous les niveaux d'automatisation).

Cette méthode d'analyse comporte plusieurs limites importantes. Les résultats doivent être interprétés comme la probabilité qu'une profession puisse être automatisée en théorie, ou comme le pourcentage approximatif de possibilité d'automatisation d'une profession donnée. Cela ne tient pas compte de l'accroissement de la demande résultant des améliorations de la productivité ni de la création de nouvelles professions en réponse aux changements technologiques. Malgré ces limites, les analyses de l'automatisation du marché du travail canadien ont employé diverses variantes de ce type de régression, ce qui permet de comparer l'analyse en cours aux résultats précédents.

Un processus gaussien est défini comme un jeu de variables aléatoires, à savoir tout nombre fini de variables aléatoires ayant une répartition gaussienne commune (Rasmussen et Williams, 2006)²⁷. Les processus gaussiens sont largement utilisés dans de nombreuses variantes d'estimation non paramétrique (Seeger, 2002). En outre, cette méthode a été appliquée pour estimer la vulnérabilité à l'automatisation des professions dans des études antérieures (Frey et Osborne, 2013; Oschinski et Wyonch, 2017).

Il est nécessaire de définir certaines professions comme entièrement automatisables ($y_i = 1$), non automatisables ($y_i = 0$) ou partiellement automatisables ($y_i = 0,5$) pour « former » le modèle. Pour élaborer une classification, j'utilise les résultats de classification en sortie des précédentes études sur l'automatisation des professions : Autor et Dorn (2013), Josten et Lordan (2019), Frey et Osborne (2013), et Oschinski et Wyonch (2017). À l'exception de celle d'Autor et Dorn (2013), qui emploie une classification binaire, ces études répartissent les professions en trois catégories. En sortie, de nombreuses professions ont obtenu une classification similaire, équivalente à celle des travaux de recherche antérieurs. Les professions jugées « à risque moyen » dans les études employant une classification ternaire ont obtenu une classification similaire dans l'ensemble d'apprentissage, en ignorant dans ce cas la classification binaire d'Autor

27 Pour en savoir plus sur l'aspect théorique et les applications des fonctions gaussiennes dans les domaines de l'apprentissage automatique et de l'analyse statistique, voir Rasmussen et Williams (2006).

et Dorn (2013). Lorsque les études précédentes ne donnaient pas de classification consensuelle, aucune catégorie n'était indiquée dans les données d'apprentissage. Cela permet de combiner efficacement les renseignements sur l'automatisation des professions qui se sont avérés cohérents avec les différentes méthodes d'analyse.

L'ensemble de données d'apprentissage est défini par $\mathcal{D} = (X, \mathbf{y})$, où X est la matrice des vecteurs d'attributs et $\mathbf{y}$ donne le libellé associé à la catégorie. Chaque élément x_{ij} de X représente le niveau de chaque compétence (j) requise, pondéré par l'importance de ladite compétence dans la profession désignée (i). Je pars du principe que $\mathcal{D} = (X, \mathbf{y})$, $\mathbf{x}_i \in \mathbb{R}^9$, $y_i \in \{0, 1\}$ est un échantillon indépendant bruité, réparti à l'identique, de la fonction latente $f: \mathbf{x} \rightarrow R$ où $\mathbf{w} = P(\mathbf{y} | f)$ indique la répartition du bruit.

À la lumière des données d'observation, $\mathcal{D} = \{(\mathbf{x}_i, y_i) \mid i = 1, 2, \dots, n\}$, nous souhaitons formuler des prédictions ($\mathbf{y}^*$) pour les nouveaux attributs d'entrée X^* qui ne figurent pas dans l'ensemble d'apprentissage $\mathcal{D}$. Pour estimer $\mathbf{y}^*$, j'emploie un processus gaussien a priori de moyenne nulle ($\mathbf{w} \sim N(0, \Sigma)$) et une méthode bayésienne générative. Cette fonction donne une pondération a priori de la probabilité pour chaque fonction possible décrivant la relation spécifiée par $\mathcal{D}$, où les plus hautes probabilités sont attribuées aux fonctions que je considère les plus vraisemblables. La probabilité d'une fonction est déterminée par sa proximité relative avec les points de données d'apprentissage et par la spécification a priori qui fixe les propriétés des fonctions prises en compte pour inférence.

Le modèle est ensuite calculé comme suit :

- (1) Introduire $\phi(\mathbf{x})$ qui cartographie le vecteur d'entrée $\mathbf{x}$ dans un espace d'attributs à n dimensions. Ensuite, $\Phi(X)$ devient l'agrégation des colonnes $\phi(\mathbf{x})$. Le modèle est défini par :

$$f(\mathbf{x}) = \phi(\mathbf{x})^T \mathbf{w}$$

- (2) Créer les conditions de répartition a priori d'après les données $\mathcal{D}$, pour obtenir une répartition a posteriori :

$$\text{posterior} = \frac{\text{likelihood} \times \text{prior}}{\text{marginal likelihood}}, \quad p(\mathbf{w} | \mathbf{y}, \Phi(X)) = \frac{p(\mathbf{y} | \Phi(X), \mathbf{w}) p(\mathbf{w})}{p(\mathbf{y} | \Phi(X))}$$

$$p(\mathbf{w} | \mathbf{y}, \Phi(X)) \sim N\left(\bar{\mathbf{w}} = \frac{1}{\sigma_n^2} A^{-1} \Phi(X) \mathbf{y}, A^{-1}\right)$$

posterior	répartition a posteriori
likelihood x prior	probabilité x répartition a priori
marginal likelihood	probabilité marginale

où $A = \sigma n^{-2} \Phi(X) \Phi(X)^T + \Sigma^{-1}$

- (3) La répartition marginale a posteriori est la répartition prédictive.

$$f_* | \mathbf{x}_*, \mathcal{D} = \int p(f_* | \phi_*, \mathbf{w}) p(\mathbf{w} | \Phi, \mathbf{y}) d\mathbf{w} \sim N\left(\frac{1}{\sigma_n^2} \phi_*^T A^{-1} \Phi \mathbf{y}, \phi_*^T A^{-1} \phi_*\right)$$

où $\Phi = \Phi(X)$ et $\phi_* = \phi(\mathbf{x}_*)$

Autre formule possible :

où $K = \Phi^T \Sigma \Phi$

$$f_* | \mathbf{x}_*, \mathcal{D} \sim N(\phi_*^T \Sigma \Phi (K + \sigma_n^2 I)^{-1} \mathbf{y}, \phi_*^T \Sigma \phi_* - \phi_*^T \Sigma \Phi (K + \sigma_n^2 I)^{-1} \Phi^T \Sigma \phi_*) \quad (5)$$

À noter : dans l'équation (5), l'espace des attributs prend toujours une forme imposant des entrées dans les matrices de type $\phi(\mathbf{x})^T \Sigma \phi(\mathbf{x}')$, où $\mathbf{x}$ et $\mathbf{x}'$ figurent soit dans l'ensemble d'apprentissage, soit dans l'ensemble d'essai.

Disons que $k(\mathbf{x}, \mathbf{x}') = \phi(\mathbf{x})^T \Sigma \phi(\mathbf{x}')$ est le noyau ou la fonction de covariance. La spécification d'une fonction de covariance implique une répartition dans différentes fonctions. Un processus gaussien est entièrement défini par sa fonction moyenne et sa fonction de covariance.

$$f(\mathbf{x}) \sim GP(m(\mathbf{x}), k(\mathbf{x}, \mathbf{x}')) \quad (6)$$

$$\mathbf{y} = f(\mathbf{x}) + \varepsilon \quad (7)$$

Pour notre modèle $f(\mathbf{x}) = \phi(\mathbf{x})^T \mathbf{w}$ avec répartition a priori $\mathbf{w} \sim N(\mathbf{0}, \Sigma)$, j'ai une moyenne et une covariance

$$E[f(\mathbf{x})] = 0 \quad (8)$$

$$E[f(\mathbf{x})f(\mathbf{x}')] = (\mathbf{x})^T \Sigma \phi(\mathbf{x}') = k(\mathbf{x}, \mathbf{x}') \quad (9)$$

À noter : la covariance entre les données de sortie est fonction des données d'entrée. Je définis le noyau comme la fonction de covariance exponentielle au carré :

$$k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{1}{2} \|\mathbf{x} - \mathbf{x}'\|^2\right) \quad (10)$$

Partant d'un bruit gaussien additif réparti à l'identique par rapport à la variance, la répartition a priori des observations du bruit devient :

$$\text{cov}(\mathbf{y}) = K(X, X) + \sigma_n^2 I. \quad (11)$$

Le paramètre σ_n^2 libre est appelé « hyperparamètre » afin de souligner qu'il s'agit d'un paramètre d'un modèle non paramétrique. Les paramètres (facteurs de pondération) du modèle paramétrique sous-jacent ont été intégrés à l'extérieur.

Pour calculer ce modèle, j'emploie le paquet de « kernlab » en R (Karatzoglou, Smola et Hornik, 2019). J'utilise plus spécifiquement la fonction « gausspr », une covariance quadratique exponentielle (noyau de fonction à base radiale = « rbfdot ») et une sélection d'hyperparamètres optimisée. La précision du modèle est estimée par une erreur de validation croisée décuplée, qui fournit une estimation d'erreur de 0,09 sur les prédictions de probabilité.

Sélection des attributs et analyse de la sensibilité

Les attributs sélectionnés comme obstacles à l'automatisation dans l'analyse principale sont précisés dans le tableau 1. Il existe, cependant, des facteurs supplémentaires susceptibles d'influer sur la probabilité d'automatisation d'une profession (tableau A1). De la même façon, les résultats seront affectés par la classification des professions donnée dans l'ensemble d'apprentissage.

Tableau A1 : Attributs faisant obstacle à l'automatisation

Perception sociale	Capacité à avoir conscience des réactions des autres et à comprendre pourquoi ils réagissent ainsi.
Originalité	Capacité à trouver des idées inhabituelles ou astucieuses concernant une situation ou un sujet donné, ou à développer des façons créatives de résoudre un problème.
Aide et assistance à d'autres personnes	Capacité à fournir une aide personnelle, des soins médicaux, un soutien émotionnel ou d'autres soins personnels à d'autres personnes, comme des collègues, des clients ou des patients.
Philosophie	Connaissance de différents systèmes philosophiques et de différentes religions, y compris leurs principes de base, leurs valeurs, leur éthique, leurs façons de penser, leurs coutumes, leurs pratiques et leurs répercussions sur la culture humaine.
Initiative	Volonté d'accepter des responsabilités et des défis.
Leadership	Volonté de diriger, de prendre le contrôle et de donner des opinions et des instructions.
Innovation	Créativité et nouvelles façons de penser permettant d'élaborer de nouvelles idées et réponses à des problèmes liés au travail.
Adaptabilité et flexibilité	Ouverture au changement (positif ou négatif) et à la variété dans le lieu de travail.
Autonomie	Développement de sa propre façon de faire les choses et capacité à s'auto-diriger avec peu de supervision, voire aucune, ainsi qu'à compter sur soi-même pour accomplir ses tâches.
Résolution de problèmes complexes	Identification de problèmes complexes et analyse des renseignements connexes pour concevoir et évaluer les options à disposition, et mettre en œuvre des solutions.
Résolution de conflits et négociation	Traitement de plaintes, résolution de litiges, de griefs et de conflits, ou autres procédures de négociation.
Réparation ou entretien de matériel électronique ou mécanique	Maintenance, réparation, étalonnage, réglage, ajustement ou mise à l'essai de machines, d'appareils, d'équipements et de pièces mobiles.
Prise de parole publique ou interprétation du sens à l'attention d'autrui	Prise de parole publique fréquente, interprétation ou explication du sens et des moyens d'utiliser les renseignements à disposition.
Définition d'objectifs et de stratégies	Établissement d'objectifs à long terme et élaboration des stratégies et actions nécessaires pour les concrétiser.

Source : base de données O*NET.

Tableau A2 : Modèles spécifiant différents attributs faisant obstacle à l'automatisation

	Analyse principale	Sélection élargie d'attributs	Activités professionnelles uniquement
Attributs	Perception sociale, originalité, aide et assistance à d'autres personnes, connaissances philosophiques, initiative, leadership, innovation, autonomie, adaptabilité et flexibilité	Perception sociale, originalité, aide et assistance à d'autres personnes, initiative, leadership, innovation, autonomie, adaptabilité et flexibilité, résolution de problèmes complexes, résolution de conflits et négociation, réparation et entretien de matériel, prise de parole publique et définition d'objectifs et de stratégies	Aide et assistance à d'autres personnes, résolution de problèmes complexes, résolution de conflits et négociation, réparation et entretien de matériel, et définition d'objectifs et de stratégies
Moyenne du risque d'automatisation des professions concernées (pourcentage non pondéré)	48,4	45,7	46,1
Corrélation de la probabilité d'automatisation estimée (sortie du modèle)			
Analyse principale	1	0,954	0,889
Sélection élargie d'attributs	0,954	1	0,921
Activités professionnelles uniquement	0,889	0,921	1

Source: Author's calculations.

Tableau A3 : Modèles spécifiant différentes classifications entrantes – Corrélation et moyenne du risque d'automatisation

	Oschinski et Wyonch	Frey et Osborne	Métaclassi- fication	Autor et Dorn	Josten et Lordan
Oschinski et Wyonch	1	0,93659	0,943999	0,636234	0,822301
Frey et Osborne		1	0,93865	0,642068	0,803695
Métaclassification			1	0,828096	0,941843
Autor et Dorn				1	0,861623
Josten et Lordan					1
Moyenne non pondérée du risque d'automatisation	0,56	0,54	0,48	0,40	0,41

Source : calculs de l'auteur.

Tableau A4 : Évolution de l'emploi par catégorie de risque (2020 à 2028)

Taux d'adoption technologique (pourcentage annuel)	2,1	1,7	4,3	7,5
Nombre d'employés				
Risque faible	490 000	522 000	254 000	-58 000
Risque moyen	-262 000	-176 000	-878 000	-1 611 000
Risque élevé	-322 000	-260 000	-753 000	-1 236 000
Emplois nets	-93 000	87 000	-1 377 000	-2 906 000
Pourcentage d'évolution				
Risque faible	3,3	3,4	1,8	-0,1
Risque moyen	-1,4	-0,9	-5,2	-9,7
Risque élevé	-1,9	-1,5	-4,5	-7,5

Source : calculs de l'auteure.

Les attributs cités comme obstacles à l'automatisation sont ceux identifiés par les études similaires incluses dans cette analyse. Une sélection de ces attributs est incluse dans l'analyse principale. Une autre analyse utilisant une sélection étendue et une incluant les activités professionnelles uniquement (à l'image d'Autor et Dorn, 2013) montrent que l'ajout ou la suppression d'attributs individuels a peu d'incidence sur les résultats agrégés (tableau A2). Cela a toutefois un effet marginal sur la classification de professions particulières dans les catégories à risque d'automatisation faible, moyen ou élevé.

Les compétences ou attributs permettant aux êtres humains de travailler avec la technologie, mais également susceptibles d'être pris en charge par des ordinateurs ou des machines, comme l'analyse de données et la reconnaissance de formes, ont un impact ambigu sur la possibilité d'automatisation. Concernant les attributs ou compétences faisant obstacle à l'automatisation, les machines/ordinateurs ne peuvent pas, par définition, être plus performants que les êtres humains. Les compétences analytiques ou numériques sont susceptibles de prendre une ampleur croissante dans de nombreuses, voire dans toutes les professions, à mesure que la conception et l'adoption de nouvelles technologies se poursuivent. Un rapport récent de l'Organisation des Nations Unies pour le développement industriel résume la relation entre les compétences humaines et les technologies d'automatisation :

« Les changements technologiques tendent à favoriser les compétences complémentaires avec les nouvelles technologies (Acemoglu, 2002; Rodrik, 2018). Même si le débat reste ouvert concernant l'ensemble de compétences requis pour exécuter les technologies de production numérique avancée, lesdites compétences devraient être orientées vers trois grandes catégories : compétences analytiques, technologiques et générales (Kupfer et coll. 2019). »

Concernant les compétences analytiques et technologiques, les êtres humains peuvent être complémentaires des machines, mais ils peuvent aussi être en concurrence avec elles. Ainsi, tant qu'ils sont complémentaires de la technologie, le fait de disposer de compétences analytiques avancées n'est pas nécessairement synonyme d'une impossibilité d'automatisation. Les personnes ayant de solides compétences analytiques et technologiques pourraient avoir une plus grande capacité d'adaptation à l'évolution des pratiques professionnelles, mais il est possible d'automatiser en grande partie de nombreuses tâches à composante analytique. L'amélioration de ces compétences dans l'ensemble de la population, indépendamment de l'âge, du revenu ou de la profession, réduirait probablement

les répercussions négatives sur l'emploi de l'adoption technologique et de son expansion. De plus, la disponibilité de travailleurs hautement qualifiés permet aux entreprises d'adopter plus facilement des technologies d'accroissement de la productivité qui augmentent le rendement sans modifier la demande de main-d'œuvre.

La présente analyse est axée sur le risque d'automatisation des professions, et non sur les gains d'efficacité que pourrait offrir l'adoption technologique ou sur sa complémentarité avec le travail humain. La méthode employée pour calculer la probabilité d'automatisation exige que les compétences et attributs inclus dans le modèle soient réservés à l'être humain, ou relèvent au moins de domaines dans lesquels les êtres humains ont des chances de rester plus performants que les machines pendant encore un certain temps.

Pour déterminer si la classification d'apprentissage entrante influe de manière significative sur les résultats de l'analyse, j'ai également estimé la probabilité qu'une profession soit automatisable d'après les classifications d'Autor et Dorn (2013), de Lordan et Jorden (2019), de Frey et Osborne (2013), et d'Oschinski et Wyonch (2017) prises indépendamment les unes des autres. De la même façon, pour déterminer si la sélection d'attributs influe de manière significative sur les résultats, j'ai fait des estimations du modèle avec différentes combinaisons d'attributs et une classification entrante similaire.

Collectivement, les résultats sont plutôt cohérents entre les différentes classifications entrantes et les sélections d'attributs explicatives. La variation de la sélection d'attributs faisant obstacle à l'automatisation donne des moyennes de probabilité d'automatisation d'environ 46 p. 100 à 48 p. 100 (tableau A2). En outre, la probabilité d'automatisation estimée est hautement corrélée dans les variantes mises à l'essai ici. De la même façon, les estimations de probabilité calculées à l'aide de classifications entrantes différentes ont donné des résultats en grande partie similaires (tableau A3). Il existe toutefois une variabilité significative des estimations propres à chaque profession. Sur les sept sélections modélisées figurant dans les tableaux A2 et A3, l'écart moyen de la probabilité d'automatisation estimée pour chaque profession s'élève à 30 p. 100 (écart maximum de 0,82; écart minimum de 0,06).

Les estimations associées à la classification combinée (libellée « métaclassification » dans le tableau A3) sont hautement corrélées à celles calculées pour chaque classification prise individuellement. La probabilité d'automatisation moyenne tombe au centre de ces estimations. De la même façon, les estimations de la probabilité d'automatisation propres à chaque profession tombent toutes entre les deux extrêmes. Face à l'incertitude qui persiste concernant la classification reflétant le mieux le potentiel d'automatisation dans le monde réel et vu le résultat plutôt commode offert par la classification combinée (dont les estimations tombent entre les estimations extrêmes des autres modèles), la « métaclassification » est employée dans l'analyse principale des effets potentiels de l'automatisation sur le marché du travail canadien.

MISE EN CORRESPONDANCE DES DONNÉES O*NET ET DES DONNÉES DU MARCHÉ DU TRAVAIL CANADIEN

L'analyse statistique estime la probabilité qu'une profession soit automatisée en fonction des attributs sélectionnés et des classifications du vecteur d'apprentissage. Les renseignements relatifs aux compétences et aux attributs de chaque profession proviennent de la base de données O*NET, parrainée par le ministère du Travail des États-Unis. Les données O*NET se fondent sur les notations données par les employeurs, les travailleurs et les analystes de professions. Les données d'O*NET sont fondées sur les codes de la Classification type des professions de 2010 (CTP2010). Pour appliquer l'information d'O*NET aux données canadiennes sur les professions, j'utilise une concordance établie par Statistique Canada et je suis la méthodologie de Frenette et Frank (2017) :

« Une concordance a été établie entre les codes à six chiffres de la CTP2010 (tirés de la base de données O*NET) et les codes à quatre chiffres de la Classification nationale des professions de 2011 (CNP2011). Cette concordance a été établie sur la base de la similarité des descriptions de professions.

Certains codes à six chiffres de la CTP2010 comportaient des sous-professions dérivées (nouvelles ou émergentes).

Puisqu'il n'existait pas d'estimations de la taille de la population pour ces sous-professions, on a présumé que la taille était la même pour toutes celles ayant un code à six chiffres de la CTP2010. Une fois les professions dérivées agrégées, le nombre de paires appariées CTP2010–CNP2011 s'établissait à 1 058. Ce nombre comprenait 495 codes uniques à quatre chiffres de la CNP2011.

Sur ces 495 codes uniques de la CNP2011, 336 ont été appariés à un code à six chiffres de la CTP2010. Pour les 159 autres codes de la CNP2011, il a fallu utiliser des codes de la CTP2010 d'un niveau supérieur, car plus d'un code à six chiffres de la CTP2010 leur correspondait (les données de l'ACS ont été utilisées à cette étape). Parmi ces 159 codes de la CNP2011, 119 ont été mis en correspondance avec des codes à cinq chiffres de la CTP2010, 38 ont été mis en correspondance avec des codes à quatre chiffres de la CTP2010, et les deux codes restants ont été mis en correspondance avec des codes à trois chiffres de la CTP2010.

À ce stade, les 495 codes uniques de la CNP2011 avaient été associés à l'information sur les compétences professionnelles exigées provenant de la base de données O*NET. Les données sur les compétences professionnelles exigées ont ensuite été couplées aux codes de la CNP2011. »