

Article

Système intelligent explicable pour l'analyse des déterminants de l'insertion professionnelle et la détection des facteurs discriminants en République démocratique du Congo

Zawadi Sirisombola Corinne^{1,2,*}, Kitondua Lubanzadio Richard¹, Claude Takenga³, Mayala Lemba Francis¹

¹ Université Pédagogique Nationale, Kinshasa, République démocratique du Congo

² Université de l'Assomption au Congo (UAC), Département d'Informatique de Gestion, Butembo, République Démocratique du Congo

³ Université Officielle de Ruwenzori (UOR), Butembo, République démocratique du Congo

*Correspondant : zawadisirisombola@uaonline.edu.cd

Abstract: Cette étude développe un système intelligent explicable pour analyser les déterminants de l'insertion professionnelle en République démocratique du Congo et repérer les facteurs associés aux disparités d'accès à l'emploi. À partir de 1961 observations de l'Office National de l'Emploi au Nord-Kivu, quatre modèles ont été comparés : régression logistique, Random Forest, Gradient Boosting et XGBoost. Les prédictions ont été interprétées avec SHAP et complétées par un score discriminatoire hybride ainsi qu'une analyse d'ablation. Les résultats indiquent que le type de contrat, la branche d'activité et le type d'entreprise dominent les facteurs individuels. La plateforme Streamlit proposée constitue un outil d'aide à la décision pour les politiques publiques d'emploi.

Mot clés: apprentissage automatique ; intelligence artificielle explicable ; SHAP ; analyse d'ablation ; insertion professionnelle.

Citation : Zawadi Sirisombola, C., Kitondua Lubanzadio, R., Takenga, C., & Mayala Lemba, F. Système intelligent explicable pour l'analyse des déterminants de l'insertion professionnelle et la détection des facteurs discriminants en République démocratique du Congo. *Etincelle*. 2026, vol. 26, no. 2. <https://doi.org/10.61532/rime262124>

Reçu: 05 juin 2026

Accepté: 27 juin 2026

Publié: 04 juillet 2026

Note de l'éditeur: Ishango-uac reste neutre en ce qui concerne les revendications juridictionnelles dans les cartes géographiques publiées et les affiliations institutionnelles des auteurs.



Copyright: © 2026 par les auteurs. Soumis pour une publication en libre accès selon les termes et conditions de la licence Creative Commons Attribution (CC BY) (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

L'insertion professionnelle est l'un des défis socio-économiques les plus sensibles pour les pays en développement. En République démocratique du Congo (RDC), la croissance démographique et l'augmentation du nombre de diplômés de l'enseignement supérieur ne s'accompagnent pas toujours d'une création suffisante d'emplois formels. Cette situation entretient le chômage, le sous-emploi et la précarité professionnelle, en particulier chez les jeunes diplômés.

Dans le même temps, les progrès de l'intelligence artificielle et de l'apprentissage automatique ont profondément renouvelé les méthodes d'analyse des phénomènes socio-économiques. Ces approches permettent d'exploiter des volumes importants de données et de repérer des relations complexes que les méthodes statistiques classiques peuvent difficilement mettre en évidence (Bishop, 2006; Mitchell, 1997).

Cependant, les systèmes prédictifs ne sont pas neutres. Ils peuvent reproduire ou amplifier les inégalités présentes dans les données historiques utilisées pour l'apprentissage, surtout lorsque certaines variables jouent le rôle de proxys sociaux ou institutionnels (Barocas & Selbst, 2016; Chouldechova, 2017; Mehrabi et al., 2021).

Face à ce risque, l'intelligence artificielle explicable (Explainable Artificial Intelligence, XAI) occupe une place centrale dans l'évaluation des modèles. Elle vise à rendre les prédictions compréhensibles pour les chercheurs, les décideurs et les utilisateurs

finiaux (Doshi-Velez & Kim, 2017; Molnar, 2022). Parmi les méthodes les plus utilisées, SHAP permet d'estimer la contribution de chaque variable à une prédiction donnée (Lundberg & Lee, 2017).

La présente étude répond à la question suivante : quels facteurs influencent le plus l'insertion professionnelle en RDC et dans quelle mesure l'IA explicable permet-elle d'identifier les variables susceptibles de refléter des inégalités d'accès à l'emploi ?

L'objectif général est de développer un système intelligent explicable capable de :

Prédire l'insertion professionnelle ;

Identifier les facteurs les plus influents ;

Expliquer les décisions algorithmiques ;

Détecter les variables susceptibles de refléter des disparités structurelles.

L'originalité de cette recherche réside dans l'intégration de trois dimensions complémentaires : la prédiction par apprentissage automatique, l'explication des prédictions par SHAP et l'évaluation des disparités au moyen d'un score discriminatoire hybride.

Cette recherche apporte des contributions méthodologiques, applicatives et scientifiques.

Sur le plan méthodologique, elle propose un score discriminatoire hybride combinant l'importance moyenne des variables issue des valeurs SHAP et les écarts statistiques observés entre groupes. Cette approche permet d'évaluer, simultanément, l'influence d'une variable sur la prédiction et son association potentielle avec des disparités structurelles dans l'accès à l'emploi.

Sur le plan applicatif, l'étude débouche sur une plateforme décisionnelle interactive développée avec Streamlit. Celle-ci intègre les modèles d'apprentissage automatique, les analyses SHAP, le calcul du score discriminatoire hybride et un module de recommandation destiné aux acteurs de l'emploi.

Sur le plan scientifique, la recherche propose une analyse explicable des déterminants de l'insertion professionnelle et la complète par une analyse d'ablation permettant de vérifier expérimentalement l'influence réelle des variables dominantes.

2. Revue de littérature

2.1. Apprentissage automatique et analyse du marché du travail

L'apprentissage automatique regroupe des méthodes qui permettent à un système informatique d'apprendre à partir des données et d'améliorer ses prédictions sans être explicitement programmé pour chaque situation (Bishop, 2006; Mitchell, 1997). Dans le domaine de l'emploi, ces méthodes sont utilisées pour analyser les compétences, prédire l'employabilité, appuyer les processus de recrutement et éclairer la gestion des ressources humaines.

Eloundou et al. (2023) montrent que les systèmes d'intelligence artificielle peuvent mettre en évidence des structures complexes dans les données relatives au marché du travail. Dans cette étude, cette perspective est mobilisée pour analyser l'insertion professionnelle à partir de données administratives congolaises.

2.2. Intelligence Artificielle Explicable

L'explicabilité est devenue un enjeu central dans les systèmes intelligents modernes. Les modèles tels que Random Forest, Gradient Boosting ou XGBoost sont souvent performants, mais leur fonctionnement interne reste difficile à comprendre pour les utilisateurs non spécialistes (Molnar, 2022; Rudin, 2019).

SHAP, proposé par Lundberg et Lee (2017), repose sur la théorie des valeurs de Shapley et attribue à chaque variable une contribution à la prédiction. Cette propriété en fait une méthode utile pour interpréter les décisions individuelles et pour classer les variables selon leur influence globale.

Les méthodes XAI renforcent ainsi la transparence, soutiennent l'audit des modèles et facilitent l'appropriation des résultats par les décideurs. Elles ne remplacent toutefois pas l'analyse sociale ou économique : elles doivent être interprétées à la lumière du contexte dans lequel les données ont été produites.

2.3. Biais algorithmiques et discrimination

Barocas et Selbst (2016) montrent que des variables apparemment neutres peuvent agir comme proxys et produire des discriminations indirectes. Cette idée est particulièrement importante dans les systèmes d'aide à la décision appliqués à l'emploi.

Chouldechova (2017) souligne, pour sa part, les tensions possibles entre performance prédictive et équité algorithmique. Une variable peut améliorer fortement la précision d'un modèle tout en renforçant des écarts préexistants entre groupes.

Raghavan et al. (2020) rappellent que les systèmes algorithmiques apprennent principalement les régularités présentes dans les données d'entraînement. Dans le domaine de l'emploi, cela impose une lecture prudente des résultats, car les modèles peuvent refléter des structures historiques plutôt que des relations causales.

Mehrabi et al. (2021) identifient les données d'entraînement, la sélection des variables et les choix de modélisation comme des sources majeures de biais. Cette littérature justifie l'usage conjoint de l'explicabilité et d'indicateurs de disparité dans la présente étude.

2.4. Cadre conceptuel de l'étude

Le cadre conceptuel de cette étude repose sur l'idée que l'insertion professionnelle résulte de l'interaction entre deux familles de facteurs :

Les facteurs individuels (niveau d'étude, expérience professionnelle, genre, âge, durée du chômage) ;

Les facteurs structurels (branche d'activité, type d'entreprise, type de contrat, province).

L'approche développée compare l'influence relative de ces facteurs à travers les modèles d'apprentissage automatique, les valeurs SHAP et un score discriminatoire hybride.

3. Méthodologie

3.1. Source des données

Les données proviennent de l'Office National de l'Emploi (ONEM) et concernent les demandeurs d'emploi enregistrés dans la province du Nord-Kivu.

Après nettoyage et prétraitement, le jeu de données final comprend 1 961 observations et plusieurs variables décrivant les caractéristiques sociodémographiques, académiques et professionnelles des individus.

3.2. Prétraitement des données

Les principales étapes de préparation des données ont été les suivantes : (i) Suppression des valeurs manquantes ; (ii) Traitement des incohérences ; (iii) Encodage des variables catégorielles ; (iv) Normalisation des variables numériques ; (v) Séparation des données en ensembles d'entraînement et de test.

Le pipeline général de traitement est illustré dans la Figure 1.

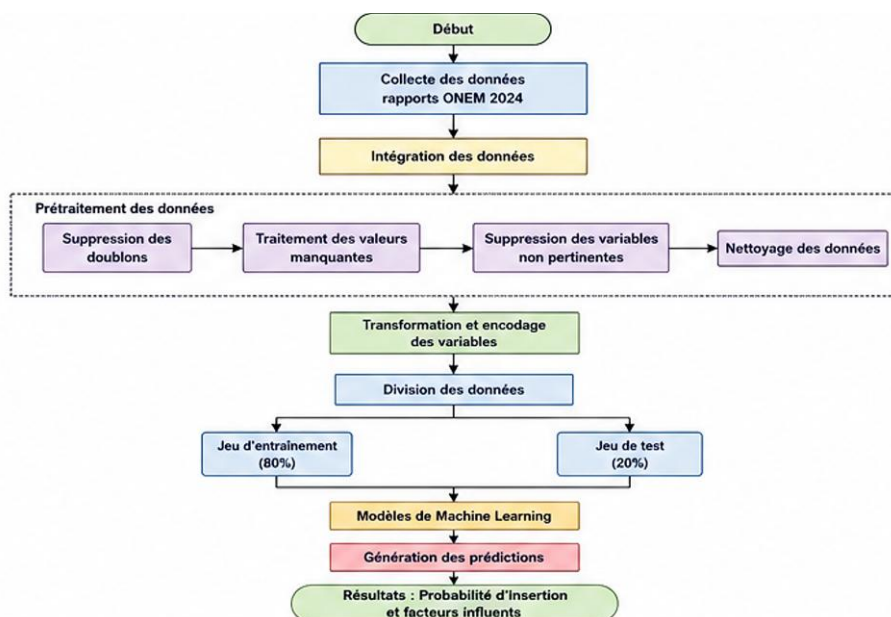


Figure 1. Pipeline global de traitement et d'analyse des données

Ce pipeline résume le cycle de traitement utilisé : préparation des données, entraînement des modèles, interprétation explicable, analyse des disparités et restitution des résultats dans la plateforme décisionnelle.

3.3. Modèles développés

Quatre algorithmes supervisés ont été entraînés et comparés : Régression logistique ; Random Forest ; Gradient Boosting ; XGBoost.

Le processus d'entraînement et d'évaluation est représenté dans la Figure 2.

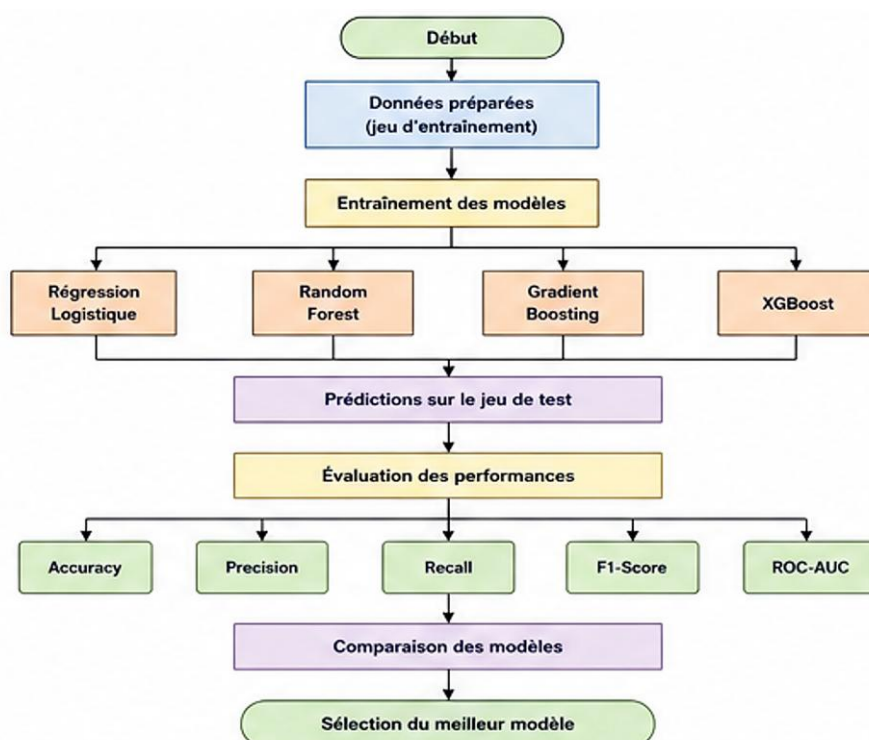


Figure 2. Processus d'entraînement et d'évaluation des modèles

Cette figure illustre le mécanisme d'apprentissage du système intelligent ainsi que le processus d'évaluation des performances prédictives des modèles utilisés dans l'étude.

3.4. Évaluation des performances

Les modèles ont été évalués à l'aide des métriques suivantes : Accuracy ; Precision ; Recall ; F1-score ; ROC-AUC. Ces indicateurs permettent d'évaluer à la fois la justesse globale des prédictions, l'équilibre entre précision et rappel, ainsi que la capacité des modèles à distinguer les classes.

3.5. Validation expérimentale

Pour évaluer les performances des modèles, les données ont été divisées en deux sous-ensembles : 80 % pour l'apprentissage et 20 % pour le test. La partition a été réalisée avec la fonction `train_test_split` de Scikit-learn, avec `test_size = 0,20` et `random_state = 42` afin d'assurer la reproductibilité. Le paramètre `stratify = y` a été utilisé pour préserver la distribution de la variable cible dans les deux sous-ensembles.

Les quatre algorithmes ont ensuite été entraînés et comparés sur l'ensemble de test à l'aide de l'accuracy, de la précision, du rappel, du F1-score et de l'aire sous la courbe ROC (ROC-AUC).

Les principaux hyperparamètres utilisés sont présentés dans le Tableau 1. Les autres paramètres ont été conservés à leurs valeurs par défaut afin de garantir une comparaison homogène entre les modèles.

Tableau 1. Principaux hyperparamètres des modèles développés

Modèle	Hyperparamètres utilisés
Régression Logistique	<code>max_iter = 500</code>
Random Forest	<code>n_estimators = 300, random_state = 42</code>
Gradient Boosting	<code>random_state = 42</code>
XGBoost	<code>eval_metric = "logloss", random_state = 42</code>

Les variables explicatives ont été préparées avant l'apprentissage. Les variables non pertinentes pour la prédiction (ID, Insertion et Insertion_num) ont été retirées. Les variables catégorielles ont ensuite été converties en valeurs numériques au moyen de LabelEncoder afin d'être exploitées par les algorithmes.

Le modèle présentant le meilleur F1-score a été retenu comme modèle de référence et sauvegardé pour son intégration dans la plateforme décisionnelle développée sous Streamlit.

Une analyse préalable de la distribution de la variable cible a permis d'évaluer l'existence d'un déséquilibre entre classes. L'usage de `stratify` lors de la partition a conservé la représentativité des classes dans les ensembles d'apprentissage et de test ; aucune technique supplémentaire de rééchantillonnage ou de pondération des classes n'a donc été appliquée.

3.6. Analyse explicable et score discriminatoire

L'interprétation des prédictions a été réalisée avec SHAP, qui fournit des explications locales et globales des décisions du modèle (Lundberg & Lee, 2017).

Un score discriminatoire hybride a été développé afin de combiner l'importance prédictive des variables et les écarts observés entre groupes :

Score discriminatoire = $\alpha \times \text{Importance_SHAP normalisée} + (1 - \alpha) \times \text{Gap_Insertion normalisé}$

où :

Importance_SHAP représente l'influence moyenne d'une variable sur les prédictions ; Gap_Insertion représente l'écart observé entre groupes pour la variable considérée.

Ce score permet d'identifier les variables qui sont à la fois importantes pour le modèle et associées aux plus fortes disparités observées dans les données.

4. Plateforme développée et fonctionnalités

4.1. Architecture générale de la plateforme

Pour rendre les modèles opérationnels, une plateforme décisionnelle intelligente a été conçue avec Streamlit. Elle regroupe, dans un environnement interactif unique, l'exploration des données, la prédiction de l'insertion professionnelle, l'interprétation des décisions algorithmiques et l'analyse des facteurs associés aux disparités structurelles.

Ce choix d'une architecture légère, modulaire et reproductible s'inscrit dans une logique d'IA adaptée aux environnements à ressources limitées. Des travaux récents ont montré qu'un pipeline d'apprentissage profond compact, combiné à un déploiement embarqué et à une notification en temps réel, peut conserver des performances élevées malgré des contraintes matérielles et opérationnelles (Mpia et al., 2025). Cette orientation soutient, dans la présente étude, le recours à une plateforme sobre, explicable et facilement réutilisable par les services publics de l'emploi.

La plateforme a été développée avec Python 3.11. L'architecture logicielle suit une logique modulaire afin de faciliter la maintenance, l'évolution et la réutilisation du système. Le fonctionnement de la plateforme suit un flux séquentiel : les données sont importées et prétraitées, les modèles produisent les prédictions, les analyses SHAP expliquent les résultats, le score discriminatoire hybride enrichit l'interprétation, puis les informations sont restituées sous forme de tableaux de bord et de recommandations.

Sur le plan expérimental, l'application exécute les prédictions individuelles et les analyses explicables en quelques secondes sur une configuration informatique standard. Les modèles entraînés sont chargés automatiquement au démarrage, ce qui réduit le temps de calcul pendant les interactions avec l'utilisateur. Cette architecture favorise une utilisation fluide et assure la reproductibilité des expérimentations. Le fonctionnement global est illustré à la Figure 3.

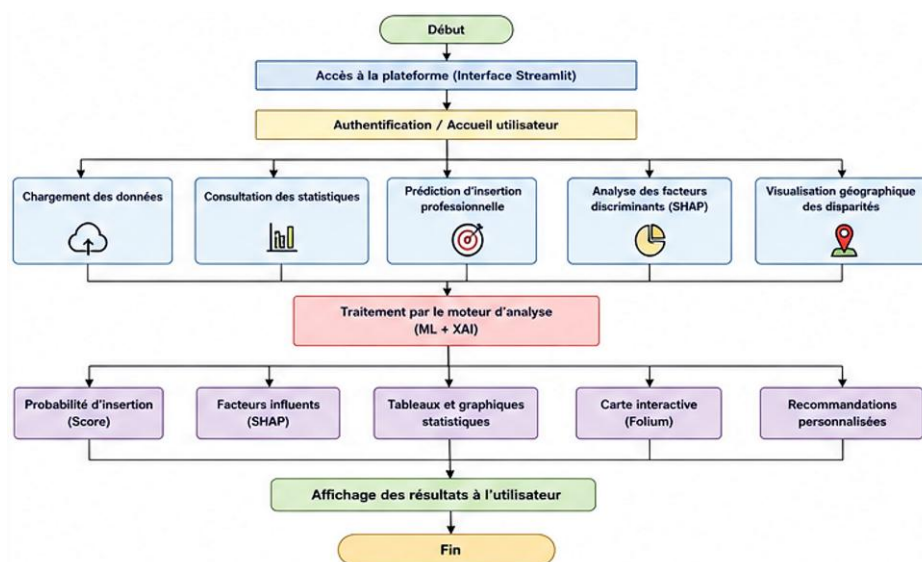


Figure 3. Architecture générale de la plateforme développée

Cette architecture assure la continuité entre l'acquisition des données, la génération des prédictions, leur explication et leur restitution aux utilisateurs.

4.2. Interface d'accueil

L'interface d'accueil fournit une vue synthétique de la base de données utilisée. Elle présente notamment le nombre total d'observations, le taux global d'insertion professionnelle et les principaux indicateurs statistiques.



Figure 4. Interface d'accueil

Cette interface constitue le point d'entrée de la plateforme et permet à l'utilisateur de comprendre la structure générale des données avant d'approfondir l'analyse.

4.3. Interface d'analyse descriptive

L'analyse descriptive permet d'explorer les principales caractéristiques de la population étudiée et d'identifier les premiers écarts entre groupes.

4.3.1. Répartition selon le genre

L'analyse de la variable genre met en évidence la représentation des hommes et des femmes dans les données.

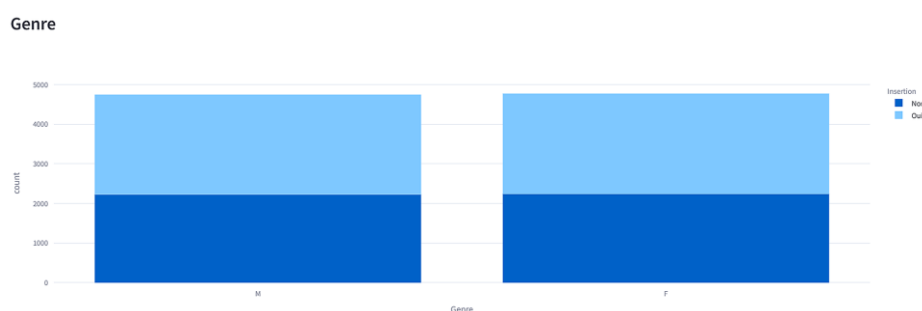


Figure 5. Répartition selon le genre

Les écarts observés selon le genre doivent être interprétés avec prudence : ils ne démontrent pas à eux seuls une discrimination, mais ils peuvent signaler des inégalités d'accès à l'emploi. Ils justifient donc l'usage d'indicateurs complémentaires, notamment le score discriminatoire hybride.

4.3.2. Répartition selon le niveau d'étude

Le niveau d'étude représente un indicateur du capital humain et peut influencer les chances d'accès à un emploi stable.

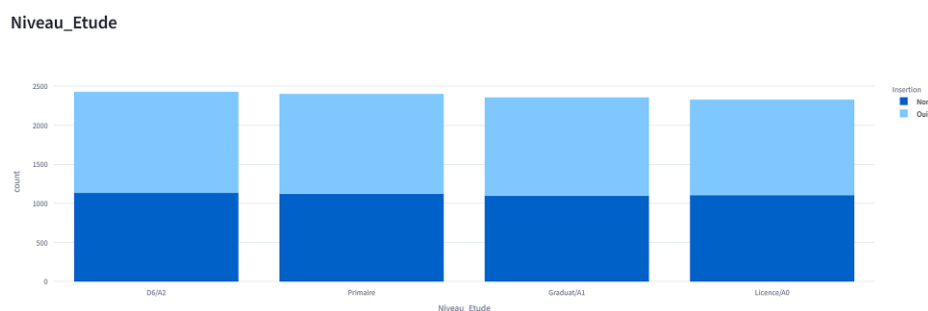


Figure 6. Répartition selon le niveau d'étude

La distribution selon le niveau d'étude suggère une association entre le capital humain et l'accès à l'emploi. Les individus ayant un niveau d'étude plus élevé semblent mieux représentés parmi les personnes insérées, tandis que les niveaux d'instruction plus faibles apparaissent davantage exposés à la non-insertion. Cette interprétation doit toutefois rester prudente, car les résultats du modèle montrent que les variables structurelles du marché du travail jouent un rôle plus déterminant.

4.3.3. Répartition selon la branche d'activité

La branche d'activité reflète la structure sectorielle du marché du travail.

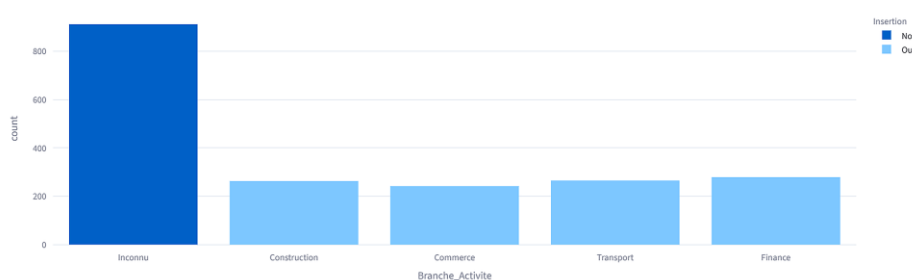


Figure 7. Répartition selon la branche d'activité

La distribution par branche d'activité montre une concentration des opportunités dans certains secteurs. Cette hétérogénéité sectorielle constitue un élément important pour comprendre les mécanismes d'insertion professionnelle.

4.4. Interface des modèles de apprentissage automatique

La plateforme permet de comparer simultanément les performances des différents modèles développés.

Tableau 2. Comparaison des performances des modèles

	Model	Accuracy	Precision	Recall	F1	ROC_AUC
0	RF	1	1	1	1	1
1	GB	1	1	1	1	1
2	LR	1	1	1	1	1
3	XGB	1	1	1	1	1

⚠ Les performances très élevées indiquent une forte relation entre certaines variables et la variable cible. Cela justifie l'usage de SHAP et de l'analyse de biais.

Les résultats indiquent des performances très élevées pour l'ensemble des algorithmes étudiés.

Ces performances doivent néanmoins être interprétées avec prudence, car certaines variables explicatives présentent une forte proximité avec la variable cible. Cette situation limite l'intérêt d'une simple comparaison prédictive et justifie l'orientation de

l’étude vers l’intelligence artificielle explicable : l’enjeu n’est pas seulement de prédire l’insertion professionnelle, mais de comprendre les facteurs qui la structurent.

4.5. Interface de prédiction individuelle

La plateforme permet d’estimer la probabilité d’insertion professionnelle pour un profil individuel.

L’utilisateur renseigne les caractéristiques du candidat — genre, niveau d’étude, province, expérience, type d’entreprise, branche d’activité, type de contrat et durée du chômage — puis le système génère une probabilité d’insertion à partir des modèles entraînés.

Tableau 3. Interface de prédiction individuelle

Modèle	Probabilité	Prédiction
0 RandomForest	0	0
1 GradientBoosting	0.00002	0
2 LogisticRegression	0.0029	0
3 XGBoost	0.0003	0

Cette fonctionnalité peut servir d’outil d’aide à la décision pour les institutions d’orientation professionnelle et les services publics de l’emploi. Elle ne doit toutefois pas être utilisée comme mécanisme automatique de décision individuelle.

4.6. Module d’Intelligence Artificielle Explicable

L’une des contributions centrales de la recherche est l’intégration d’un module d’explication fondé sur SHAP.

Cas d’une probabilité d’insertion faible

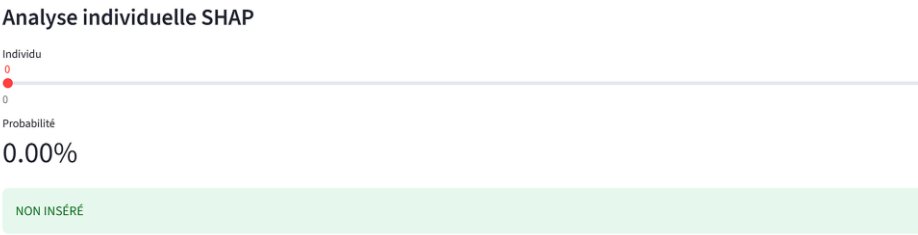


Figure 8. Analyse SHAP individuelle (0 %)

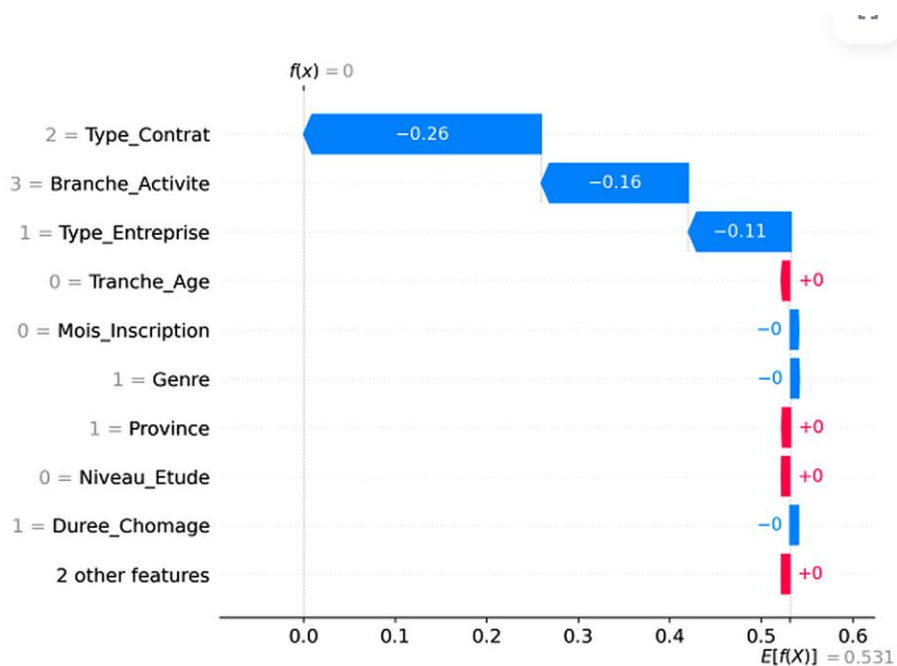


Figure 9. Interface XAI dans le cas où la probabilité est de 0%

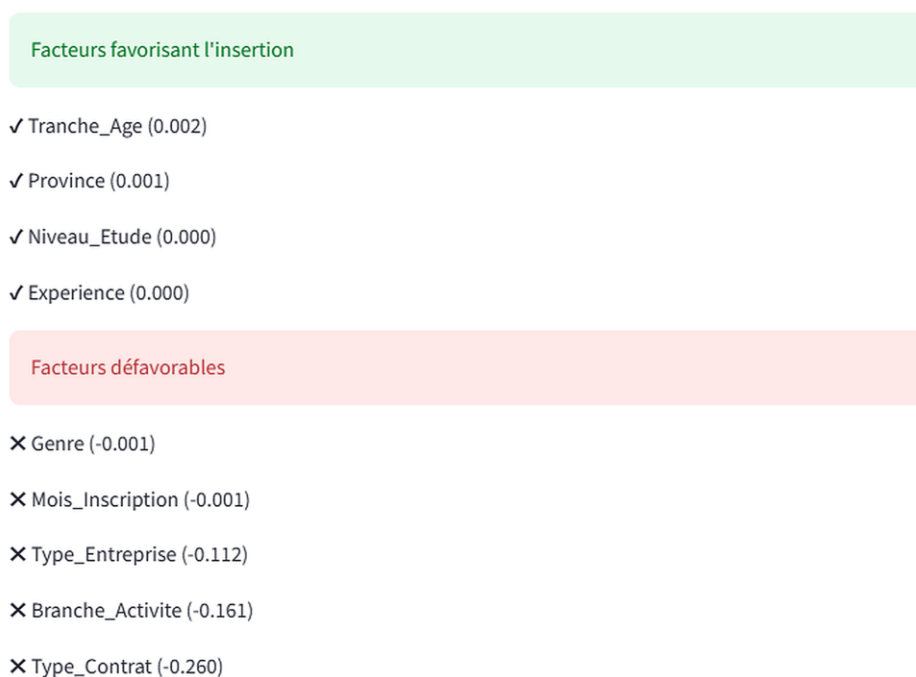


Figure 10. Facteurs favorisant et défavorisant l'insertion

Cas d'une probabilité d'insertion forte

Analyse individuelle SHAP

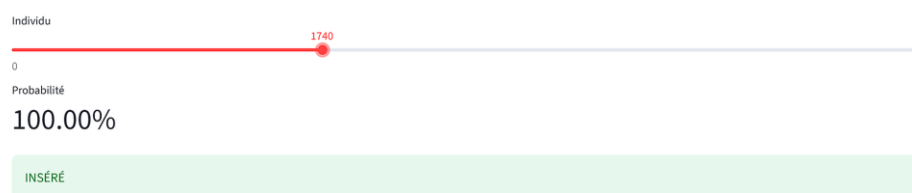


Figure 11. Analyse individuelle SHAP (100%)

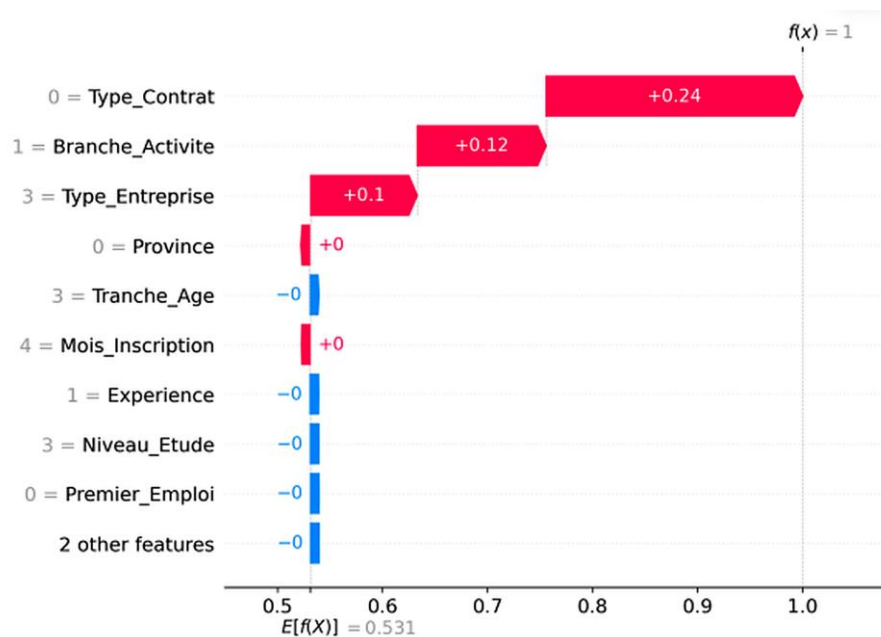


Figure 12. Interface XAI dans le cas où la probabilité est de 100%

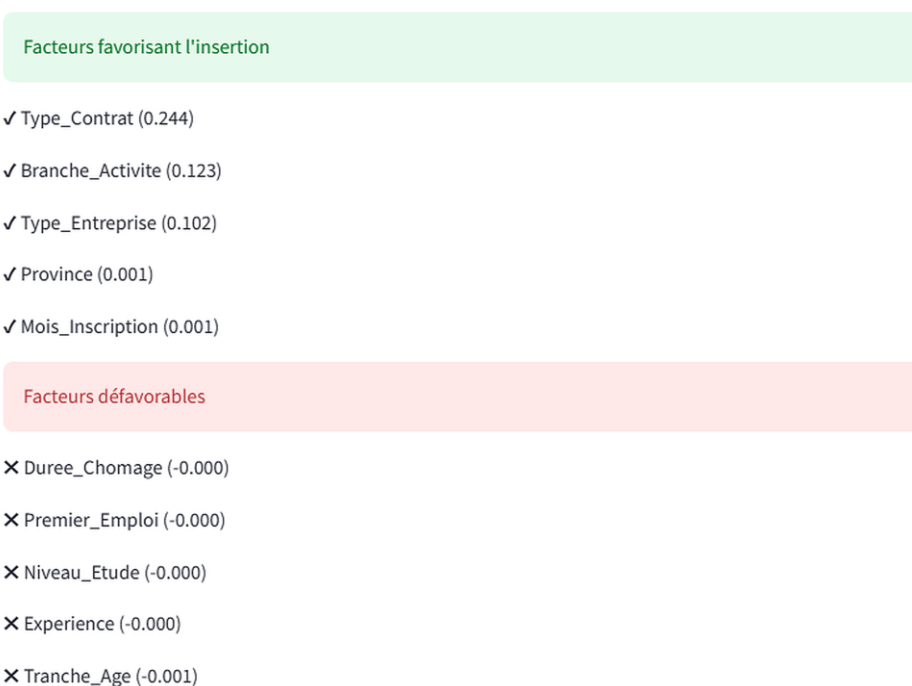


Figure 13. Facteurs favorisant et défavorisant l'insertion

Ces visualisations expliquent les facteurs qui contribuent à chaque prédiction et permettent de distinguer les variables favorables et défavorables selon le profil analysé.

4.7. Analyse SHAP globale

L'analyse globale vise à identifier les variables ayant la plus forte influence moyenne sur les prédictions.

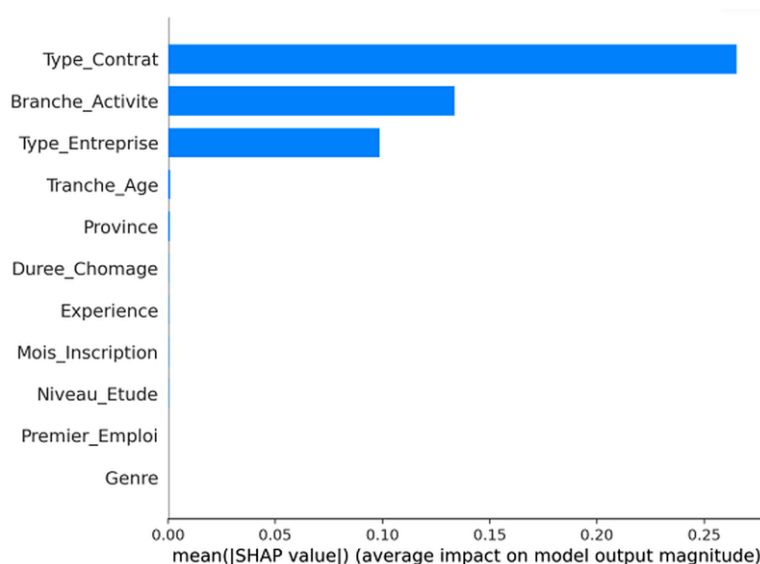


Figure 14. Analyse SHAP globale

Les résultats indiquent que :

Type_Contrat est la variable la plus influente ;

Branche_Activite occupe la deuxième position ;

Type_Entreprise apparaît en troisième position.

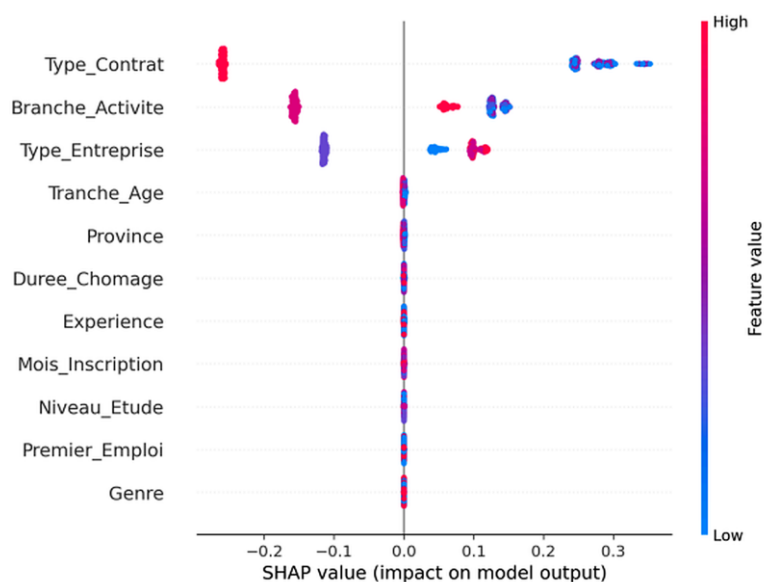


Figure 15. SHAP Summary Plot

Cette visualisation confirme le poids dominant des variables structurelles du marché du travail.

4.8. Interface de score discriminatoire hybride

Afin d'identifier les variables associées aux plus fortes disparités observées, la plateforme intègre le score discriminatoire hybride proposé dans cette étude.

Score discriminatoire hybride (SHAP + inégalités)

	Variable	SHAP_mean	Gap_statistique	Score_discriminatoire
8	Type_Contrat	0.2657	100	26.5712
9	Branche_Activite	0.1324	100	13.2448
10	Type_Entreprise	0.0983	100	9.8317
3	Tranche_Age	0.0011	4.0483	0.0044
2	Province	0.0007	2.0476	0.0014
0	Mois_Inscription	0.0002	4.8061	0.0011
5	Duree_Chomage	0.0003	2.0589	0.0006
6	Experience	0.0003	1.1975	0.0003
4	Niveau_Etude	0.0002	0.8607	0.0002
7	Premier_Emploi	0.0001	1.1975	0.0002

Figure 16. Score discriminatoire hybride

Le score combine l'importance moyenne SHAP et les écarts observés entre groupes.

Cette approche permet de repérer les variables qui peuvent refléter des mécanismes d'inégalités structurelles.

4.9. Interface de recommandation

La plateforme propose enfin des recommandations fondées sur les secteurs et les profils présentant les meilleures probabilités d'insertion.

Secteurs recommandés ⇄

Branche_Activite	Insertion_num
Commerce	1
Construction	1
Finance	1
Transport	1
Inconnu	0

Secteur optimal : Commerce

Figure 17. Interface de recommandation

Cette fonctionnalité traduit les résultats du modèle en informations directement exploitables par les acteurs de l'emploi.

5. Résultats et discussion

5.1. Performances globales des modèles

Les quatre algorithmes étudiés présentent des performances très élevées.

Tableau 4. Résultats comparatifs des modèles

	Model	Accuracy	Precision	Recall	F1	ROC_AUC
0	RF	1	1	1	1	1
1	GB	1	1	1	1	1
2	LR	1	1	1	1	1
3	XGB	1	1	1	1	1

⚠ Les performances très élevées indiquent une forte relation entre certaines variables et la variable cible. Cela justifie l'usage de SHAP et de l'analyse de biais.

Les métriques d'accuracy, de précision, de rappel, de F1-score et de ROC-AUC atteignent des valeurs proches de 100 %.

Ces performances s'expliquent principalement par la présence de variables fortement liées à la variable cible, notamment :

Type_Contrat ;
Branche_Activite ;
Type_Entreprise.

Dans ce contexte, l'intérêt scientifique principal ne réside pas uniquement dans la prédiction, mais dans la compréhension des mécanismes qui expliquent les décisions produites par les modèles.

Cette observation rejoint les recommandations de Mehrabi et al. (2021), selon lesquelles l'évaluation des systèmes intelligents doit dépasser la seule mesure de la précision et intégrer des dimensions d'explicabilité, d'équité et d'audit.

5.2. Résultats de l'analyse explicable

5.2.1. Importance globale des variables

Les résultats SHAP montrent que les variables les plus influentes sont :

Type_Contrat ;
Branche_Activite ;
Type_Entreprise ;
Niveau_Etude ;
Experience ;
Duree_Chomage ;
Genre.

Les facteurs structurels dominent nettement les facteurs individuels.

Cette observation suggère que l'accès à l'emploi dépend principalement des caractéristiques institutionnelles du marché du travail, de la segmentation sectorielle et de la disponibilité des opportunités économiques.

5.2.2. Analyse locale des prédictions

Les analyses SHAP individuelles montrent que les facteurs favorables ou défavorables à l'insertion varient selon les profils.

Les mécanismes d'accès à l'emploi apparaissent ainsi comme le résultat d'interactions entre plusieurs variables plutôt que comme l'effet isolé d'un seul facteur.

Cette capacité à produire des explications locales constitue l'un des principaux apports des approches XAI.

5.3. Analyse du score discriminatoire hybride

Les résultats du score discriminatoire hybride montrent une hiérarchie particulièrement marquée.

Tableau 5. Résultats du score discriminatoire hybride

Score discriminatoire hybride (SHAP + inégalités)

	Variable	SHAP_mean	Gap_statistique	Score_discriminatoire
8	Type_Contrat	0.2657	100	26.5712
9	Branche_Activite	0.1324	100	13.2448
10	Type_Entreprise	0.0983	100	9.8317
3	Tranche_Age	0.0011	4.0483	0.0044
2	Province	0.0007	2.0476	0.0014
0	Mois_Inscription	0.0002	4.8061	0.0011
5	Duree_Chomage	0.0003	2.0589	0.0006
6	Experience	0.0003	1.1975	0.0003
4	Niveau_Etude	0.0002	0.8607	0.0002
7	Premier_Emploi	0.0001	1.1975	0.0002

Les scores les plus élevés concernent Type_Contrat, Branche_Activite et Type_Entreprise. À l'inverse, les variables individuelles présentent des scores plus faibles.

Ces résultats indiquent que les disparités observées dans l'accès à l'emploi semblent principalement liées à la structure du marché du travail.

5.4. Analyse d'ablation des variables dominantes

Afin d'évaluer l'influence réelle des principales variables identifiées par SHAP, une analyse d'ablation a été réalisée. Cette méthode consiste à supprimer progressivement les variables les plus importantes, puis à réentraîner les modèles afin d'observer l'évolution de leurs performances.

Quatre scénarios expérimentaux ont été considérés :

Scénario 0 : utilisation de toutes les variables ;

Scénario 1 : suppression de Type_Contrat ;

Scénario 2 : suppression de Type_Contrat et Branche_Activite ;

Scénario 3 : suppression de Type_Contrat, Branche_Activite et Type_Entreprise.

Tableau 6. Résultats de l'analyse d'ablation

	Scenario	Modele	Nb_variables	Accuracy	Precision	Recall	F1	ROC_AUC
0	Aucune	RF	10	1	1	1	1	1
1	Aucune	GB	10	1	1	1	1	1
2	Aucune	LR	10	1	1	1	1	1
3	Aucune	XGB	10	1	1	1	1	1
4	Sans_Type_Contrat	RF	9	1	1	1	1	1
5	Sans_Type_Contrat	GB	9	1	1	1	1	1
6	Sans_Type_Contrat	LR	9	0.8397	1	0.7	0.8235	0.8255
7	Sans_Type_Contrat	XGB	9	1	1	1	1	1
8	Sans_Type_Contrat_Branche	RF	8	1	1	1	1	1
9	Sans_Type_Contrat_Branche	GB	8	1	1	1	1	1
10	Sans_Type_Contrat_Branche	LR	8	0.7023	0.7627	0.6429	0.6977	0.6429
11	Sans_Type_Contrat_Branche	XGB	8	1	1	1	1	1
12	Sans_Type_Contrat_Branche_Entreprise	RF	7	0.4885	0.5205	0.5429	0.5315	0.4591
13	Sans_Type_Contrat_Branche_Entreprise	GB	7	0.4936	0.5209	0.6524	0.5793	0.4991
14	Sans_Type_Contrat_Branche_Entreprise	LR	7	0.5369	0.5387	0.9286	0.6818	0.5413
15	Sans_Type_Contrat_Branche_Entreprise	XGB	7	0.4911	0.5225	0.5524	0.537	0.4763

Les résultats mettent en évidence des comportements différents selon les familles d'algorithmes.

Dans le scénario de référence, c'est-à-dire sans suppression de variables, les quatre modèles atteignent des performances parfaites, avec des F1-scores égaux à 1. Ce résultat traduit une excellente capacité de classification sur le jeu de données étudié, mais il appelle également une interprétation prudente.

Lorsque Type_Contrat est supprimée, la régression logistique présente une diminution sensible de performance ($F1 = 0,8235$), tandis que Random Forest, Gradient Boosting et XGBoost conservent des performances maximales. Cette différence montre que les modèles fondés sur les arbres exploitent des interactions non linéaires entre plusieurs variables et compensent partiellement l'absence de cette variable.

La suppression simultanée de Type_Contrat et Branche_Activite accentue cette tendance. Le F1-score de la régression logistique baisse à 0,6977, alors que les modèles fondés sur les arbres maintiennent encore des performances optimales.

Lorsque Type_Entreprise est également retirée, les performances de tous les modèles diminuent fortement. Les F1-scores atteignent respectivement 0,5315 pour Random Forest, 0,5793 pour Gradient Boosting, 0,6818 pour la régression logistique et 0,5370 pour XGBoost.

Ces résultats montrent que les performances initiales ne reposent pas sur une seule variable dominante, mais sur la combinaison de plusieurs variables structurelles du marché du travail. Ils confirment que Type_Contrat, Branche_Activite et Type_Entreprise constituent conjointement les principaux déterminants de l'insertion professionnelle identifiés par les modèles.

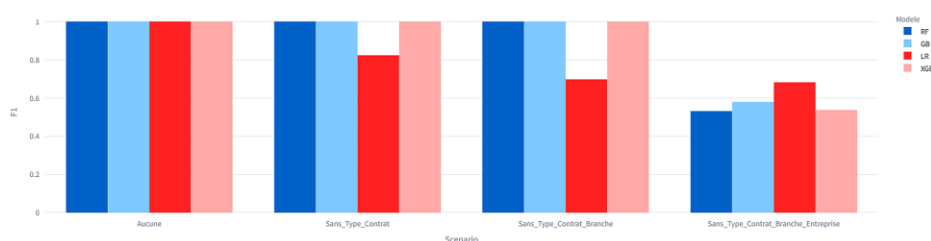


Figure 18. Évolution du F1-score selon les scénarios d'ablation

5.5. Discussion des résultats

5.5.1. Comparaison avec les travaux existants

Les résultats obtenus sont cohérents avec les travaux récents consacrés à l'application de l'intelligence artificielle au marché du travail. À l'instar d'Eloundou et al. (2023), cette étude montre que l'apprentissage automatique peut identifier des déterminants de l'insertion professionnelle à partir de données administratives. L'analyse SHAP confirme que le type de contrat, la branche d'activité et le type d'entreprise sont les variables les plus influentes dans les décisions du modèle.

Ces résultats rejoignent également Raghavan et al. (2020), qui montrent que les systèmes algorithmiques apprennent les régularités présentes dans les données historiques. Les variables identifiées dans cette étude reflètent donc les structures observées du marché du travail plutôt que des relations strictement causales.

Les travaux de Barocas et Selbst (2016), Chouldechova (2017) et Mehrabi et al. (2021) soulignent que certaines variables apparemment neutres peuvent agir comme proxys et reproduire des inégalités préexistantes. Les résultats de cette étude vont dans le même sens, puisque plusieurs variables structurelles apparaissent à la fois fortement prédictives et susceptibles de refléter des disparités dans l'accès à l'emploi.

Tableau 7. Confrontation des résultats avec la littérature

Auteur(s)	Principaux résultats de la littérature	Résultats obtenus dans cette étude
Eloundou et al. (2023)	Les systèmes d'intelligence artificielle per-mettent d'identifier des structures complexes dans les données du marché du travail et de mettre en évidence les facteurs influençant l'emploi.	Les analyses SHAP ont permis d'identifier le Type_Contrat, la Branche_Activite et le Type_Entreprise comme les principaux déterminants de l'insertion professionnelle.
Barocas & Selbst (2016)	Des variables apparemment neutres peuvent agir comme des variables proxy produisant des discriminations indirectes.	Les variables Branche_Activite et Province présentent les scores discriminatoires les plus élevés et peuvent refléter des inégalités structurelles du marché du travail.
Raghavan et al. (2020)	Les systèmes algorithmiques reproduisent souvent les structures historiques présentes dans les données utilisées pour leur apprentissage.	Les variables les plus influentes correspondent aux structures existantes du marché du travail congolais, notamment la segmentation sectorielle et les caractéristiques de l'emploi formel.
Chouldechova (2017)	Il existe une tension entre performance prédictive et équité algorithmique.	Les modèles atteignent des performances très élevées, tandis que les analyses SHAP mettent en évidence l'influence dominante

Fabris et al. (2025)	Les systèmes de recrutement automatisés peuvent reproduire des biais historiques et nécessitent des mécanismes d'audit d'explicabilité. L'approche combinant SHAP et score discriminatoire hybride permet d'identifier les variables influentes et les facteurs susceptibles de refléter des inégalités structurelles.
Mehrabani et al. (2021)	Les biais algorithmiques proviennent souvent des données d'entraînement, des variables utilisées et des choix de modélisation. Le score discriminatoire hybride montre que certaines inégalités observées dans les données sont partiellement reproduites par le modèle.

L'analyse d'ablation renforce les conclusions issues de SHAP. La suppression progressive des variables les plus influentes entraîne une dégradation des performances prédictives, ce qui montre que les résultats initiaux proviennent de l'effet combiné de plusieurs variables structurelles. Cette analyse limite le risque d'attribuer les performances élevées à une seule variable fortement corrélée avec la cible et confirme la pertinence de l'approche explicable adoptée.

5.5.2 Interprétation dans le contexte socio-économique congolais

L'interprétation des résultats doit être replacée dans le contexte du marché du travail congolais. En RDC, l'accès à l'emploi dépend fortement de la structure du tissu économique, marqué par la prédominance du secteur informel, la disponibilité limitée des emplois salariés et la concentration des opportunités dans certains secteurs d'activité.

Dans ce contexte, le poids du type de contrat, de la branche d'activité et du type d'entreprise est cohérent avec les réalités du marché de l'emploi. Les caractéristiques institutionnelles des employeurs semblent exercer une influence plus forte que les caractéristiques individuelles telles que le niveau d'étude ou l'expérience. Les politiques d'emploi devraient donc agir à la fois sur les compétences individuelles et sur les mécanismes structurels de création d'emplois.

5.5.3 Limites de l'étude

Cette étude présente plusieurs limites. Premièrement, les données proviennent exclusivement de l'ONEM et concernent les demandeurs d'emploi enregistrés dans la province du Nord-Kivu. Les résultats ne peuvent donc pas être généralisés à l'ensemble du territoire national sans validation complémentaire.

Deuxièmement, certaines variables explicatives, notamment le type de contrat et le type d'entreprise, présentent une forte proximité avec la variable cible. Cette situation peut contribuer aux performances très élevées observées et impose une interprétation prudente.

Troisièmement, les données utilisées sont historiques et peuvent contenir des biais reflétant les conditions du marché du travail. Les modèles développés doivent donc être considérés comme des outils d'aide à l'analyse et non comme des mécanismes automatisés de décision individuelle.

5.5.4 Implications scientifiques et politiques publiques

Au-delà de la performance prédictive, cette recherche montre l'intérêt des approches d'IA explicable pour l'analyse des phénomènes socio-économiques. L'association de l'apprentissage automatique, de SHAP et du score discriminatoire hybride améliore la transparence des modèles et facilite l'identification des facteurs associés aux disparités structurelles.

Sur le plan opérationnel, la plateforme développée peut appuyer les institutions chargées de l'emploi, de la formation professionnelle et de la planification des politiques d'insertion. Les résultats peuvent aider à orienter les actions vers les secteurs où

les difficultés d'accès à l'emploi sont les plus fortes et à soutenir des politiques publiques fondées sur les données.

5.6. Validation des hypothèses

Les résultats montrent que l'hypothèse principale est validée, que HS1 est validée, que HS2 est partiellement validée et que HS3 est validée.

L'approche fondée sur SHAP démontre sa capacité à identifier les facteurs influençant l'accès à l'emploi et à détecter certains mécanismes susceptibles de refléter des disparités structurelles.

Tableau 8. Validation des hypothèses

Hypothèse	Formulation	Méthode d'analyse	Résultats observés	Décision
Hypothèse principale	Les techniques d'IA explicables permettent d'identifier les facteurs discriminants influençant l'accès à l'emploi.	Analyse SHAP globale et locale	Les variables les plus influentes sont : Type_Contrat, Branche_Activite, Type_Entreprise. Hiérarchisation stable des variables dans tous les modèles.	Validée
HS1	Les facteurs structurels du marché du travail influencent davantage l'accès à l'emploi que les facteurs individuels.	Comparaison des contributions SHAP (variables structurelles vs individuelles)	Les variables structurelles présentent des valeurs SHAP significativement plus élevées que les variables individuelles.	Validée
HS2	Le niveau d'étude et la localisation influencent significativement l'insertion, mais moins que les facteurs structurels.	Analyse SHAP + score discriminatoire hybride	Niveau d'étude contributif mais secondaire ; Province = 2.05 contre Branche_Activite = 100.	Partiellement validée
HS3	SHAP améliore l'interprétabilité des modèles et permet de détecter des mécanismes de discrimination indirecte.	SHAP global + SHAP local + score discriminatoire	Explication individuelle des prédictions + identification d'écarts importants entre groupes (ex : genre vs branche d'activité).	Validée

6. Conclusion

Cette recherche avait pour objectif de développer un système intelligent explicable destiné à analyser les déterminants de l'insertion professionnelle en République démocratique du Congo.

Les résultats montrent que les modèles d'apprentissage automatique obtiennent des performances prédictives très élevées. L'apport principal de l'étude ne réside toutefois pas dans cette performance seule, mais dans l'intégration de l'intelligence artificielle explicable pour comprendre les mécanismes sous-jacents.

L'analyse SHAP a révélé que les facteurs les plus influents sont le type de contrat, la branche d'activité et le type d'entreprise.

Ces résultats indiquent que l'insertion professionnelle dépend davantage des caractéristiques structurelles du marché du travail que des seules caractéristiques individuelles.

Le score discriminatoire hybride a permis d'identifier les variables associées aux disparités observées dans les données et d'enrichir l'interprétation des prédictions.

Ainsi, l'IA explicable apparaît non seulement comme un outil de transparence, mais aussi comme un instrument d'analyse des inégalités dans l'accès à l'emploi.

Enfin, l'analyse d'ablation confirme la robustesse des résultats. La suppression progressive des variables identifiées par SHAP entraîne une diminution importante des performances, ce qui montre que la capacité prédictive des modèles repose sur la combinaison de plusieurs facteurs structurels. Cette expérimentation renforce la validité des conclusions relatives aux déterminants de l'insertion professionnelle et répond aux interrogations méthodologiques soulevées par les performances très élevées des modèles.

Contributions : Conceptualisation, Z.S.C. ; méthodologie, Z.S.C.; validation, Z.S.C., K.L.R., C.T., & M.L.F.; investigation, Z.S.C., K.L.R., C.T., & M.L.F.; ressources, Z.S.C., K.L.R., C.T., & M.L.F.; traitement des données, Z.S.C., K.L.R., C.T., & M.L.F.; écrire le manuscrit, Z.S.C., K.L.R., C.T., & M.L.F.; visualisation, Z.S.C., K.L.R., C.T., & M.L.F.; supervision, Z.S.C., K.L.R., C.T., & M.L.F.; correction du manuscrit, Z.S.C., K.L.R., C.T., & M.L.F. Les auteurs ont lu et approuvé la version publiée de ce manuscrit.

Sponsor financier : Cette recherche n'a reçu aucun soutien financier.

Disponibilité des données : Les données ne sont pas disponibles.

Remerciement : Non applicable.

Conflits d'intérêt : Les auteurs déclarent aucun conflit d'intérêt.

Références

1. Banque mondiale. (2021). World development report 2021: Data for better lives. World Bank. <https://doi.org/10.1596/978-1-4648-1600-0>
2. Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732. <https://doi.org/10.15779/Z38BG31>
3. Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer.
4. Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
5. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). Association for Computing Machinery. <https://doi.org/10.1145/2939672.2939785>
6. Chouldechova, A. (2017). Fair prediction with disparate impact: A study of bias in recidivism prediction instruments. *Big Data*, 5(2), 153–163. <https://doi.org/10.1089/big.2016.0047>
7. Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv*. <https://arxiv.org/abs/1702.08608>
8. Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2023). GPTs are GPTs: An early look at the labor market impact potential of large language models. *arXiv*. <https://arxiv.org/abs/2303.10130>
9. Fabris, A., et al. (2025). Fairness and bias in algorithmic hiring: A multidisciplinary survey. *ACM Computing Surveys*. <https://doi.org/10.1145/3696457>
10. Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5), 1189–1232. <https://doi.org/10.1214/aos/1013203451>
11. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
12. Han, J., Kamber, M., & Pei, J. (2011). *Data mining: Concepts and techniques* (3rd ed.). Morgan Kaufmann.
13. Leskovec, J., Rajaraman, A., & Ullman, J. D. (2020). *Mining of massive datasets* (3rd ed.). Cambridge University Press.
14. Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems 30* (pp. 4765–4774).
15. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), Article 115. <https://doi.org/10.1145/3457607>

16. Mitchell, T. M. (1997). Machine learning. McGraw-Hill.
17. Mpia, H. N., Kavalami, M. M., Lusenge, G. K., Ushindi, K. P., Kyalengekana, D.-D. K., & Cswaka, O. M. (2025). Lightweight deep learning system for infant-cry recognition with real-time notification in resource-constrained environments. *Machine Learning with Applications*, 22, Article 100791. <https://doi.org/10.1016/j.mlwa.2025.100791>
18. Molnar, C. (2022). Interpretable machine learning: A guide for making black box models explainable (2nd ed.). <https://christophm.github.io/interpretable-ml-book/>
19. OCDE. (2019). Artificial intelligence in society. OECD Publishing. <https://doi.org/10.1787/eedfee77-en>
20. Office National de l'Emploi. (s. d.). Rapports statistiques annuels sur l'emploi et l'insertion professionnelle au Nord-Kivu. ONEM.
21. Organisation internationale du Travail. (2023). World employment and social outlook: Trends 2023. International Labour Office.
22. Pearl, J. (2009). Causality: Models, reasoning, and inference (2nd ed.). Cambridge University Press.
23. Programme des Nations Unies pour le Développement. (2023). Rapport sur le développement humain. PNUD.
24. Raghavan, M., Barocas, S., Kleinberg, J., & Levy, K. (2020). Mitigating bias in algorithmic hiring: Evaluating claims and practices. In *Proceedings of the Conference on Fairness, Accountability, and Transparency* (pp. 469–481). Association for Computing Machinery. <https://doi.org/10.1145/3351095.3372828>
25. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?" Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 1135–1144). Association for Computing Machinery. <https://doi.org/10.1145/2939672.2939778>
26. Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206–215. <https://doi.org/10.1038/s42256-019-0048-x>
27. UNESCO. (2022). Global education monitoring report 2022: Non-state actors in education: Who chooses? Who loses? UNESCO.