

De la recherche en intelligence artificielle pour la santé au module pédagogique : un transfert de connaissances appliquées

Sandie Cabon^{1*}, Régine Le Bouquin-Jeannès¹

¹Univ Rennes, Inserm, LTSI - UMR 1099, F-35000 Rennes, France

*Auteur de correspondance : sandie.cabon@univ-rennes.fr

Résumé : *Au cours de l'année 2023-2024, un module intitulé "Développement de systèmes à base d'apprentissage pour la santé" a été proposé dans le cadre de la 3ème année de cycle ingénieur de l'Ecole Supérieure d'Ingénieurs de Rennes et du projet ReDHI. Ce module visait à transmettre des concepts avancés d'intelligence artificielle appliquée au domaine de la santé, tout en s'appuyant sur des travaux de recherche récents (gestion qualité des données, interprétabilité des modèles...) et, en particulier, ceux portés par le Laboratoire Traitement du Signal et de l'Image de l'Université de Rennes. Il était composé de trois parties : un cours magistral de 3 heures, deux séances de travaux dirigés d'1 heure 30 et de deux séances de travaux pratiques de 2 heures. A la fin du module, les étudiants ont été consultés. Ils ont apprécié le fait de travailler avec des méthodes récentes sur des sujets de recherche actuels. La durée allouée au module a été cependant jugée trop courte. Le nombre de séances de travaux pratiques et dirigés sera augmenté dès cette année pour adapter le rythme.*

Mots clés : intelligence artificielle pour la santé, dispositif pédagogique, transfert de connaissances appliquées, retour d'expérience.

Abstract: From Artificial Intelligence Research in Healthcare to a Teaching Module: A Transfer of Applied Knowledge

In the academic year 2023–2024, a module entitled Learning-Based Systems Development for Health was introduced in the third year of the engineering program at the École Supérieure d'Ingénieurs de Rennes, within the ReDHI project. The course aimed to teach advanced artificial intelligence concepts applied to healthcare, drawing on recent research—including data quality management and model interpretability—particularly from the Signal and Image Processing Laboratory at the University of Rennes. The module consisted of a 3-hour lecture, two 1.5-

hour tutorials, and two 2-hour practical sessions. Student feedback highlighted the value of working with state-of-the-art methods on current research topics, while also noting that the allocated time was too short. Consequently, the number of tutorials and practical sessions will be increased in the upcoming year to better match the learning pace.

Keywords: artificial intelligence for healthcare, teaching module, applied knowledge transfer, feedback from experience.

1 Introduction

L'enseignement supérieur a pour ambition de préparer les étudiants à relever les défis technologiques et scientifiques qu'ils rencontreront dans leurs parcours professionnels [1]. Dans un contexte d'évolution rapide et constante des pratiques et des outils numériques, cette mission peut s'avérer complexe. Ce défi est marqué dans un domaine comme l'intelligence artificielle (IA). L'engouement pour cette technologie la rend aujourd'hui omniprésente et incontournable [2]. L'enseignement de l'IA a ainsi été rapidement intégré aux cursus scientifiques et techniques en France. Cependant, peu de formations abordent les défis liés à sa mise en œuvre pratique : gestion de la qualité des données, besoin d'explicabilité, techniques d'évaluation.... Cela s'explique par le fait que ce sont des problématiques auxquelles les chercheurs se confrontent actuellement et que les solutions proposées ne sont ni totalement déployées ni communément admises. En parallèle, la réglementation sur ces systèmes évolue rapidement, comme en témoigne l'adoption récente de l'AI Act au niveau européen, qui fixe de nouvelles exigences impliquant directement la conception et l'évaluation de ce type d'outils [3].

Dans ce contexte, l'application de l'IA à la santé, notamment pour le développement de systèmes d'aide à la décision, constitue un terrain idéal pour mettre en lumière les défis à surmonter afin de concevoir des outils fiables. Cela permet également aux étudiants de développer un regard avisé sur cette technologie. C'est précisément cet objectif qui a guidé la conception du module "Développement de systèmes à base d'apprentissage pour la santé", proposé dans le cadre de la 3ème année de cycle ingénieur de l'Ecole Supérieure d'Ingénieurs de Rennes, dans la spécialité Technologies de l'Information pour la Santé. Ce module avait quatre objectifs :

- Identifier toutes les étapes de conception d'un système intégrant des méthodes par apprentissage pour la santé ;
- Former aux méthodes et bonnes pratiques qui interviennent à chacune de ces étapes ;
- Transmettre la capacité d'analyser et de présenter les performances des méthodes par apprentissage ;
- Faire le lien entre les activités de recherche actuelles en santé et l'enseignement.

Ce module a été proposé dans le cadre du projet ReDHI, lauréat de l'Appel à Manifestation d'Intérêts Compétences et Métiers d'Avenir en Santé Numérique en 2023. C'est un projet au croisement des filières des sciences et technologies du numérique et de la santé qui fédère 3

établissements Rennais : l'Université de Rennes, l'Ecole des Hautes Etudes en Santé Publique et le Centre Hospitalier Universitaire. Son objectif est d'attirer et de former à un niveau Bac + 5, des ingénieurs en Technologies de l'Information pour la Santé, des étudiants de Master de mathématiques et applications, orienté simulation numérique pour la santé, capables de répondre aux besoins du monde professionnel : dans l'élaboration de solutions technologiques innovantes en santé, dans le déploiement de systèmes d'information complexes et dans la mise en œuvre de dispositifs médicaux de nouvelle génération.

À travers cet article, nous analysons la manière dont ce module a tenté de relever ces défis, en commençant par présenter son contenu, puis en nous appuyant sur les retours des étudiants et sur les enseignements tirés.

2 Structure et contenu

Le module s'est articulé autour de trois parties : « Processus de développement pas à pas », « A vous d'analyser », « A vous de coder ».

2.1 Processus de développement pas à pas

Cette partie se présentait sous la forme d'un cours magistral de trois heures. Le principe était de parcourir, de manière exhaustive, 6 principales étapes du cycle de développement d'un système à base d'apprentissage pour l'aide à la décision en santé, en partant de l'objectif clinique jusqu'au déploiement (Figure 1).

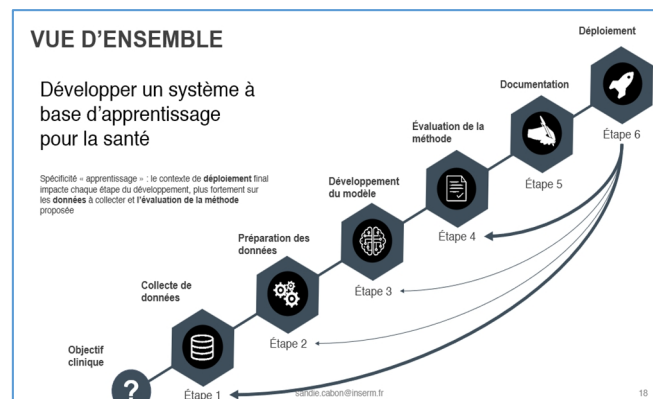


Figure 1. Vue d'ensemble des étapes de développement d'un système d'aide à la décision clinique basé sur l'apprentissage

Pour chacune des étapes, un ensemble de considérations réglementaires, méthodologiques et pratiques étaient présentées. Chaque partie se concluait par une diapositive sous forme de liste de cases à cocher (Figure 2), reprenant les points évoqués et pouvant alors servir de support pour des développements futurs et pour des analyses critiques de ce type de système. Le contenu de ce cours est basé sur toutes les étapes entreprises pour le développement de nouvelles solutions de surveillance des nouveau-nés prématurés, entreprises au Laboratoire Traitement du Signal et de l'Image (LTSI). Il s'appuyait sur trois principaux projets de recherche :

- Le projet européen Digi-NewB [4] visant à développer de nouvelles solutions de surveillance des nouveau-nés prématurés ;

- Le projet européen TEF-Health [5] ayant pour objectif de créer un réseau européen de test et d'expérimentation des dispositifs médicaux à base d'intelligence artificielle ;
- Le projet Improved pour lequel l'entrepôt de données de santé du CHU de Rennes a été exploité afin d'améliorer une solution de dépistage du cancer de la vessie [6].

COLLECTE DE DONNÉES	
Conception de l'étude	
DONNÉES	☒ Prévoir un jeu pour l'évaluation interne et un jeu pour l'évaluation externe (ou plusieurs) ☒ Définir la quantité de données / la durée d'une étude ☒ Prévenir les biais de genre, de localisation, d'ethnicité, socio-économiques, temporels
ACQUISITION	☒ Définir les méthodes d'acquisition de la donnée (système, logiciel...) ☒ Prendre en compte la diversité des systèmes, des réglages et des environnements d'acquisition ☒ Définir le plan de gestion des données incluant les lieux de stockage des données pour l'étude (plateforme sécurisée, HDS), la gestion des accès...
ANNOTATION	☒ Définir une procédure pour récolter la vérité terrain (= gold standard =)

sandie.cabon@inserra.fr

108

Figure 2. Exemple de diapositive pour la partie « Collecte de données » regroupant les différents points d'attention évoqués

2.2 A vous d'analyser

Cette partie a été conduite en Travaux Dirigés (TD). Deux séances d'1h30 ont été consacrées à l'analyse critique d'articles scientifiques présentant des systèmes d'aide à la décision clinique [7]. L'idée était de sensibiliser les étudiants aux différentes facettes de l'évaluation d'un système à base d'apprentissage mentionnées dans l'AI Act (Exactitude, Robustesse, Transparence, Cybersécurité, Éthique, Contrôle humain et Qualité des données). Les étudiants ont été regroupés par 4. Chaque groupe s'est vu attribuer un article et une grille de lecture vide. Les étudiants ont été invités à :

- Définir les questions à se poser derrière chaque critère d'évaluation, en faisant le lien avec le contenu de la 1ère partie du module.
- Analyser l'article au regard de ces questions ;
- Préparer une synthèse orale ;
- Présenter le travail réalisé aux autres groupes.

En cours de la première séance de TD, la grille complétée a été fournie pour assurer de la direction prise dans les réflexions (Figure 3)

Critère	Exemples de questions associées
Objectif clinique et méthode(s) proposée(s) description globale de l'étude	Quel objectif ? Quelle population est visée ? Quel type de modélisation ? Quelle chaîne de traitement ? Quelles hypothèse(s) pour résoudre le problème ? ...
Exactitude	Quelle(s) mesure(s) d'exactitude sont remontées ? Sont-elles cohérentes avec l'objectif clinique ? Est-ce qu'elles sont standards ou spécialement développées ? Comment est définie « la vérité terrain » ? ...
Robustesse technique	Est-ce que le système fonctionne dans des conditions variées ? Est-ce que le système sait faire face aux situations inattendues (e.g. données manquantes) ? Est-ce que les résultats sont stables ? ...
Transparence	Est-ce que la méthode est claire ? Est-ce que le chemin qui mène à une décision est intelligible ? Quels éléments permettent d'expliquer ...

Figure 3. Sous-partie de la grille de lecture guidant l'analyse des articles.

2.3 A vous de coder

La dernière phase du module s'est déroulée lors de Travaux Pratiques (TP). Deux séances consécutives de 2 heures, axées sur des manipulations concrètes en langage Python, ont été proposées. Les étudiants étaient en binômes. Pour chaque séance, un ensemble de ressources était fourni :

- Texte contenant les instructions et des annexes guidant l'interprétation des figures générées.
- Programmes « à compléter » en langage Python avec des instructions en commentaire et des liens vers la documentation associée aux fonctions à compléter.
- Des indications sur le temps à allouer pour chaque partie.

Cette structure permettait de concentrer le temps sur l'analyse des résultats tout en initiant à l'utilisation des fonctions communément utilisées pour les générer. Les deux séances faisaient écho à des éléments évoqués pendant le cours et discutés pendant les TD. Un compte-rendu était attendu. Suite aux séances, un mail regroupant les points forts et les erreurs courantes observées dans les comptes-rendus a été envoyé au groupe.

2.3.1 Importance de la préparation des données pour le développement d'un outil de prédiction de la survie des patients souffrant d'insuffisance cardiaque

Le premier sujet venait illustrer l'importance d'utiliser des méthodes de préparation des données adaptées et de savoir se questionner sur la meilleure méthode à choisir. Il s'agissait d'observer concrètement l'impact d'une gestion mal adaptée (et adaptée) des données manquantes sur les performances d'un modèle de classification. Ce TP a été construit à partir de l'étude de Chicco et al. dont les données avaient été rendues accessibles par les auteurs sur la plateforme Kaggle [8]. Cette étude portait sur la prédiction de la survie des patients souffrant d'insuffisance cardiaque à partir de données clinico-biologiques (Figure 4).

Variable	Description
age	Age du patient
anemia	Diminution des globules rouges ou de l'hémoglobine
creatinine_phosphokinase	Taux de l'enzyme CPK dans le sang (mcg/L)
diabetes	Présence de diabète
ejection_fraction	Pourcentage de sang quittant le cœur à chaque contraction
high_blood_pressure	Présence d'hypertension
platelets	Plaquettes dans le sang (kiloplaquettes/mL)
serum_creatinine	Taux de créatinine sérique dans le sang (mg/dL)
serum_sodium	Taux de sodium sérique dans le sang (mEq/L)
sex	Sexe du patient
smoking	Consommation de tabac
time	Période de suivi (jours)
DEATH_EVENT	Décès du patient pendant la période de suivi

Figure 4. Variables clinico-biologiques utilisées pour la prédiction de la survie.

En amont, le jeu de données a été artificiellement modifié afin d'y intégrer des données manquantes, avec un mécanisme non complètement aléatoire (donnée manquante en fonction des valeurs des autres variables). En effet, lorsque l'on traite des données réelles, les jeux de données ne sont souvent pas complets. Cela se vérifie, en particulier, dans le cadre de réutilisation de la donnée pour un usage secondaire (ex : exploitation de données provenant d'un entrepôt de données de santé). Pour gérer ces données manquantes, plusieurs stratégies existent : analyse en cas-complet, imputation univariée par la médiane ou équivalent, imputation multivariée par une approche de K-plus-proches voisins. Appliquer une mauvaise méthode d'imputation peut artificiellement ajouter un biais au jeu de données qui peut se transmettre au modèle de prédictions lors de l'apprentissage.

Dans ce TP, les étudiants ont d'abord pu caractériser le jeu de données à la fois au travers d'analyses statistiques simples (calcul de la moyenne, médiane pour les variables continues, comptage pour les variables binaires) et graphiques (boxplots, graphes de corrélation). Ensuite, trois méthodes d'imputation des données manquantes ont été appliquées avant d'entraîner trois modèles Random Forest à prédire le risque de décès chez les patients atteints d'insuffisance cardiaque. Les performances de ces modèles (concordance équilibrée, sensibilité et spécificité) ont ensuite été comparées, montrant alors que la méthode d'imputation multivariée était la plus adéquate pour la résolution de ce problème.

Tableau 1 Tableau de comparaison des performances des modèles après application de trois méthodes d'imputation, complété par les étudiants.

Métrique	Modèle Cas-complet	Modèle après imputation par la médiane	Modèle après imputation par k-plus proche voisin	Modèle à partir du jeu sans données manquantes
concordance équilibrée	0.72 +/- 0.05	0.79 +/- 0.02	0.82 +/- 0.07	0.82 +/- 0.05
sensibilité	0.77 +/- 0.02	0.77 +/- 0.08	0.79 +/- 0.08	0.81 +/- 0.04
spécificité	0.68 +/- 0.19	0.80 +/- 0.08	0.86 +/- 0.07	0.83 +/- 0.08

Pour aider dans l'interprétation des métriques et de ce qu'elles transmettent comme information sur les performances d'un modèle (ex : « le modèle détecte tous les patients qui sont décédés mais renvoie beaucoup de faux positifs » se traduit par une sensibilité élevée et une spécificité faible), une annexe était fournie.

2.3.2 Expliquer le comportement d'un modèle prédisant le risque de décès chez des patients insuffisants cardiaques

Le deuxième sujet faisait suite au premier, mais pouvait être réalisé indépendamment. Le principe est qu'une fois que nous avons obtenu un modèle suffisamment performant, il est

nécessaire d'essayer de comprendre son fonctionnement. L'interprétabilité des modèles revêt une importance cruciale dans le contexte clinique, notamment pour garantir la compréhension des décisions prises, favoriser la confiance des praticiens, et assurer une prise de décision éclairée. Dans le cadre de la santé, où des choix peuvent influencer directement la vie des patients, la transparence des modèles est une exigence essentielle. L'interprétabilité se réfère à la capacité à comprendre et interpréter les décisions prises par un modèle.

Au cours de ce TP, les étudiants tentaient d'expliquer comment le modèle Random Forest, entraîné lors du TP précédent, prédisait le décès chez les patients atteints d'insuffisance cardiaque. Ce modèle avait été évalué avec une concordance équilibrée de 0.82 ± 0.05 , une sensibilité de 0.81 ± 0.04 et une spécificité de 0.83 ± 0.08 .

D'abord, ils étaient invités à observer les 20 arbres de décision composant la forêt aléatoire et à constater la difficulté d'éclaircir le processus de prise de décision à partir de ces observations. Comme il est difficile d'observer le comportement de 20 arbres simultanément, des méthodes permettant d'expliquer les décisions du modèle de manière globale et locale ont ensuite été étudiées : l'importance des variables, les graphes de dépendance partielle et les valeurs de Shap.

La première méthode était l'importance des variables données intrinsèquement par le Random Forest (Figure 5). Pour celle-ci, les importances des variables sont calculées en évaluant la somme de leur impact sur la qualité des divisions effectuées à chaque nœud de chaque arbre. Cet impact est mesuré par la diminution de l'impureté, qui peut être quantifiée à l'aide de mesures telles que l'indice de Gini ou l'entropie (utilisée lors de l'apprentissage du modèle). Dans la même veine, mais avec une dimension plus interprétable des valeurs d'importance renvoyées, la permutation d'importance évalue l'importance d'une variable en permutant aléatoirement ses valeurs et en mesurant l'impact sur la performance du modèle (mesurée par une métrique de performance spécifique, comme la précision, le rappel, etc.). Une valeur élevée indique que la permutation aléatoire des valeurs de cette variable a un impact significatif sur la métrique de performance évaluée, suggérant ainsi que la variable joue un rôle crucial dans les prédictions du modèle. En revanche, une valeur proche de zéro suggère une moindre influence de la variable sur la performance du modèle.

Les graphiques de dépendance partielle (Partial Dependence Plot, PDP) offrent un moyen de visualiser l'influence globale d'une ou deux variables sur les prédictions d'un modèle (Figure 6). Ces graphiques fournissent un aperçu de la relation entre ces variables et les prédictions du modèle. Pour cela, une grille de valeurs couvrant la plage de la variable (ou des deux) est créée, et des prédictions sont effectuées en utilisant chaque valeur de la grille à la variable étudiée pour toutes (ou une partie) des observations du jeu de données, tandis que les autres variables restent constantes. La moyenne des prédictions obtenues est ensuite calculée et utilisée pour générer le graphique de dépendance partielle.

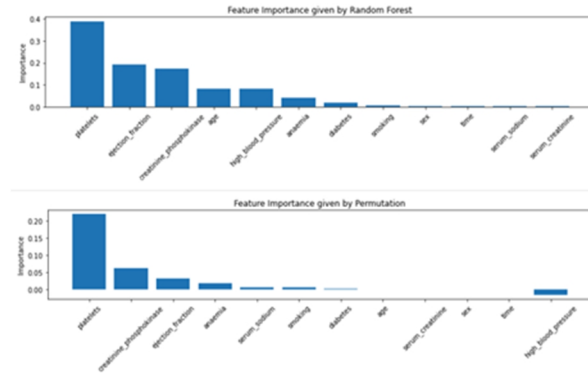


Figure 5. Importance des variables par Random Forest et permutation.

L'interprétation du graphique se fait en observant comment la prédiction évolue en fonction des valeurs de la variable ou des deux variables. Si le graphique montre une relation linéaire ou non linéaire, ou s'il y a des seuils critiques, cela peut fournir des informations importantes sur la manière dont le modèle utilise ces caractéristiques pour faire des prédictions. Les analyses ont été réalisées grâce à la librairie *pdp*.

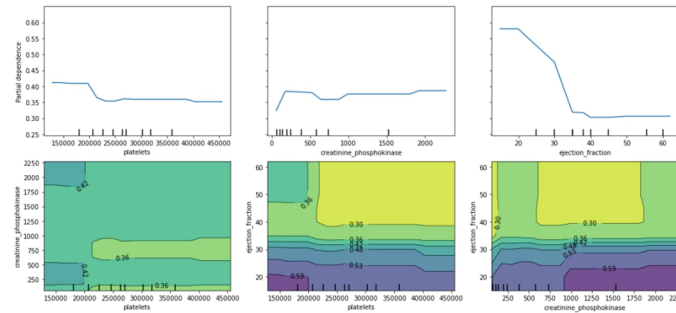


Figure 6. Graphes de dépendance partielle obtenus pour les trois variables les plus importantes.

Les valeurs de Shapley, issues de la théorie des jeux, évaluent l'importance de chaque variable dans un modèle prédictif en quantifiant leur contribution marginale à la prédiction. Dans le cadre de l'apprentissage automatique, les valeurs des variables d'une observation sont considérées comme les membres d'une coalition. Les valeurs de Shapley permettent de déterminer équitablement la répartition du "gain" (i.e., la valeur prédite) entre ces variables, éclairant ainsi sur l'importance de la contribution de chaque variable à la prédiction (Figure 7). Cela s'inspire du scénario où une coalition de joueurs collabore pour atteindre un bénéfice commun, et les valeurs de Shapley offrent une solution pour répartir équitablement les bénéfices entre les "joueurs", qu'ils soient des valeurs de variables individuelles ou des ensembles de valeurs de variables.

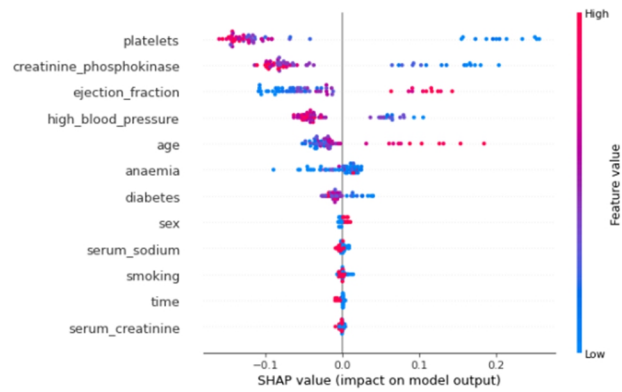


Figure 7. Visualisation "Beeswarm" des valeurs de Shapley.

Jusqu'alors, seulement des méthodes d'analyse du comportement global du modèle avaient été étudiées. Il existe aussi des méthodes qui permettent cette analyse à un niveau local. Les valeurs de Shapley ont aussi été analysées pour des observations ciblées (Figure 8).

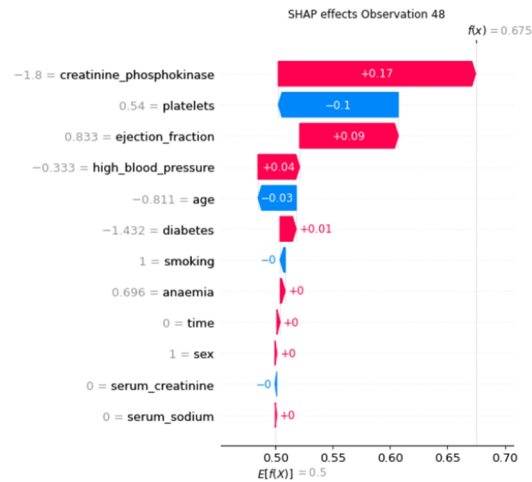


Figure 8. Analyse des valeurs de Shapley pour un vrai positif produit par le modèle.

A la fin de ce TP, les élèves étaient en mesure de donner des explications sur le fonctionnement du modèle et de proposer un modèle plus compact basé uniquement sur les 3 variables ayant montré l'importance la plus forte et aussi performant. L'objectif était de leur montrer qu'une analyse fine du comportement d'un modèle peut aider à proposer des solutions simples, plus interprétables pour les cliniciens.

3 Les retours des étudiants

Suite à ce module, un questionnaire a été envoyé aux étudiants. Treize étudiants sur seize ont répondu (Figure 9). Ce taux de retour indique la volonté des étudiants de participer à l'évaluation du contenu proposé.

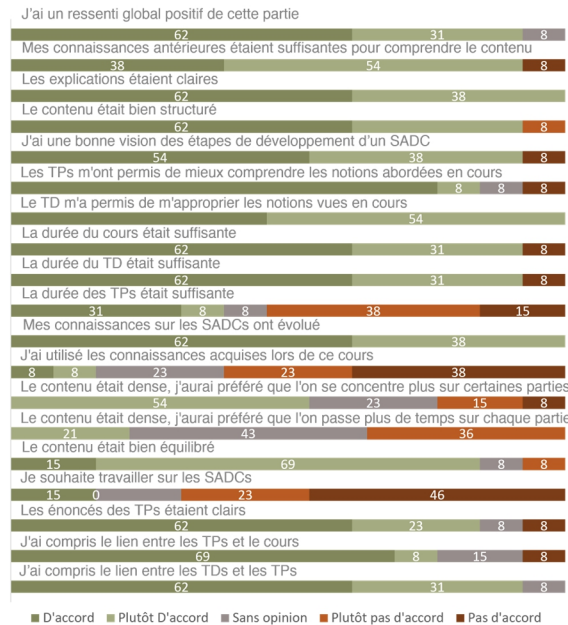


Figure 9. Evaluation par les étudiants.

Plusieurs questions ouvertes ont aussi été posées :

- Quels sont les points positifs à retenir pour cette partie ?
- Quels sont les points négatifs à retenir et améliorations possibles pour cette partie ?
- Quelles sont vos suggestions d'amélioration ?

3.1 Les points forts soulignés

Les retours des étudiants révèlent une adhésion forte à cette approche. Ils ont particulièrement apprécié :

- La pertinence du contenu : les exemples cliniques et les cas d'étude concrets ont permis de donner du sens aux apprentissages ;
- La qualité des TP : considérés comme les plus réussis de leur formation, les TP ont été salués pour leur structure, leurs consignes claires et les ressources annexes fournies ;
- Les liens avec la recherche : ce module a permis de se projeter dans des problématiques actuelles, tout en acquérant des compétences directement transposables.

3.2 Les points à améliorer

Malgré ces succès, plusieurs limites ont été identifiées :

- Le rythme intense : enchaîner 4 heures de TP ou couvrir de nombreux concepts en un temps limité a parfois rendu l'assimilation difficile ;
- La charge cognitive : les supports de cours, riches et détaillés, ont été perçus comme trop volumineux ;

- Le besoin de modularité : une suggestion fréquente a été d'entrelacer davantage cours et pratique pour permettre des pauses et des approfondissements progressifs.

4 Les améliorations prévues

En s'appuyant sur ces retours, plusieurs pistes d'ajustement se dessinent pour améliorer ce module :

- Prévoir plus de temps : augmenter les temps dédiés aux TD et TP ;
- Favoriser l'interactivité : introduire davantage de discussions et d'échanges pour maintenir l'attention et enrichir l'apprentissage ;
- Adapter la communication : expliquer que la densité du cours est liée à la volonté de leur fournir un guide exhaustif, à consulter plus tard et dont le contenu n'a pas vocation à être appris par cœur.

Ces ajustements permettront de mieux aligner les besoins des étudiants avec les ambitions du module.

5 Conclusion

Ce module a démontré qu'il est possible de transformer des problématiques de recherche actuelles en un format pédagogique accessible. En particulier, les travaux pratiques basés sur une application concrète et actuelle ont été bien reçus. Les étudiants ont acquis des compétences techniques, et ont également été initiés à des réflexions critiques sur l'utilisation de l'IA en santé. L'expérience de ce module souligne aussi l'importance d'un processus itératif dans la conception pédagogique : en intégrant les retours des participants, le contenu pourra être ajusté pour mieux répondre à leurs besoins. C'est un résultat prometteur pour renforcer les ponts entre recherche et enseignement.

Remerciements

Ce projet a bénéficié d'une aide de l'Etat gérée par l'Agence Nationale de la Recherche au titre de France2030 portant la référence « ANR-23-CMAS-0039 ».

Références

- [1] OCDE, « Études économiques de l'OCDE : France 2021 », Éditions OCDE (2021), <https://doi.org/10.1787/80013359-fr>.
- [2] Babaniyazovich, A. A. Shamsiya A. "The impact of artificial intelligence on human life in the future." International Multidisciplinary Journal for Research & Development, Vol 10,10 (2023).
- [3] Regulation 2024/1689. Regulation (EU) No 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonized rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act) Text with EEA relevance.

- [4] European Union's Horizon 2020 research and innovation program under grant agreement no. 689260 (Digi-NewB project). <https://www.digi-newb.eu/>
- [5] European Union's Digital Europe program under grant agreement no. 101100700 (TEF-Health pro-ject).<https://tefhealth.eu/>
- [6] Cabon S, Brihi S, Fezzani R, Pierre-Jean M, Cuggia M, Bouzillé G. "Combining a risk factor score designed from electronic health records with a digital cytology image scoring system to improve bladder-cancer detection". *Journal of Medical Internet Research*. 07/11/2024:56946 (forthcoming/in press), DOI: 10.2196/56946
- [7] Poiron A, Cabon S, Cuggia M. "How Trueness of Clinical Decision Support Systems Based on Machine Learning Is Assessed?". *Stud Health Technol Inform*. 2024 Aug 22;316:813-817. DOI 10.3233/SHTI240535.
- [8] Davide Chicco, Giuseppe Jurman: "Machine learning can predict survival of patients with heart failure from serum creatinine and ejection fraction alone." *BMC Medical Informatics and Decision Making* 20, 16 (2020)

Biographie des auteurs

Sandie Cabon : Sandie Cabon est Ingénieure de recherche à l'Institut National de la Santé et de la Recherche Médicale (Inserm). Elle exerce son activité au Laboratoire Traitement du Signal et de l'Image (LTSI UMR INSERM).

Régine Le Bouquin-Jeannès : Régine Le Bouquin-Jeannès est Professeure des Universités à l'Université de Rennes. Elle enseigne à l'École Nationale d'Ingénieurs de Rennes (ESIR) et exerce son activité de recherche au Laboratoire Traitement du Signal et de l'Image (LTSI UMR INSERM).