

AVRIL 2024

DONNÉES DE SANTÉ ARTIFICIELLES : ANALYSE ET PISTES DE RÉFLEXION

Livre blanc coordonné par :

Pr. ALLASSONNIÈRE
Stéphanie



Dr. FRAYSSE
Jean-Louis

BOTdesign

REMERCIEMENTS

La rédaction de ce livre blanc est une entreprise collégiale.
Nous remercions chaleureusement :

Pr. Stéphanie ALLASSONNIÈRE, Professeure de Mathématiques Appliquées et Vice-Présidence Valorisation et partenariats industriels, Université Paris-Cité, Chaire et directrice adjointe de PR[AI]RIE, Co-fondatrice de Sonio

Mme Manon DE FALLOIS, Adjointe à la cheffe du service de la santé, CNIL

Mr. Stanley DURRLEMAN, Directeur de recherche Inria à l'Institut du Cerveau, Co-fondateur et CEO de Qairnel (la clinique du Docteur Mémoto), Chaire PR[AI]RIE

Dr. Fabrice FERRÉ, PhD, Anesthésiste, CHU de Toulouse

Mr. Marco FIORINI, Directeur Général, Filière Intelligence Artificielle et Cancers

Dr. Jean-Louis FRAYSSE, Co-fondateur, BOTdesign

Mr. David GRUSON, Fondateur, ETHIK-IA

Mr. Jérôme KALIFA, Président et fondateur de Let it Care

Dr. Yann Maël LE DOUARIN, Conseiller médical à la DGOS et chef du département Santé et Transformation numérique, Ministère du Travail de la Santé et des Solidarités

Pr. Marie-France MAMZER, PU-PH, Université Paris Cité

M. Thierry MARCHAL, Président, Secrétaire Général, Avicenna Alliance et Chief Technologist Santé, EMEA, Ansys, Co-fondateur de Biomed In Silico France, Membre des groupes d'experts EMA et eSanté, Commission Européenne

Mr. Jean-Baptiste MASSON, Chercheur à l'Institut Pasteur, Chaire PR[AI]RIE, Co-fondateur d'AVATAR MÉDICAL

Mr. Hervé NABARETTE, Directeur adjoint Affaires Publiques, AFM-Téléthon

Mme Raphaëlle PARKER, Principal Scientist, The Janssen Pharmaceutical Companies of Johnson & Johnson

Pr. Raphaël PORCHER, PU-PH Université Paris Cité, Chaire PR[AI]RIE

Dr. Camille SCHURTZ, Responsable de l'accélération des Process réglementaires et de l'accès au marché, Agence de l'Innovation en Santé

Mme Sylvie TROY, Directrice Médical Adjointe et RWD/RWE Lead France, Pfizer

Mr. Vinh-Phuc LUU, Responsable du pôle épidémiologie et innovation, Filière Intelligence Artificielle et Cancers

Dr. Sarah ZOHAR, Directrice de Recherche, Inserm, Responsable de l'équipe de recherche "HeKA" - Inserm, Inria, Université Paris Cité

Pour ce travail studieux, sympathique et contributif.

Nous souhaitons également remercier France Asso Santé, Dr. Lise Alter (AIS), Mme Gaëlle Bouvet (BOTdesign), Mme Corinne Colignon (HAS), M. Florent Della Valle (CNIL), M. Jérôme Lang (IRIT), Mme Virginie Lasserre (Janssen), Mme Floriane Pelon (HAS), Mme Isabelle Ryl (Chaire PR[AI]RIE), Pr. Antoine Tesnières (PariSanté Campus), M. Olivier Thuillart (BOTdesign) et M. Félicien Vallet (CNIL) pour leurs précieuses contributions.

REMERCIEMENTS



AD (Augmentation des Données) : Dans le domaine de l'intelligence artificielle, le processus d'augmentation de données accroît la quantité de données d'entraînement par la création de nouvelles données à partir des données existantes.

AMM (Autorisation de Mise sur le Marché) : Accord donné à un titulaire des droits d'exploitation d'un médicament fabriqué industriellement pour qu'il puisse le commercialiser.

Anonymisation des données : Méthode rendant impossible l'identification d'une personne à partir d'un jeu de données

AIPD (Analyse d'Impact relative à la Protection des Données) : Outil permettant de construire un traitement conforme au RGPD et respectueux de la vie privée. Elle concerne les traitements de données personnelles qui sont susceptibles d'engendrer un risque élevé pour les droits et libertés des personnes concernées.

Apprentissage automatique : L'apprentissage automatique est un champ d'étude de l'intelligence artificielle qui vise à donner aux machines la capacité d'« apprendre » à partir de données, via des modèles mathématiques.

Avatardisation : Transformation des données en conservant leurs propriétés statistiques pour éviter la réidentification.

Bras contrôle ou bras témoin : Lors d'un essai clinique, groupe de participants qui sont soumis à la nouvelle méthode à l'étude visant à prévenir, à détecter, à traiter ou à contrôler la maladie.

CEPD (Comité Européen sur la Protection des Données) : Autorité de contrôle indépendante des institutions européennes (par exemple la Commission européenne) sur la protection des données.

Classification : Méthode statistique permettant d'attribuer une étiquette de groupe à des données qui sont non étiquetées.

Clustering : Méthode statistique qui permet le partitionnement des données en sous-groupes homogènes.

CNIL (Commission Nationale de l'Informatique et des Libertés) : Elle accompagne les professionnels dans leur mise en conformité et aide les particuliers à maîtriser leurs données personnelles et exercer leurs droits.

GLOSSAIRE

Cohortes virtuelles : Données stockées, traitées ou échangées sous forme numérique.

Cohortes synthétiques : Données réelles issues de l'agrégation de données collectées antérieurement

Cohortes artificielles : Données générées par des algorithmes ou des processus automatisés.

CPNPC (Cancer du Poumon Non à Petites Cellules) : Il s'agit du type de cancer du poumon le plus fréquent, qui représente 85-90 % de l'ensemble des cancers du poumon.

CPP (Comités de Protection des Personnes) : Ils sont chargés d'émettre un avis préalable sur les conditions de validité de toute recherche impliquant la personne humaine, au regard des critères définis par l'article L 1123-7 du Code de la Santé Publique.

Deep learning : L'apprentissage profond est un procédé d'apprentissage automatique utilisant des réseaux de neurones possédants plusieurs couches de neurones cachées. Ces algorithmes possédant de très nombreux paramètres, ils demandent un nombre très important de données afin d'être entraînés.

Déterministe : Relatif au déterminisme, à la doctrine philosophique suivant laquelle tous les événements sont liés et déterminés par la chaîne des événements antérieurs.

DM (Dispositif Médical) : Produit de santé destiné, par son fabricant, à être utilisé à des fins de diagnostic, de prévention, de contrôle, de traitement ou d'atténuation d'une maladie ou d'une blessure.

DMN (Dispositif Médical Numérique) : Les dispositifs médicaux font partie intégrante de la prise en charge médicale, ceux qui intègrent une fonction numérique peuvent générer un grand nombre de données de vie réelle et ouvrent la voie à une médecine plus personnalisée.

Données de vie réelle : Les données en vie réelle proviennent de sources multiples et reflètent le quotidien « réel » des patients et des médecins : la pratique de soins, la prise en charge, ou encore l'impact de la maladie et du traitement sur la vie courante.

Données HDLSS (High Dimension Low Sample Size) : Ensemble de données dont le nombre d'observation est très inférieur à la dimension de ses attributs. Exemple : base de données d'imagerie scanner abdominal dont la cohorte contiendrait 50 patients, chacun représenté par une image à 100 000 pixels.

Données multimodales : Les données multimodales se réfèrent aux ensembles de données contenant plusieurs modes ou sources de données, tels que le texte, l'audio, la vidéo et les images.

Dossier Patient Numérique (EHR – Electronic Health Record) : Fichier conservé sur support électronique contenant des informations sur la santé et les soins d'un patient tout au long de sa vie.

EFPIA (European Federation of Pharmaceutical Industries Associations) : Ses missions sont de promouvoir la recherche et le développement pharmaceutique en Europe, ainsi que la création d'un environnement économique favorable, réglementaire et politique, permettant de répondre aux besoins de santé et aux attentes croissantes des patients.

EMA (European Medicines Agency) : Contribue à protéger et à promouvoir la santé humaine et animale en évaluant et en contrôlant les médicaments au sein de l'Union européenne (UE) et de l'Espace économique européen.

Équation de Navier-Stokes : Équation de la mécanique des fluides décrivant le mouvement des fluides de type gaz ou la majeure partie des liquides.

Équations paramétriques : En mathématiques, une équation paramétrique est une équation dont les inconnues à estimer sont en nombre fini. A opposer à une équation non-paramétrique où le nombre de degrés de liberté est donc infini.

FDA (Food and Drug Administration) : Institution américaine chargée de la surveillance des denrées alimentaires et des médicaments. Elle autorise notamment la commercialisation de produits pharmaceutiques aux États-Unis.

Federated Learning ou apprentissage fédéré : Un paradigme d'apprentissage dans lequel plusieurs entités entraînent collaborativement un modèle d'IA sans mise en commun de leurs données respectives.

France 2030 : Le plan « France 2030 », doté de 54 milliards d'euros déployés sur 5 ans, vise à développer la compétitivité industrielle et les technologies d'avenir, dont la moitié des financements sont destinés à des acteurs émergents, et la moitié aux actions de décarbonation.

GLOSSAIRE

GAN based (réseaux adversariaux génératifs) : En apprentissage utilisant les réseaux de neurones, les réseaux antagonistes génératifs, parfois aussi appelés réseaux adversariaux génératifs, sont une classe d'algorithmes d'apprentissage permettant de générer des données à partir d'une compétition entre un générateur proposant des données les plus pertinentes possibles et un classifieur qui cherche à séparer les vraies données de ces données générées.

Généralisabilité : Propriété des modèles mathématiques permettant de généraliser ses résultats à une classe d'observations qu'il n'a pas rencontré préalablement dans sa phase de calibration.

HAS (Haute Autorité de Santé) : La HAS promeut les bonnes pratiques et le bon usage des soins auprès des usagers. Elle participe à l'information du grand public et à améliorer la qualité de l'information médicale.

IA (Intelligence artificielle) : L'intelligence artificielle (ou IA) est de plus en plus présente dans notre quotidien, notamment au travers de nouveaux produits ou services. Elle repose cependant sur des algorithmes gourmands en données, souvent personnelles, et son usage nécessite le respect de certaines précautions.

IA Act : L'IA Act (Artificial Intelligence Act) est un règlement qui vise à encadrer et favoriser le développement et la commercialisation des systèmes d'IA en Union européenne.

ICH : Conseil international d'harmonisation des exigences techniques pour l'enregistrement des médicaments à usage humain.

Inférence (statistique) : Ensemble des méthodes mathématiques et algorithmiques permettant d'induire les caractéristiques d'un groupe général à partir de celles d'un échantillon observé.

IRB (Institutional Review Board) : Institutional Review Board est un comité indépendant de personnes, composé de médecins, de scientifiques et de non scientifiques, et dont la responsabilité est d'assurer la protection des droits, de la sécurité, et du bien-être des sujets humains participant à des essais.

IRM (images de résonance magnétiques) : Ensemble des techniques permettant d'obtenir des images à partir de la résonance magnétique nucléaire.

Jumeaux artificiels ou numériques : Un Jumeau Numérique ou digital twin, est une réplique virtuelle d'un objet physique.

GLOSSAIRE

K-Nearest Neighbors (k plus proches voisins) : Méthode de classification basée sur la proximité d'une observation à ses plus proches voisins lui transférant leurs caractéristiques.

LEEM (Les Entreprises du Médicament) : Il s'agit d'un syndicat professionnel français du milieu pharmaceutique et un lobby qui s'est substitué en 2002 au Syndicat national de l'industrie pharmaceutique (SNIP).

LLM (Large Language Model) : Grands modèles de langage basés sur l'estimation de très grands réseaux de neurones mimant les relations probabilistes des mots entre eux dans des phrases ou des paragraphes.

Machine Learning (ML) : apprentissage machine : méthodes mathématiques permettant de calibrer des modèles à partir de l'observations de données. Ces modèles donnent aux ordinateurs la possibilité d'apprendre à partir de données à résoudre des tâches définies par l'utilisateur sans avoir été programmés explicitement.

Metaverse : Espace virtuel dans lequel on interagit de manière totalement immersive, grâce à des avatars.

Norme ISO 13485 : L'ISO 13485:2016 énonce les exigences relatives au système de management de la qualité lorsqu'un organisme doit démontrer son aptitude à fournir régulièrement des dispositifs médicaux et des services associés conformes aux exigences des clients et aux exigences réglementaires applicables.

Omniverse : Un ensemble d'univers supposés coexistants.

Outlier : Observation qui est distante du reste de la population observée.

Phase 1 (essai clinique) : Les essais de phase I correspondent le plus souvent à la première administration d'un médicament à l'homme. Les essais de phase I/II sont une variante des essais de phase I, ils permettent une évaluation préliminaire de l'efficacité à la dose sélectionnée ou bien de tester des combinaisons de médicaments.

Phase 2 (essai clinique) : Les essais de phase II, ont pour objectif de confirmer l'activité clinique préliminaire et/ou pharmacologique du médicament à la dose recommandée à l'issue de la phase I. Un nombre limité de malades est inclus dans ces essais (40 à 80 en moyenne). Certains essais de phase II comparent deux traitements. La durée d'une phase II est généralement de deux à trois ans, dépendant de la pathologie sélectionnée et du nombre de malades.

GLOSSAIRE

Phase 3 (essai clinique) : Les essais comparatifs sont destinés à comparer le nouveau médicament à un traitement standard afin de déterminer son efficacité. Les essais de phase III incluent plusieurs centaines, voire plusieurs milliers de malades, et durent d'ordinaire au moins quatre à cinq ans, selon la pathologie et l'effet attendu.

Phase 4 (essai clinique) : Après leur commercialisation, les médicaments continuent à faire l'objet d'un suivi strict à long terme, dit post-AMM, afin d'identifier tout effet secondaire grave et/ou inattendu dû à son administration. On parle de pharmacovigilance. Les essais de phase IV peuvent aussi être destinés à évaluer ce nouveau médicament approuvé dans des conditions d'administration différentes, par exemple la fréquence d'administration, le nombre de cures, la durée de la perfusion...

PK-PD : La modélisation pharmacocinétique et pharmacodynamique (fait intégralement partie du processus de développement des médicaments).

Pseudonymisation des données : La pseudonymisation des données est une technique de protection des données, qui consiste à traiter des données de telle sorte qu'il ne soit pas possible de les attribuer à une personne spécifique sans avoir recours à des informations additionnelles. Concrètement, il s'agit de remplacer les identificateurs personnels réels (noms, prénoms, emails, adresses, numéros de téléphone, etc.) avec des pseudonymes.

Réalité mixte (MR) : Technologie qui consiste à ajouter des éléments virtuels à ce que nous voyons et de pouvoir interagir physiquement avec ces éléments. La MR est un mélange entre réalité augmentée et réalité virtuelle. Pour son utilisation, vous devrez utiliser un casque en Réalité Augmentée, sachant que la position de l'utilisateur est calculée en temps réel, il pourra interagir avec les éléments virtuels par des gestes ou des manettes.

Réalité Étendue (XR) : Représente toutes les technologies qui créent des environnements et des objets générés par ordinateur. C'est-à-dire, la réalité augmentée (AR), la réalité mixte (MR) ou la réalité virtuelle (VR).

Réseaux de neurones : également connus sous le nom de réseaux de neurones artificiels (ANN) ou réseaux de neurones simulés (SNN) constituent un sous-ensemble de l'apprentissage automatique et sont au cœur des algorithmes de l'apprentissage en profondeur. Leur nom et leur structure sont inspirés du cerveau humain, imitant la manière dont les neurones biologiques s'envoient des signaux.

RGPD (Règlement Général de Protection des Données) : Texte réglementaire européen qui encadre le traitement des données de manière égalitaire sur tout le territoire de l'Union européenne (UE).

RWD (Real World Data) : Analyse de données de vie réelle. Ce sont des données relatives à l'état de santé et au mode de vie en général de patients réels. Elles sont recueillies dans des real world settings, ce qui signifie qu'elles sont obtenues à partir de situations réelles (et non à partir de formulaires ou de rapports de laboratoire). Par exemple, les habitudes sportives, la profession et le cadre de vie.

Segmentation : Délimitation de la frontière des objets dans une image.

Singleton : En mathématiques, ensemble constitué d'un seul élément.

SoC (Standard of Care) : Traitement reconnu par les experts médicaux comme le plus approprié pour un certain type de maladie dans un contexte particulier et largement utilisé par les professionnels de la santé. Également appelés meilleures pratiques, soins médicaux standard, meilleure thérapie disponible et thérapie standard.

Supervision humaine ou garantie humaine : Supervision humaine ou garantie humaine est un label reconnu à l'échelle française, européenne et même internationale, le principe de Garantie Humaine assure le développement éthique des intelligences artificielles concourant à la santé, en établissant des points de supervision humaine tout au long de leur évolution.

SVM (Machine à Vecteur de Support) : Méthode de clustering basé sur la création de frontières entre les populations à partir d'un petit nombre d'observations pivots.

VAE (Variational Auto-Encoders) : Modèle mathématique permettant l'encodage de données dans un espace représentatif de petite dimension ainsi que le décodage de points de cet espace, générant ainsi de nouvelles observations.

Voxel : Pixel (Picture Élément) tri-dimensionnel (Volume Élément).

PRÉFACE

Il découle de notre responsabilité collective d'utiliser toutes les méthodes éthiques pour garantir la santé de nos concitoyens actuels et futurs. Il relève de notre éthique personnelle d'embrasser toute technologie susceptible d'accomplir cette mission, même si nous ne la comprenons pas encore ou qu'elle menace notre méthode de travail qui nous est si familière. Certains nous diront qu'il ne faut pas aller trop vite ; mais pouvons-nous ralentir quand les patients meurent ? Toutefois, nous ne pouvons en aucun cas nous précipiter et risquer de mettre ces patients en danger. Aussi, nous devons accueillir avec enthousiasme toutes les critiques constructives et veiller à leur apporter des réponses satisfaisantes.

Le premier quart de ce vingt et unième siècle a vu émerger et se développer de nombreuses méthodes numériques telles que les statistiques prédictives, la modélisation et simulation numérique et donc l'intelligence artificielle. Elles sont largement exploitées pour drastiquement améliorer la sécurité des voitures, avions et autres centrales énergétiques. Pourquoi en serait-il autrement pour la santé ?

Il y a juste mille ans, Avicenna, un médecin et scientifique Perse, publiait « Le Canon de la Médecine » ; il fut le premier à suggérer de tester tout nouveau traitement sur un groupe de patients tout en gardant un groupe de contrôle : les premiers tests cliniques. Au cours des dix derniers siècles, la médecine a prudemment mais systématiquement adopté les nouvelles technologies afin de réduire les risques encourus par ces cobayes humains et ainsi accélérer l'innovation médicale. Ces attitudes progressistes ont mené à des révolutions majeures, souvent apparues dans nos contrées, telles que la chirurgie moderne, la pharmacie et ses vaccins. Nous sommes à l'aube d'une nouvelle révolution médicale qui nous mènera rapidement à la santé personnalisée.

Nous disposons d'une masse incommensurable de données que nous pouvons désormais exploiter en développant des algorithmes plus pertinents. Nous avons la sagesse d'adopter des régulations pour la protection des données personnelles et nous établissons des protocoles pour établir la validité et la fiabilité de ces données. Bien au-delà de l'expérience et la mémoire du médecin particulier, il est désormais possible de comparer une nouvelle situation à l'immensité de la mémoire collective et deviner l'évolution probable d'un patient via une intelligence artificielle des données (Data Driven). La maîtrise, certes relative, de la physique, chimie et biologie du corps humain nous offre des modèles numériques de plus en plus performants pour prédire l'évolution d'une pathologie via une intelligence artificielle de la connaissance (knowledge driven). L'intelligence artificielle combine ces approches, donnée et connaissance, non pas pour remplacer les médecins et infirmiers ou infirmières, mais pour les assister dans ces tâches de routine et de mémoire et leur permettre de se consacrer à élaborer le meilleur traitement pour le patient, la quintessence du génie humain.

Comme expliqué dans ce livre blanc, nos chercheurs et chercheuses développent également de multiples approches pour augmenter ces données et modèles, les anonymiser pour respecter les patients, dans le but de tester tout nouveau traitement sur des cohortes potentiellement infinies de patients, désormais virtuels, bien avant d'entamer des tests cliniques traditionnels. Il serait absurde de se priver de ces tests cliniques dits 'in silico', par différence avec les tests 'in vivo', réalisés sur des cohortes de patients virtuels, pouvant inclure des patients extrêmes, des maladies rares voire des minorités (enfants, femmes enceintes, minorités ethniques, etc.). Les tests cliniques in silico vont ainsi apporter des informations très utiles, plus rapidement et à moindre coût, qui vont drastiquement réduire les risques lors des tests cliniques traditionnels.

Aussi, nous remercions les auteurs de cet ouvrage pour cette remarquable synthèse et implorons les autorités d'y accorder toute leur attention, afin de garder une porte largement ouverte tout en conservant un esprit critique pour défier ces approches. Ainsi, nos concitoyens et concitoyennes pourront bénéficier pleinement et rapidement des nouveaux traitements et notre industrie de la santé jouir d'une régulation leur permettant de tester rapidement et efficacement leurs innovations médicales sans devoir s'exporter dans des régions plus ouvertes à ces méthodes numériques.

Nous sommes dans une course mondiale où nous devons agir vite, mais prudemment.



Thierry Marchal

Président, Secrétaire General, Avicenna Alliance

Chief Technologist Santé, EMEA, Ansys

Co-fondateur de Biomed In Silico France

Membre des groupes d'experts EMA et eSanté,
Commission Européenne

SOMMAIRE

Contexte	1
Justification scientifique	2
Usage des cohortes de patients artificiels	3
Protection des données et génération de cohortes de patients artificiels	4
Preuves et méthodologie pour la validation des produits de santé à l'aide de données artificielles	5
Enjeux et garanties éthiques	6
Conclusion	7
Références bibliographiques	8

CONTEXTE

01

La recherche et le développement au service de la santé s'accélèrent, elles sont marquées par une augmentation rapide et continue des découvertes, des avancées technologiques et des thérapeutiques. L'enjeu est de construire un plan de développement permettant de démontrer l'intérêt des nouvelles technologies et leur apport par rapport à l'arsenal disponible et permettre ainsi aux réelles innovations d'être détectées et rendues disponibles.

Plusieurs facteurs contribuent à cette accélération :

- La collaboration et le partage des données à l'échelle mondiale qui facilitent la circulation rapide d'informations et d'idées ;
- L'augmentation importante des capacités de calcul ;
- Le recours à l'intelligence artificielle (IA) ;
- Les financements et investissements dans le cadre de France 2030 notamment ;
- Les voies d'accès au remboursement dérogatoires qui visent à accélérer la mise à disposition des technologies de santé tout en assurant la mise en place des preuves nécessaires à l'évaluation de l'intérêt et de leur apport dans l'arsenal de soins. La HAS conduit en particulier des travaux pour définir quelles peuvent être les conditions méthodologiques acceptables lors d'un accès anticipé pour les médicaments ;
- La recherche et les essais cliniques qui évoluent pour devenir plus efficaces.

Par ailleurs, notre système de santé doit faire face à des défis structurels lourds (vieillesse de la population, hausse des maladies chroniques, inégalités d'accès aux soins, *etc*).

Dans sa publication « Essais cliniques 2030 » en mars 2022, l'organisation professionnelle des entreprises du médicament (LEEM) pointe les enjeux, perspectives et challenges afin que la France reste compétitive à l'échelle mondiale de la mise sur le marché de nouvelles molécules, et qui peuvent être étendus aux dispositifs médicaux (DM). Le recours grandissant à la médecine personnalisée, la difficulté de recrutement des patients et leur rétention dans les études, la méthodologie des essais se complexifiant, la concurrence accrue avec les autres pays, et le coût que représente la mise en place de ces études cliniques, interrogent sur l'adaptation des structures et systèmes de façon à ce que la France continue de participer au développement des molécules de demain.^{1 2}

1 "Les essais cliniques apportent un premier accès à l'innovation pour les patients, et notamment pour ceux atteints de pathologies graves comme le cancer (42 % des essais initiés), mais également pour les maladies auto-immunes (18%) ou encore pour les maladies du système nerveux central (13%)." Source LEEM.

2 Si, pour démontrer la preuve de l'efficacité d'un médicament, l'essai randomisé en double aveugle reste le gold standard, donc à privilégier, la HAS introduit la possibilité d'intégrer des données moins consolidées à condition qu'elles permettent la comparaison avec les traitements disponibles. En effet, seule la comparaison permet de se prononcer sur la valeur ajoutée d'un nouveau traitement. L'objectif est de permettre l'accès au remboursement de produits immatures, tout en maintenant un niveau d'exigence de qualité acceptable. La nouvelle doctrine d'évaluation de la CT s'ouvre ainsi aux données de comparaison indirecte de bonne qualité méthodologique ou encore à celles issues de groupe contrôle, à condition qu'elles soient expliquées et justifiées en amont par l'industriel.

Pour répondre à ces défis, l'essor du numérique et de l'IA permet d'envisager, dans un avenir proche, la mise en œuvre de nouveaux schémas de protocoles et d'essais cliniques qui viendront compléter l'arsenal méthodologique en utilisant, par exemple, des données de vie réelle ou des bras de patients artificiels. En effet, en s'appuyant sur des données médicales réelles, collectées dans le cadre des soins, de la prévention ou de la recherche, des outils de Machine Learning (ML) permettent de générer des données dites virtuelles ou artificielles à partir de données médicales réelles. Ces outils, dès lors qu'ils seront validés, pourraient permettre de compenser les difficultés de recrutement des patients réels dans les bras de contrôle grâce à des patients générés artificiellement.

On parle ainsi de cohortes augmentées qui sont constituées, d'une part, de patients réels et, d'autre part, de patients artificiels. Les données de ces patients artificiels sont générées de toute pièce par des modèles de Machine Learning s'appuyant sur les données réelles de la cohorte, provenant des patients recrutés pour l'étude concernée. Ces cohortes augmentées sont à distinguer des cohortes dites synthétiques (tels les bras de contrôle synthétiques) qui, elles, sont composées de données réelles « recyclées » provenant de patients inclus dans des cohortes antérieures. Ces cohortes doivent être fiables, représentatives des patients sources et ces derniers ne doivent pas être ré-identifiables.

Nous ferons donc la différence dans ce document entre (1) les cohortes synthétiques qui impliquent des données collectées antérieurement, réutilisées et colligées pour la création d'une nouvelle cohorte et (2) les cohortes artificielles qui augmentent le nombre de patients réels recrutés pour l'étude à mener.

Les données artificielles peuvent présenter d'autres intérêts pour la recherche dans le domaine de la santé en permettant de réaliser des expérimentations, telles que le développement de nouvelles méthodes d'analyses ou encore des entraînements et/ou tests de modèles d'IA, sans avoir recours à des données sensibles puisque n'appartenant pas à des patients réels.

Ainsi, les données artificielles pourraient contribuer à faciliter le développement de modèles permettant :

- D'assister les professionnels de santé dans la pose de diagnostic ;
- D'aider à la personnalisation des traitements (prise de décision thérapeutique) ;
- D'identifier les populations de patients sensibles à un traitement donné ;
- De soutenir la formation médicale (simulations pour apprentissage et entraînement) ;

- De simuler la propagation d'épidémies et test de stratégies de prévention et contrôle ;
- De tester la robustesse et la sécurité des systèmes informatiques et des applications numériques (simuler des attaques potentielles et des failles de sécurité) ;
- De mettre en œuvre des essais cliniques avec des patients artificiels pour évaluer l'efficacité et la sécurité de nouvelles molécules ou dispositifs médicaux.

Pour des raisons éthiques liées à la confidentialité des données sensibles que sont les données de santé, et donc à la protection de la vie privée notamment par le Règlement Général sur la Protection des Données au sein de l'Union européenne (RGPD), l'accès aux données de patients nécessaires au développement de ces modèles est très complexe et limité. Les données artificielles, générées à partir de données réelles dans des environnements sécurisés très contraignants, mais ne provenant pas de personnes physiques, pourraient, elles, être facilement mises à disposition des chercheurs sur des plateformes plus accessibles.

L'utilisation raisonnée des données artificielles dans le domaine de la santé offrirait des avantages précieux tout en limitant certains risques associés à l'utilisation de données réelles. Leur « utilisation en routine » nécessite une validation rigoureuse par des experts (professionnels de santé, mathématiciens, patients) pour s'assurer de leur fiabilité, c'est-à-dire de leur capacité à reproduire fidèlement les données réelles. Leur utilisation implique en effet d'en maîtriser les risques spécifiques et de définir les périmètres d'utilisation possibles.

Il n'existe actuellement pas de recommandations émanant d'agences ou d'autorités de régulation définissant les critères d'acceptabilité de cohortes de patients artificiels pour l'évaluation des produits de santé ou des dispositifs médicaux, le concept même de recours à des cohortes de patients artificiels dans cet objectif ne faisant pas consensus. Les expériences d'utilisation de ce type de cohortes dans les études soumises aux agences sont d'ailleurs, à ce stade, limitées.

En lisant ce livre blanc, nous vous proposons de mieux comprendre le processus de création des patients artificiels, les différents champs d'application possibles et les limites de ces cohortes de patients, mais surtout de décrire les processus de validation obtenir pour garantir la fiabilité et la représentativité de ces patients d'un genre nouveau. Cette validité est indispensable avant d'envisager que ces données puissent être utilisées par les autorités sanitaires lors de la mise sur le marché des produits de santé, de leur évaluation dans le cadre de démarches de remboursement mais également pour apporter des garanties aux utilisateurs de ces technologies de santé.

L'innovation n'ayant de valeur que si elle est synonyme de bénéfice pour les usagers, nous espérons que ce livre blanc y contribuera.

Bonne lecture.



La mise en œuvre de cohortes de patients artificiels à partir d'IA générative confirme l'efficacité de la collaboration des équipes publiques et privées dans les domaines de la recherche, la médecine, l'éthiques, les régulateurs les associations de patients et les industriels. Ces collaborations agiles et structurées permettent d'accélérer la validation et la mise à disposition de technologies très innovantes dans le cadre de France 2030.

Pr. Stéphanie Allassonnière

Professeure de Mathématiques Appliquées et
Vice-Présidence Valorisation et partenariats industriels,
Université Paris-Cité,
Chaire et directrice adjointe de PR[AI]RIE
Co-fondatrice de Sonio



**JUSTIFICATIONS
SCIENTIFIQUES**

02

Nous vous proposons maintenant de rentrer dans le cœur du sujet.

L'augmentation des données (AD) est l'art d'augmenter la taille d'un ensemble de données d'intérêt en créant des données artificielles annotées ayant les mêmes caractéristiques que la population d'origine sans toutefois reproduire les données réelles à l'identique. C'est ce que nous appelons la « généralisabilité » des modèles.

Cette méthode permet également d'envisager le rééquilibrage du nombre d'échantillons par classe en suréchantillonnant les classes minoritaires. Cela permet de traiter des cas où une sous-population n'est pas assez représentée dans une étude, comme des populations fragiles par exemple (femmes enceintes, enfants, malades souffrants de maladie rares, etc).

En complément des données collectées en vie réelle, dans des situations dans lesquelles le recrutement est difficile, générer numériquement des données dites artificielles peut être un levier. Ces données artificielles sont des informations ou des ensembles de données créés numériquement pour simuler des caractéristiques et des structures similaires à celles des données réelles.

Elles sont générées à l'aide d'algorithmes implémentant des modélisations mathématiques. Il existe deux types de modèles.

Les premiers sont les modèles dits mécanistiques. Ces modèles mathématiques consistent en la concaténation et l'interaction d'équations connues de la mécanique, de la chimie ou d'autres disciplines qui permettent de décrire un phénomène. Ces équations sont déterministes en cela qu'il n'y a pas d'aléa. Il en va ainsi des équations de Navier-Stokes, des équations aux dérivées partielles non linéaires qui décrivent le mouvement des fluides newtoniens. Ces équations permettent de générer des sorties qui sont pour la plupart du temps des quantités physiques caractérisant le phénomène étudié. Dans le cas des essais cliniques, de telles méthodes sont employées par Nova Discovery notamment qui, par la recherche fine dans la littérature des équations impliquées dans un processus pharmacologique, sont en mesure de mimer les résultats d'un essai clinique. Cela peut néanmoins impliquer de nombreuses équations avec de nombreux paramètres dont il est important de trouver dans la littérature les valeurs pivot.

Les autres modèles mathématiques utilisés sont basés sur des modèles statistiques. Leur but est de proposer des équations mimant des schémas observés dans les données réelles ou représentant la distribution de probabilité de ces données dans un espace mathématique. Ces approches peuvent utiliser des équations mécanistiques, par exemple les équations PK-PD, c'est-à-dire la modélisation pharmacocinétique (Pharmaco-Kinetic) et pharmacodynamique (Pharmaco-Dynamic) mais y ajoutent de l'aléa permettant une grande souplesse en particulier pour

prendre en compte des échelles populationnelles et individuelles dans le même modèle. On parle alors de modèles à effets mixtes. Ils sont à privilégier lorsqu'il n'existe pas d'équation régissant le phénomène étudié (exemple la taille des enfants entre 0 et 20 ans n'est générée par aucune équation et pourtant elle est très bien décrite dans les carnets de santé).

Ces approches, basées sur des modèles dit d'Intelligence Artificielle Générative permettent le rééchantillonnage : après une phase de calibration de ces modèles (apprentissage), l'utilisateur peut générer de nouvelles données, différentes des données d'apprentissage, mais reproduisant les caractéristiques populationnelles observées et capturées. Ces modèles peuvent impliquer un niveau plus ou moins complexe d'équations explicites ou implicites (par exemple en utilisant un réseau de neurones). De telles méthodes sont employées par Quinten par exemple.

À partir de ces modèles, nous sommes en mesure de décrire un phénomène à partir de l'échantillon d'une population qui est observé. Ils apportent une autre perspective d'utilisation : permettre de créer de nouvelles observations non vues dans l'échantillon initial, mais qui conservent les mêmes caractéristiques. Ces modèles sont dits génératifs. La création de données supplémentaires s'appelle l'augmentation de données.

Il existe de manière plus générale plusieurs techniques pour permettre de générer des données artificielles. Elles peuvent être divisées en trois niveaux.

L'utilisation de l'IA dans le domaine de la santé présente des challenges particulièrement complexes inhérents aux types de données sur lesquelles elle s'appuie (données sensibles, sparses, hétérogènes, quantités souvent limitées, ...).

Cependant, les efforts de recherche de ces dernières années qui s'attèlent à répondre à ces défis spécifiques, permettent aujourd'hui d'entrevoir un potentiel vertigineux, entre autres et en particulier, dans le domaine des études cliniques, portant des promesses d'études moins chères mais surtout plus justes, plus inclusives et plus efficaces.

Le recours à des données artificielles pour le développement de nouveaux médicaments est une ambition particulièrement audacieuse et sensible qui nécessite un effort commun et soutenu d'acteurs divers et complémentaires afin d'avancer de la manière la plus efficace possible.



Raphaëlle Parker

Principal Scientist, The Janssen Pharmaceutical Companies
of Johnson & Johnson



I. Une donnée d'un seul patient produisant une ou plusieurs autres données similaires

Un exemple simple est celui des images. Il s'agit d'appliquer des transformations simples telles que l'ajout de bruit (modifications aléatoires des niveaux de gris des pixels), le floutage (appelé convolution), le zoom ou des déformations telles que des translations et/ou rotations. On attribue ensuite le label de l'image initiale aux images créées.

Bien que ces techniques d'augmentation se soient révélées très utiles, elles restent fortement dépendantes des données et donc limitées. Certaines transformations peuvent en effet ne pas être informatives ou même induire des biais. Prenons l'exemple d'un chiffre représentant un 6 qui donne un 9 lorsqu'il est trop tourné ou d'un 4 qui peut ressembler à un 9 si on floute « trop » la forme (voir figure ci-dessous). La difficulté d'évaluation de la pertinence des données augmentées s'accroît en fonction de la complexité des données d'origine et peut nécessiter l'intervention d'un expert qui évalue le degré de pertinence des transformations proposées. Cette évaluation peut aussi être difficile techniquement, car l'expertise est confrontée à des limites (penser aux déformations admissibles d'un foie sain difficiles à caractériser). Un autre exemple santé sera le cas d'un nodule qui pourrait être confondu avec une structure mixte tel qu'un kyste hémorragique s'il était trop bruité.

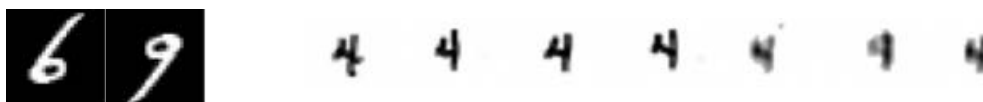


Fig 1. Exemples de modifications sujettes à erreurs.

Ces techniques sont donc très prisées, car simples et très efficaces pour l'imagerie. Elles sont toutefois plus difficiles à appliquer pour d'autres types de données, en particulier des données tabulaires ou des données génétiques.

II. Une donnée d'un patient produisant une donnée artificielle modifiée par utilisation de ces voisins identifiés dans une population : le bruitage par moyennage des voisins.

Cette technique a été mise en lumière par la société Octopize qui a proposé une méthode dite d'avatardisation.

Remarque : Initialement prévue pour anonymiser les données, cette méthode pourrait s'envisager pour créer également des données artificielles. Il s'agirait alors de ne pas produire une seule donnée par observation, mais un petit nombre pour chacune. Cela étant, produire plus de données pourrait remettre en cause l'aspect anonymisant de la cohorte produite. Cette question reste ouverte à ce jour.

Le principe est de représenter les données dans un espace mathématique dans lequel il est possible de faire des moyennes locales. Les données tabulaires quantitatives sont tout naturellement déjà dans un espace de ce type. Pour les données qualitatives (c'est-à-dire des données catégorielles issues de choix multiples) ou des données plus complexes (par exemple des images), il est nécessaire de définir correctement la notion de moyenne locale. En effet la moyenne de deux images considérées comme des tableaux de chiffres ne fait aucun sens, car produirait une image floue avec deux fois les éléments présents (penser à l'exemple de la moyenne de deux visages produisant une image avec quatre yeux, deux nez et deux bouches). Pour les images, il est par exemple nécessaire de recourir à des notions de modèles déformables [69] généralisant la notion de moyenne aux formes contenues dans les images.

Une fois cette notion de moyenne locale définie mathématiquement, chaque donnée sera alors modifiée par l'influence de k de ses plus proches voisins. Le nombre de voisins utilisés (k) définit le niveau de bruitage.

Outre le fait que ces méthodes peuvent demander un modèle mathématique complexe, ces méthodes sont parfois difficiles dans le cas de cohortes de petite taille où chaque individu est un « singleton », c'est-à-dire unique au regard des autres. Par ailleurs, les « outliers » de populations (données isolées du reste des observations) plus grandes peuvent aussi disparaître par ce phénomène qui aura tendance à concentrer la population autour de données « moyennes ».

Pour pallier ce problème, des méthodes de suréchantillonnage, visant à augmenter le nombre d'échantillons de classes minoritaires et à prendre en compte leur difficulté d'apprentissage, ont été développées. La technique de suréchantillonnage minoritaire synthétique (SMOTE), introduite pour la première fois par l'équipe de D. K. Lakovidis [6], implique l'interpolation de points de données appartenant à des classes minoritaires dans leur espace de caractéristiques.

Cette approche a été étendue dans d'autres travaux où les auteurs ont proposé de suréchantillonner (produire de nouveaux échantillons non observés) près de la limite de décision (zone où les sous-groupes sont séparés par l'algorithme) en utilisant divers algorithmes de classifications (soit l'algorithme k plus proches voisins (k-Nearest Neighbor) [7], soit une machine à vecteur de support [8]), se concentrant ainsi sur les échantillons qui sont potentiellement mal classifiés afin d'affiner la méthode. D'autres méthodes ont également été proposées [9], [10], cependant, celles-ci sont difficiles à adapter à des données de hautes dimensions [11], [12]."

Remarque : la solution d'anonymisation proposée par Octopize a fait l'objet d'une évaluation par les services de la CNIL qui ont conclu à l'absence d'obstacles à la démonstration de la conformité aux critères européens relatifs à l'anonymisation des données sous réserve d'un paramétrage adéquat. De plus, une mise à jour des lignes directrices européennes en matière d'anonymisation est en cours de réécriture par le groupe de travail « article 29 ».

III. Une population de patients caractérisée globalement par l'identification de ses caractéristiques et reproduites en créant de nouveaux individus les partageant (jumeaux numériques notamment)

i. Mécanistique

L'utilisation des propriétés des systèmes réels dits mécanistiques et incluant par exemple des connaissances physiques, chimiques et biologiques permet de prédire des propriétés d'un organe et préparer les décisions médicales. Les simulations cardiaques [21-27] sont une réalisation majeure dans ce domaine. Ces équations sont souvent paramétriques et le choix des paramètres influe beaucoup sur la qualité du modèle. Ces paramètres représentant souvent des quantités physiques, il devient possible de faire une calibration de ces modèles en comparant leurs résultats à la réalité, via des données observées de l'organe, ainsi qu'aux données de la littérature. Une fois calibrés, ces modèles fournissent un jumeau d'une observation. On parle souvent de jumeau numérique, qui s'applique également dans le cadre aléatoire présenté ci-dessous. Il fournit une représentation digitale de l'organe, du patient, d'un phénomène (flux aux urgences par exemple).

De nombreux jumeaux mécanistiques simulés sont adaptés à des applications chirurgicales portant sur divers organes tels que le foie ou le cerveau, et peuvent également inclure les propriétés physiques des matériaux chirurgicaux des interventions chirurgicales elles-mêmes. Des extensions de ces simulations utilisant la réalité mixte (XR) apparaissent rapidement, à la fois dans la procédure de planification chirurgicale avec la réalité virtuelle et dans la salle d'opération avec la

réalité augmentée. Les initiatives majeures des grandes entreprises technologiques pourraient s'étendre aux applications médicales. Le Metaverse [38] fait allusion aux applications médicales par le biais de partenariats et d'annonces commerciales (principalement sur les opérations chirurgicales) et, très probablement à plus court terme, l'Omniverse [39], qui fusionne des simulations physiques approximatives, un rendu avancé, de nombreuses initiatives du Monai [40] pour fournir un ensemble d'applications médicales.

ii. Modèles statistiques d'apprentissage sans Deep Learning

1) Méthodes de type ABC : simulation de données selon un modèle calibré au fur et à mesure par confrontation au réel

Les méthodes de calcul bayésien approximatif (ABC), également connues sous le nom de techniques sans vraisemblance (*likelihood free methods*), sont apparues au cours des trente dernières années comme l'approche la plus satisfaisante des problèmes de vraisemblance incalculables.

Elles offrent une solution presque automatique aux difficultés rencontrées avec les modèles qui sont complexes, mais à partir desquels il est possible d'effectuer des simulations. Elles ont été proposées pour la première fois en génétique des populations par Tavaré *et al.* (1997) [91], qui ont introduit les méthodes ABC en tant que technique contournant le calcul de la fonction de vraisemblance par le biais d'une simulation à partir de la distribution correspondante du modèle statistique génératif. Toutefois, ces méthodes souffraient dans une certaine mesure de difficultés de calibration qui les rendaient plutôt instables dans leur mise en œuvre et donc peu utilisées. Diverses améliorations et des extensions apportées à l'algorithme ABC original ont permis de rendre ces méthodes plus robustes et sont toujours au cœur de nouvelles recherches [70, 71, 72].

Le principe est simple : un modèle génératif paramétrique est construit pour mimer un phénomène, c'est-à-dire que des paramètres sont simulés *a priori* puis des observations sont générées à partir d'eux. Les données artificielles les plus proches des données observées (réelles) pointent vers un ensemble de paramètres initiaux meilleurs que les autres. C'est la partie calibration. Ensuite, le modèle est en mesure de générer des données selon cette configuration.

2) Inférences basées sur des simulations et inférences amorties

Récemment, les méthodes ABC se sont vues augmentées par les avancées en Machine Learning ou apprentissage automatique. Ainsi les inférences basées sur des données simulées et l'amortissement des inférences viennent accélérer l'analyse et la génération de données.

Le principe fondamental de ces méthodes est une utilisation de données simulées afin de créer et optimiser la procédure de génération de données virtuelles. Elle nécessite d'avoir déjà un modèle génératif de données, même approximatif a priori qui sera optimisé au cours de l'apprentissage de l'algorithme.

Les extensions de ces approches seront alimentées par les progrès récents des inférences basées sur la simulation [68]. Bien qu'elles reposent sur des principes similaires à ceux des approches ABC, elles visent à faciliter la procédure d'inférence en entraînant les réseaux neuronaux sur des simulations afin de pouvoir les exécuter directement et plus rapidement sur des données expérimentales.

3) Méthodes type estimation d'atlas

La particularité des méthodes de type estimation d'atlas est qu'elles se concentrent sur un postulat : une population peut être représentée par un élément moyen autour duquel varient les données. À partir de cette hypothèse, les populations sont décrites à deux échelles : une échelle populationnelle, par la moyenne et la variance interindividuelle, et à l'échelle individuelle, par la variation spécifique de chaque individu par rapport à cette moyenne. Les moyennes et variances (pouvant être des quantités complexes : des images, des formes, des réseaux de régulation pour les moyennes et des variations de formes, de structures pour les variances) sont des quantités inconnues qu'il faut estimer à partir des observations. Il ne s'agit plus de comparer des simulations a priori aux données réelles, mais de maximiser la vraisemblance des observations qui va pointer vers un (ou un ensemble de) paramètre optimal. C'est la partie apprentissage ou calibration du modèle. Une fois cette étape faite, les modèles peuvent générer des données artificielles.

Ces techniques ont été utilisées à de nombreuses reprises pour estimer des éléments caractéristiques de populations ainsi que sa variabilité « normale » tant pour des données de formes, d'images³ que pour des processus longitudinaux⁴.

³ www.deformetrica.org

⁴ <https://disease-progression-modelling.github.io/pages/main.html>

Ceci a permis de générer des cohortes de patients artificiels en grand nombre avec différentes modalités (voici un exemple⁵ où un million de sujets ont été générés, chacun avec des scores cognitifs, son atrophie corticale et hippocampique et sa dynamique métabolique). Ces avancées sont au cœur de la technologie de détection anticipée des maladies liées à des pertes de mémoires exploitée par www.docteurmemo.fr [73].



Stanley Durrleman

Directeur de recherche Inria à l'Institut du Cerveau,
Co-fondateur et CEO de Qairnel (la clinique du Docteur
Mémo)
Chaire PR[AI]RIE

**Les modèles de progression de
maladies génèrent des données
artificielles dont on montre
maintenant qu'elles permettent
d'augmenter la puissance
statistique des essais cliniques en
les combinant avec les données
réelles observées.**



⁵ <https://project.inria.fr/digitalbrain>

iii. Techniques de l'apprentissage profond ou Deep Learning

1) GAN based

L'augmentation récente des performances des modèles génératifs tels que les réseaux adversaires génératifs (GAN) [13] ou les autoencodeurs variationnels (VAE) [14], [15] en a fait des modèles très attrayants pour l'analyse des données. Les GAN ont déjà été largement utilisés dans de nombreux domaines d'application [16], [17], [18], [19], [20], y compris en médecine [21]. Par exemple, les GAN ont été utilisés sur des images de résonance magnétique (IRM) [22], [23], de tomodensitométrie (CT) [24], [25], de radiographie [26], [27], [28], de tomographie par émission de positrons (PET) [29], de spectroscopie de masse [30], de dermoscopie [31] ou de mammographie [32], [33] et ont donné des résultats prometteurs.

Néanmoins, la plupart de ces études portaient sur un ensemble d'entraînement assez important (plus de 1 000 échantillons d'entraînement) ou sur des données de dimension assez faible, alors que dans les applications médicales quotidiennes, il reste très difficile de rassembler des cohortes aussi importantes de patients étiquetés. Par conséquent, à ce jour, le cas des données à haute dimension combinées à une taille d'échantillon très faible reste peu exploré.

2) VAE (Variational Auto-Encoders) based

Comparés aux GAN, les VAE n'avaient suscité qu'un moindre intérêt pour l'augmentation de données et avaient été principalement utilisés pour des applications vocales [34], [35], [36]. L'utilisation de ces modèles génératifs sur des données médicales pour des tâches de classification [37], [38] ou de segmentation [39], [40], [41] se développe. Mal maîtrisés, ces modèles peuvent produire des échantillons flous et imprécis. Cet effet indésirable était encore plus marqué lorsqu'ils étaient entraînés avec un petit nombre d'échantillons, ce qui les rendait très difficiles à utiliser pour effectuer l'augmentation de données dans le cadre d'une dimension (très) élevée à faible taille d'échantillon (HDLSS, pour High Dimension Low Sample Size). Des travaux récents démontrent que les VAE peuvent être utilisées pour l'augmentation des données de manière fiable, même dans le contexte des données HDLSS, à condition d'apporter une certaine modélisation de l'espace latent et de modifier la manière dont sont générées les données.

La plateforme ORIGA de BOTdesign exploite la puissance des VAE en partenariat avec Pr. Stéphanie Allassonnière. Des publications existantes et en cours viennent appuyer la pertinence de l'utilisation de cette technologie dans la génération des données artificielles appliquées à la santé.

3) Denoising diffusion

Les modèles de diffusion à partir d'un modèle de bruit utilisent des équations différentielles stochastiques (SDE) et un processus d'ajout puis de suppression itérative du bruit pour générer de nouvelles images. Le modèle transforme essentiellement la tâche de génération d'images en une tâche de débruitage par le biais d'un processus de diffusion. Bien que les applications médicales soient plus récentes [17-21], ces modèles présentent le même niveau d'efficacité que la génération d'images généraliste.

4) Futur de l'utilisation de ces modèles de génération de données

La numérisation à grande échelle des dossiers des patients a ouvert la voie à la création de dossiers virtuels efficaces. Diverses approches ont permis de synthétiser ces rapports ou Dossiers Patients Numérique (EHR, Electronic Health Record) en utilisant principalement des GAN [41-44] et, plus récemment, des modèles de diffusion [45,46] et des VAE [47,48]. Les débats sont très actifs [49,50] sur les caractéristiques pertinentes à explorer pour garantir l'utilité médicale de ces données générées. Lorsque ces données sont générées dans le contexte d'études longitudinales, de nombreuses approches [44,51-54] (s'appuyant grosso modo sur les trois mêmes architectures principales) ont vu le jour ces dix dernières années. Elles doivent être mises en relation avec des approches statistiques destinées à modéliser les données longitudinales [55-57] pour asseoir leur précision. Au-delà de ces initiatives, les récentes percées dans le domaine des modèles de génération de texte [58-61], aussi appelés Large Language Model (LLM), sont appelées à jouer un rôle important en fournissant aux données artificielles des mises à jour quasi hebdomadaires sur les propriétés de ces modèles et les possibilités de réglage fin.

La plupart des modèles génératifs actuels présentent des limites liées au fait qu'ils génèrent principalement des types uniques de données cliniques pour un patient (imagerie, rapports médicaux, plus généralement l'ensemble des EHR). Il existe encore peu d'approches complètes de synthèse multimodale qui couplent ces différents types de données, ce qui poserait d'énormes problèmes pour l'évaluation de leurs propriétés pertinentes. Cependant, de récentes avancées techniques adaptées pour l'instant à l'analyse médicale intègrent désormais des données multimodales. Med-PALM [62] est un modèle génératif multimodal qui code et analyse les données biomédicales, y compris le langage clinique, l'imagerie et la génomique.

Le modèle JNF-VAE [90] utilise des VAE et des flots normalisant pour proposer une représentation des données multimodales en une seule population et conjointement des représentations de chaque modalité conditionnellement à toute autre combinaison des autres. De nombreuses initiatives [63-66] intègrent désormais des multiples modalités dans le processus d'analyse des données. Il est probable que similaires au modèle de fondation [67] domineront bientôt ces initiatives avec les défis habituels associés à la propagation d'erreurs imprévisibles, aux anomalies en matière de protection de la vie privée et à l'absence d'analyse des données.

Différentes méthodes peuvent donc être utilisées pour créer des données artificielles. La typologie des données initiales, l'objectif poursuivi et l'expertise des équipes guideront le choix de la technologie à retenir pour augmenter les données.



Jean-Baptiste Masson
Chercheur à l'Institut Pasteur
Chaire PR[AI]RIE
Co-fondateur d'AVATAR MEDICAL

Les cohortes artificielles offriront la possibilité d'étendre le domaine des inférences basées sur les simulations au domaine médicale. Ainsi des méthodes développées sur simulations pourront être testées et améliorées sur les données médicales réelles.



USAGE DES COHORTES DE PATIENTS ARTIFICIELS

03



**Toward regulatory framework for
the use of in silico trials and virtual
cohorts as clinical evidence**

Sarah Zohar

Directrice de Recherche, Inserm
Responsable de l'équipe de recherche "HeKA" - Inserm, Inria,
Université Paris Cité



Dans le chapitre précédent, nous avons décrit la possibilité d'augmenter des données d'imagerie, tabulaires et génétiques de façon transversale, longitudinale ou multimodale qui seront utilisées pour les différentes applications précitées.

Nous allons aborder successivement ces trois thématiques.

I. Utilisation des données artificielles pour l'accélération de la recherche clinique et la mise sur le marché de produits de santé (médicaments et Dispositifs Médicaux (DM))

Les études cliniques sont des études scientifiques visant à documenter l'efficacité et la tolérance d'un médicament ou d'un dispositif médical.

Les études cliniques sont indispensables à l'évaluation d'un nouveau produit par les autorités de santé, qu'il s'agisse de l'accès au marché ou de l'obtention d'un remboursement, le cas échéant, mais également pour éclairer les utilisateurs. La possibilité de mise en place d'études en France représente une chance à l'échelle nationale et individuelle, le bras expérimental étant présumé aussi ou plus efficace que le bras contrôle. Le bras contrôle devrait être systématiquement le « Standard of Care » (SoC), c'est-à-dire le soin administré en pratique courante selon les guidelines du moment ou Bonnes Pratiques Cliniques (BPC) ; si le besoin médical est non couvert, il pourra s'agir d'un bras placebo. Ces études sont encadrées, soumises à autorisations et répondent à une éthique pour ne pas engendrer de perte de chance pour les patients.

L'inclusion d'un nombre suffisant de patients dans l'essai est un des facteurs qui garantit la puissance statistique et la robustesse des résultats, sans garantir la pertinence clinique de ces résultats. La taille d'effet attendue est également un élément clé dans un protocole : plus elle est importante et moins il faudra de patients. Environ 80% des essais cliniques ne parviendraient pas à recruter le nombre de patients nécessaires dans les délais attendus et 55% seraient arrêtés prématurément faute de recrutement patients.⁶

Les difficultés de recrutement et de maintien des patients dans les essais sont particulièrement aiguës pour différentes raisons.

Dans la recherche clinique, la génération de données artificielles permettrait d'aider au design pertinent des essais cliniques, de constituer des cohortes de patients artificiels qui pourraient renforcer les bras contrôles d'études de phase 3 et/ou les patients inclus dans les phases 2.

⁶ Desai M. Recruitment and retention of participants in clinical studies: Critical issues and challenges. *Perspect Clin Res.* 2020 Apr-Jun;11(2):51-53. doi: 10.4103/picr.PICR_6_20. Epub 2020 May 6. PMID: 32670827; PMCID: PMC7342339.

Dans la recherche clinique, la génération de données artificielles permettrait d'aider au design pertinent des essais cliniques, de constituer des cohortes de patients artificiels qui pourraient renforcer les bras contrôles d'études de phase 3 et/ou les patients inclus dans les phases 2.

Actuellement, des essais ayant recourt à des patients synthétiques (et non encore artificiels – voir Lexique) ont déjà été autorisés par les autorités de régulation (FDA ou *Food and Drug Administration*, EMA ou *European Medicines Agency*) pour la mise en œuvre d'essais cliniques.

On peut citer, par exemple, l'Alecensa® (alectinib)⁷, développé par le laboratoire Roche, indiqué en monothérapie dans le traitement du cancer du poumon non à petites cellules (CPNPC).

Ce médicament avait reçu une approbation conditionnelle par l'EMA en 2017. En constituant un bras de contrôle actif synthétique (67 patients recevant le ceritinib), la compagnie a pu fournir à l'EMA des preuves d'efficacité par rapport à un traitement de référence. Le niveau de preuve a été reconnu comme robuste et a permis la mise sur le marché de l'alectinib 18 mois plut tôt que s'il avait fallu attendre les données d'une étude clinique « traditionnelle ».

Le rôle du bras contrôle synthétique dans l'obtention de l'AMM conditionnel devra être précisé.

L'utilisation de données artificielles pourrait jouer un rôle identique à ces données synthétiques, avec l'avantage d'être temporellement équivalentes au bras traité.

On peut également envisager que ces données artificielles puissent un jour compléter les bras actifs des phases 3, permettant d'augmenter la puissance statistique des tests comparatifs de deux bras, et les patients. Dans les phases 4, les données artificielles pourraient aussi permettre de rendre plus « visibles » des signaux faibles.

Nous avons listé ci-dessous des situations qui nous ont semblé être des cas d'usage où les données artificielles pourraient permettre de lever certains verrous. Cette liste, non exhaustive, montre déjà les multiples enjeux auxquels ces patients artificiels pourraient répondre.

⁷ <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7218288/#CIT0009>

i. Usages

1) Favoriser la mise en place d'études (faciliter le recrutement des patients)

Dans de nombreuses situations, le recrutement d'un nombre suffisant de patients nécessaire pour mener une étude d'envergure (phase 3) présente un véritable défi. Un exemple évident est celui des essais dans le domaine des maladies rares, car les patients sont, par définition, rares, mais cette situation est par ailleurs de plus en plus fréquente avec les pratiques de médecine de précision qui visent à traiter une population très précise de patients caractérisée par un ensemble de biomarqueurs et donc définis par des critères d'inclusion et d'exclusion très stricts. Cette difficulté de recrutement se retrouve aussi dans les essais pédiatriques, qui nécessitent des phases de recrutement toujours très longues.

Par ailleurs, la génération de patients artificiels pourrait aussi permettre d'enrichir une cohorte avec des patients sous-représentés, permettant ainsi d'augmenter la diversité de ces cohortes. Cela concernerait des populations difficilement accessibles (certaines ethnicités) ou fragiles (femmes enceintes, personnes âgées / polymédiquées / immunodéprimées / en rémission de cancer, etc.). Or, l'inclusion de ces patients et l'analyse de leurs caractéristiques au sein des essais constituent un enjeu majeur pour permettre le développement d'une médecine prédictive et personnalisée. Cependant, l'inclusion de ces populations dites « vulnérables » ne peut intervenir qu'à certaines conditions décrites dans les articles L.1121-5 et suivants du Code de la Santé Publique.

2) Accélérer la mise sur le marché de nouveaux traitements

En diminuant le nombre de patients à recruter dans une première étude, si la robustesse de ces nouvelles méthodologies est établie, la génération de patients artificiels permettrait de réduire le coût et la durée des études.

Les fonds dégagés pourraient alors être réinvestis dans la recherche de nouveaux traitements innovants et le temps gagné permettrait de rendre ces traitements accessibles aux patients plus rapidement.

3) Faciliter les inclusions dans un essai (notamment lorsque les difficultés d'inclusion dans le bras contrôle sont connues)

Afin de démontrer l'efficacité ou la supériorité d'un traitement il faut pouvoir démontrer une différence d'effet par rapport à une stratégie de référence (selon les cas, absence de traitement, traitement de référence), la constitution d'un bras

contrôle dans les essais est donc incontournable. Cela pose toujours des questions éthiques délicates quant aux traitements concomitants autorisés afin de ne pas biaiser les analyses.

Dans certaines études, des bras « d'échappement » sont prévus pour permettre aux patients qui souffrent trop d'être passés sous traitement. Cela présente un risque pour l'étude, car si le nombre de patients arrivant au terme de cette période placebo est insuffisant, le critère d'évaluation principal ne pourra pas être évalué avec la puissance statistique nécessaire.

Sous réserve de la validation de ces méthodes, la perspective de pouvoir remplacer les patients réels sous placebo par des patients artificiels pourrait résoudre ces dilemmes. Cela pourrait aussi permettre de comparer l'évolution du bras traité par un bras non traité sur une plus longue période, ce qui pourrait être très informatif.

De plus, si les patients, comme les médecins acceptant ces études, n'avaient plus cette appréhension de potentiellement recevoir un traitement non actif, cela faciliterait de facto le recrutement.

Il faut noter cependant, que les patients sachant qu'ils reçoivent forcément un traitement auront potentiellement un effet placebo plus conséquent. Pour cette raison, il sera sans doute préférable de ne pas substituer complètement le bras placebo par un bras entièrement composé de patients artificiels.



Sylvie Troy

Directrice Médicale Adjointe et RWD/RWE Lead France
Pfizer

Les données virtuelles sont naissantes mais pourraient modifier l'organisation et le développement de l'innovation dans les industries de santé, en raccourcissant le temps de développement des nouvelles molécules en particulier dans des pathologies rares



ii. Points d'attention pour l'utilisation de données artificielles dans le cadre de la recherche clinique

1) Définir ce dont doit être composé un patient artificiel

Tout au long de ce livre blanc, nous parlons de patients artificiels, mais ce dont est composé un patient artificiel reste à définir, et dépendra probablement des situations.

Il faudra, en concertation avec les experts pertinents, les paramètres à générer (potentiellement conditionné par le nombre maximal de dimensions qui pourront être gérées par le modèle utilisé), le format de ces paramètres, le nombre de time points nécessaires...

Le plus gros défi résidera sans doute dans la génération des potentiels événements indésirables.

2) Pertinence des données

La première exigence, dont l'importance est évidente, est de s'assurer que les jeux de données utilisés pour entraîner le modèle qui permettra la génération de patients artificiels représente bien la population à augmenter. Ces jeux de données doivent représenter :

a) La pathologie étudiée

Comme évoqué, la médecine de précision requiert de plus en plus de se concentrer sur des sous-catégories de patients au sein même d'une pathologie. Il faudra donc s'assurer que les patients réels correspondent bien à ces populations précises.

b) Les spécificités de la population à augmenter

De la même manière, s'il s'agit d'enrichir la diversité d'une cohorte, il faudra trouver des jeux de données représentatifs des spécificités physiologiques de ces populations (ethnies, patient immunodéprimé, femme enceinte, enfant/adolescent, personnes âgées,...). Cela représentera un défi certain étant donné que ces données peuvent s'avérer rares et qu'elles devront, en sus, toujours représenter la pathologie étudiée.

La création d'une cohorte de patients artificiels à partir de données de patients réels, suivis dans le cadre de l'anesthésie, confirme le potentiel de l'intelligence artificielle dans les processus de soins et de recherche clinique. En tant que professionnels de la santé et scientifiques, nous travaillons à la validation de ces technologies en routine pour les mettre à disposition de nos confrères et des patients les plus fragiles. La collaboration entre les domaines mathématiques, statistiques et médicaux pour développer des outils d'intelligence artificielle fiables, sécurisés et éthiques préfigure l'avènement d'une médecine toujours plus personnalisée et humaine



Dr. Fabrice Ferré
PhD, Anesthésiste
CHU de Toulouse



c) L'époque au cours de laquelle l'essai est mené

Sur le long terme, il peut y avoir une évolution naturelle de certaines maladies (exemples de la diminution de la fréquence des infections vs augmentation des terrains allergiques, diminution non expliquée de la fréquence de certaines formes de spondylarthrites,...) et sur le court terme, on peut faire face à des situations particulières, telle la pandémie COVID pour donner un exemple parlant (contexte infectieux, mais aussi impact sur la santé mentale).

3) La conformité des données sources avec le dernier état des connaissances scientifiques

Comme évoqué dans le point précédent, l'époque et l'environnement au cours de laquelle se déroule l'essai peut impacter l'efficacité du traitement évalué (surestimé ou sous-estimé). De même, les traitements concomitants tolérés dans certaines études varient au fil des époques. Les bases de données utilisées pour générer les patients artificiels devront être le plus proche possible du contexte dans lequel se déroule l'essai.

Lors de l'évaluation du brodalumab, un nombre anormalement accru de suicides avait été constaté au cours des phases 3. Biologiquement, la raison pour laquelle l'inhibition du récepteur à l'IL-17 pourrait augmenter les pensées suicidaires ainsi que les passages à l'acte était très énigmatique. Une hypothèse émise a été qu'en parallèle de cette étude, les États-Unis ont traversé la crise des subprimes. Pour autant, il n'a pas été possible de formellement démontrer que cela expliquait cet effet indésirable sévère. Cet exemple est extrême, et n'aurait pas pu être anticipé, mais il permet d'illustrer l'importance d'avoir un bras comparateur au plus proche du contexte au cours duquel a lieu l'essai.

4) Les critères d'éligibilité

Comme mentionné dans les points précédents, les données réelles utilisées pour générer des données synthétiques devront être représentatives de la population étudiée. Une des difficultés rencontrées en ce sens, sera de s'assurer que ces patients respecteront les critères d'éligibilité des études : sévérité de la maladie, traitements antérieurs, biomarkers spécifiques, etc. Certains de ces critères seront peut-être difficiles à identifier dans les bases de données utilisées pour générer des données artificielles.

5) La qualité de la base de données de référence

Il sera nécessaire de calibrer la quantité de données nécessaire et suffisante pour capter les signaux faibles dans la population à augmenter. Il peut s'avérer difficile de trouver une base suffisante permettant une représentativité de la variabilité interindividuelle. Ces analyses préliminaires devront permettre de s'orienter vers le choix optimal entre bras synthétique (et donc rétrospectif) et bras augmenté dans le cas de la possibilité d'une représentativité suffisante de la population cible.

Ces points d'attention seront complétés dans le chapitre preuves et méthodologie pour la validation des produits de santé.



Hervé Nabarette

Directeur adjoint Affaires Publiques
AFM-Téléthon

La création raisonnée de données et patients artificiels peut contribuer dans l'avenir à diminuer le nombre de patient nécessaire dans le bras comparateur des essais cliniques, et à accélérer ces essais, en particulier dans les maladies rares. Elle doit contribuer à encourager et pas désinciter les partages et mises en commun de données réelles. Suite à la maturation scientifique en cours sur le sujet, les agences réglementaires devront organiser les ateliers et colloques utiles, puis mettre en place des recommandations et avis scientifiques destinés aux développeurs de thérapeutiques.



II. Utilisation des données artificielles pour l'entraînement des algorithmes d'intelligence artificielle (calibration)

La calibration, appelée estimation, des modèles mathématiques à l'aide du Machine Learning ou du Deep Learning nécessite des données en quantité suffisante, fiables et de bonne qualité, ce qui implique qu'elles soient suffisamment représentatives de la population d'étude. Ceci sous-entend souvent des données multicentriques difficilement mobilisables, induisant des démarches consommatrices de temps et de moyens, lorsqu'il s'agit de données de santé qui ne peuvent pas facilement être utilisées en dehors de leur établissement de collecte.

Afin d'obtenir des données multicentriques, l'émergence de l'apprentissage fédéré (Federated Learning) est plein de promesses. Il permet d'entraîner des modèles sans centraliser les données sur un même serveur et donc de les conserver en local. Cependant, il est pour l'instant encore un sujet de recherche et son déploiement n'est pas encore une évidence à court terme. En effet, l'apprentissage fédéré peut présenter des risques de sécurité importants sur les modèles appris avec par exemple un risque d'empoisonnement des données⁸. Le recours aux données artificielles permettrait de créer des bases de données locales, artificielles, non ré-identifiables, qui pourraient être colligées pour former une base de données artificielles multicentrique, représentative. Cela représenterait un gain de temps et d'argent conséquent permettant une mise à disposition plus rapide des nouveaux outils d'aide à la décision pour les soignants et les patients.

Une autre situation où ces données pourraient apporter un gain de temps et d'expertise : si l'on considère, par exemple, que pour entraîner un algorithme de reconnaissance d'images en radiologie, il faut collecter et annoter (par plusieurs radiologues expérimentés) 10 000 clichés, l'utilisation de données artificielles permettrait de ne collecter qu'une partie de cette base qui serait ensuite augmentée. Il s'agirait de ne collecter par exemple que 2 000 clichés à annoter, puis d'augmenter ces images correctement annotées pour arriver aux 10 000 recherchées. La collecte de ces données peut être longue et coûteuse, en particulier à cause d'une contractualisation difficile et des clauses de partage de propriété intellectuelle encore mal encadrées pour ces usages. De même, l'annotation par des experts est un coût pour le développement de ces nouveaux dispositifs. Le gain que représente la diminution des données à collecter et annoter (2 000 plutôt que 10 000), nous montre un nouvel intérêt des méthodes d'augmentation de données.

L'industrie du diagnostic est également concernée par les données artificielles. Ainsi, les GAN ont été utilisés avec succès par des chercheurs américains pour améliorer la prévision du diagnostic du diabète de type 2⁹.

⁸ <https://cybersecurity.springeropen.com/articles/10.1186/s42400-021-00105-6>

⁹ <https://pubmed.ncbi.nlm.nih.gov/37873778/>



Vinh-Phuc Luu

Responsable du pôle Épidémiologie et Innovation
Filière Intelligence Artificielle et Cancers

Les cohortes de patients virtuels sont un exemple emblématique de comment l'IA peut innover dans la génération des éléments de preuve et être complémentaire de la recherche « classique », y compris dans les situations où les patients et les données sur leur maladie sont rares.



PROTECTION DES DONNÉES ET GÉNÉRATION DE COHORTES DE PATIENTS ARTIFICIELS

04

Les algorithmes permettant de générer des cohortes de patients artificiels se nourrissent de données de « patients réels » collectées dans le cadre de la prise en charge ou de précédentes recherches. Lorsqu'elles présentent un caractère personnel, le Règlement général sur la protection des données (RGPD)¹⁰ et la loi n° 78-17 du 6 janvier 1978 modifiée (dite « loi informatique et libertés ») auront vocation à s'appliquer, sous réserve que ces traitements relèvent du champ d'application matériel et territorial du règlement et de la loi.

Par ailleurs, selon le contexte d'utilisation de ces systèmes, d'autres dispositions sectorielles issues du droit national (Code de la santé publique et Code pénal) et européen (notamment les règlements sur les essais cliniques¹¹ et les dispositifs médicaux¹²) pourraient avoir également vocation à s'appliquer.

Enfin, à ce paysage réglementaire s'ajouteront prochainement les dispositions du règlement européen sur l'intelligence artificielle dit « IA act » une fois celui-ci entré en vigueur et plus particulièrement celles applicables aux systèmes d'intelligence artificielle à haut risque. Il en ira ainsi notamment des systèmes intégrés dans certains dispositifs médicaux. Avant leur mise sur le marché de l'Union européenne ou en service, la conformité de ces systèmes devra être évaluée afin de démontrer qu'ils respectent certaines exigences obligatoires (par exemple, la qualité des données, la documentation et la traçabilité, la transparence, le contrôle humain, l'exactitude, la cybersécurité et la robustesse). Cette évaluation devra être répétée en cas de modification substantielle du système ou de sa finalité. Par ailleurs, les fournisseurs de systèmes d'intelligence artificielle à haut risque devront mettre en œuvre des systèmes de gestion de la qualité et des risques afin de garantir qu'ils respectent les nouvelles exigences et de minimiser les risques pour les utilisateurs et les personnes concernées, même après la mise sur le marché du produit.

Le présent chapitre n'aura toutefois vocation à aborder que les questionnements en lien avec la réglementation applicable en matière de traitements de données à caractère personnel.

¹⁰ Règlement (UE) 2016/679 du Parlement européen et du Conseil du 27 avril 2016, relatif à la protection des personnes physiques à l'égard du traitement des données à caractère personnel et à la libre circulation de ces données, et abrogeant la directive 95/46/CE (règlement général sur la protection des données).

¹¹ Règlement (UE) 536/2014 du Parlement européen et du Conseil du 16 avril 2014 relatif aux essais cliniques de médicaments à usage humain et abrogeant la directive 2001/20/CE

¹² Il en va ainsi notamment du règlement (UE) 2017/745 du Parlement européen et du Conseil du 5 avril 2017 relatif aux dispositifs médicaux et du règlement (UE) 2017/746 du Parlement européen et du Conseil du 5 avril 2017 relatif aux dispositifs médicaux de diagnostic in vitro. Il est à noter que certaines dispositions nationales spécifiques sont également prévues dans le code de la santé publique (articles L.1125-1 et suivants).



Les cohortes virtuelles sont un des leviers pour accélérer le développement des innovations technologiques. Ce livre blanc ouvre des pistes pour impulser les conditions de leur validation.

Corinne Colignon

Cheffe de service Mission Numérique en Santé
Haute Autorité de Santé



I. La qualification des données d'apprentissage et des données générées

Afin d'identifier le cadre juridique applicable, il convient de qualifier les données d'apprentissage, le modèle d'intelligence artificielle et les données générées.

i. La notion de traitements de données à caractère personnel

D'une part, un « traitement » désigne toute opération portant sur des données personnelles telle que l'enregistrement, la structuration, la conservation, l'adaptation, la modification, l'extraction, la consultation, l'utilisation, la diffusion ou la mise à disposition. A ce titre, l'entraînement d'un modèle d'intelligence artificielle constitue ainsi un traitement au sens de l'article 4 du RGPD.

D'autre part, une « donnée personnelle » est, selon cette même disposition, « toute information se rapportant à une personne physique identifiée ou identifiable ». Cette personne « peut être identifiée, directement ou indirectement, notamment par référence à un identifiant, tel qu'un nom, un numéro d'identification, des données de localisation, un identifiant en ligne, ou à un ou plusieurs éléments spécifiques propres à son identité physique, physiologique, génétique, psychique, économique, culturelle ou sociale ». Ces données peuvent être directement identifiantes (rattachées au prénom et au nom de la personne) ou indirectement identifiantes (rattachées à un numéro d'ordre ou à un code alphanumérique, il s'agit alors de données pseudonymisées¹⁵).

¹⁵ Définie à l'article 4 du RGPD, la pseudonymisation est « le traitement de données à caractère personnel de telle façon que celles-ci ne puissent plus être attribuées à une personne concernée précise sans avoir recours à des informations supplémentaires, pour autant que ces informations supplémentaires soient conservées séparément et soumises à des mesures techniques et organisationnelles afin de garantir que les données à caractère personnel ne sont pas attribuées à une personne physique identifiée ou identifiable ».

Le RGPD accorde une protection particulière à certaines catégories de données mentionnées au sein de son article 9 (aussi appelées « données sensibles »). Il en va ainsi des données de santé¹⁶ définies à l'article 4 du RGPD comme « *les données à caractère personnel relatives à la santé physique ou mentale d'une personne physique, y compris la prestation de services de soins de santé, qui révèlent des informations sur l'état de santé de cette personne* ». Le traitement de ces données est, par principe, interdit, sauf ce que puisse être mobilisée l'une des exceptions prévues par le RGPD ou la loi « informatique et libertés » telles que le recueil du consentement explicite de la personne concernée, la nécessité de traiter les données à des fins de prise en charge sanitaire ou de recherche scientifique. Ce régime juridique s'explique par le fait que le traitement de ces catégories de données particulières pourrait engendrer des risques importants pour les libertés et les droits des personnes concernées.

ii. La notion d'anonymisation

L'anonymisation est un traitement de données personnelles produisant des informations qu'il n'est pas possible, par des « *moyens raisonnables* », de rattacher à des personnes identifiées ou identifiables. La législation relative à la protection des données reste applicable s'agissant du processus d'anonymisation, mais ne s'applique plus aux données issues de ce processus dès lors que leur caractère anonyme a été démontré.

Le processus d'anonymisation ayant pour effet de réduire l'information présente dans le jeu de données d'origine, il convient de le construire de sorte à préserver le plus possible les informations de forte valeur tout en acceptant d'atténuer les autres. Pour être utile, ce processus est donc susceptible d'être très différent en fonction de la nature du jeu de données d'origine et des besoins métier.

En pratique, deux familles de techniques d'anonymisation peuvent être combinées à cette fin :

- La perturbation aléatoire (ou randomisation) consiste à modifier les attributs dans un jeu de données de telle sorte qu'elles soient suffisamment incertaines pour en altérer l'exactitude. Cette incertitude affaiblit donc le lien entre une donnée et la personne physique à laquelle elle se rapporte, tout en conservant les propriétés statistiques globales du jeu de données. Cette famille de techniques permet de protéger le jeu de données du risque d'inférence. Par exemple : ajouter un bruit aléatoire dans les données collectées ($\pm$ deux centimètres à la taille des patients) ; permuter les valeurs d'un attribut entre les patients (échanger la donnée « âge » entre deux patients choisis de manière aléatoire) ;

¹⁶ Pour en savoir plus, voir la fiche pratique publiée sur le site web de la CNIL.

- La généralisation consiste à modifier l'échelle ou l'ordre de grandeur des attributs d'un jeu de données, afin de rendre ces dernières communes à un groupe plus large de personnes. Elle contribue ainsi à rendre les informations moins spécifiques à des personnes tout en préservant leur utilité. Cette famille de techniques permet d'éviter l'individualisation d'un jeu de données. Elle limite également les possibles corrélations du jeu de données avec d'autres (cf. ci-dessous). L'agrégation de données (comptage, moyenne, etc.) et les techniques conférant au jeu de données les propriétés de k-anonymat, l-diversité et t-proximité sont des exemples de méthodes de généralisation.

La seule mobilisation de ces techniques ne permet pas à elle seule de conclure qu'un jeu de données est anonyme. À titre d'exemple, les systèmes d'intelligence artificielle sont constitués en principe de « modèles » statistiques entraînés sur des données réelles (ultima ratio), qui peuvent notamment être des données de santé pseudonymes. Ces modèles comportent donc eux-mêmes des informations issues des données dont le système d'intelligence artificielle est nourri et sont donc susceptibles de permettre une identification des personnes présentes dans les bases de données d'apprentissage. Ils sont ainsi potentiellement sujets à des attaques telles que les attaques par reconstruction ou les attaques par inférence d'attribut ou d'appartenance¹⁷ qui permettent de remonter aux individus. Il n'est donc a priori pas possible de considérer par défaut qu'un modèle d'intelligence artificielle est anonyme.

En 2014¹⁸, les autorités de protection des données européennes ont dégagé trois critères qui, s'ils sont conjointement réunis, permettent de conclure que les données concernées sont anonymes :

1. **L'individualisation** : il ne doit pas être possible d'isoler un individu dans le jeu de données ;
2. **La corrélation** : il ne doit pas être possible de relier les données avec un autre jeu de données distinct concernant un même individu ;
3. **L'inférence** : il ne doit pas être possible de déduire, de façon quasi certaine, de nouvelles informations sur un individu.

Le processus d'anonymisation doit permettre de produire un jeu de données conforme aux trois critères susmentionnés. Cette conformité devra être documentée et démontrable.

Toutefois, si l'un des critères n'est pas satisfait, il convient de procéder à un examen plus approfondi afin d'apprécier si le RGPD et la loi « informatique et libertés » sont

¹⁷ Voir cette publication récente de l'institut des standards américains NIST : <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-2e2023.pdf>

¹⁸ Pour en savoir plus, voir les lignes directrices. Celles-ci sont en cours de mise à jour mais les trois critères exposés dans le présent document demeurent valides.

Ainsi, si ces trois critères ne sont pas parfaitement remplis, le responsable de traitement qui souhaite rendre un jeu de données anonyme doit démontrer, via une évaluation approfondie des risques d'identification, que la probabilité d'identification d'une personne à l'aide de moyens raisonnablement susceptibles d'être utilisés, par lui-même ou par toute autre personne, est négligeable¹⁹.

En conclusion, pour identifier le cadre juridique applicable au cas des cohortes constituées de patients virtuels engendrés à travers des techniques reposant sur l'intelligence artificielle, il convient de distinguer trois questions :

- Celle des données d'apprentissage : il s'agira le plus souvent de données personnelles. Toutefois, dans certains cas, au regard de la finalité poursuivie, il est concevable que les données nécessaires à l'entraînement du modèle d'intelligence artificielle ne contiennent aucune donnée identifiante et soient de granularité suffisamment grossière pour qu'il soit possible de démontrer qu'elles sont anonymes ;
- Celle des données des patient virtuels engendrées par des modèles d'intelligence artificielle : une analyse, au regard des trois critères mentionnés ci-dessus ou de manière ad hoc, devrait être menée pour démontrer que les données des patients virtuels ainsi constituées ne reproduisent pas des données de patients réels ;
- Celle du modèle d'intelligence artificielle lui-même qui, s'il est entraîné à partir de données personnelles, peut ne pas être anonyme et devrait faire l'objet d'une même analyse.

Dans l'hypothèse où il s'agirait de données personnelles, le cadre juridique suivant aura vocation à s'appliquer²⁰, en sus des éventuelles réglementations sectorielles mentionnées en introduction de ce chapitre.

II. Le respect des principes clés « informatique et libertés » à tous les stades du cycle de vie de l'algorithme

Les phases de développement et de déploiement d'un système d'intelligence artificielle constituent des traitements de données personnelles distincts. Par exemple, l'entraînement d'un modèle visant à produire des données de patients virtuels fait partie de la phase de développement, alors que son utilisation pour la génération de cohortes de patients virtuels dans un contexte hospitalier se fait dans une phase de déploiement. Comme décrit précédemment, il est possible qu'aucune de ces phases ne constitue de traitement de données personnelles, juste l'une d'entre elle ou les deux. Bien souvent et en première analyse, s'agissant de la constitution de cohortes de patients virtuels, un traitement de données sera mis en œuvre pour l'apprentissage du modèle de génération et les traitements relatifs à l'utilisation du modèle et à la création de données de patients

¹⁹ Pour en savoir plus, voir la [fiche pratique](#) publiée par la CNIL.

²⁰ Pour en savoir plus, voir la [fiche pratique](#) publiée par la CNIL.

virtuels seront considérés comme anonymes, sous réserve de respecter les conditions précédemment énoncées.

La conformité de chacun de ces traitements devra être documentée par leur responsable de traitement i.e. l'organisme qui, seul ou conjointement, détermine les finalités - les objectifs - et les moyens - la manière de parvenir aux objectifs - du traitement : le fournisseur de l'algorithme²¹ lors de son développement et son utilisateur lors de son déploiement, sauf en cas d'utilisation personnelle non commerciale²². Le respect de ces principes, qui peut parfois poser des questions nouvelles, permettra de générer des données fiables et de qualité. Leur prise en compte devra être anticipée le plus en amont possible, dans une logique de protection des données par défaut et dès la conception²³. Pour ce faire, les fournisseurs d'algorithmes pourront notamment être accompagnés par la Commission nationale de l'informatique et des libertés (CNIL) qui a d'ores et déjà publié un certain nombre de fiches pratiques²⁴ et un plan d'action sur l'intelligence artificielle²⁵.

i. Le principe de licéité du traitement

Comme tout traitement de données personnelles, la constitution et l'utilisation d'une base de données pour l'apprentissage de systèmes d'intelligence artificielle doit se fonder sur l'une des six bases légales prévues par l'article 6 du RGPD pour pouvoir être légalement mises en œuvre. Parmi les fondements juridiques susceptibles d'être mobilisés figurent le consentement, l'obligation légale, l'exécution d'un contrat, l'exécution de la mission d'intérêt public, la sauvegarde des intérêts vitaux, la poursuite des intérêts légitimes²⁶.

Les traitements de données personnelles mis en œuvre dans le cadre du développement et du déploiement de l'algorithme devront donc être fondés en toutes circonstances sur l'une des bases légales mentionnées ci-dessus. Il est à noter que le choix de la base légale retenue, qui obéit à des conditions spécifiques, a des conséquences sur les modalités d'exercice des droits des personnes concernées.

Par ailleurs, si le développement de l'algorithme implique de réutiliser des données déjà collectées (dans le cadre de la prise en charge médicale ou d'une précédente recherche), il conviendra de s'assurer que la base de données initiale a été

21 La personne physique ou morale, l'autorité publique, une agence ou un autre organisme qui développe ou fait développer un système d'IA en vue de le mettre sur le marché ou de le mettre en service sous son propre nom ou sa propre marque, que ce soit contre rémunération ou gratuitement.

22 Pour en savoir plus, voir la fiche pratique publiée par la CNIL.

23 Pour en savoir plus, voir la fiche pratique publiée par la CNIL.

24 Afin d'aider les organismes à développer des solutions respectueuses de la vie privée, la CNIL a publié de nombreuses fiches dont un guide d'auto-évaluation pour les systèmes d'intelligence artificielle.

25 Pour en savoir plus, voir le plan d'action.

26 Pour en savoir plus, voir la fiche pratique publiée par la CNIL concernant les bases légales et les spécificités liées au recours à l'intelligence artificielle.

régulièrement constituée et mise en œuvre (par exemple en s'assurant que le traitement initial a fait, si nécessaire, l'objet d'une formalité adéquate et que la durée de conservation de la base de données initiale n'a pas expiré).

ii. La poursuite d'une finalité déterminée, explicite et légitime

Tout traitement de données personnelles doit s'inscrire dans un objectif suffisamment déterminé et précis. Cette finalité doit être légitime au regard des missions du responsable de traitement. En pratique, dans le cadre d'une recherche en santé, cette finalité correspond à la question scientifique précise à laquelle entend répondre le projet. Elle est matérialisée au sein de l'objectif principal ainsi que, le cas échéant, des objectifs secondaires mentionnés dans le protocole de l'étude.

De la même manière, la constitution d'une base de données contenant des données personnelles pour le développement d'un système d'intelligence artificielle constitue un traitement de données personnelles qui doit poursuivre une finalité déterminée, explicite et légitime. Afin d'aider les responsables de traitement à définir cet objectif selon leurs cas d'usage, et notamment en fonction de l'identification ou non de l'usage opérationnel du système lors de la phase de développement, la CNIL a d'ailleurs publié une fiche pratique²⁷.

Enfin, il est à noter qu'en cas de réutilisation de données déjà collectées, il appartient également au responsable de traitement de vérifier que la finalité du nouveau traitement est compatible avec celle du traitement initial, étant entendu que la réutilisation des données à des fins de recherche scientifique est présumée compatible lorsque le traitement présente certaines garanties (considérant 50 et article 5.1.b) du RGPD ainsi que l'article 4.2° de la loi « informatique et libertés »)²⁸.

²⁷ Pour en savoir plus, voir la [fiche pratique](#) publiée par la CNIL.

²⁸ Pour en savoir plus, voir la [fiche pratique](#) publiée par la CNIL.



Pr. Marie-France Mamzer
PU-PH, Université Paris Cité

**Un éclairage multidisciplinaire
indispensable à la compréhension
des enjeux scientifiques, juridiques
et éthiques des essais cliniques de
demain**



iii. Un traitement de données exactes et nécessaires

Les données traitées doivent être pertinentes, adéquates et limitées à ce qui est strictement nécessaire par rapport à l'objectif poursuivi. Cette règle, dénommée « principe de minimisation », est prévue à l'article 5 du RGPD. Dans le cadre d'un projet de recherche, le traitement de chaque catégorie de données doit être justifié scientifiquement (par exemple à travers des publications scientifiques existantes) ou en lien avec l'un des objectifs ou l'un des critères d'évaluation de l'étude. Ainsi, le traitement du numéro de sécurité sociale des patients dans le cadre d'un projet de recherche visant à développer un algorithme générateur de cohortes de patients virtuels n'apparaîtrait pas conforme au principe de minimisation.

Si l'utilisation de quantités importantes de données est au cœur du développement et de l'utilisation des systèmes d'intelligence artificielle, le principe de minimisation ne constitue cependant pas un obstacle à la réalisation de ces traitements, tout comme il ne fut pas un frein au développement des entrepôts de données de santé autorisés par la CNIL depuis 2017.

Les données nécessaires au développement de l'algorithme devront, par ailleurs, être exactes et, si nécessaires, tenues à jour par le responsable de traitement.

iv. Une conservation des données limitée

Les données personnelles doivent être traitées pendant une durée n'excédant pas celle nécessaire au regard des finalités pour lesquelles elles ont été collectées, conformément aux dispositions de l'article 5 du RGPD. Elles ne sont conservées en « base active », c'est-à-dire le temps d'être analysées, que le temps nécessaire à la réalisation de l'objectif poursuivi. Elles doivent être par la suite détruites, anonymisées ou archivées dans le respect des obligations légales applicables.

Cette nécessité de définir une durée de conservation pour les données utilisées par un traitement ne fait pas obstacle à la mise en œuvre des traitements d'intelligence artificielle (notamment en cas d'apprentissage continu) dès lors que la durée retenue, qui peut être relativement longue, est justifiée.

Il est ainsi nécessaire de fixer une durée de conservation des données utilisées pour le développement du système d'intelligence artificielle :

- Pour la phase de développement : la conservation des données doit ainsi faire l'objet d'une planification en amont et d'un suivi dans le temps ;

- Pour la maintenance, la surveillance ou l'amélioration du produit: lorsque les données n'ont plus à être accessibles pour les tâches quotidiennes des personnes en charge du développement du système d'intelligence artificielle, elles doivent en principe être supprimées. Elles peuvent toutefois être conservées pour la maintenance et la surveillance du produit ou son amélioration si des garanties sont mises en œuvre (support cloisonné, restriction des accès aux seules personnes habilitées, etc).

La conservation des données d'apprentissage peut permettre d'effectuer des audits et faciliter la mesure de certains biais. Dans ces cas, une conservation prolongée des données peut être justifiée. Cette conservation doit être limitée aux données nécessaires, et s'accompagner de mesures de sécurité renforcées.

v. La mise en place de mesures techniques et organisationnelles appropriées et la réalisation d'une analyse d'impact

Le responsable de traitement doit mettre en œuvre des mesures techniques et organisationnelles appropriées afin de garantir un niveau de sécurité adapté aux risques identifiés (tels que l'accès illégitime, la modification non désirée ou la disparition des données)²⁹. En général, les mesures de sécurité adaptées au traitement de données de santé des patients se matérialisent par l'utilisation d'une plateforme sécurisée du type de celles proposées par les hébergeurs de données de santé³⁰ ou mises en place au sein d'entrepôts de données de santé³¹.

Par ailleurs, le développement de systèmes d'intelligence artificielle nécessite, dans certains cas, la réalisation d'une analyse d'impact relative à la protection des données (AIPD). Cette dernière est obligatoire si le traitement envisagé est susceptible d'engendrer un risque élevé pour les droits et libertés des personnes physiques (article 35 du RGPD), ce qui est susceptible d'être le cas s'agissant de la constitution d'une base de données de santé pour l'apprentissage d'un système d'intelligence artificielle.

Le Comité européen de la protection des données (CEPD) a identifié neuf critères permettant d'aider les responsables de traitement à déterminer si une AIPD est requise tels que la collecte de données sensibles, la collecte de données de personnes vulnérables, la collecte de données à large échelle ou l'usage innovant. Tout traitement remplissant au moins deux critères de cette liste sera présumé soumis à l'obligation de réaliser une AIPD.

²⁹ Pour en savoir plus, voir la [guide pratique](#) publié par la CNIL.

³⁰ Pour en savoir plus, voir la [page](#) consacrée aux hébergeurs de données de santé publiée par l'Agence du numérique en santé

³¹ Pour en savoir plus, voir le [référentiel relatif aux entrepôts de données de santé](#) de la CNIL

En pratique, un ou plusieurs traitements de données de santé (considérées comme des données sensibles) de patients (considérés comme des personnes vulnérables au sens des lignes directrices CEPD) à des fins de développement d'un système d'intelligence artificielle permettant de générer des cohortes de patients virtuels (pouvant être considéré comme un « usage innovant » au regard des connaissances technologiques actuelles) remplissent au moins deux des neuf critères du CEPD. Une AIPD (qui peut porter sur un ensemble d'opérations de traitement similaires qui présentent des risques élevés similaires) devra donc être réalisée par le responsable de traitement. Elle permettra de cartographier et d'évaluer les risques du traitement sur la protection des données personnelles et d'établir un plan d'action pour les réduire à un niveau acceptable³².

vi. La transparence, la loyauté et le respect des droits des personnes

Les personnes dont les données sont réutilisées (aussi appelées « personnes concernées ») doivent être informées et être en mesure d'exercer leurs droits.

1) Les modalités d'information

Plusieurs dispositions sont applicables en matière d'information. Certaines sont spécifiques au contexte dans lequel les données seront traitées.

Le RGPD impose une information des personnes concernées concise, transparente, compréhensible et aisément accessible en des termes clairs et simples. Les articles 13 (collecte directe) et 14 du RGPD (collecte indirecte ou lorsque la recherche implique la réutilisation de données d'une base existante) prévoient les modalités d'information et les éléments devant figurer dans la note d'information (identité du responsable de traitement, finalités du traitement, destinataires des données, durée de conservation des données. etc³³). En présence d'une prise de décision automatisée, les personnes doivent être informées, lors de la collecte de leurs données et à tout moment sur leur demande, de l'existence d'une telle décision, de la logique sous-jacente ainsi que de l'importance et des conséquences prévues de cette décision. Cette information doit être adaptée à chaque catégorie de personnes concernées par le traitement, notamment en présence de mineurs ou de personnes vulnérables afin que celle-ci soit la plus transparente possible.

³² Pour en savoir plus, voir la [fiche pratique](#) publiée par la CNIL concernant l'AIPD et les spécificités liées au [recours à l'intelligence artificielle](#).

³³ Par ailleurs, en présence d'une prise de décision automatisée, les personnes doivent être informées, lors de la collecte de leurs données et à tout moment sur leur demande, de l'existence d'une telle décision, de la logique sous-jacente ainsi que de l'importance et des conséquences prévues de cette décision.

S'agissant plus spécifiquement des recherches, études ou évaluations dans le domaine de la santé, l'article 69 de la loi « informatique et libertés » prévoit une information individuelle des personnes concernées conformément au RGPD. Il en va ainsi notamment en cas de collecte directe des données auprès des personnes concernées.

Des données issues du soin ou d'une précédente étude peuvent faire l'objet d'une réutilisation sans qu'il soit procédé à une nouvelle information individuelle des personnes notamment lorsque l'information délivrée lors de la collecte initiale des données prévoit la possibilité de réutiliser les données et renvoie à un dispositif spécifique d'information auquel les personnes concernées pourront se reporter préalablement à la mise en œuvre de chaque nouveau traitement de données (par exemple : un site web ou une page dédiée appelé « portail de transparence » sur lequel serait présenté chaque projet de recherche mené sur les données collectées dans le cadre de l'information initiale). Cette méthode est vivement recommandée par la CNIL car elle constitue un moyen simple d'informer les personnes concernées sur l'ensemble des projets de recherche ultérieurs. Elle nécessite néanmoins d'avoir anticipé la possibilité de réutilisation dès la conception du traitement initial en créant une page web dédiée à l'information.

Dans l'hypothèse où une nouvelle information individuelle ne pourrait être délivrée et que la réutilisation des données n'aurait pas été anticipée au moyen d'un portail de transparence, trois exceptions à la fourniture d'une information individuelle sont prévues par l'article 14.5.b) du RGPD pour les traitements mis en œuvre à des fins de recherche scientifique. Il en va ainsi lorsque la fourniture d'une information individuelle :

- Se révèle impossible ;
- Ou exigerait des efforts disproportionnés de la part du responsable de traitement au regard de l'ancienneté des données ou du nombre de personnes concernées ;
- Ou est susceptible de compromettre ou de rendre impossible la réalisation des objectifs du traitement.

Des mesures appropriées doivent alors être mises en œuvre par le responsable de traitement afin de protéger les droits et libertés des personnes concernées. Dans ce cas, le RGPD, éclairé par les lignes directrices adoptées par le CEPD, prévoit que des mesures appropriées doivent être prises pour protéger les droits et libertés des personnes concernées, notamment par la mise en place systématique d'une information collective via des canaux de communication appropriés au regard du contexte de l'étude (a minima, sur le site web du responsable de traitement)³⁴.

³⁴ Pour en savoir plus, voir la [fiche pratique](#) publiée sur le site web de la CNIL.

Par ailleurs, dans l'hypothèse où des cohortes de participants virtuels seraient générées dans le cadre de recherches impliquant la personne humaine, d'essais cliniques, d'investigations cliniques ou d'études de performance, certaines dispositions spécifiques issues du Code de la santé publique et des règlements européens relatifs aux essais cliniques, aux dispositifs médicaux³⁴ (y compris de diagnostic in vitro³⁵) auraient vocation à s'appliquer concernant le contenu des documents d'information et leurs destinataires.

Par exemple, les participants à une recherche impliquant la personne humaine ont, en application de l'article L. 1122-1 du code de la santé publique, le droit d'être informés de façon intelligible sur la méthodologie de la recherche dont fait partie intégrante le recours à un algorithme pour générer des cohortes virtuelles. Dans le cadre de son appréciation des conditions de validité de la recherche, il appartiendrait alors au comité de protection des personnes de se prononcer notamment sur la méthodologie envisagée ainsi que sur l'adéquation, l'exhaustivité et l'intelligibilité des documents d'information conformément aux dispositions de l'article L. 1123-7 du code de la santé publique. Quant aux patients dont les données issues de la prise en charge ou d'une précédente recherche seraient réutilisées dans le cadre d'un groupe comparateur d'une recherche impliquant la personne humaine, ils devraient être, en l'état du droit positif, informés selon les mêmes modalités que les autres participants et leur consentement, lorsqu'il est requis, devrait être recueilli, en l'absence de dispositions spécifiques prévues par le code de la santé publique et les règlements européens applicables aux essais cliniques et aux dispositifs médicaux. Il est à noter que des travaux, auxquels la CNIL est associée, sont en cours à ce sujet pour adapter ce cadre juridique aux nouvelles méthodologies de recherche.

Enfin, les dispositions de l'article L. 4001-3 au sein du code de la santé publique encadrent l'usage, pour un acte de prévention, de diagnostic ou de soin, de dispositifs médicaux comportant un traitement de données algorithmique dont l'apprentissage a été réalisé à partir de données massives. Dans cette hypothèse, le professionnel de santé doit informer la personne concernée de l'usage d'un traitement algorithmique et, le cas échéant, de l'interprétation qui en résulte. Par ailleurs, le professionnel de santé doit également être informé du traitement de données (et notamment des données utilisées et des résultats). Enfin, les concepteurs du traitement algorithmique doivent s'assurer de l'explicabilité de son fonctionnement pour les utilisateurs. La nature des dispositifs médicaux concernés et leurs modalités d'utilisation doivent être précisés par un arrêté du ministre chargé de la santé établi après avis de la Haute autorité de santé et de la CNIL. Il n'a, pour l'heure, pas été publié, ne permettant pas de déterminer si le dispositif permettant de générer des cohortes de patients virtuels pourrait être soumis à ces dispositions.

³⁴ Règlement (UE) 2017/745 du Parlement européen et du Conseil du 5 avril 2017 relatif aux dispositifs médicaux

³⁵ Règlement (UE) 2017/746 du Parlement européen et du Conseil du 5 avril 2017 relatif aux dispositifs médicaux de diagnostic in vitro et abrogeant la directive 98/79/CE et la décision 2010/227/UE de la Commission

2) Les modalités d'exercice des droits des personnes concernées

Le responsable de traitement doit garantir aux personnes dont les données sont traitées la possibilité d'exercer leurs droits d'accès (article 15 du RGPD), de rectification (article 16 du RGPD), d'opposition (article 21 du RGPD), d'effacement (article 17 du RGPD), de limitation (article 18 du RGPD) et de portabilité (article 20 du RGPD) dans les conditions prévues par le RGPD et la loi « Informatique et libertés »³⁶.

Par ailleurs, l'article 22 du RGPD prévoit que, par principe, les personnes ont le droit de ne pas faire l'objet d'une décision fondée exclusivement sur un traitement automatisé, y compris le profilage, produisant des effets juridiques les concernant ou les affectant de manière significative de façon similaire. Certaines exceptions sont prévues, notamment lorsque la décision est fondée sur le consentement explicite des personnes. Dans cette hypothèse, le responsable de traitement doit mettre en œuvre des mesures appropriées pour la sauvegarde des droits et libertés et des intérêts légitimes des personnes concernées. Outre les obligations de transparence susmentionnées, toute personne ayant fait l'objet d'une telle décision peut demander qu'une personne humaine intervienne, notamment afin d'obtenir un réexamen de sa situation, d'exprimer son propre point de vue, d'obtenir une explication sur la décision prise ou de contester la décision. Enfin, des décisions individuelles automatisées ne peuvent être fondées sur des données de santé, sauf cas particuliers (recueil du consentement de la personne concernée permettant de déroger au principe d'interdiction de traitement de telles données ou traitement nécessaire pour des motifs d'intérêt public dans le domaine de la santé publique). Dans ce cas, des mesures appropriées pour la sauvegarde des droits et libertés et des intérêts légitimes des personnes concernées doivent être mises en place.

Ces droits constituent une protection essentielle pour les personnes concernées, en leur permettant de ne pas subir les conséquences d'un système automatisé sans avoir la possibilité de comprendre et, si nécessaire, de s'opposer à des traitements de données qui les concernent. En pratique, ces droits trouvent à s'appliquer tout au long du cycle de vie du système d'intelligence artificielle.

³⁶ Pour en savoir plus, voir les [fiches pratiques](#) publiées par la CNIL.

III. Les éventuelles formalités applicables conformément à la loi « informatique et libertés »

Les éventuelles formalités différeront en fonction du cycle de vie et du contexte, l'utilisation de l'algorithme permettant de générer les cohortes virtuelles. À date, ces questions sont, pour l'heure, assez récentes, et n'ont pas encore été arbitrées par le collège de la CNIL.

i. Dans le cadre du développement de l'algorithme

Le traitement de données de santé à des fins de développement d'un algorithme destiné à générer des cohortes de patients virtuels dans un contexte médical pourrait constituer une recherche, une étude ou une évaluation dans le domaine de la santé relevant de la « loi informatique et libertés ». Conformément aux dispositions de l'article 66 de la loi « informatique et libertés », les recherches scientifiques en santé nécessitant un traitement de données de santé, doivent, sauf exceptions prévues par la loi, faire l'objet d'une déclaration de conformité à une méthodologie standard de traitement de données décrite dans un référentiel publié par la CNIL appelé « méthodologie de référence »³⁷. Par exception, les recherches qui ne sont pas conformes à ces référentiels en raison d'une sensibilité particulière (traitements de données d'infractions par exemple) doivent être autorisées par la CNIL, après avis du comité compétent.

En l'absence de conformité à une méthodologie de référence, la CNIL peut également délivrer une décision unique après avis du comité compétent en application des dispositions de l'article 66 IV de la loi « informatique et libertés ». Par ce mécanisme, la CNIL peut autoriser, grâce à une seule décision, un nombre important de traitements pendant une durée déterminée répondant à une même finalité, portant sur les mêmes catégories de données et ayant des catégories de destinataires identiques. Afin que le comité compétent puisse se prononcer sur la pertinence scientifique du projet, ces traitements doivent être mis en œuvre sur la base d'une méthodologie standardisée établie par le responsable de traitement. Le recours à ce mécanisme peut notamment se justifier par la nécessité de mettre en œuvre une importante volumétrie de traitements. Il pourrait donc être exploré dans l'hypothèse où le développement de l'algorithme impliquerait la mise en œuvre de nombreux traitements de données à caractère personnel.

³⁷ Pour en savoir plus, voir la [rubrique](#) consacrée aux méthodologies de référence publiée sur le site web de la CNIL.

En tout état de cause, et quelle que soit la formalité à réaliser auprès de la CNIL, le protocole scientifique élaboré par le responsable de traitement devrait notamment :

- Décrire la finalité du traitement de données personnelles, la nature des données traitées, la proportionnalité des techniques choisies, l'impact et les avantages du recours à ce système d'intelligence artificielle ;
- Documenter la méthodologie (hypothèses effectuées sur les données d'entraînement, réalisation d'une étude des biais) pour faire en sorte d'avoir un traitement fiable, robuste dans le temps (pour minimiser les risques de dérives) et démontrer que le traitement envisagé poursuit une finalité d'intérêt public.

À l'issue de cette phase de développement, une évaluation scientifique du système devrait être effectuée avant que celui-ci puisse être déployé.

ii. Dans le cadre du déploiement du système d'intelligence artificielle

1) La génération de nouvelles cohortes de patients virtuels dans le cadre de projets de recherche

Dans cette hypothèse, le traitement de données nécessaire à la génération de nouvelles cohortes dans le cadre d'un essai clinique pourrait s'inscrire dans le cadre de la formalité préalable relative à cet essai (déclaration de conformité à une méthodologie de référence ou, en l'absence de conformité, dépôt d'une demande d'autorisation auprès de la CNIL).

2) La génération de nouvelles cohortes de patients virtuels dans le cadre de la prise en charge des patients

En cas d'utilisation du système d'intelligence artificielle par un professionnel de santé ou l'établissement qui prend en charge le patient, à des fins de diagnostic, de l'administration de soins ou de traitements, le traitement pourrait s'inscrire dans l'une des exceptions à formalités préalables prévues à l'article 65 de la loi « informatique et libertés », plus particulièrement, l'élaboration des diagnostics médicaux, ou encore l'administration de soins ou de traitement.

En tout état de cause, quels que soient les usages dans le cadre de la phase de déploiement, il conviendra de vérifier si le système d'intelligence artificielle satisfait bien les besoins pour lesquels il a été conçu. Par ailleurs, et comme exposé plus haut, dans l'hypothèse où le jeu de données initial comporterait des données à caractère personnel, il apparaît nécessaire d'analyser le statut du modèle d'intelligence artificielle et des données générées par l'algorithme au regard notamment des critères européens relatifs à l'anonymisation.



Manon de Fallois

Adjointe à la cheffe du service de la santé, CNIL

Le développement de ces systèmes d'intelligence artificielle s'accompagne d'enjeux émergents en matière de protection des données. La CNIL souhaite soutenir leur essor en accompagnant leurs fournisseurs afin qu'ils soient déployés dans le respect de la vie privée des personnes concernées



PREUVES ET MÉTHODOLOGIE POUR LA VALIDATION DES PRODUITS DE SANTÉ À L'AIDE DE DONNÉES ARTIFICIELLES

05



Marco Fiorini

Directeur Général du projet de Filière Intelligence
Artificielle et Cancers

Au regard de l'exponentialité de la production de données, de la qualité des capteurs, de la puissance d'analyse...se préparer à l'usage des données pour simuler tous les paramètres habituels de la recherche clinique ou de la recherche en vraie vie n'est pas une option. Nous devons avancer en appréciant honnêtement les potentiels et limites de ces technologies réunies pour participer à une course mondiale. Je suis persuadé que c'est par des livres blancs comme celui-ci, articulés à des projets réels, que nous y parviendrons.

”

La littérature foisonnante nous montre que les modèles mathématiques peuvent avoir de nombreuses utilisations en santé. Cela peut concerner des aspects très techniques (tester la robustesse et la sécurité d'outils informatiques) ou plus opérationnels (assister des chirurgiens lors de la planification d'une intervention, développer des modèles diagnostiques ou pronostiques pour construire des systèmes d'aide à la décision, créer des jumeaux numériques d'établissement permettant d'anticiper les flux de patients en son sein). Nous avons vu aussi précédemment que les données artificielles peuvent aussi avoir des applications importantes en termes de développement de produits de santé, à des phases très précoces de découverte de nouveaux médicaments (par exemple, [75, 76]), mais aussi pour l'évaluation de leur efficacité et de leur sécurité, ou en personnalisant les stratégies de traitement. Ces dernières utilisations en sont actuellement à leurs débuts, et il est important d'envisager les preuves et les méthodes de validation à mettre en œuvre pour que l'utilisation de ces données virtuelles soit acceptable par les différents acteurs concernés. La confiance apportée par la validation de ces méthodes est nécessaire si l'on souhaite un jour utiliser les données de cohortes artificielles ou partiellement artificielles pour l'accès aux nouveaux traitements.

Il est aussi important de faire la distinction, comme proposé par la FDA, entre la validation analytique des algorithmes et leur validation clinique dans des conditions d'utilisation en vie réelle. Il y a souvent une confusion entre la validation analytique qui comprend une validation sur des données (cela peut inclure de multiples sources de données de provenance différentes) et la validation clinique qui peut comprendre des essais cliniques prospectifs ou sur des RWD (« Real World Data ») une fois que l'algorithme (via une app ou autre) est utilisé et génère lui-même des données. Dans les deux cas, les méthodes et les designs d'étude sont différents ainsi que les critères de validation.

Il n'existe actuellement pas de recommandations émanant d'agences ou d'autorités de régulation définissant les critères d'acceptabilité de cohortes de patients artificiels pour l'évaluation des produits de santé ou des dispositifs médicaux. On trouve des références à des études de simulations, à partir de modèles mécanistiques, dans les recommandations ICH E11, dans la partie ICH guideline E11A sur l'extrapolation pédiatrique [77], et l'EMA et l'EFPIA ont coorganisé un séminaire sur la modélisation et la simulation³⁸, il y a maintenant plus de 10 ans. Cette thématique est toujours en cours de réflexion à l'EMA³⁹. Le guide méthodologique de la HAS⁴⁰ pour le développement clinique des dispositifs médicaux, mis à jour en 2021, évoquait aussi les études in silico. Il concluait toutefois qu'à ce jour, les études in silico devaient être

38 European Medicines Agency-European Federation of Pharmaceutical Industries and Associations modelling and simulation workshop, 2011 ; <https://www.ema.europa.eu/en/events/european-medicines-agency-european-federation-pharmaceutical-industries-associations-modelling>

39 <https://www.ema.europa.eu/en/human-regulatory-overview/research-and-development/innovation-medicines>

40 https://www.has-sante.fr/upload/docs/application/pdf/2013-11/guide_methodologique_pour_le_developpement_clinique_des_dispositifs_medicaux.pdf

considérées comme des outils complémentaires aux méthodologies existantes lorsqu'il existe un modèle physiopathologique et qu'elles devaient être réservées aux études de faisabilité. Enfin, un cadre d'évaluation de la crédibilité des modèles mécanistiques a été développé récemment par un groupe de travail coordonné par l'EMA, Modelling and Simulation Working Party, pour l'utilisation d'études in silico reposant des modèles mécanistiques [78].

En résumé, ce type de données commence à se développer et les agences et autorités compétentes ont été à ce jour peu sollicitées sur des données de ce type. Dès lors, aucun des documents ou actes actuellement disponibles ne donne d'attendus précis pour l'acceptabilité des méthodes permettant de générer des données artificielles et leur utilisation pour l'évaluation de l'efficacité et de la sécurité de produits de santé ou de dispositifs médicaux, comme souligné dans un *position paper* du consortium INSILICO WORLD⁴¹.

Au contraire, il existe un grand nombre de travaux concernant l'utilisation de bras de contrôles synthétiques - ou externes - construits à partir de données observationnelles [79-82], y compris émanant d'agences comme la HAS et la commission de la transparence [83], et même des recommandations de la FDA⁴².

L'EMA et le réseau des directeurs d'agence nationale du médicament ont dévoilé en décembre dernier leur plan de travail sur l'intelligence artificielle visant à élaborer une stratégie européenne. Il s'agit du plan de travail destiné à mieux encadrer l'utilisation de l'intelligence artificielle par les régulateurs nationaux.

I. Études in-silico à partir de modèles mécanistiques

Comme décrit au chapitre 2, la génération de données de cohortes artificielles peut reposer sur des approches mécanistiques, reposant sur la modélisation des systèmes physiques, chimiques et biologiques considérés, ou plutôt sur des méthodes d'apprentissage, visant à reproduire des caractéristiques similaires à celles d'une population observée. Pour ce type d'approche, un cadre d'évaluation de la crédibilité des modèles mécanistiques a aussi été proposé [78].

41 Toward good simulation practice: best practices for the use of computational modelling & simulation in the regulatory process of biomedical products. Version: R6, 06-05-2023. <https://insilico.world/sito/wp-content/uploads/2023/06/Position-Paper-GSP-R6.pdf>

42 <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/considerations-design-and-conduct-externally-controlled-trials-drug-and-biological-products>

Les principes énoncés dans ces recommandations peuvent aussi s'appliquer aux études utilisant des données artificielles. Ce groupe d'auteurs a proposé une approche en sept points :

1. Description de la question d'intérêt et contexte d'utilisation ;
2. Définition des critères d'acceptabilité du modèle pour cette question d'intérêt et ce contexte d'utilisation ;
3. Description de l'impact réglementaire (pour les soumissions à des agences comme l'EMA) ;
4. Analyse des conséquences des décisions ;
5. Description de la crédibilité du modèle et des algorithmes (vérification et validation) ;
6. Applicabilité et incertitude ;
7. Prise de décision éclairée par le modèle.

Plusieurs points importants pour définir les méthodes et preuves à mettre en œuvre peuvent donc être identifiés : les critères d'acceptabilité pour une question particulière, qui comprend entre autres la pertinence théorique des modèles utilisés, mais aussi la qualité de leur implémentation informatique (et en particulier les processus d'assurance-qualité mis en œuvre), les méthodes de validation utilisées, les aspects concernant les données utilisées pour développer le modèle et le valider, et une étude d'impact des différentes sources d'incertitude. Les éléments de crédibilité du modèle concernent en premier lieu la capacité du modèle à (re)produire des données qui correspondent bien à ce qui aurait été observé si des patients réels avaient été inclus. Les aspects détaillés dans la liste ci-dessous, reprenant les propositions de Musuamba et coll. [78], devraient ainsi être considérés pour valider ce type d'approche, en les adaptant si nécessaire au contexte des données artificielles et aux méthodes utilisées pour les générer :

Détailler les modèles de génération des données

- Approches utilisées pour générer les données artificielles, que ce soient des modèles mécanistiques ou des modèles d'apprentissage, tels que décrits au chapitre 2 ;
- Évaluation de l'adéquation du modèle en lien avec les objectifs (question d'intérêt ou cas d'usage), notamment si le groupe de patients artificiel a pour objectif d'améliorer le développement d'algorithmes, d'augmenter un groupe contrôle d'un essai, voire éventuellement de le remplacer ; ces points seront réabordés plus bas.

Détailler l'implémentation, en lien avec les réglementations de la norme ISO 13485 :

- Méthodes utilisées pour vérifier le code et résultats des actions de vérification ;
- Aspects liés à la transparence de l'algorithme (et des méthodes de développement et validation) : description des méthodes utilisées, disponibilité de l'algorithme, de son implémentation, etc. ;
- Démarche d'assurance qualité et respect de la norme ISO 13485 en particulier ;
- Éléments attestant de la robustesse de l'algorithme (à des perturbations dans les données, attaques contradictoires [*adversarial attacks*], etc.), à la fois d'un point de vue théorique et pratique.

Détailler les méthodes de validation :

- Capacité de l'algorithme à reproduire des observations déjà obtenues (distributions de probabilité de populations identiques pour les données initiales et augmentées), et sur lesquelles il n'a pas été entraîné (même distribution de probabilité pour une base test) ;
- Étude de la validation externe de modèles de prédiction, dans la population d'intérêt et sur l'objet d'intérêt, par exemple en prenant 50% des données d'un bras comparateur, en le complétant artificiellement puis en regardant si l'on obtient le même résultat en comparant le bras traité et le bras témoins originel et le bras traité avec le bras augmenté ;
- Métriques classiques (R2, discrimination, calibration, etc.) ou utilisation de méthodes plus spécifiques (par exemple, [84]).

Détailler les données utilisées :

- Pour le développement de l'algorithme permettant de générer données artificielles ;
- Pour sa calibration et validation ;
- Pour le test sur des bases externes.

En 2023, la Commission Européenne a publié dans le cluster santé d'Horizon Europe un appel à projets afin de proposer des méthodes de simulations, modélisation et essais cliniques *in silico* avec un but d'être compatible et accepté d'un point de vue réglementaire pour les petits effectifs maladies rares, pédiatriques, etc.). Deux projets ont été financés et ont démarré en 2024, ERAMET⁴³ et INVENTS⁴⁴. Dans les prochaines années, des méthodes seront donc disponibles dans ce cadre.

43 <https://cordis.europa.eu/project/id/101137141>

44 <https://ecrin.org/projects/invents>



Raphaël Porcher

PU-PH Université Paris Cité
Chaire PR[AI]RIE

L'utilisation de données de patients artificiels va vraisemblablement se développer dans les études cliniques. Il appartient encore aux scientifiques d'identifier les conditions et les limites de leur utilisation, et d'apporter la preuve de leur fiabilité et de leur sécurité pour une utilisation bénéfice des patients et du système de santé.

”

II. Modélisation à l'aide d'une approche bayésienne sans génération de données artificielles

Un exemple d'utilisation de modélisation, mais sans génération de données artificielles, pour l'évaluation de l'efficacité d'un médicament concerne le développement secukinumab, un anti IL-17A, pour le psoriasis en pédiatrie [85]. Bien que cet exemple n'ait pas utilisé la génération de cohortes artificielles, il repose sur une approche prédictive à partir d'une méta-analyse bayésienne, qui pourrait être utilisée pour générer des données dans une approche similaire aux méthodes de type ABC ou d'estimation d'atlas décrites au chapitre 2, et il nous a paru pertinent de le développer, dans la mesure où il s'agit d'un des rares exemples d'évaluation d'un dossier reposant sur une modélisation non mécanistique par les autorités de santé. Le médicament, initialement développé pour le psoriasis modéré à sévère de l'adulte, a été autorisé au Japon (2014), en Europe (2015) et aux États-Unis (2015). Il est aussi prescrit dans la spondylarthrite ankylosante, le rhumatisme psoriasique, et d'autres maladies inflammatoires, avec plus de 500 000 patients traités dans le monde.

Pour le développement en pédiatrie, deux essais randomisés ont été conduits :

- Un essai dans le psoriasis sévère, avec trois bras de traitement : secukinumab 150 mg, secukinumab 300 mg, comparateur actif (n = 160 prévus, n = 162 inclus).
- Un essai dans le psoriasis modéré, avec deux bras de traitement : secukinumab 150 mg, secukinumab 300 mg (n = 120 prévus, n = 84 inclus), mais sans comparateur actif (ni placebo).

Il a été en outre prévu une étude de modélisation pour extrapoler l'efficacité de secukinumab dans le psoriasis modéré de l'enfant et de l'adolescent en utilisant un modèle bayésien de méta-analyse prédictive. L'idée est de modéliser le phénomène en utilisant de l'information *a priori* d'une autre population. Le cadre bayésien de modélisation permet de formaliser l'incorporation de ces données issues d'une autre population dans une distribution de probabilités dite *a priori*. Le modèle combine alors les données observées dans la population d'intérêt avec la distribution *a priori*, pour donner une distribution de l'effet du traitement dite *a posteriori*. Cette approche utilise donc des prédictions issues d'autres groupes pour proposer une mise à jour de l'*a priori* au regard de ce qui est observé. Cependant, elle ne repose pas sur la génération de données artificielles pour construire un groupe contrôle pour l'essai dans le psoriasis modéré de l'enfant.

L'extrapolation de l'adulte vers l'enfant répondait ici en particulier à plusieurs éléments-clés rappelés dans les recommandations ICH-E11A [77], à savoir :

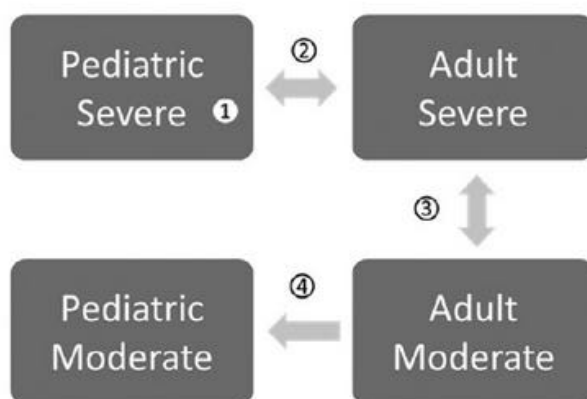
1. Une pharmacologie similaire chez l'enfant et l'adulte ;
2. Une physiopathologie et une progression de la maladie similaires chez l'enfant et l'adulte ;
3. Une réponse thérapeutique similaire chez l'enfant et l'adulte.

Il est à noter que le développement de nouvelles formes galéniques pour les enfants de moins de six ans quand les règles PK-PD diffèrent est un vrai enjeu.

Il avait en outre été considéré que le secukinumab, utilisé chez un nombre important d'adultes dans le monde, avait un profil de sécurité rassurant.

L'extrapolation du psoriasis pédiatrique sévère au psoriasis pédiatrique modéré a été faite en utilisant un modèle en 4 étapes, décrit sur la figure ci-dessous. L'étape 4 consistait en particulier en un modèle joint bayésien tenant compte de toutes les données (adulte et enfant). Le dossier a finalement été soumis à l'EMA en 2019 et à la FDA en 2020, et l'autorisation pour les formes pédiatriques modérées a été obtenue à la fois en Europe (EMA, 2020) et aux Etats-Unis (2021). Toutefois en France, le premier avis de la Commission de la Transparence⁴⁵, n'a été favorable que pour le traitement du psoriasis de l'enfant et de l'adolescent en échec d'au moins deux traitements et en cas de forme étendue et/ou un retentissement psychosocial important. Le service médical rendu était jugé insuffisant dans les autres formes. Cet avis a été récemment réévalué à la suite des résultats d'une post-inscription, avec un avis favorable au remboursement dans l'indication de son AMM et un service médical rendu considéré comme important⁴⁶. On peut noter que dans l'avis de novembre 2021, le modèle de méta-analyse prédictive bayésienne permettant d'estimer indirectement l'effet contre placebo avait été jugé insuffisamment robuste.

En conclusion, même si dans cet exemple aucune donnée de patients artificiels n'a été utilisée et que la Commission de la Transparence n'a pas retenu la méta-analyse prédictive estimant indirectement l'effet contre placebo, cette dernière analyse a été jugée suffisamment convaincante par l'EMA et la FDA. Il est donc permis d'envisager un cas d'usage de cohortes artificielles permettant d'augmenter les données d'essais cliniques, sous réserve que la méthodologie employée soit considérée comme suffisamment robuste.



Copyright You et al. Clin Pharmacol Ther. 2022 Mar;111(3):697-704.

45 https://www.has-sante.fr/upload/docs/evamed/CT-18841_COSENTYX_enfant_PIC_EI_AvisDef_CT18841&18848.pdf (3 nov. 2021)

46 https://www.has-sante.fr/jcms/p_3426148/fr/cosentyx-secukinumab-psoriasis-en-plaques-chez-l-enfant-et-l-adolescent (15 mars 2023)



Jérôme Kalifa

Président et fondateur de Let it Care

Il est cruel de devoir abandonner le développement d'un médicament prometteur sur une maladie rare, faute de faisabilité d'un essai clinique avec un nombre suffisant de patients. Les données artificielles promettent une réponse en complétant les essais avec des patients virtuels offrant la preuve statistique requise.

”

III. Questions d'évaluation soulevées par les données artificielles

Au-delà des aspects développés au I., l'utilisation de données artificielle soulève des questions spécifiques, pour lesquelles la pratique permettra d'identifier les prérequis ou les recommandations adéquates. Plusieurs points peuvent être dès à présent être considérés.

De nombreuses méthodes présentées au chapitre 2 ont montré leur capacité à reproduire des données dont les caractéristiques, en termes de distribution de probabilité, sont bien similaires à celles utilisées pour l'entraînement des modèles. Se pose alors en premier lieu la question de savoir quels niveau et quantité de preuves sont nécessaires pour considérer qu'une approche est suffisamment fiable et validée pour pouvoir être utilisée en pratique. Deux approches différentes peuvent être envisagées. Soit la validation d'une méthode dès lors qu'elle a fait ses preuves dans un certain nombre de situations (notamment dans la situation d'application concernée, par exemple l'indication ciblée, ou une maladie proche), soit la qualification par un organisme compétent. Cela pourrait dépendre aussi des types de critères de jugement utilisés, des objectifs de l'utilisation de données artificielles, et de l'indication, en tenant compte entre autres de la prévalence de la maladie (maladies rares), de sa proximité avec d'autres pathologies déjà modélisées, de l'existence de modèles physiologique validés, et de la population cible (comme une indication ciblant population fragile, par exemple).

Il est ainsi possible d'envisager des niveaux d'exigence différents en fonction des contextes d'utilisation. Les approches d'augmentation de données pour entraîner des algorithmes d'analyse d'images, par exemple, ou de prédiction pour des applications diagnostiques ou pronostiques nécessitent peut-être moins d'apport de preuve de la capacité de l'algorithme à générer des données similaires à celles observées que celles consistant à augmenter les données du groupe contrôle – voire du groupe expérimental – d'un essai clinique avec des données artificielles. En effet, dans le premier cas, les mesures de performance des algorithmes dans un jeu de données de test pourraient suffire à quantifier l'intérêt des données artificielles, ce d'autant que des données indépendantes de test seront nécessaires, alors que dans le deuxième cas, les données artificielles devraient représenter ce qui pourrait être observé si d'autres participants avaient été recrutés dans l'étude.

Le troisième élément concerne la question de la taille des jeux de données, qui comprend la taille minimale des jeux de données d'entraînement nécessaires même si non suffisante, pour qu'un algorithme soit performant, mais aussi la taille maximale d'un jeu de données artificielles qui puisse raisonnablement être utilisé pour cet entraînement. Ici aussi, ces tailles maximales pourraient varier selon le contexte d'utilisation. Il s'agira de proposer des méthodes pour les évaluer de manière fiable et robuste. La représentativité de la base de données est également un facteur primordial.

Le quatrième élément, enfin, est en rapport avec la validation clinique de l'approche, si ces données artificielles sont utilisées dans le cadre d'un système clinique apprenant d'aide à la décision (diagnostique, thérapeutique, prévention...) et donc apparenté à un dispositif médical numérique (DMN) à usage professionnel. Dans ce cas, il est important de ne pas seulement évaluer l'algorithme par lui-même, mais aussi l'ensemble des systèmes avec lesquels il va interagir en termes de logiciel, système de soins et humain. Dernièrement, la HAS⁴⁷ a publié un guide à usage des professionnels utilisant des DMN⁴⁸. Même si ce guide ne traite pas particulièrement des données simulées, il peut servir de première base de réflexion.

Pour apporter des réponses à ces questions, il semble primordial que des exemples d'utilisation de données artificielles dans les différents contextes soient publiés. À cet égard, des travaux montrant la capacité de données artificielles à reproduire des situations réelles sont importants [86, 87, 88, 89]. Il faudra aussi compléter ces applications par des travaux plus méthodologiques permettant d'évaluer les limites d'utilisation de ces approches prometteuses, et les jalons à mettre en place pour permettre des premières utilisations. Les éléments cités plus haut peuvent servir de prérequis pour produire et utiliser des cohortes artificielles ou augmentées dans lesquelles nous pouvons avoir confiance, et aider l'écosystème à s'emparer de ces approches. À terme, cela devrait permettre aux autorités compétentes et aux agences d'évaluation de fixer des règles claires d'utilisation et d'acceptabilité de ces cohortes artificielles, notamment pour des demandes d'accès au marché ou de remboursement, sachant qu'elles devront offrir le même niveau de garantie que les méthodes de références actuelles. De plus, les données artificielles pourraient être prises en compte dans la modélisation des parcours de soin par exemple. Ce livre blanc permet ainsi de poser les premières pierres de cet édifice.

Il conviendra de se référer au tout nouvel article The GSP (Good Simulation Practices) Task Force, publié le 24 février 2024 par Avicenna Alliance, proposant :

- Une meilleure définition du champ d'application, en recherchant activement l'engagement le plus large possible avec toutes les organisations représentant les parties prenantes (universités, industrie, patients, régulateurs et payeurs) ;
- Une diffusion plus large des travaux réalisés jusqu'à présent ;
- Un processus de consensus par l'intermédiaire de la communauté de pratique In Silico Word.

47 48 https://www.has-sante.fr/jcms/p_3363066/fr/dispositifs-medicaux-numeriques-a-usage-professionnel

L'augmentation des cohortes par IA est un levier formidable pour apporter des solutions dans des domaines thérapeutiques encore trop souvent peu ou mal adressés (maladies rares ou d'évolution lente), ou pour des populations peu évaluées (pédiatrie, gériatrie). Ce livre blanc ouvre la réflexion sur la démonstration de la valeur et les conditions d'utilisation de ces outils pour une utilisation en recherche clinique en toute confiance. C'est une étape indispensable pour permettre aux patients de bénéficier des traitements, efficaces et sûrs, dont ils ont besoin.



Camille Schurtz

Responsable process réglementaires et accès au marché
Agence de l'Innovation en Santé
Copilote du Groupe de travail AIS/ F-CRIN (L'évolution des méthodologies d'essais cliniques: nouveaux outils, nouveaux usages et conditions de recours)



ENJEUX ET GARANTIES ÉTHIQUES

06



David Gruson
Fondateur, ETHIK-IA

Ce Livre Blanc sur les données artificielles éclaire une transformation majeure en cours de notre système de santé qui permettra notamment d'accélérer considérablement l'innovation thérapeutique. Il intègre également la régulation éthique by design en recommandant la mise en oeuvre de dispositifs de garantie humaine par les professionnels de santé et les représentants des patients, en cohérence avec la loi de bioéthique française de 2021 et le nouveau règlement européen sur l'IA.

”

Le recours à des cohortes de patients virtuels constitue une opportunité majeure pour le progrès de la recherche médicale et l'amélioration de la qualité de la prise en charge des patients. Cette innovation essentielle des méthodes d'apprentissage et de méthodologie clinique soulève néanmoins quelques interrogations sur le plan éthique, également en termes de robustesse comme déjà évoqué et doivent être positionnées dans une logique de régulation positive en liaison avec l'utilisation de modèles d'IA pour la création de ces patients artificiels. Il apparaît donc important de prendre en considération plusieurs points clés.

I. Supervision humaine des cohortes de patients artificiels (en référence à l'article 14 de l'IA Act)

Sur le plan éthique, si la qualité et le choix des données restent centraux, la validation des modèles mathématiques et des algorithmes l'est tout autant. Comme tout traitement de données, la mise en œuvre de cohortes de patients artificiels nécessite une supervision humaine. Ce principe de supervision humaine (« Garantie Humaine » ou « contrôle humain ») renvoie à la nécessité de ne pas abandonner toute autonomie d'action ou de décision humaine dans un contexte de diffusion de plus en plus rapide de l'intelligence artificielle.

Concrètement, des experts devront être chargés de valider que les cohortes créées sont bien conformes aux groupes de patients initiaux. Organisées sous la forme de collèges de garantie humaine, ces mesures de contrôle permettront de mieux comprendre la phase de modélisation des populations des patients artificiels et veiller ainsi à ce qu'elle soit la moins biaisée et la plus fiable possible.

Cette surveillance continue pourra s'effectuer dès la phase de création de la cohorte, mais également a posteriori durant toute la phase d'utilisation de ces bases. Les experts pourront par exemple superviser et valider les données collectées, mais également surveiller, aider au calibrage et à l'ajustement des modèles mathématiques et des outils algorithmiques utilisés dans la création des patients afin que la cohorte de patients artificiels soit la plus réaliste possible. Une validation d'un échantillon aléatoire de patients pourra être demandée afin de vérifier que tout patient artificiel pourrait être un patient réel. Enfin, il est absolument impératif qu'une supervision humaine soit opérée sur les résultats des utilisations effectuées à partir de ces cohortes de patients artificiels afin de veiller à la détection de potentielles erreurs ou scénario inattendu.

La mise en œuvre de ces principes de garantie humaine pourra intégrer, en plus des cliniciens de la discipline concernée, des représentants des patients de la [HN1] cohorte ou d'une indication similaire afin d'avoir une validation complémentaire. Les patients réels de la cohorte augmentés pourraient ainsi être mieux informés du traitement statistique de leurs données, aidant à l'adoption de ces technologies par le grand public et à l'ouverture des données de santé dans le cadre de la recherche et de l'innovation.

II. Vers un nouveau principe incontournable de l'éthique de la recherche impliquant la personne humaine et une révision des normes éthiques ?

Les enjeux éthiques du développement de ces nouvelles pratiques débordent la seule mise en balance de l'accélération des progrès à moindre coût financier d'un côté avec la nécessité de protéger les données de santé et les droits et libertés des personnes y afférant, question qui sera d'ailleurs sans doute réglée avec les progrès attendus en matière d'anonymisation. Une fois les questions scientifiques et techniques d'intégrité scientifique et de fiabilité résolues, et donc l'équivalence des résultats garantis, l'éventail des possibilités offertes en matière de remplacement des personnes humaines dans les pratiques interventionnelles de recherche apparaissent comme de formidables outils de protection des personnes. Ainsi, à l'instar de la réglementation sur l'expérimentation animale, il pourrait rapidement devenir possible d'imposer une étape de réflexion initiale sur le bien-fondé scientifique, éthique et sociétal de la participation de patients réels aux essais cliniques, et a fortiori de la participation de volontaires sains. Cette réflexion, qui est déjà imposée par les principes internationaux de l'éthique de la recherche et reprise dans la loi française, pour les personnes en situation de particulière vulnérabilité comme les enfants, les femmes enceintes, les personnes les plus âgées ou les personnes sans couverture sociale, se traduit trop souvent par l'exclusion de ces populations des essais cliniques ciblant la population générale, et la difficulté de financer les essais spécifiques pourtant incontournables. Il s'ensuit une iniquité d'accès à l'innovation, ou a minima une inégalité des conditions de sécurité de cet accès, lorsqu'il est offert hors de l'AMM, au prix d'une majoration des risques et sous la responsabilité des prescripteurs.

Une politique plus large et plus juste de remplacement des personnes humaines pourrait donc être mise en place à moyen terme, chaque fois qu'elle sera techniquement possible, dans l'esprit de diminuer le nombre des personnes exposées à des risques, même minimes, et/ou à des contraintes, sans préjudice de leur accès à l'innovation thérapeutique. Dans le cas des essais de phase III, la plus-value de la participation de témoins réels est à l'évidence à mettre en balance avec le manque d'appétence des participants pour le bras contrôle, et donc les difficultés à les maintenir dans l'essai. Ces réticences, qui peuvent au maximum se manifester par un refus d'emblée d'entrer dans l'essai ou une rétractation dès la connaissance d'une randomisation dans le bras contrôle, peuvent s'accompagner d'une demande claire d'accès au principe actif testé, dès lors que celui-ci est prometteur. L'incertitude de son efficacité, qui justifie la randomisation pour les scientifiques, est plutôt perçue par les patients comme un obstacle infondé à l'accès au produit, perçu comme une perte de chance. À l'inverse, le remplacement de certains patients par des patients virtuels dans les bras des produits testés permettrait de réduire l'exposition des personnes à des risques, tout en réalisant des essais et l'accès aux résultats, et donc soit l'abandon plus rapide des produits dangereux ou inutiles, soit l'accès à des traitements validés.

Une telle politique de réduction des interventions portant atteinte à l'intégrité du corps des personnes humaines sans bénéfice thérapeutique direct pour elles, ou dans l'objectif de réduire les risques d'une intervention dont ils auraient spécifiquement besoin, devrait s'appliquer tout autant au contexte pédagogique ou de la sécurisation des soins individuels. Ainsi, les techniques de productions de patients virtuels, jumeaux numériques de patients réels, permettent déjà aux chirurgiens ou autres médecins interventionnels un apprentissage des gestes les plus invasifs, sans risque physique pour les patients et dans le respect de leur dignité.

Bien sûr, de telles modifications des pratiques doivent être accompagnées non seulement de référentiels scientifiques validés, mais aussi d'une rénovation des référentiels de l'éthique de la recherche. D'ailleurs au-delà d'une possible obligation de réflexion sur la réduction du nombre de personnes exposées aux essais cliniques sans intérêt pour elle, de nouveaux devoirs d'information sont à réfléchir, notamment vis-à-vis des personnes qui recevront des traitements validés dans ces nouvelles conditions. Une obligation renforcée de pharmacovigilance post-AMM pourrait, en complément des collèges de garantie humaine, contribuer à sécuriser et à valider ces techniques. Enfin, une révision des guidelines pour les comités d'éthique de la recherche (CPP en France ou IRB), de même qu'une formation de leurs membres semble nécessaire. Une réflexion parallèle sur le statut des cohortes virtuelles une fois créées, de leurs périmètres de réutilisabilité, et des conditions de cette réutilisabilité pourrait être envisagé dans un triple objectif d'intérêt public, de parcimonie, et de prise en compte des enjeux environnementaux. Dans cet objectif, il est d'ailleurs probable que les cohortes de patients virtuels pour la recherche contribuent encore à diminuer le recours à l'expérimentation animale.

Par ailleurs, il est nécessaire dans le cadre d'une recherche rétrospective de faire valider le SIA pour faire la preuve de l'efficacité des cohortes artificielles avant qu'elles ne soient utilisées en vie réelle pour de nouvelles recherches (recherches prospectives). Cela implique notamment que le développement de l'algorithme se fasse dans le cadre d'un protocole scientifique rigoureux, évalué par un comité éthique et respectueux de la vie privée des personnes concernées. Cette preuve de concept doit être publiée pour une évaluation par les pairs.



Yann-Maël Le Douarin

Conseiller médical à la DGOS et chef du département Santé
et Transformation numérique
Ministère du Travail de la Santé et des Solidarités

**Les données artificielles offrent
une opportunité prometteuse pour
notre système de santé et
pourraient jouer un rôle essentiel
dans la structuration des parcours
de soins à long terme. Elles
incarnent de manière tangible les
convergences entre les
mathématiques, la médecine et le
numérique.**

”

07

CONCLUSION

Ce livre blanc nous montre combien les données artificielles pourraient à terme jouer un rôle important dans la recherche clinique et l'optimisation des algorithmes de santé notamment.

Il s'agit bien sûr d'enjeux diagnostiques, pronostiques et thérapeutiques pour les patients et les professionnels de santé, mais également d'un enjeu de compétitivité et de valeur ajoutée pour la France et pour l'Europe.

La génération de données artificielles est un exemple d'utilisation secondaire des données de santé telle que décrite dans le rapport « Fédérer les acteurs de l'écosystème pour libérer l'utilisation secondaire des données de santé » rédigé par Mr Jérôme Marchand-Arvier, Pr Stéphanie Allassonnière, Mr Aymeril Hoang et Dr Anne-Sophie Jannot.⁴⁹

Dans ce document, nous avons montré qu'il existe déjà des fondations pour ouvrir les travaux nécessaires à la définition et à la mise en œuvre d'une validation de ces « patients d'un genre nouveau ». Il s'agit notamment de la mise en place des collèges de garantie humaine (pour valider la fiabilité et la sécurité des données produites par l'intelligence artificielle) et bien sûr leur reconnaissance méthodologique et scientifique pour la recherche in silico.

L'Agence de l'Innovation en Santé ainsi que l'Agence Nationale de la Recherche notamment pilotent actuellement des groupes de travail sur l'usage des nouvelles méthodologies de recherche clinique pour stimuler l'attractivité de la France dans ce domaine stratégique porté par France 2030. D'autre part le rapport interministériel pour « libérer l'utilisation secondaire des données de santé » publié par l'IGAS en janvier 2024 ainsi que les recommandations remises au Président de la République en mars par la Commission de l'Intelligence Artificielles permettent également de faire avancer rapidement ces sujets.

Ce livre blanc montre que tous les acteurs sont et se sentent concernés et comprennent le potentiel de ces nouvelles technologies, mais aussi le besoin de régulation. Nous espérons avoir apporté le début d'un éclairage sur ce sujet d'avenir.

⁴⁹ https://sante.gouv.fr/IMG/pdf/rapport_donnees_de_sante.pdf



Dr. Jean-Louis Fraysse

Co-fondateur de BOTdesign

Membre du groupe éthique de la Délégation du Numérique en Santé et du groupe de réflexion sur ces nouvelles méthodologies de recherche clinique de l'AIS et du F-Crin

Dans son avis du 2 avril 2024 le Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé (CCNE) rappelle « l'impérieuse nécessité de protéger les participants aux essais cliniques ». Les nouvelles approches de recherches médicales et notamment les patients artificiels générés à partir d'un petit nombre de patients réels contribueront à cet objectif.



RÉFÉRENCES BIBLIOGRAPHIQUES

08

1. Rebuffi, S.-A. et al. Data Augmentation Can Improve Robustness. Preprint at <https://doi.org/10.48550/arXiv.2111.05328> (2021).
2. Kingma, D. P. & Welling, M. Auto-Encoding Variational Bayes. (2013).
3. Goodfellow, I. J. et al. Generative Adversarial Networks. Preprint at <https://doi.org/10.48550/arXiv.1406.2661> (2014).
4. Sohl-Dickstein, J., Weiss, E. A., Maheswaranathan, N. & Ganguli, S. Deep Unsupervised Learning using Nonequilibrium Thermodynamics. ArXiv150303585 Cond-Mat Q-Bio Stat (2015).
5. Hu, Q., Li, H. & Zhang, J. Domain-Adaptive 3D Medical Image Synthesis: An Efficient Unsupervised Approach. in Medical Image Computing and Computer Assisted Intervention – MICCAI 2022 (eds. Wang, L., Dou, Q., Fletcher, P. T., Speidel, S. & Li, S.) 495–504 (Springer Nature Switzerland, 2022). doi:10.1007/978-3-031-16446-0_47.
6. Diamantis, D. E., Gatoula, P. & Iakovidis, D. K. EndoVAE: Generating Endoscopic Images with a Variational Autoencoder. in 2022 IEEE 14th Image, Video, and Multidimensional Signal Processing Workshop (IVMSP) 1–5 (2022). doi:10.1109/IVMSP54334.2022.9816329.
7. Barbano, R., Arridge, S., Jin, B. & Tanno, R. Chapter 26 - Uncertainty quantification in medical image synthesis. in Biomedical Image Synthesis and Simulation (eds. Burgos, N. & Svoboda, D.) 601–641 (Academic Press, 2022). doi:10.1016/B978-0-12-824349-7.00033-5.
8. Chadebec C, Thibeau-Sutre E, Burgos N, and Allasonnière S. Data Augmentation in High Dimensional Low Sample Size Setting Using a Geometry-Based Variational Autoencoder. IEEE Pattern Analysis and Machine Intelligence..
9. Rhodes, C. The Man With 100,000 Brains: AI's Big Donation to Science. NVIDIA Blog <https://blogs.nvidia.com/blog/2022/05/30/ai-brain-images-kcl/> (2022).
10. Nie, D. et al. Medical Image Synthesis with Deep Convolutional Adversarial Networks. IEEE Trans. Biomed. Eng. 65, 2720–2730 (2018).
11. Wolterink, J. M. et al. Deep MR to CT Synthesis using Unpaired Data. ArXiv170801155 Cs (2017).
12. Hu, Y. et al. Freehand Ultrasound Image Simulation with Spatially-Conditioned Generative Adversarial Networks. in vol. 10555 105–115 (2017).
13. Chuquicusma, M. J. M., Hussein, S., Burt, J. & Bagci, U. How to Fool Radiologists with Generative Adversarial Networks? A Visual Turing Test for Lung Cancer Diagnosis. Preprint at <https://doi.org/10.48550/arXiv.1710.09762> (2018).
14. Baur, C., Albarqouni, S. & Navab, N. MelanoGANs: High Resolution Skin Lesion Synthesis with GANs. Preprint at <https://doi.org/10.48550/arXiv.1804.04338> (2018).
15. Wolterink, J. M., Leiner, T. & Isgum, I. Blood Vessel Geometry Synthesis using Generative Adversarial Networks. Preprint at <https://doi.org/10.48550/arXiv.1804.04381> (2018).
16. Skandarani, Y., Jodoin, P.-M. & Lalande, A. GANs for Medical Image Synthesis: An Empirical Study. Preprint at <https://doi.org/10.48550/arXiv.2105.05318> (2021).
17. Khader, F. et al. Medical Diffusion: Denoising Diffusion Probabilistic Models for 3D Medical Image Generation. Preprint at <https://doi.org/10.48550/arXiv.2211.03364> (2023).

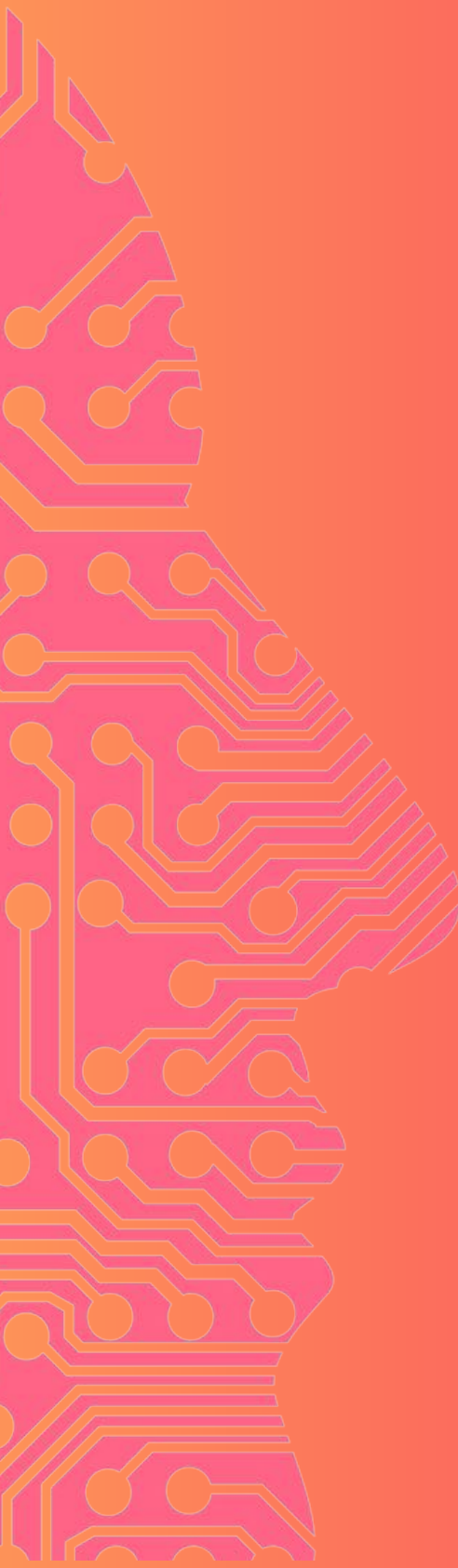
- 18.Kazerouni, A. et al. Diffusion models in medical imaging: A comprehensive survey. *Med. Image Anal.* 88, 102846 (2023).
- 19.Pan, S. et al. 2D medical image synthesis using transformer-based denoising diffusion probabilistic model. *Phys. Med. Biol.* 68, 105004 (2023).
- 20.Dorjsembe, Z., Pao, H.-K., Odonchimed, S. & Xiao, F. Conditional Diffusion Models for Semantic 3D Medical Image Synthesis. Preprint at <https://doi.org/10.48550/arXiv.2305.18453> (2023).
- 21.Weber, T., Ingrisch, M., Bischl, B. & Rügamer, D. Cascaded Latent Diffusion Models for High-Resolution Chest X-ray Synthesis. Preprint at <https://doi.org/10.48550/arXiv.2303.11224> (2023).
- 22.Peirlinck, M. et al. Precision medicine in human heart modeling. *Biomech. Model. Mechanobiol.* 20, 803–831 (2021).
- 23.Applied Sciences³⁵ | Free Full-Text | MUSIC: Cardiac Imaging, Modelling and Visualisation Software for Diagnosis and Therapy. <https://www.mdpi.com/2076-3417/12/12/6145>.
- 24.Nakatani, Y. et al. Preoperative personalization of atrial fibrillation ablation strategy to prevent esophageal injury: Impact of changes in esophageal position. *J. Cardiovasc. Electrophysiol.* 33, 908–916 (2022).
- 25.Banus, J., Lorenzi, M., Camara, O. & Sermesant, M. Biophysics-based statistical learning: Application to heart and brain interactions. *Med. Image Anal.* 72, 102089 (2021).
- 26.Levine, S. et al. Dassault Systèmes' Living Heart Project. in *Modelling Congenital Heart Disease: Engineering a Patient-specific Therapy* (eds. Butera, G., Schievano, S., Biglino, G. & McElhinney, D. B.) 245–259 (Springer International Publishing, 2022). doi:10.1007/978-3-030-88892-3_25.
- 27.Segars, W. P., Veress, A. I., Sturgeon, G. M. & Samei, E. Incorporation of the Living Heart Model Into the 4-D XCAT Phantom for Cardiac Imaging Research. *IEEE Trans. Radiat. Plasma Med. Sci.* 3, 54–60 (2019).
- 28.Kaboudian, A., Cherry, E. M. & Fenton, F. H. Real-time interactive simulations of large-scale systems on personal computers and cell phones: Toward patient-specific heart modeling and other applications. *Sci. Adv.* 5, eaav6019 (2019).
- 29.Peterlík, I., Duriez, C. & Cotin, S. Modeling and Real-Time Simulation of a Vascularized Liver Tissue. in *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2012* (eds. Ayache, N., Delingette, H., Golland, P. & Mori, K.) 50–57 (Springer, 2012). doi:10.1007/978-3-642-33415-3_7.
- 30.Plantefève, R., Haouchine, N., Radoux, J.-P. & Cotin, S. Automatic Alignment of Pre and Intraoperative Data Using Anatomical Landmarks for Augmented Laparoscopic Liver Surgery. in *Biomedical Simulation* (eds. Bello, F. & Cotin, S.) 58–66 (Springer International Publishing, 2014). doi:10.1007/978-3-319-12057-7_7.
- 31.Mendizabal, A., Márquez-Neila, P. & Cotin, S. Simulation of hyperelastic materials in real-time using deep learning. *Med. Image Anal.* 59, 101569 (2020).

32. Brunet, J.-N. et al. Physics-Based Deep Neural Network for Augmented Reality During Liver Surgery. in *Medical Image Computing and Computer Assisted Intervention – MICCAI 2019* (eds. Shen, D. et al.) 137–145 (Springer International Publishing, 2019). doi:10.1007/978-3-030-32254-0_16.
33. Pulkkinen, A., Werner, B., Martin, E. & Hynynen, K. Numerical simulations of clinical focused ultrasound functional neurosurgery. *Phys. Med. Biol.* 59, 1679 (2014).
34. Jaroudi, R., Åström, F., Johansson, B. T. & Baravdish, G. Numerical simulations in 3-dimensions of reaction–diffusion models for brain tumour growth. *Int. J. Comput. Math.* 97, 1151–1169 (2020).
35. Santaniello, S., Gale, J. T. & Sarma, S. V. Systems approaches to optimizing deep brain stimulation therapies in Parkinson’s disease. *WIREs Syst. Biol. Med.* 10, e1421 (2018).
36. Johannis, P. et al.³⁵ The strength of surgical knots involves a critical interplay between friction and elastoplasticity. *Sci. Adv.* 9, eadg8861 (2023).
37. Boussès, Y., Brulat-Bouchard, N., Bouchard, P.-O., Abouelleil, H. & Tillier, Y. Theoretical prediction of dental composites yield stress and flexural modulus based on filler volume ratio. *Dent. Mater.* 36, 97–107 (2020).
38. Digital connection in the metaverse | Meta. Digital connection in the metaverse <https://about.meta.com/metaverse/>.
39. NVIDIA Omniverse. NVIDIA <https://www.nvidia.com/en-us/omniverse/>.
40. Cardoso, M. J. et al. MONAI: An open-source framework for deep learning in healthcare. Preprint at <https://doi.org/10.48550/arXiv.2211.02701> (2022).
41. Baowaly, M. K., Lin, C.-C., Liu, C.-L. & Chen, K.-T. Synthesizing electronic health records using improved generative adversarial networks. *J. Am. Med. Inform. Assoc. JAMIA* 26, 228–241 (2018).
42. Ghosheh, G., Li, J. & Zhu, T. A review of Generative Adversarial Networks for Electronic Health Records: applications, evaluation measures and data sources. Preprint at <https://doi.org/10.48550/arXiv.2203.07018> (2022).
43. Li, R. et al. Improving an Electronic Health Record–Based Clinical Prediction Model Under Label Deficiency: Network-Based Generative Adversarial Semisupervised Approach. *JMIR Med. Inform.* 11, e47862 (2023).
44. Li, J., Cairns, B. J., Li, J. & Zhu, T. Generating synthetic mixed-type longitudinal electronic health records for artificial intelligent applications. *Npj Digit. Med.* 6, 1–18 (2023).
45. Kotelnikov, A., Baranchuk, D., Rubachev, I. & Babenko, A. TabDDPM: Modelling Tabular Data with Diffusion Models. Preprint at <http://arxiv.org/abs/2209.15421> (2022).
46. He, H., Zhao, S., Xi, Y. & Ho, J. C. MedDiff: Generating Electronic Health Records using Accelerated Denoising Diffusion Model. Preprint at <https://doi.org/10.48550/arXiv.2302.04355> (2023).
47. Nikolentzos, G., Vazirgiannis, M., Xypolopoulos, C., Lingman, M. & Brandt, E. G. Synthetic electronic health records generated with variational graph autoencoders. *Npj Digit. Med.* 6, 1–12 (2023).

- 48.Liao, W. et al. Dual autoencoders modeling of electronic health records for adverse drug event preventability prediction. *Intell.-Based Med.* 6, 100077 (2022).
- 49.Muller, E., Zheng, X. & Hayes, J. Evaluation of the Synthetic Electronic Health Records. Preprint at <https://doi.org/10.48550/arXiv.2210.08655> (2022).
- 50.Yan, C. et al. A Multifaceted Benchmarking of Synthetic Electronic Health Record Generation Models. *Nat. Commun.* 13, 7609 (2022).
- 51.Kuo, N. I.-H., Jorm, L. & Barbieri, S. Synthetic Health-related Longitudinal Data with Mixed-type Variables Generated using Diffusion Models. Preprint at <https://doi.org/10.48550/arXiv.2303.12281> (2023).
- 52.Biswal, S. et al. EVA: Generating Longitudinal Electronic Health Records Using Conditional Variational Autoencoders. Preprint at <https://doi.org/10.48550/arXiv.2012.10020> (2020).
- 53.Lee, D. et al. ³⁵Generating sequential electronic health records using dual adversarial autoencoder. *J. Am. Med. Inform. Assoc.* 27, 1411–1419 (2020).
- 54.Theodorou, B., Xiao, C. & Sun, J. Synthesize High-dimensional Longitudinal Electronic Health Records via Hierarchical Autoregressive Language Model. Preprint at <https://doi.org/10.48550/arXiv.2304.02169> (2023).
- 55.Koval, I. et al. Forecasting individual progression trajectories in Huntington disease enables more powered clinical trials. *Sci. Rep.* 12, 18928 (2022).
- 56.Sauty, B. & Durrleman, S. Riemannian Metric Learning for Progression Modeling of Longitudinal Datasets. in 2022 IEEE 19th International Symposium on Biomedical Imaging (ISBI) 1–5 (2022). doi:10.1109/ISBI52829.2022.9761641.
- 57.Schiratti, J.-B., Allasonnière, S., Colliot, O. & Durrleman, S. Mixed-effects model for the spatiotemporal analysis of longitudinal manifold-valued data. in 5th MICCAI Workshop on Mathematical Foundations of Computational Anatomy (2015).
- 58.Shahriar, S. & Hayawi, K. Let's have a chat! A Conversation with ChatGPT: Technology, Applications, and Limitations. Preprint at <https://doi.org/10.48550/arXiv.2302.13817> (2023).
- 59.Liu, Y. et al. Summary of ChatGPT-Related Research and Perspective Towards the Future of Large Language Models. *Meta-Radiol.* 1, 100017 (2023).
- 60.OpenAI. GPT-4 Technical Report. Preprint at <https://doi.org/10.48550/arXiv.2303.08774> (2023).
- 61.Touvron, H. et al. LLaMA: Open and Efficient Foundation Language Models. Preprint at <https://doi.org/10.48550/arXiv.2302.13971> (2023).
- 62.Tu, T. et al. Towards Generalist Biomedical AI. Preprint at <https://doi.org/10.48550/arXiv.2307.14334> (2023).
- 63.Kline, A. et al. Multimodal machine learning in precision health: A scoping review. *Npj Digit. Med.* 5, 1–14 (2022).
- 64.Behrad, F. & Saniee Abadeh, M. An overview of deep learning methods for multimodal medical data mining. *Expert Syst. Appl.* 200, 117006 (2022).
- 65.Qiu, J. et al. Multimodal Representation Learning of Cardiovascular Magnetic Resonance Imaging. Preprint at <https://doi.org/10.48550/arXiv.2304.07675> (2023).
- 66.Belyaeva, A. et al. Multimodal LLMs for health grounded in individual-specific data. Preprint at <https://doi.org/10.48550/arXiv.2307.09018> (2023).

67. Bommasani, R. et al. On the Opportunities and Risks of Foundation Models. *ArXiv210807258 Cs* (2021).
68. Cranmer, K., Brehmer, J. & Louppe, G. The frontier of simulation-based inference. *Proc. Natl. Acad. Sci.* 117, 30055–30062 (2020).
69. Mirza Faisal Beg, Michael I. Miller, Alain Trounev, Laurent Younes. Computing Large Deformation Metric Mappings via Geodesic Flows of Diffeomorphisms *International Journal of Computer Vision* 61(2):139-157
70. Jean-Michel Marin, Pierre Pudlo, Christian P. Robert, Robin J. Ryder. Approximate Bayesian computational methods. *Stat Comput* (2012) 22:1167–1180
71. Jesús Murga-Moreno, Sònia Casillas, Antonio Barbadilla, Lawrence Uricchio and David Enard. An efficient and robust ABC approach to infer the rate and strength of adaptation.
72. <https://www.biorxiv.org/content/biorxiv/early/2023/09/28/2023.08.29.555322.full.pdf>
73. Louis Raynal, Jean-Michel Marin, Pierre Pudlo, Mathieu Ribatet, Christian P Robert, Arnaud Estoup. ABC random forests for Bayesian parameter inference. *Bioinformatics*, Volume 35, Issue 10, May 2019, Pages 1720–1728,
74. Etienne Maheux, Igor Koval, Juliette Ortholand, Colin Birkenbihl, Damiano Archetti, Vincent Bouteloup, Stéphane Epelbaum, Carole Dufouil, Martin Hofmann-Apitius & Stanley Durrleman. Forecasting individual progression trajectories in Alzheimer's disease. *Nature Communications* 14, 761 (2023)
75. Liu G, Catacutan DB, Rathod K, Swanson K, Jin W, Mohammed JC, et al. Deep learning-guided discovery of an antibiotic targeting *Acinetobacter baumannii*. *Nat Chem Biol.* 2023;19: 1342-50.
76. Wong F, Zheng EJ, Valeri JA, Donghia NM, Anahtar MN, Omori S, et al. Discovery of a structural class of antibiotics with explainable deep learning. *Nature.* 2023.
77. European Medicines Agency. ICH guideline E11A on pediatric extrapolation. 2022.
78. Musuamba FT, Skottheim Rusten I, Lesage R, Russo G, Bursi R, Emili L, et al. Scientific and regulatory evaluation of mechanistic in silico drug and disease models in drug development: Building model credibility. *CPT Pharmacometrics Syst Pharmacol.* 2021;10: 804-25.
79. Azizi Z, Zheng C, Mosquera L, Pilote L, El Emam K, Collaborators G-F. Can synthetic data be a proxy for real clinical trial data? A validation study. *BMJ Open.* 2021;11: e043497.
80. Thorlund K, Dron L, Park JJH, Mills EJ. Synthetic and External Controls in Clinical Trials - A Primer for Researchers. *Clin Epidemiol.* 2020;12: 457-67.
81. Lambert J, Lengline E, Porcher R, Thiebaut R, Zohar S, Chevret S. Enriching single-arm clinical trials with external controls: possibilities and pitfalls. *Blood Adv.* 2023;7: 5680-90.
82. Bakker E, Plueschke K, Jonker CJ, Kurz X, Starokozhko V, Mol PGM. Contribution of Real-World Evidence in European Medicines Agency's Regulatory Decision Making. *Clin Pharmacol Ther.* 2023;113: 135-51.

83. Vanier A, Fernandez J, Kelley S, Alter L, Semenzato P, Alberti C, et al. Rapid access to innovative medicinal products while ensuring relevant health technology assessment. Position of the French National Authority for Health. *BMJ Evid Based Med*. 2024;29: 1-5.
84. Jacob E, Perrillat-Mercerot A, Palgen JL, L'Hostis A, Ceres N, Boissel JP, et al. Empirical methods for the validation of time-to-event mathematical models taking into account uncertainty and variability: application to EGFR + lung adenocarcinoma. *BMC Bioinformatics*. 2023;24: 331.
85. You R, Weber S, Bieth B, Vandemeulebroecke M. Innovative Pediatric Development for Secukinumab in Psoriasis: Faster Patient Access, Reduction of Patients on Control. *Clin Pharmacol Ther*. 2022;111: 697-704.
86. Neal ML, Trister AD, Cloke T, Sodt R, Ahn S, Baldock AL, et al. Discriminating survival outcomes in patients with glioblastoma using a simulation-based, patient-specific response metric. *PLoS One*. 2013;8: e51951.
87. Switchenko JM, Heeke AL, Pan TC, Read WL. The use of a predictive statistical model to make a virtual control arm for a clinical trial. *PLoS One*. 2019;14: e0221336.
88. Guillaudeau M, Rousseau O, Petot J, Bennis Z, Dein CA, Goronflot T, et al. Patient-centric synthetic data generation, no reason to risk re-identification in biomedical data analysis. *NPJ Digit Med*. 2023;6: 37.
89. Creemers JHA, Ankan A, Roes KCB, Schroder G, Mehra N, Figdor CG, et al. In silico cancer immunotherapy trials uncover the consequences of therapy-specific response patterns for clinical trial design and outcome. *Nat Commun*. 2023;14: 2348.
90. Senellart A, Chadebec C, Allasonnière S. Improving Multimodal Joint Variational Autoencoders through Normalizing Flows and Correlation Analysis. <https://arxiv.org/abs/2305.11832>
91. Tavaré, S., Balding, D., Griffith, R., Donnelly, P.: Inferring coalescence times from DNA sequence data. *Genetics* 145(2), 505–518 (1997)



MERCI

Pour plus d'informations :

stephanie.allassonniere@u-paris.fr

jlfraysse@botdesign.net