



Enquête sur l'Intelligence Artificielle dans le Département MathNum : Référentiel, Positionnement et Dynamique

Frédéric Garcia, DR INRAE, Unité Mathématiques et Informatique Appliquées de Toulouse (MIAT)

Vincent Guigue, PR AgroParisTech, UMR Mathématiques et Informatique Appliquées de Paris-Saclay (MIA-PS)

Septembre 2025 – version diffusable

Enquête sur l'Intelligence Artificielle dans le Département MathNum : Référentiel, Positionnement et Dynamique

Frédéric Garcia¹ et Vincent Guigue^{2,3}

¹MIAT, INRAE
²MLA-PS, INRAE
³AgroParisTech

Novembre 2025, version partielle diffusable

Résumé

Les objectifs de ce rapport sont de proposer un référentiel permettant de décrypter les différents courants dans le domaine très large, foisonnant et dynamique de l'Intelligence Artificielle puis de situer les unités MathNum dans ce cadre. La vague actuelle de l'IA, principalement centrée sur l'apprentissage automatique, a un impact fort sur de nombreuses disciplines à la fois à l'intérieur et à l'extérieur de l'IA (e.g. simulation numérique, modélisation dans les domaines applicatifs). Le référentiel établi doit nous permettre de mieux cerner la nature de ces bouleversements dans différents domaines. L'objectif devient alors non seulement de positionner les unités par rapport à cette transition mais également d'analyser la vision des unités pour le futur.

Keywords: Intelligence Artificielle, Statistiques, Machine Learning, Deep-learning, IA générative, IA Hybride

■ Table des matières

1	Introduction	3
2	Brève histoire de l'IA	4
3	Éléments d'un référentiel	4
3.1	Modalités de données et problèmes génériques associés	5
	(Tab) Données tabulaires	
	(ST) Séries temporelles & capteurs	
	(SpT) Données spatiales, séries spatio-temporelles	
	(PM) Paroles et musique	
	(Vis) Image, (Vid) vidéo	
	(Tel) Télémétrie	
	(Spm) Spectrométrie	
	(TAL) Texte	
	(Exp) Données d'expertise	
	(Om) Données omiques	
	(Rec) Données d'interaction & recommandation	
	(Enq) Données d'enquête	
	(Par) Données de structuration IA	
3.2	Approches et modélisation en Intelligence Artificielle . . .	7
	(IAS) Logique et IA symbolique	
	(KB) Ingénierie des connaissances et ontologies	
	(RI) Indexation et Recherche d'Information	
	Apprentissage statistique	
	(ML1) Modélisation statistique	
	(ML2) Formulation régularisée	
	(ML3) Ensembles, forêts	
	(ML4) Clustering (non supervisé)	
	Deep-learning	
	(DL1) Perceptrons, Embeddings	
	(DL2) CNN, RNN, GNN, Transformer	
	(DL3) Modèles pré-entraînés	
	(DL4) IA Générative	
	(DL5) PINNs	
	(CTR) Automatique & Contrôle	
	Apprentissage par renforcement	
	(RL1) Processus de Markov (discret)	
	(RL2) Modélisation continue en RL	
	(RL3) Modèles pré-entraînés (AlphaFold)	
	(Seq) Approches séquentielles	
	Optimisation, satisfaction de contraintes et recherche opérationnelle	
	(OPT) Optimisation continue	
	(RO) Recherche Opérationnelle	
	(GR) Graphes	
	Simulation	
	(SIM1) Modélisation Mécaniste	
	(SIM2) Simulation multi-agents	
3.3	De la théorie à la pratique	10
	Recherche amont	
	Au cœur de l'IA	
	Utilisateurs des outils d'IA	
	Utilisateurs finaux	
3.4	Modalités multiples, problématique composite et mots clés émergents	13
	IA générative	
	Modèles de fondation	
	Approches hybrides & systèmes dynamiques	
	Jumeaux numériques	
	Robotique	
3.5	Outils de développement et moyens de calcul	14
	Développement	
	Calcul	
3.6	Groupes bibliographiques	14
4	Positionnement des unités MathNum	15
4.1	BIOSP	15
4.2	LISC	16
4.3	ITAP	16
4.4	TSCF	17
4.5	MISTEA	17
4.6	MIAT	18
4.7	MAIAGE	18

4.8	MIA-PS	19
4.9	OPAALÉ	20
4.10	TETIS	20
5	Conclusion	21
A	Premiers résultats	22
B	Guide d'entretien	22
B.1	Applications & Modalités	22
B.2	De la théorie à la pratique	22
B.3	Types d'IA, cadre de développement	22
B.4	Bibliographie	25
B.5	Prospective	25

1. Introduction

Comme l'informatique, l'intelligence artificielle (IA) est un domaine de recherche que l'on peut encore qualifier de récent. Depuis ses origines il y a 70 ans, le périmètre de l'IA a régulièrement évolué, ce champ de recherche foisonnant de nombreuses propositions et la définition même du terme faisant débat, récupéré successivement par différentes sous-communautés au gré des courants et des modes. Le domaine de l'Intelligence Artificielle est marqué depuis 15 ans par une tendance forte des GAFAM/GAMMA (Google, Apple, Facebook/Meta, Amazon, Microsoft, ...) à préempter l'innovation à travers des débauchages massifs de chercheurs et une centralisation des moyens de calcul mondiaux. Cette tendance a notamment eu pour conséquence la redéfinition du terme *IA* au milieu des années 2010 vers les techniques de machine-learning et plus précisément de deep-learning, sous l'impulsion du trio J. Bengio, Y. LeCun et G. Hinton. Ce cadre, déjà très mouvant, est aujourd'hui percuté par l'ouverture au public de chatGPT qui a bouleversé les usages autour de ces outils et déplacé –de nouveau– le centre de gravité de la communication vers la notion d'IA générative.

Ces dernières étapes, aussi marquantes soient-elles, ne doivent toutefois pas faire oublier des dizaines d'années d'avancées pertinentes de la communauté scientifique dans d'autres sous-domaines de l'intelligence artificielle, comme le raisonnement logique, la représentation des connaissances, les méthodes heuristiques, etc., et qui correspondent à autant de compétences progressivement développées au sein du département MathNum.

Il est alors apparu nécessaire pour le département MathNum d'Inrae de faire un "point d'étape" sur la place de l'IA dans les recherches conduites au sein des unités du département, en considérant bien sûr la définition la plus large possible de l'IA, pour ne pas mettre de côté des techniques ne relevant pas du machine learning ou de l'IA générative, mais qui restent pourtant critiques dans le développement des méthodes d'intégration et d'analyse des données et modèles considérés au sein des unités MathNum, et plus largement des systèmes informatiques modernes. Dans le même temps, nous nous interrogerons sur la dynamique des domaines de recherche en IA : la dernière vague est très large et propose parfois de nouvelles approches très efficaces mais le plus souvent, elle correspond aussi à des outils complémentaires ou de nouveaux modules intégrables dans des chaînes historiques. Parfois aussi, évidemment, les nouvelles propositions sont simplement orthogonales aux domaines de recherche des équipes.

Ce rapport, dont la rédaction nous a été confiée par la direction de MathNum, vise ainsi en premier lieu à proposer un référentiel décrivant les nombreuses facettes de l'IA et permettant de positionner les recherches menées et les compétences présentes au sein des différentes unités du département. Mais au-delà du positionnement, nous avons également tenté de décrire les directions actuellement prises par les unités et leurs équipes, entre leur expertise historique et la prise en compte des techniques émergentes, en particulier du deep-learning et de l'IA générative. Enfin, nous avons tenté d'aborder une dimension supplémentaire, concernant le ressenti par rapport à cette vague actuelle de l'IA, et des transformations qu'elle entraîne à la fois dans les questions et dans les pratiques de recherche. Les références mentionnées dans le rapport sont

très généralistes, elles étaient le référentiel mais ne concernent pas directement l'état de l'art.

Pour ce faire, nous avons fait le choix de baser ce rapport sur une série d'entretiens avec les unités MathNum. Ces entretiens ont été conduits de manière semi-directive, avec un groupe représentatif du laboratoire (entre deux et six personnes selon les unités). Le guide d'entretien utilisé est rapporté en section B. Les entretiens ont toutefois été conduits de manière assez libre, en laissant un certain nombre de questions dériver de manière à explorer au mieux les ressentis et les associations d'idées des uns et des autres, autour du thème de l'intelligence artificielle.

2. Brève histoire de l'IA

Le foisonnement des approches en IA, qui s'est fait en plusieurs vagues successives (Systèmes Experts, Machine-Learning, Big-Data, Deep-Learning, IA génératives, ...) depuis des dizaines d'années, semble rendre difficile la définition de l'*Intelligence Artificielle*. Néanmoins, revenir aux origines du terme permet de donner du sens et de la cohérence à cette profusion de méthodes. L'IA, c'est donc à la base une discipline scientifique qui se construit avec un projet clair : élaborer les principes, concevoir et étudier les algorithmes associés qui permettent le développement opérationnel de machines informatiques "intelligentes", comme le proposait Minsky, capables de faire des choses qui nécessiteraient de l'intelligence si elles étaient faites par des humains ("*AI : the science of making machines do things that would require intelligence if done by men*"). Ou, comme le proposait Turing avec son test, capables de se rendre indistinguables d'un humain testé simultanément. Si au départ l'IA reposait explicitement sur une démarche interdisciplinaire mixant psychologie, mathématiques, informatique, automatique, électronique etc., avec le temps l'IA est devenue progressivement une discipline basée essentiellement sur les algorithmes et l'informatique, laissant à la robotique également en plein développement le soin de résoudre la question de la corporalité, de l'inscription matérielle de la machine dans la réalité. Russel et Norvig pouvaient écrire "*An intelligent agent is anything that can be viewed as perceiving its environment through sensors and acting rationally upon that environment through effectors*", la mission de l'IA est bien de traduire de manière algorithmique ces tâches de perception et de décision rationnelle.

Dans ce cadre, un problème d'IA repose toujours sur un schéma $x \rightarrow (A) \rightarrow y$, où (A) est un algorithme, une suite d'instructions et d'opérations permettant de résoudre une classe de problèmes, x étant une description formelle d'une instance de cette classe de problèmes (les données d'entrées de l'algorithme), et y la réponse élaborée par le calcul (la sortie de l'algorithme). Les types de problèmes à résoudre ont historiquement structuré fortement les sous-communautés de l'IA. On peut ainsi classiquement considérer quatre champs historiques au sein de l'IA : *decision and planning*, *knowledge representation and management*, *modeling, simulation and optimisation*, et enfin *machine-learning*. Evidemment, il y a toujours eu des recouvrements entre ces catégories, mais globalement elles tiennent encore. En *decision and planning*, le problème posé est du type : quoi faire pour atteindre un but ? La sortie y fait référence à un choix, une décision, un plan d'action, un comportement à adopter. Les entrées x décrivent le but, l'environnement, les choix et actions possibles, et leurs

effets sur l'environnement. Pour *knowledge representation and management*, le problème porte plus généralement sur la question de la manipulation des concepts, des catégories, et des raisonnements logiques que l'on peut faire dessus. Les entrées x et les sorties y peuvent être des descriptions formelles de cas, d'objets, d'entités, mais aussi de connaissances structurées sous formes de règles, de théories, de graphes, etc. Avec *modeling, simulation and optimisation*, l'IA a introduit des méthodes originales de modélisation largement basées sur les interactions entre entités (SMA, GA, DEVS, ...), et a développé des approches originales de résolution de problème exploitant des heuristiques pour construire des solutions à des instances de très grandes dimensions, de manière exacte ou approchée, avec des contraintes de ressources (CSP, MDP, ...). La dernière catégorie, le *machine-learning*, peut être vue comme une méta-catégorie qui traverse potentiellement les trois autres. Il s'agit de concevoir des algorithmes qui vont "apprendre" à résoudre des problèmes sur la base de cas observés, d'expériences, ces problèmes pouvant être des problèmes de décision, de planification, d'optimisation, de raisonnement, etc. Historiquement, le *machine-learning*, ou apprentissage automatique, a plutôt été influencé à ses débuts par les communautés *electrical engineering* et *cybernetics*, avec les premiers travaux sur le perceptron et les réseaux de neurones, et jusque dans les années 2000, cette branche était plutôt à part de la communauté mainstream de l'IA, qui était plutôt orientée symbolique. Il y avait néanmoins une seconde branche apprentissage symbolique, qui mettait l'accent sur le raisonnement et les objets structurés, en lien fort avec la représentation des connaissances. A partir des années 1990, l'arrivée des données massives et l'augmentation des capacités de calcul a permis l'émergence et le développement de méthodes efficaces à la croisée de l'apprentissage et des statistiques (CART, Random Forests, SVM, ...). Mais ce sont les méthodes d'apprentissage profond, utilisant des réseaux de neurones possédant plusieurs couches de neurones cachés qui ont à partir des années 2010 réellement révolutionné le domaine du *machine learning*, par leur capacité à apprendre des tâches de plus en plus complexes sur la base de jeux de données de tailles de plus en plus importantes. Cela a clairement transformé le périmètre de la communauté IA au sens large, qui a donc intégré tout le champ du machine learning et de l'apprentissage statistique, mais également tout l'apprentissage à base de réseaux de neurones. Selon chatGPT, à qui nous posons la question, *30-40% of IJCAI papers in recent years are either directly related to deep learning or use it as a significant component of the research*.

3. Eléments d'un référentiel

Pour décrire un référentiel qui permettrait de positionner les recherches et les compétences en IA des unités du département, les types de problèmes à résoudre sont bien sûr importants à considérer, comme nous l'avons expliqué ci-dessus. Mais il nous est apparu également intéressant, sur la base du schéma générique $x \rightarrow (A) \rightarrow y$, de considérer les types de données d'entrées ou de sorties, décrivant les problèmes à résoudre, qui peuvent largement différer d'une sous-communauté de l'IA à l'autre. Il en va de même des types de modèles utilisés pour représenter ces données x et y , ou pour les exploiter au sein de (A) . Souvent, ces communautés sont liées à un domaine applicatif, que nous évoquerons. Types

de problèmes, types de données, de modèles, domaines applicatifs se croisent souvent, et notre référentiel ne visera pas l'exhaustivité, mais plutôt la bonne description et représentativité de ce qui se fait dans le département MathNum. Quelques références générales décrivent la pluralité des approches aujourd'hui disponibles en IA [27], [28].

3.1. Modalités de données et problèmes génériques associés

Traiter des images, des documents textuels, des séquences temporelles, des données tabulaires ou des processus définis par des experts implique des chaînes de traitements spécifiques et des classes d'algorithmes particulières, pour des problèmes généralement spécifiques.

La liste que nous avons envisagée est la suivante :

(Tab) Données tabulaires Les données les plus classiques en apprentissage automatique (Machine-Learning) sont les données tabulaires : issues de n'importe quel domaine applicatif, les ingénieurs et chercheurs collectent souvent les données d'expériences dans un tableur ; chaque ligne est alors une projection d'une instance sur un ensemble de caractéristiques observées [8].

Les variables descriptives peuvent-être numériques ou catégorielles (rarement ordinales). Les problématiques associées à ces données relèvent principalement de la **classification** et de la **régression** sur des valeurs continues, plus rarement de l'ordonnancement. Sur le volet non-supervisé (**clustering**), l'enjeu est de regrouper les individus par : ce type d'approche est typiquement très délicat du fait du grand nombre de degrés de liberté, il est généralement nécessaire d'introduire de la connaissance experte pour orienter le clustering dans une direction pertinente. Les données manquantes requièrent également un traitement particulier.

Extraire les régularités à partir de grands échantillons d'observations pour faire des prédictions est le cœur de l'analyse de données. Identifier ces régularités à partir de *petits échantillons* est un défi spécifique : le problème est alors quasiment impossible avec une approche agnostique. Parmi les différentes manières d'introduire des connaissances a priori, la plus courante consiste à introduire des hypothèses basées sur des modèles probabilistes : cela permet à la fois de guider l'apprentissage et d'obtenir des garanties théoriques sur la solution, en particulier des bornes et intervalles de confiance pour déterminer quand est ce qu'un résultat est significatif. Les modèles entraînés et les approches génératives sont un autre moyen de faire face à ce défi spécifique.

(ST) Séries temporelles & capteurs Les séries temporelles sont présentes dans de nombreux domaines applicatifs : séries financières, capteurs de données physiologiques (Température, EEG, ECG), données météorologiques (température, anémomètre, pluviométrie)... Tous les *capteurs* envoyant régulièrement des données génèrent des séries temporelles [16]. Les problématiques récurrentes sont la **prédiction** des valeurs futures et l'**interpolation** de valeurs intermédiaires. Mais il est aussi possible de classer des séries temporelles et de détecter des **anomalies** dans les flux de données, ce qui ouvre encore un peu plus les possibilités applicatives.

(SpT) Données spatiales, séries spatio-temporelles Les données spatiales (région, parcelle, POI –Point of Interest–...) sont

critiques dans de nombreuses applications.

Mais les données spatiales sont souvent combinées avec des aspects temporels pour former des *réseaux de capteurs*. Lorsque le réseau est déployé sur un territoire, nous parlerons de séries spatio-temporelles (e.g. données météorologiques), lorsqu'il est posé sur une machine ou un système complexe, nous irons plutôt vers le domaine de la *maintenance prédictive* et de l'usine 4.0. Dans le domaine agricole, cette modalité de données va permettre le développement de modèles de prédiction de rendement ou plus globalement de jumeaux numériques (sur lesquels nous reviendrons ultérieurement).

Le domaine des systèmes dynamiques relève historiquement de la simulation avec des EDP mais une tendance récente consiste à entraîner des modèles de machine-learning sur des données issues de simulations numériques (et directement sur des mesures) pour gagner en temps de calcul, en précision ou en pouvoir de généralisation selon les cas. Ces données spécifiques sont à la base des PINNs (*Physics Informed Neural Networks*).

(PM) Paroles et musique Les données audio sont des séries temporelles mais avant tout des données sémantiques : leur traitement est suffisamment spécifique pour qu'une communauté scientifique se soit bâtie autour. Les enjeux autour de la parole sont la **transcription** (segmentation et la reconnaissance des mots), la **séparation des sources** lorsque plusieurs personnes parlent simultanément [2]. Plus marginalement, des applications existent en reconnaissance du locuteur voire en diagnostic à partir de la voix (e.g. fatigue). La tâche la plus importante est de loin la transcription qui mêle traitement du signal et analyse sémantique issue du texte pour valider l'enchaînement des syllabes puis des mots.

Le domaine de l'informatique musicale travaille beaucoup sur la **classification** des morceaux, la construction de nomenclature et l'intégration des systèmes de recommandation. Dans toutes ces applications, cette communauté est en pointe sur les questions d'explicabilité qui sont critiques pour justifier les catégorisations et rendre les systèmes acceptables par les utilisateurs.

Cette communauté projette souvent les données dans un espace temps-fréquence qui met en évidence des caractéristiques discriminantes intéressantes. Cette projection donne à ces données les dimensions d'une image et offre l'opportunité d'exploiter et de spécialiser des outils issus de la communauté image. La question de la **génération** a pris beaucoup de place ces dernières années. Côté parole, l'enjeu est notamment de transformer les interfaces auxquelles nous sommes habitués ; côté musique, l'aide à la composition ou à l'accompagnement de morceaux devient peu à peu une réalité.

(Vis) Image, (Vid) vidéo A l'instar des données textuelles, la vision rassemble une communauté scientifique spécifique dite *Computer Vision* que l'on traduit parfois par *Perception*. Il est bon de rappeler que l'architecture *deep-learning* proposée en 2012 est à l'origine un système de reconnaissance d'objets dans les images. L'évolution récente de ce domaine suit la même dynamique que le texte en termes de performance et d'accessibilité. L'avènement des modèles de fondation fait actuellement converger les communautés texte et image : le terme de *données sémantiques* permet de regrouper texte, image et voix.

Les problématiques classiques en image sont la **détection** d'objets, la **localisation** de ces objets (via des boîtes englobantes)

et la **segmentation** des objets au niveau des pixels [25].

Il faut souligner la multiplication des capteurs d'imagerie avec les satellites, la télémétrie, les drones, les différentes techniques d'imagerie médicale, les caméras des véhicules semi-autonomes, les caméras de surveillance... Il a été démontré que les modèles d'analyse bénéficient les uns des autres en particulier en exploitant le pré-entraînement sur lequel nous reviendrons dans la section suivante.

La vidéo est une combinaison d'analyse d'images et de séries temporelles : l'enjeu est de combiner la détection et le **suivi** d'objet ou la détection de rupture. La vidéo peut aussi intégrer une dimension audio : la combinaison de l'analyse des différentes modalités de données correspond au domaine des modèles de fondation.

(Tel) Télémétrie La télémétrie est un nom générique pour désigner des technique de mesure des variables à distance (comme la température, la pression, la vitesse, les vibrations etc.). Ce terme regroupe aussi bien les capteurs de maintenance prédictive que certaines mesures issues de satellites avec de nombreuses applications dans le domaine de l'agriculture [13]. Le terme est cependant structurant pour certaines communautés scientifiques : nous le conservons malgré les superpositions larges avec d'autres modalités mentionnées.

(Spm) Spectrométrie Les données de *spectrométrie* sont aussi des données de perception mais pas forcément dans le domaine visible. L'enjeu est souvent de **déterminer la composition** chimique ou les **propriétés physiques** d'un échantillon en mesurant l'absorption, l'émission, ou la réflexion de la lumière sur une certaine gamme de longueurs d'onde.

(TAL) Texte La communauté scientifique autour de l'analyse automatique du texte est l'une des plus importantes de l'analyse de données mais, comme pour l'image, elle est aussi un peu à part. Les chercheurs en *Traitement Automatique de la Langue Naturelle* (TALN), ou *Natural Language Processing* (NLP) en anglais, sont souvent dans des laboratoires dédiés (ou au moins des équipes spécifiques), ils ont leur propre réseau de conférences et revues (notamment estampillées ACL). Cet isolement relatif s'explique par la spécificité des données et des méthodes associées et la taille des corpus, historiquement beaucoup plus volumineux que dans d'autres domaines [3]. Les mots sont des entités discrètes, porteuses de sens, comportant de nombreuses flexions : manipuler cette modalité conduit à des formulations en très grande dimension où les heuristiques jouent (ou ont joué) un rôle prépondérant. Les applications en texte sont variées : **indexation & moteur de recherche**, **classification** de documents (e.g. pertinence, fouille d'opinion), **extraction d'information** pour la construction de bases de connaissances (extraction & identification des entités nommées, des relations)...

Comme pour l'image, la maturité des outils impacte le périmètre des communautés scientifiques : beaucoup plus d'acteurs peuvent se saisir de ces outils et un continuum dense apparaît depuis les linguistes modélisant les propriétés des textes, les statisticiens et informaticiens travaillant sur les modèles et architectures, les développeurs recombinaient des modules existants jusqu'aux utilisateurs qui peuvent aujourd'hui appliquer des modèles très complexes beaucoup plus facilement.

L'ouverture des outils en texte, image et son a permis de construire des ponts entre les communautés. Les modèles ont

bénéficié de ces échanges et l'on assiste aujourd'hui à une convergence vers des modèles de fondation tirant parti de toutes les modalités pour améliorer les performances d'un large spectre d'applications.

(Exp) Données d'expertise A l'origine du terme *Intelligence Artificielle* en 1956, l'idée était principalement de mimer le comportement d'un expert et le développement d'un système automatique passait donc par la collecte de données d'expertise à insérer manuellement dans un algorithme comme un arbre ou un système expert.

Ces systèmes sont encore beaucoup utilisés lorsque les données d'expertise sont disponibles (e.g. clé de détermination en biologie) et lorsque que la collecte d'échantillons est difficile ou impossible (e.g. fuites de gaz, défaillance de systèmes critiques...).

L'exploitation de ce type de données nécessite souvent des outils issus de différentes communautés : des outils de **logique** et de **raisonnement** pour généraliser les **connaissances** à différentes situations, des outils bases de données pour stocker et requêter ces connaissances, des outils de textes pour **extraire les connaissances** des textes et mieux modéliser la sémantique des connaissances [23].

Notons que les ontologies peuvent aussi être considérées comme des données d'expertise en entrée d'applications de plus haut niveau.

Les données d'expertise sont également une porte d'entrée pour la modélisation : depuis les outils de **simulation numérique** jusqu'aux **systèmes multi-agents**, l'ancrage des modèles dans la réalité dépend aussi de données d'expertise.

(Omq) Données omiques Les données omiques regroupent différentes modalités présentes dans les sciences de la vie à différentes échelles, en particulier *génomique*, *transcriptomique* et *protéomique*. Les données omiques fournissent les informations moléculaires fondamentales qui, par leur interaction et leur régulation, conduisent à l'expression des traits phénotypiques [19]. L'analyse de ces données est donc critique pour comprendre des maladies ou optimiser des processus biologiques. Les problématiques générales associées relèvent de la **classification** (identification) ou du **clustering** et de la **régression** (scoring).

Les passerelles entre la génomique et le TALN sont historiques puisqu'il s'agit de séquences d'objets discrets dans les deux cas. La taille des vocabulaires et les erreurs induites par les processus d'acquisition et la présence de nombreux variants sont cependant très spécifiques et mènent à des chaînes de traitements originales. Une fois les séquences codantes identifiées, la communauté scientifique travaille sur des problèmes de quantification et de lien avec les traits phénotypiques qui ne relèvent plus forcément de données séquentielles.

La communauté a été fortement impactée par les outils d'analyse qui ont permis d'accéder à de nouvelles données de manière massive. L'*intégration multi-omique* est une tendance forte consistant à considérer les données à toutes les échelles en même temps pour mieux identifier les caractères d'intérêt. Une partie de la communauté omique a été impactée par l'outil *AlphaFold* (2018) qui permet de prédire la structure 3D d'une protéine à partir de sa séquence, ouvrant la voie à une prédiction bien plus fine des propriétés associées à ces chaînes.

(Rec) Données d'interaction & recommandation Les clics et les notes d'utilisateurs sur le web 2.0 (i.e. interactifs et plus seulement accessibles en lecture) forment une modalité de données spécifiques souvent associées aux **systèmes de recommandation** [12]. L'enjeu est d'exploiter cette source pour suggérer à l'utilisateur des choix susceptibles de lui plaire. Lorsque ces suggestions reposent sur l'apprentissage de profils utilisateurs, nous parlerons de **filtrage collaboratif**.

Ces algorithmes forment aujourd'hui la clé de voûte du modèle économique d'Internet à travers la publicité en ligne mais impactent aussi de nombreuses interfaces web et logicielles. Il s'agit aussi d'outils qui sont utiles à l'optimisation des interfaces au sens beaucoup plus large : la collecte d'interactions sur de l'électroménager, des systèmes embarqués dans les voitures ou de la domotique permet de mieux cerner les catégories d'utilisateurs et d'améliorer itérativement les outils qui leur sont proposés.

Suivant la tendance générale, ces données d'interaction sont aujourd'hui souvent combinées à du texte, de l'image ou d'autres modalités pour enrichir les profils, affiner les prédictions ou mieux expliquer les suggestions.

(Eng) Données d'enquête L'analyse des données d'enquête constitue un outil historique pour comprendre les forces et les faiblesses d'un système complexe (e.g. un réseau de transport en commun) ou de quelque chose de subjectif (e.g. opinion, élection). Ces données sont assez spécifiques : la construction des questions de l'enquête requiert une expertise et les échantillons collectés sont souvent de taille très limitée et peuvent nécessiter un travail de ré-équilibre (de type sondage d'opinion). La collecte est également souvent faite sous forme d'enregistrement audio ou vidéo.

Ces données étaient historiquement essentiellement traitées manuellement et donc très à la marge du périmètre de l'intelligence artificielle. Les outils de l'IA sont aujourd'hui largement utilisés pour la transcription et, de plus en plus, pour l'analyse des données transcrites.

(Par) Données de structuration IA Cette catégories au nom abstrait vise à regrouper des données particulières : celles qui décrivent un problème à résoudre. Sur un plateau de jeu (échec, go, ...), il faut décrire les règles du jeu et la nature du terrain. Dans les problèmes de **planification** (e.g. Sudoku), il faut indiquer les éléments disponibles et les contraintes [1]. Les données fournies permettent ainsi la structuration d'un solveur.

Ce type de données correspond en général à des **problèmes séquentiels ou combinatoires** ; dans les jeux à deux joueurs, les décisions des joueurs sont cohérentes d'un tour à l'autre et correspondent à leur stratégie ; en **robotique**, chaque action élémentaire correspond à un mouvement global motivé par une tâche. De nombreuses applications complexes entrent dans cette définition où l'enjeu est donc de générer des séquences de décisions cohérentes par rapport à un système où les actions sont spécifiques et limitées par un ensemble de contraintes. Ces données de structure s'accompagnent de plus en plus d'observations partielles ou complètes de certaines instances. Dans l'exemple du jeu d'échec : vu sous l'angle des données d'instances, l'idée est d'apprendre à jouer à partir des parties enregistrées précédemment en particulier à l'aide d'algorithmes d'**apprentissage par renforcement** et de **processus de décision de Markov** ; vu sous l'angle combinatoire, l'idée est

de trouver la séquence d'actions optimale parmi une combinaison gigantesque de possibilités en prenant en compte les contraintes correspondant aux règles du jeu et en maximisant la probabilité de victoire. Il s'agit alors plus de **recherche opérationnelle** ou d'**optimisation sous contrainte** (SAT).

Cette ouverture permet d'intégrer ici les problèmes de **planification** et de **logistique**, de la gestion des entrepôts ou des emplois du temps à la construction d'itinéraires optimaux.

3.2. Approches et modélisation en Intelligence Artificielle

Toutes les données listées dans la section précédente ont leurs spécificités. Parfois il est possible de partir d'une approche générique et simplement adapter les entrées et/ou les sorties pour y faire face ; parfois, des architectures appropriées sont nécessaires pour relever ces challenges.

Nous proposons de construire un deuxième axe de notre référentiel de l'intelligence artificielle centré sur grandes familles d'approches, qui correspondent souvent à des communautés scientifiques déterminées. Identifier et regrouper ces approches est un travail qui a déjà été effectué par différents acteurs dont le Groupe de Recherche (GdR) RADIA – Raisonement, Apprentissage, et Décision en Intelligence Artificielle – [28] ou d'autres chercheurs renommés [27].

La grille d'analyse est *assez déséquilibrée* entre les approches historiques (e.g. IA symbolique) qui sont regroupées dans de vastes catégories et les approches plus récentes (e.g. deep-learning) qui sont découpées plus finement. C'est un choix assumé pour comprendre plus finement le positionnement des chercheurs par rapport à ces nouveaux outils, pour décrire les hybridations possibles entre approches historiques et machine-learning.

(IAS) Logique et IA symbolique Depuis l'émergence de l'intelligence artificielle, les chercheurs travaillent autour de la notion de causalité là où les statistiques sont principalement centrées sur la notion beaucoup plus lâche de corrélation. La communauté scientifique travaillant dans le domaine de la logique formelle ou de la logique floue et de l'apprentissage de grammaire est au cœur de l'IA symbolique [10], [17], [24]. Alors que l'IA statistique est très majoritairement fondée sur la notion de corrélation, l'IA symbolique est plus ancrée dans la causalité : les décisions sont prises sur des critères logiques. Ces systèmes sont donc (plus) transparents et explicables, ils sont adaptés non seulement au diagnostic sur diverses problématiques mais aussi à la vérification de systèmes. Une des difficultés de développement de ces approches réside dans les pré-requis sur la *qualité des données* : il est très délicat de gérer des contradictions (e.g. des exemples erronés ou incohérents) dans les données à analyser. Une des forces de ces systèmes repose sur leur efficacité dans le déploiement : une fois appris, un ensemble de règles s'implémente très efficacement dans n'importe quel langage de programmation.

Mots-clés : Décision dans l'incertain ; Fonctions de croyance ; Logique floue ; Inférence logique ; Causalité

Modèles et outils : Frequent Item Set ; Arbre de décision ; Prolog ; Grammaire ; Automate fini ; Machine de Turing ; Système expert

(KB) Ingénierie des connaissances et ontologies L'ingénierie des connaissances regroupe la communauté scientifique travaillant sur la collecte mais surtout la valorisation des connaissances expertes à travers des formats de stockage des données

adaptés à des faits et des règles qui permettent ensuite de faire de l'inférence logique pour raisonner et généraliser. Une ontologie désigne l'ensemble des composants formant l'un de ces systèmes. Comme nous l'avons évoqué précédemment, cette communauté est liée à celle du texte pour l'extraction de connaissances et le peuplement des bases de données d'une part et pour l'alignement des ontologies qui nécessite de lier des catégories synonymes ou proches [4].

Cette communauté est intimement liée à celle travaillant sur les outils de logique, elle utilise très largement les mots-clés évoqués ci-dessus.

Mots-clés : Système expert ; Lexique ; Thésaurus ; Taxonomie ; Ontologie ; base de connaissances ; Logique de Description ; Web Sémantique

Modèles et outils : Prolog ; RDF ; OWL ; SparQL ; Protégé ; Dublin Core

(RI) Indexation et Recherche d'Information Le domaine de la recherche d'information est critique dans le monde moderne pour l'accès à l'information. L'enjeu est d'indexer des sources de données (souvent des documents) puis d'être en mesure d'ordonner ces sources par rapport à une requête de l'utilisateur. La masse des données à traiter est souvent colossale, tout comme les contraintes opérationnelles comme le temps de réponse [7].

Ce domaine ne se limite pas aux données textuelles : les SIG *Systèmes d'Information Géographiques* sont par exemple très importants dans le domaine de l'agriculture moderne et présentent des défis spécifiques à la fois en termes d'indexation et de réponse aux requêtes utilisateurs.

Mots-clés : Recherche d'information ; Indexation ; Base de données ; Moteur de recherche

Modèles et outils : BM25 ; Terrier ; Elastic Search

Apprentissage statistique Techniquement, le terme **Machine-Learning** (ML) se traduit par apprentissage automatique et englobe les IA symbolique et numérique [6]. Dans la pratique, le ML est quasiment synonyme d'apprentissage statistique. Ce domaine vise à construire automatiquement un algorithme de classification ou de régression (prédiction de valeurs continues) à partir de l'observation d'un ensemble de couples (entrée, sortie). Cette définition est néanmoins très large et va regrouper différentes tendances que nous allons détailler.

(ML1) La modélisation statistique consiste à choisir une loi de probabilité adaptée au type de problème visé puis à optimiser les paramètres de cette loi sur les données observées, ce qui pose plusieurs défis. Il faut de l'expertise pour choisir les bonnes lois de probabilité puis des compétences en optimisation pour identifier les paramètres optimaux. Enfin, il faut être en mesure de donner des garanties à l'utilisateur du système sur la qualité de l'estimateur, son pouvoir de généralisation, la convergence du processus d'optimisation etc...

Mots-clés : Approche bayésienne ; Intervalle de confiance
Modèles et outils : Maximum de vraisemblance ; MAP ; BFGS ; EM ; Stochastic Block Model

(ML2) D'autres chercheurs formulent l'apprentissage comme un problème de **minimisation d'un coût**, souvent composite, qui prend en compte plusieurs facteurs : bien prédire les étiquettes observées, éviter les solutions trop complexes qui risquent d'apprendre les observations par cœur, intégrer

le contexte... L'idée est d'optimiser des fonctions linéaires ou non-linéaires sur ces critères pour converger vers une solution exploitable. Cette catégorie est très proche de la précédente mais moins centrée sur les lois de probabilités.

Mots-clés : Régularisation ; Noyaux

Modèles et outils : SVM ; Generalized Linear Model ;

(ML3) La problématique peut aussi directement être abordée sous l'angle **algorithmique** avec des analyses sur les plus proches voisins ou la construction automatique d'arbre de décision. Une tendance forte du machine-learning consiste à considérer des **approches ensemblistes** mettant en parallèle plusieurs modèles. Cette catégorie est encore une fois très proche des deux précédentes mais elle est centrée sur des techniques qui proposent des performances de premier plan sur les données tabulaires, y compris sans réglage particulier. Ces outils sont utilisés par les chercheurs de l'IA mais aussi par de nombreux chercheurs de différents domaines qui peuvent se saisir facilement de ces outils. **Mots-clés :** Ensembles

Modèles et outils : Arbre de décision ; forêt aléatoire ; Bagging ; Boosting ; Plus-Proches-Voisins

(ML4) Une partie de la communauté se focalise spécifiquement sur des données non-supervisées où l'enjeu est de regrouper les instances qui se ressemblent. Le problème du **clustering** est classique (la supervision étant souvent chère) mais difficile : il faut souvent trouver une manière d'introduire des informations du domaine dans le problème pour réduire les degrés de liberté et avoir une chance de trouver des solutions pertinentes.

Mots-clés : Clustering

Modèles et outils : Frequent-Item-Set ; K-means ; DBScan ; EM

Deep-learning Le deep-learning est le nom donné aux réseaux de neurones depuis 2012, il s'agit techniquement d'un modèle particulier de machine-learning [18], [22]. Dans le détail, ce modèle permet de définir des distances particulièrement pertinentes entre des concepts (mots, phrase, phonème, image, ...) à partir de nouveaux paradigmes d'apprentissage (e.g. apprentissage de représentation, apprentissage contrastif, auto-supervision). La vague du deep learning s'est d'abord concentrée sur les données *sémantiques* (texte, parole & image) où les gains en performance ont été stupéfiants. Ces techniques, combinées à la maturité des outils de machine learning après 2015 ont ouvert la possibilité de traiter des données textuelles ou images à une communauté beaucoup plus large. Les (grands) modèles de langue ont amplifié ce mouvement, à la fois sur l'augmentation des performances et sur l'ouverture après 2020.

Le deep-learning propose aujourd'hui des solutions très performantes dans de nombreux domaines applicatifs (génomique, séries temporelles et spatio-temporelles, ...) mais la solution est loin d'être universelle. D'une part, les réseaux de neurones restent en retrait dans différents domaines, d'autre part le mot clé *réseau de neurones* cache en réalité une myriade d'architectures très différentes et toutes largement adaptables. La force de cette famille d'approches réside justement dans la plasticité de l'architecture qui permet l'introduction de connaissances expertes ou la prise en compte de différentes modalités de données ; le revers de la médaille est que le développement d'une architecture pour une application est bien plus coûteux qu'un modèle de machine-learning comme les forêts aléatoires par

exemple.

Des architectures spécifiques ont émergé ces dernières années qui ont un impact très fort sur l'ensemble des sciences des données : d'une part, les modèles génératifs qui permettent de construire des textes (**Language Model**), des images mais aussi des données de tous types avec des applications en data-augmentation ; d'autre part, un apprentissage en deux temps avec de l'auto-supervision, des modèles intermédiaires sauvegardés et disponibles qu'il est possible d'adapter à différentes tâches.

Nous proposons de distinguer différents cas d'usage spécifiques du deep-learning :

(DL1) Un modèle fortement non-linéaire pour aller au-delà des forêts aléatoires sur des jeux de données classiques. L'intérêt de l'architecture est alors souvent la capacité d'apprentissage de représentation (*embeddings*) sur les variables discrètes (e.g. utilisateurs d'un système de recommandation, régions géographiques).

Mots-clés : Apprentissage de représentation, embeddings

Modèles et outils : Perceptron, MLP (Multi-Layer-Perceptron)

(DL2) Jouer avec différents types d'architectures : à convolution, récurrent, sur des graphes (GNN), Transformer, multiplier les critères d'apprentissage. Nous regroupons dans cette catégorie les chercheurs qui travaillent sur les architectures et exploitent la flexibilité du cadre proposé par les librairies actuelles de réseaux de neurones comme pytorch ou tensorflow.

Mots-clés : convolution ; fonction de coût

Modèles et outils : CNN ; GNN ; RNN ; Transformer

(DL3) L'apparition des modèles pré-entraînés est une révolution dans le domaine de l'IA à plusieurs titres : il s'agit de limiter les coûts et de faciliter la reproductibilité en proposant des modèles déjà entraînés sur des données. Il s'agit aussi de pouvoir exploiter des architectures complexes sur des jeux de données très petits en exploitant le *fine-tuning* (BERT, YOLO, ...). Enfin, il s'agit aussi de récupérer des outils opérationnels directement applicables sur les données (e.g. VGG sur les classes d'Imagenet ou chatGPT pour les applications de NLP). Pour les modèles les plus versatiles, appris sur des masses de données gigantesques, on parle aujourd'hui de modèles de fondation.

Mots-clés : modèles de langue, de vision ; modèles de fondation ; fine-tuning ; few-shot

Modèles et outils : VGG ; ResNet ; YOLO ; U-Net ; GPT ; Llama

(DL4) Nous discutons du mot-clé *IA générative* plus loin dans le rapport. L'idée de cette section est de regrouper les chercheurs travaillant sur des approches génératives à la fois en lien avec les modèles génératifs de statistique ou l'augmentation de données. L'émergence des VAE (Variational Auto-Encoder) en 2014 a ouvert la voie à de nombreux travaux à l'interface avec les statistiques. En parallèle, une des solutions pour rendre les modèles plus robustes est de les apprendre sur plus de données. Comme cela revient très cher, une astuce consiste à générer des données proches de celles observées pour bénéficier de l'étiquetage tout en introduisant des variations crédibles à même de renforcer le modèle.

Mots-clés : Modèles à variables latentes ; augmentation de données

Modèles et outils : Auto-Encoder ; VAE ; GAN (Generative Adversarial Networks) ; Modèle de diffusion

(DL5) Le deep learning a un impact considérable sur la modélisation des systèmes dynamiques avec différentes architectures. Certains modèles sont appris sur des simulations à base d'équations différentielles pour économiser du temps de calcul, d'autres architectures miment le comportement d'une équation différentielle tout en permettant plus de flexibilité sur les conditions initiales. Enfin, de nombreuses approches hybrides combinent des équations différentielles et des approches neuronales pour l'explication des résidus.

Mots-clés : PINNs ; Hybridation

Modèles et outils : RNN ; NeuralODE

(CTR) Automatique & Contrôle Les algorithmes de contrôle des actionneurs en automatique reposent historiquement sur des classes de fonctions mathématiques qui sont franchement hors du périmètre de l'IA (PI -régulateur Proportionnel Intégral-, PID, ...) mais l'intégration de données d'observation permet d'ajuster ces stratégies pour tenir compte de situations plus complexes et contextualisées. Avec la prise en compte des données la problématique devient de la décision séquentielle. Cette classe de problèmes correspond à des algorithmes spécifiques et à la communauté des processus décisionnels de Markov et de l'apprentissage par renforcement décrite juste après.

Mots-clés : Automatique ; Contrôle flou

Modèles et outils : PI ; PID ; Contrôleur ; MDP ; RL

Apprentissage par renforcement Les algorithmes de machine-learning prennent en général des décisions indépendantes les unes des autres. Par exemple, le fait d'étiqueter un mail comme un spam ou pas est une décision indépendante de celle prise pour le mail précédent. A l'inverse, de nombreux processus nécessitent des décisions coordonnées séquentiellement ; par exemple, l'affinage d'un diagnostic médical avec des décisions à prendre concernant les examens à réaliser, les décisions permettant de contrôler un véhicule, un robot ou un avatar dans un monde virtuel. Ce processus de décision revient à établir une planification des tâches, une stratégie pour résoudre des jeux classiques (Echec, Go) mais aussi des problèmes bien plus concrets [15], [20].

Comme pour le deep-learning, nous considérons trois catégories distinctes dans la manière de travailler sur ces outils.

(RL1) La première catégorie regroupe la modélisation des problèmes séquentiels et les outils historiques à base d'états discrets. On parle souvent de processus de Markov. Le fait d'englober largement la modélisation à état discret rend cette catégorie assez proche de celle que nous avons appelée **(OPT2)**.

Mots-clés : Processus décisionnels de Markov (MDP) ; POMDP ; Planification ; modèle à états

Modèles et outils : Apprentissage par renforcement ; Q-Learning ; Bandits

(RL2) Les algorithmes de deep-learning permettent de passer à une modélisation continue des états et des actions [26]. De nombreuses propositions ont émergé ces dernières années dans ce domaine. Des algorithmes reposent aussi sur une interaction avec les utilisateurs, nous parlons alors de RLHF *Reinforcement Learning with Human Feedback*.

Mots-clés : Apprentissage de représentation ; RLHF

Modèles et outils : DQN (deep Q-network) ; PPO (Proximal Policy Optimization) ; AC (Actor Critic) ;

(RL3) La dernière catégorie regroupe les usages reposant sur des modèles pré-entraînés. Dans le domaine de l'apprentissage par renforcement, il s'agit pour l'instant essentiellement du modèle AlphaFold sur la prédiction de structure 3D des protéines.

Mots-clés : Modèles pré-entraînés

Modèles et outils : AlphaFold

(Seq) Approches séquentielles Les données séquentielles (données textuelles, séries temporelles, décisions séquentielles, etc.) ont toujours posé un défi spécifique à la communauté qui a développé une panoplie d'outils spécifiques au fil des années. Les premières propositions ont reposé sur des extractions & détectons de motifs caractéristiques, puis sur des automates de type chaînes de Markov cachées dans les années 90, puis des CRF *Conditional Random Fields* dans les années 2000. Les réseaux de neurones ont profondément changé ce domaine en proposant des modélisations à état continu (e.g. RNN -Recurrent Neural Network-), basées sur l'apprentissage de représentation.

Cette section se superpose largement avec d'autres sections **(RL1),(DL2)** mais correspond à une communauté qui développe et utilise ces approches.

Mots-clés : Séquence; Processus de Markov

Modèles et outils : AR; ARMA; Chaîne de Markov (cachée); CRF *Conditional Random Field*; RNN *Recurrent Neural Network*

Optimisation, satisfaction de contraintes et recherche opérationnelle L'optimisation est un sujet transverse et donc un axe assez particulier dans les thématiques listées dans cette section. En effet, la plupart des problématiques évoquées précédemment se traitent de la manière suivante : (1) analyse d'un problème concret; (2) formulation mathématique générale; (3) optimisation et identification d'une solution. Ainsi, la plupart des communautés ont recours à des algorithmes d'optimisation. Certains chercheurs font de l'optimisation leur domaine de recherche : selon leur positionnement, ils peuvent être plutôt à l'extérieur du périmètre de l'IA, dans les mathématiques appliquées (e.g. optimisation convexe ou non-convexe).

(OPT) Lorsque ces algorithmes d'optimisation continue sont adaptés pour les besoins d'un modèle statistique ou d'une architecture neuronale, les chercheurs travaillent alors très en amont mais bien en Intelligence Artificielle.

Mots-clés : Optimisation continue

Modèles et outils : BFGS; Descente de Gradient

(RO) Une classe particulière d'algorithmes d'optimisation est une partie intégrante de l'IA : les algorithmes de recherche opérationnelle ou d'optimisation combinatoire, d'optimisation sous-contraintes [5], [14]. Le domaine de la satisfaction de contraintes (SAT ou CSP en anglais *Constraint Satisfaction Problem*) vise à trouver une solution optimale à un problème combinatoire tout en respectant un certain nombre de règles et restrictions. Les exemples classiques de ce type de problème sont la gestion d'emplois du temps au niveau d'une structure (avec des contraintes personnelles, de salles, ...) ou tout simplement un Sudoku (où les contraintes sont les règles du jeu). Ce domaine peut être relié à la planification et l'apprentissage par renforcement : de nombreux problèmes combinatoires sont également des problèmes de décision séquentielle.

Mots-clés : Planification, SAT/CSP, combinatoire

Modèles et outils : Programmation dynamique; Simplex;

Cplex; Algorithmes évolutionnaires

(GR) Graphes Le graphe est un outil générique qui permet de représenter de nombreuses situations dans différents domaines applicatifs et d'introduire ensuite des algorithmes génériques pour optimiser des classes entières de problèmes. Cette communauté scientifique est souvent à cheval soit du côté de la recherche opérationnelle, soit du côté de la logique et de la modélisation des connaissances ou en encore du côté du clustering. Ainsi, cette case ne sera probablement jamais cochée seule. Nous avons cependant choisi de conserver une catégorie à part car cet outil de modélisation fédère une réelle communauté scientifique.

Le graphe est un outil de modélisation vraiment versatile, utile des systèmes de recommandation à l'optimisation de trajet en cartographie (plus court chemin, voyageur de commerce, ...) en passant par l'analyse des réseaux sociaux.

Mots-clés : Graphes; Réseaux; Clustering

Modèles et outils : Détection de communautés; Voyageur de commerce; Plus court chemin; GNN; SBM

Simulation La simulation consiste à construire un modèle pour une situation ou un système donné. Nous distinguons ici deux types de simulation : la modélisation mécaniste d'une part et la simulation des systèmes complexes d'autre part.

(SIM1) La simulation numérique est historiquement en dehors du périmètre de l'IA. Il s'agit de construire un modèle mécaniste d'un système à base d'équations pour interpoler et prédire différentes situations. Les systèmes dynamiques sont modélisés à l'aide d'équations différentielles. La plupart des systèmes modernes étant également décrits par des données, nous assistons aujourd'hui à une multiplication des approches hybrides qui nous poussent à inclure la simulation dans notre périmètre.

Mots-clés : Simulation; Modélisation Mécaniste; PINNs

Modèles et outils : Eléments finis; EDP (équations aux dérivées partielles); Systèmes d'équations

(SIM2) Les systèmes multi-agents (SMA), ou plus largement les approches systèmes, sont ancrés dans la logique de simulation, mais avec des outils plus informatiques que mathématiques. Un système complexe est composé d'une multitude d'objets en interactions avec beaucoup de comportements stochastiques : la prédiction de l'état du système repose sur une modélisation fine des agents, du contexte et des interactions sur la base des connaissances des experts du domaine [11]. Ce domaine est évidemment en interaction avec les systèmes de décision séquentielle pour rendre les agents cohérents dans le temps.

Mots-clés : Simulation; Agents; Interactions; modélisation markovienne

Modèles et outils : Systèmes multi-Agents; NetLogo; GAMA; DEVS

3.3. De la théorie à la pratique

En considérant les différents types de données et les problématiques associées comme un périmètre, nous proposons d'explorer ici différentes *profondeurs*. Entre la recherche amont sur les théories sous-jacentes de l'intelligence artificielle, les modèles probabilistes ou l'optimisation combinatoire et l'utilisation des outils aujourd'hui accessibles à tous, nous avons essayé d'identifier des profils significatifs.

[illegible]

FIGURE 1. Référentiel à deux axes : modalités des données en ligne et techniques d'analyse en colonne. Dans l'enquête associée à ce rapport nous souhaitons projeter les effectifs des équipes dans ce tableau. Le troisième axe, *de la théorie à la pratique*, prend la forme d'une division en 3 de chaque case : U pour désigner les *Utilisateurs* ; D pour désigner les *Développeurs* et T pour désigner les *Théoriciens* qui travaillent sur ces techniques. Enfin, la dernière case rouge doit nous permettre de mesurer la dynamique : elle permet de renseigner si la case est une prospective de l'équipe (P) ou plutôt en régression dans l'équipe (R).

Il est important de comprendre la difficulté de cet exercice à l'heure où les profils de chercheurs se sont largement diversifiés en IA et où nous observons un continuum de positionnements scientifiques entre la théorie et la pratique, souvent à cheval entre différents domaines. La frontière entre les statistiques et l'optimisation est totalement perméable.

Recherche amont L'IA, dans le périmètre très large décrit précédemment, repose sur différents domaines scientifiques, en particulier en mathématiques : les probabilités, l'optimisation ou le calcul différentiel. Sur le plan informatique, la programmation objet a profondément fait évoluer la manière d'implémenter la simulation et le développement des bibliothèques d'outils tandis que le calcul massivement parallèle est un sujet à part entière, à cheval entre l'algorithmique et le développement de matériel. Citons enfin le domaine de la visualisation d'information (infoviz) et les bases de données.

Ces outils sont importants voire critiques pour l'IA mais ces domaines ne rentrent pas à proprement parler dans l'IA. Une communauté existe tout de même dans l'IA pour l'adaptation de ces outils aux problématiques spécifiques de l'IA.

En **statistique**, ces chercheurs à l'interface de l'IA restent généralement centrés sur l'étude des propriétés théoriques des lois de probabilité mais ils dérivent ces résultats pour des applications concrètes où les propriétés deviennent des intervalles de confiances directement exploitables par des experts métier.

En **machine-learning** ou **deep-learning** cette interface concerne les algorithmes d'optimisation, les architectures de calcul distribué (en particulier le calcul sur GPU) et l'architecture logicielle (pour l'industrialisation, le déploiement en embarqué -téléphone ou autre-, les plateformes de **deep-learning** comme PyTorch ou TensorFlow).

Dans le domaine de l'**ingénierie des connaissances**, l'interface se situe au niveau des systèmes de bases de données qui supportent les bases de connaissances. Du côté du web sémantique, il s'agit de faire le lien entre les protocoles réseaux de haut-niveau et le raisonnement sur des formats structurés de connaissances.

Au cœur de l'IA Un titre si large va forcément conduire non pas à un niveau de théoricit  ciblé mais plutôt à un intervalle sur le spectre. Il est cependant difficile de proposer des catégories plus restreintes tant les chercheurs ont tendance à travailler à différents niveaux selon leurs projets.

Les **statisticiens** mêlent une connaissance des probabilités, une capacité à optimiser les paramètres des approches et une affinit  pour le domaine applicatif vis . Ce type de recherche permet à la fois d'introduire des a priori sur les problèmes traités pour guider la solution et d'obtenir des garanties fortes sur modèles développés. Ce type de recherche est très bien adapté aux petits échantillons de données. Ce niveau fait directement  cho aux techniques de mod lisation statistique d crites dans (ML1).

Les **informaticiens du machine-learning** inversent souvent le processus : d'abord comprendre en profondeur une probl matique m tier pour ensuite chercher des solutions dans une panoplie d'outils bien plus large que les lois de probabilit s. L'enjeu reste d'obtenir des garanties sur la solution : ces garanties peuvent  tre th oriques, comme en A1,   travers des d monstrations sur les formulations. Ces garanties peuvent aussi  tre empiriques   travers des campagnes d'exp riences

et d'ablation montrant l'int r t de l'approche sur une classe de probl mes illustr e par plusieurs bases de donn es.

En **deep-learning**, une fois l'architecture modulaire fiabilis e et op rationnelle, l'enjeu est de d velopper des modules ayant de bonnes propri t s pour l'optimisation (e.g. batch norm, softmax, activation RELU, ...) et/ou pour des applications g n riques. La combinaison des modules, la d finition de co ts sp cifiques, l'introduction de repr sentations interm diaires et de contraintes sur les donn es font le quotidien du chercheur sur les r seaux de neurones.

Dans le domaine de l'**IA Symbolique**, les chercheurs en logique, grammaire ou graphes travaillent sur des descriptions formelles et des outils de raisonnement parfois  loign s des probl matiques concr tes. Mais la plupart des chercheurs de cette communaut  travaillent aussi sur l'application des th ories pr c dentes   des cas d'usage r els.

En **recherche op rationnelle**, comme dans l'exemple pr c dent, il est possible de travailler sur des formulations permettant de transformer un probl me concret en une instance traitable par un solveur g n rique (type CPLEX) ou sur des algorithmes de programmation dynamique ou des algorithmes g n tiques d di s   une t che tout en restant au centre de l'IA.

En **ing nierie des connaissances**, les chercheurs travaillent sur les formats de stockage, d'interrogation et d' change des connaissances. L'extraction d'information, pour la construction des bases de connaissances, et la fusion de diff rents formats (au niveau s mantique) sont autant de ponts avec les autres communaut s de l'IA.

Utilisateurs des outils d'IA Les outils de l'IA en g n ral et du machine-learning en particulier se sont d velopp s mais aussi fortement consolid s au cours de la derni re d cennie : les biblioth ques disponibles, au premier rang desquelles **scikit-learn**, permettent aujourd'hui d'effectuer des op rations tr s avanc es avec des comp tences minimales en programmation. En parall le, la plupart des  tudiants, en particulier issus de classes pr paratoires, arrivent aujourd'hui avec un bagage en programmation suffisant. Il est donc aujourd'hui possible pour presque n'importe quel ing nieur (a fortiori chercheur) d'effectuer du travail en intelligence artificielle. Nous parlerons d'un acc s *utilisateur*.

En **machine-learning**, les ing nieurs peuvent se saisir des outils et travailler sur des donn es r elles tr s rapidement. M me l'optimisation des hyper-param tres devient accessible avec des outils comme **optuna**.

L'exploitation des **r seaux de neurones** est nettement plus technique que celles des algorithmes de machine learning mais n anmoins accessible. En particulier, certains mod les de **deep-learning**, en texte ou en vision, sont pr -entra n s et disponibles en librairie (github ou huggingface par exemple) : ils sont directement op rationnels sur certains probl mes.

Dans le cadre des **syst mes multi-agents**, des plateformes comme **GAMA** permettent   quasiment n'importe quel informaticien de d velopper des syst mes complexes.

En **ing nierie des connaissances**, des interfaces d' dition des ontologies comme **prot g ** permettent aussi d'exploiter le travail de la communaut  scientifique plus efficacement.

Utilisateurs finaux La plupart des syst mes d'IA sont d velopp s pour r pondre aux besoins des utilisateurs finaux. Ce niveau d'usage correspond typiquement   un utilisateur

lambda de *chatGPT* ; un lecteur exploitant un moteur de recherche bibliothécaire basé sur le Dublin Core ; un médecin utilisant des outils d'aide au diagnostic sur de l'imagerie médicale. Ces utilisateurs sont directement au contact de l'IA, ils doivent être formés, même de manière minimaliste, à ces systèmes pour connaître leurs forces et faiblesses et, ainsi, pouvoir en tirer le meilleur.

3.4. Modalités multiples, problématique composite et mots clés émergents

IA générative L'IA générative est à la fois ancienne et très récente, créant une ambiguïté sur le terme. Historiquement, les approches de modélisation statistique tentent de construire un modèle capable de générer les données observées : maximiser la vraisemblance d'une observation consiste à trouver le paramétrage du modèle qui va maximiser la probabilité que ce modèle puisse générer cette observation. On parle alors de modèle génératif. De la même manière, n'importe quel modèle de prédiction de série temporelle va *générer* la suite d'un processus, mais on parlera plus de processus auto-régressif dans ce cas de figure.

A l'aube des années 2010, les premières architectures profondes de réseaux de neurones souffrent de problèmes d'optimisation chronique et la solution de l'époque a consisté à apprendre successivement des couches capables de régénérer (parfois en enlevant du bruit) les données d'entrée : on parle alors de *Stacked Denoising Auto-Encoder*.

L'émergence moderne du terme *IA générative* a lieu en 2014 avec l'introduction de l'architecture GAN (Generative Adversarial Network) qui permet de générer algorithmiquement des images d'une crédibilité jamais vue auparavant. Le terme *IA générative* va être repris pour décrire les architectures de génération d'image (e.g. modèles de diffusion), de texte (e.g. modèles de langue) ou de la parole : des modalités sémantiques qui ont du sens pour les utilisateurs du système.

En parallèle, une classe de modèles spécifiques, les auto-encodeurs variationnels, mêle modélisation statistique et architecture neuronale générative. Ces architectures sont développées dans différentes communautés avec des visées sur de nombreux domaines applicatifs : elles forment l'une des tendances récentes de l'IA générative.

Aujourd'hui, l'IA générative a acquis une très grande notoriété dans le sillage de *chatGPT* et beaucoup d'architectures reprennent a posteriori cette étiquette : les modèles de prédiction de série temporelle, les auto-encodeurs et bien d'autres approches rentrent dans le giron de l'IA générative, parfois de manière raisonnable et parfois en jouant vraiment sur les mots. Les différents courants de l'IA générative s'ignorent parfois ou tentent d'accaparer le terme entretenant une certaine confusion sur le périmètre réel du mot-clé.

Pour que la question *Etes-vous impacté par l'IA générative ?* soit pertinente, il faut donc d'abord bien expliquer les contours de la définition retenue.

Modèles de fondation Une des grandes tendances récentes (après 2020) consiste à construire des modèles à cheval entre les modalités de données pour améliorer à la fois la compréhension des données (sémantique, désambiguïsation) et les interfaces utilisateur autour de ces données (dialogue en langage naturel) tout en ouvrant de nouvelles perspectives applicatives comme la réponse aux questions (*question-answering*) ou la

génération d'images. La fusion entre le texte et l'image a donné naissance aux modèles de fondation qui tentent maintenant d'intégrer des données, les signaux audio, la vidéo voire les séries temporelles.

Le paradigme de l'apprentissage de représentation (projection des données dans un espace vectoriel latent) simplifie beaucoup la conception des systèmes multimodaux : s'il est très difficile d'aligner du texte et de l'image dans l'espace d'origine, il est aisé de définir des opérations mathématiques combinant des vecteurs représentant ces objets hétérogènes.

Approches hybrides & systèmes dynamiques Une tendance très forte ces dernières années (à partir 2020) consiste à construire des approches hybrides, typiquement des réseaux de neurones dont l'architecture s'apparente à la décomposition d'une équation différentielle. L'idée est de combiner les avantages des deux approches : exploiter plus de données, ne pas régler explicitement des conditions aux limites inconnues et obtenir des prédictions fiables.

Ces modèles sont dits hybrides car ils créent un pont entre la modélisation mécaniste d'une part, à base d'équations physiques et en particulier des équations différentielles pour les systèmes dynamiques, et l'analyse de données d'autre part. Les apports de l'hybridation peuvent être très différents d'une application à l'autre : parfois, il s'agit de gagner en temps de calcul avec les approches orientées données au détriment de la précision, parfois il s'agit d'une nouvelle manière d'aborder le problèmes des conditions limites inconnues et des problèmes mal posés.

On parlera dans ce domaine de PINNs (Physical Informed Neural Networks) pour les architectures neuronales apprises sur des modèles d'EDP ou de NeuralODE (Neural Ordinary Differential Equation) pour les modèles mimant les EDP et plus ancrés dans les données réelles.

Jumeaux numériques Les jumeaux numériques visent à modéliser des systèmes complexes pour mieux les comprendre et à optimiser des processus en interaction avec des données de terrain et des actionneurs, en tenant compte de contextes riches : il s'agit donc d'un mot-clé applicable à de nombreux domaines et de nombreux projets (la notion de complexité n'étant pas très précise). La multiplication des projets de jumeaux ces dernières années a poussé la communauté à distinguer trois types d'approches :

1. **Modèle numérique/Digital modeling** : il s'agit de modéliser, hors ligne, un système complexe en utilisant diverses données enregistrées sur le système et son contexte.
2. **Ombre numérique/Digital shadow** : il s'agit de brancher un modèle numérique sur des données captées en temps-réel. L'enjeu est de fournir des propositions d'action ou de lever des alarmes.
3. **Jumeau numérique/Digital twin** : dans son acception récente, cette dénomination est réservée aux modèles numériques qui sont à la fois branchés en temps-réel sur les données du système mais également raccordés aux actionneurs pour effectivement avoir un impact sur le système. Ce type de projet contient donc un volet de commande des systèmes.

Le dénominateur commun aux jumeaux numériques est en général la multiplicité (éventuellement l'hétérogénéité) des sources de données et des prédicteurs attendus (qui tra-

duisent l'observation d'un système et d'un contexte complexes). Par exemple, un jumeau numérique d'exploitation agricole va prendre en compte des informations sur la nature des sols, la météo, la configuration de l'exploitation, les techniques utilisées sur le site mais aussi potentiellement des capteurs présents à des endroits clés voire des modèles d'interactions avec la faune et la flore environnantes. Les sorties seraient également multiples au niveau de rendements, de consommation d'énergies ou d'intrants...

A leur apparition, les jumeaux numériques étaient principalement basés sur des approches mécanistes en exploitant les connaissances d'experts du métier. Récemment, l'expansion de l'analyse de données, la complémentarité des approches voire leur hybridation a fait évoluer de nombreux projets de jumeaux numériques vers le machine learning, en particulier avec les modèles hybrides mentionnés ci-dessus.

Robotique La robotique implique à la fois une dimension matérielle très importante, un certain nombre d'algorithmes de contrôle pour respecter des consignes elles-mêmes souvent issues d'algorithmes de décision séquentielle. Nous sommes donc face à un problème composite particulièrement ambitieux et les roboticiens ont coutume de considérer que les algorithmes de décision qui fonctionnent dans un environnement virtuel échouent généralement lors du passage dans le réel à cause des interactions avec des actionneurs imparfaits ou des contextes complexes.

En plus des algorithmes de décision séquentielle (apprentissage par renforcement pour déterminer la séquence d'actions permettant de réaliser une tâche), les robots modernes reposent énormément sur les algorithmes de vision (et d'analyse radar/lidar) pour la perception de leur environnement. Ainsi la robotique concentre un ensemble de problématiques à la fois à l'intérieur et à l'extérieur du domaine de l'intelligence artificielle... mais avec de plus en plus de connexions au fur et à mesure que l'IA gagne en maturité et en robustesse pour être exploitée sur le terrain.

Le courant de l'IA en robotique bénéficie aussi de l'évolution des algorithmes de contrôle des actionneurs vers l'analyse de données. De nouveau, la maturité des approches joue un rôle clé : on est passé en quelques années de théories inapplicables sur le terrain combinées à des actionneurs analogiques assez sensibles à des stratégies où les données peuvent –dans certains cas– apporter robustesse et flexibilité qui se combinent parfaitement avec des actionneurs numériques bien plus faciles à contrôler.

3.5. Outils de développement et moyens de calcul

Les outils de développement et les moyens de calcul sont des marqueurs forts des différentes communautés scientifiques de l'IA :

- R chez les statisticiens,
- Python chez les informaticiens,
- parfois du C ou du JAVA pour les modélisateurs,
- des besoins en CPU multiples pour la simulation, en particulier dans les modélisations par éléments finis (HPC *High Performance Computing*).
- des besoins en GPU multiples pour les réseaux de neurones.

Développement Comme nous l'avons mentionné à plusieurs reprises, l'évolution des plateformes logicielles est l'un des facteurs clés de l'expansion de l'IA moderne. Les outils du machine-learning ont progressé en 10 ans de manière remarquable : les outils sont plus efficaces et surtout plus robustes et plus accessibles, n'importe quel ingénieur peut les prendre en main. Jouer avec le GPU d'un ordinateur est à la portée de tous.

Les outils ont tous progressé mais l'environnement Python est celui qui a le plus bénéficié de l'attention des GAFAM : la plupart des bibliothèques ont une interface dans ce langage dans le domaine du machine-learning mais aussi au-delà (simulation, multi-agents, gestion des connaissances,...).

Il y a donc une tendance à aller vers le Python pour différentes communautés, tendance renforcée par la formation des lycéens et étudiants qui penche fortement vers ce langage.

Calcul Si tout le monde bénéficie des avancées dans le domaine logiciel, les besoins en calcul dépendent fortement des domaines de recherche.

Les techniques de machine-learning sont plus accessibles que jamais au niveau du développement mais aussi du calcul : les approches ensemblistes sur les arbres sont implémentées de manière efficace mais sont aussi très peu sensibles au paramétrage. Finalement, un simple PC permet d'aller très loin dans les expériences et d'exploiter l'état de l'art du domaine.

Les réseaux de neurones appartiennent au domaine du machine-learning mais la complexité de ces architectures se traduit par des besoins spécifiques : il faut au moins un voire plusieurs GPUs pour traiter du texte ou de l'image avec du deep-learning. Le matériel est cher, gourmand en énergie et demande quelques compétences spécifiques pour l'administration : il s'agit d'un verrou important. Cette barrière est toutefois aujourd'hui moins contraignante : les machines se sont démocratisées, les pilotes sont plus matures et, surtout, des centres de calcul se sont construits aux niveaux local et national et sont accessibles à la communauté.

Les besoins de calcul en simulation (notamment sur les systèmes dynamiques et le spatio-temporel) sont très importants. Il s'agit de paralléliser des modélisations en éléments finis qui tournent historiquement sur des clusters de CPU mais se déploient de plus en plus sur les GPU pour bénéficier de l'explosion de puissance des architectures récentes développées dans le sillage de NVidia.

3.6. Groupes bibliographiques

Les communautés bibliographiques sont à la fois marquées, perméables et mouvantes. Les communautés sont structurées autour des approches, autour des types de données et autour des domaines applicatifs ; la communauté IA ayant beaucoup grossi ces dernières décennies, des communautés structurées existent à l'intersection entre certains types de données et certains domaines applicatifs (e.g. les données textuelles dans le domaine médical).

Au-delà des sources bibliographiques, chaque communauté possède évidemment ses codes, son organisation et un vocabulaire spécifique. La plupart des communautés scientifiques placent les journaux bien au-dessus des conférences mais l'informatique fait exception et dans le cas du machine-learning en particulier, les publications les plus sélectives sont des conférences.

On peut distinguer de très nombreuses communautés bibliographiques dans l'IA :

- les statistiques : *Journal of the American Statistical Association (JASA)*, *Annals of Statistics*, *Conf. on Learning Theory*, *AISTATS*, ...
- l'IA Symbolique : *AAAI*, *IJCAI*, *International Joint Conference on AI*, *ECAI*, *FUZZ-IEEE (IEEE Int. Conf. on Fuzzy Systems)*
- La gestion des connaissances : *ISWC (Int. Semantic Web Conf.)*, *KDD (Knowledge Discovery and Data Mining)*
- le machine-learning/deep learning : *NeurIPS*, *Neural Information Processing System*, *ICML*, *Int. Conf. on Machine Learning*, *ECML*, *ICLR (Int. Conf. on Learning Representation)*
- les simulations multi-agents : *AAMAS (Int. Conf. on Autonomous Agents and Multi-Agent Systems)*
- l'optimisation discrète : *CP (Int. Conf. on Principles and Practice of Constraint Programming)*, *SODA (ACM-SIAM Symposium on Discrete Algorithms)*
- Les algorithmes évolutionnaires : *GECCO (Genetic and Evolutionary Computation Conference)*
- diverses communautés centrées sur les modalités de données : le texte (conférences ACL : ACL, EACL, NAACL, EMNLP), l'image (ECCV, ICCV, CVPR), les données biologiques (*Biometrika*), les séries temporelles (ICASP), la robotique (IROS)...

Les sources bibliographiques sont données à titre d'exemples, elles indiquent bien l'existence d'une communauté thématique mais ne sont en aucun cas exhaustives.

INRAE développe des approches théoriques en tant qu'institut de recherche mais dispose aussi de domaines applicatifs privilégiés autour de l'agronomie et des sciences de la vie en général. La plupart des équipes émergent ainsi dans plusieurs communautés bibliographiques.

Nous verrons que l'ensemble des communautés scientifiques sont impactées par la vague actuelle en IA : nouveau matériel, nouveaux modèles, nouvelles performances de référence, parfois nouveaux processus pour aborder une problématique : toutes les communautés sont touchées. Avec toujours les mêmes questions : qu'est-ce qui relève du phénomène de mode et qu'est-ce qui relève du changement profond ? Que vaut mon expertise historique dans le domaine ? Quelles sont les nouvelles clés de publication dans les supports de référence et combien cela coûte-t-il ?

Une autre évolution notable est l'ouverture des communautés, directement liée à l'ouverture des outils de l'IA : les codes sont plus accessibles, les modalités de données deviennent plus accessibles, les modèles complexes deviennent envisageables. Les chercheurs d'un domaine peuvent élargir leur horizon de recherche à moindre frais. Les différents niveaux de théoricités deviennent un continuum : les théoriciens peuvent aller jusqu'à l'expérimentation et, réciproquement, le temps gagné dans les expérimentations permet aux praticiens de se lancer dans des justifications plus théoriques.

4. Positionnement des unités MathNum

Nous proposons ici une description rapide du positionnement des unités MathNum dans le référentiel proposé. Les travaux de ces unités sont pour la plupart conduits en collabo-

ration avec des équipes d'autres départements d'INRAE, mais ces liens souvent étroits ne sont pas décrits ci-dessous.

4.1. BIOSP

Contexte : L'unité BioSP regroupe deux équipes : l'équipe Recherche, composée de 14 chercheurs permanents, 7 doctorants et 1 post-doctorant, ainsi que l'équipe OPE (Opérationnelle à INRAE pour la Plateforme d'Épidémiologie-surveillance en santé végétale), constituée de 4 ingénieurs. Ses travaux se concentrent sur les domaines de la statistique, des systèmes dynamiques et de l'écologie-épidémiologie.

Données : Les données manipulées sont principalement de nature climatique et environnementale, constituées majoritairement de séries temporelles et spatio-temporelles (mesurées en différents points). Il y a une tendance croissante à traiter des ensembles de données volumineux, notamment plusieurs téra-octets de données climatiques. Les données de ré-analyse (*regrid*) sont particulièrement précises et denses. Une grande partie des données concerne des informations spatiales et l'occupation des sols, avec quelques images et une abondance de données textuelles. Sur la plateforme d'épidémiologie-surveillance, les données textuelles proviennent de sources médiatiques et scientifiques (presse, publications), en complément d'une veille scientifique et de la fouille d'URLs.

Méthodes : Les méthodes employées incluent la simulation stochastique pour l'analyse des séries temporelles et spatio-temporelles, ainsi que la modélisation par équations aux dérivées partielles (EDP). L'équipe utilise des modèles prédictifs tels que les forêts aléatoires, XGBoost et les réseaux de neurones. Plus récemment, elle s'est tournée vers les GAN et les modèles de diffusion. Il y a une volonté d'intégrer des approches hybrides combinant statistiques spatiales et apprentissage automatique.

L'équipe travaille également sur des jeux de données de petite taille et s'intéresse aux événements rares et extrêmes, en s'appuyant sur la théorie des valeurs extrêmes. Elle explore de plus en plus le deep learning et l'IA générative pour la génération de champs de précipitations (modèles de diffusion) et la modélisation d'événements extrêmes. Enfin, elle s'intéresse à la fouille et au résumé de texte, utilisant pour cela des modèles de langue.

Calcul : L'infrastructure de calcul comprend un cluster doté de 600 cœurs dédié à la simulation, ainsi qu'un GPU pour expérimenter les approches de deep-learning. L'équipe bénéficie également de l'accès à CollabIA. Pour l'instant, cette configuration est suffisante et il n'est pas nécessaire d'accéder à Jean Zay.

Bibliographie : Conférences et publications : l'équipe participe aux conférences des domaines des statistiques, de la géostatistique et de la géophysique. Ces événements ont progressivement intégré des thématiques liées à l'IA et au machine-learning, avec des sessions entières dédiées et de nombreuses expériences présentées sur ces sujets.

Perspective : Bien que l'IA générative soit probablement sur-estimée dans certains domaines, elle ne représente pas une simple tendance passagère : il est essentiel de l'adopter rapidement. Il s'agit de préserver les compétences traditionnelles tout en intégrant les nouveaux outils. Cette transition est rapide et

exigeante, nécessitant des recrutements (deux sont en cours), mais la concurrence est intense.

L'équipe utilise les formations FIDLE du CNRS, avec pour objectif que chacun puisse intégrer les outils en fonction de ses besoins : certains peuvent être utilisateurs, d'autres développeurs ou concepteurs, selon leur parcours et les opportunités liées à leur spécialité.

4.2. LISC

Contexte : Le laboratoire, composé de 17 membres dont 8 permanents, est organisé autour de trois grands axes de recherche. Le premier axe porte sur la modélisation des systèmes complexes, avec un intérêt particulier pour les modèles à base d'agents pour étudier l'émergence de comportements collectifs qui ne peuvent être prédits à partir des seuls comportements individuels. Le second axe est consacré aux systèmes dynamiques, qu'ils soient contrôlés ou non. Enfin, un axe transversal aborde les aspects méthodologiques, notamment le fléau de la dimensionalité, l'approximation des fonctions de réponse et l'analyse des données.

Données : Les systèmes complexes reposent principalement sur la modélisation à base d'agents, avec peu de données disponibles. En revanche pour les systèmes dynamiques, des données spatiales, récemment issues d'images satellites, sont utilisées.

Méthodes : La modélisation des systèmes complexes repose sur des agents (systèmes multi-agents). Elle inclut l'approximation de surfaces de réponse et l'utilisation, ainsi que l'adaptation, des machines à vecteurs de support (SVM). Des architectures neuronales sont également employées pour le clustering, afin d'étendre la capacité des SVM à se positionner par rapport à la frontière d'un sous-espace. Les algorithmes génétiques sont utilisés pour la calibration des modèles, vue comme un problème d'optimisation, en complément des méthodes de type ABC (Approximate Bayesian Computation). Les équations aux dérivées partielles (EDP) sont utilisées pour modéliser la dynamique. Bien que les modèles hybrides n'aient pas encore été développés, un intérêt croissant se manifeste pour les PINNs (Physically Informed Neural Networks).

Calcul : La majorité des projets sont développés directement en Java, Smalltalk et surtout Python, après avoir utilisé des langages comme Matlab. Les systèmes multi-agents (SMA) s'appuient également sur des plateformes telles que Capsis, Maelia et Gama. Le GPU, notamment via Pycuda, est exploité pour les modèles d'agents. Actuellement, les calculs sont réalisés localement, sans avoir encore recours au calculateur Jean Zay (Genci).

Bibliographie : Les conférences sur les systèmes complexes, telles que l'ICCS et l'ECCS, ainsi que l'atelier "Contagion on Complex Social Systems" (CCSS) sont des références majeures. Concernant les conférences sur les systèmes multi-agents, l'intelligence artificielle y est davantage présente aujourd'hui, bien que sans révolution majeure. En revanche, dans les conférences dédiées à l'IA et à la modélisation écologique, à la physique, l'IA axée sur les données occupe une place de plus en plus importante.

Perspective : En tant que petite unité, le laboratoire a exprimé le besoin de se développer, particulièrement après les

départs consécutifs à la fusion Inra-Irstea et la perte de certaines compétences, notamment en théorie de la viabilité.

Des perspectives émergent pour renforcer le développement des agents grâce à l'apprentissage par renforcement, ainsi que pour explorer l'hybridation dans la simulation de processus dynamiques à l'aide des PINNs (Physics Informed Neural Networks).

4.3. ITAP

Contexte : L'unité est sous les tutelles d'INRAE et de l'Institut Agro Montpellier, elle est centrée sur la métrologie environnementale, couvrant toutes les étapes de la chaîne, de la mesure des signaux à la prise de décision, et cela à différentes échelles. Le projet implique quatre équipes (Évaluation environnementale et sociale (ELSA); Modélisation et décision agro-environnementale (DeMo); Capteurs optiques pour les milieux complexes (COMIC); Procédés Environnement Pesticides Santé (PEPS)) sous la tutelle de deux départements de recherche INRAE : AgroEcoSystem et MathNum. Environ 80 personnes y participent, dont 30 titulaires, avec une proportion importante de contrats à durée déterminée. Bien que la viticulture soit une spécialité, les thématiques abordées sont variées et couvrent un large éventail.

Données : Les données traitées incluent des données spatiales, issues des systèmes d'information géographique (SIG), des images et des séries temporelles. Cela comprend des images hyperspectrales (3D), ainsi que quelques images RGB. Les données proviennent soit de capteurs, soit d'extractions de bases de données, avec une prédominance de données quantitatives. Il y a également une petite part de données d'experts.

Méthodes : Des évolutions récentes ont marqué l'unité : la composante logique floue s'est éteinte suite au départ en retraite du scientifique concerné par cette thématique, il n'y a plus non plus de composante gestion des connaissances. En matière de données spectrales, des méthodes de machine-learning sont employées, mais sans recours au deep learning pour l'instant. Parmi les techniques utilisées, on trouve l'apprentissage par transfert, la gestion de données hétérogènes, des approches semi-supervisées, ainsi que l'apprentissage de descripteurs, d'encodeurs et l'analyse en composantes principales (ACP). L'analyse multivariée, notamment en chimio-métrie, est aussi présente. L'analyse du cycle de vie (ACV) et l'utilisation de nombreux modèles mécanistes sont courantes. Le deep-learning est tout juste en phase de démarrage. Bien que l'équipe soit principalement composée d'utilisateurs du deep-learning, quatre chercheurs possèdent déjà une certaine expertise de ces approches.

Calcul : Les scientifiques d'ITAP font des développements en Python, Matlab, un peu de R. Une partie de l'unité ne code pas. Ils se positionnent en utilisateurs de méthodes statistiques. L'unité déplore un manque de connaissance sur les ressources de calcul, avec lesquelles elle n'est pas familiarisée.

Bibliographie : Les revues des domaines d'application intègrent de plus en plus fréquemment des travaux méthodologiques en machine learning, notamment dans les secteurs de l'optique et de la chimie. Parmi les conférences importantes, on peut citer la SETAC (Annual Meeting of the Society of Environmental Toxicology and Chemistry). Les revues et conférences

spécialisées dans l'agriculture numérique, comme la European Conference on Precision Agriculture, sont également des références pour l'unité.

Perspective : Développer des outils pour les données spectrales en s'appuyant sur des techniques d'augmentation de données. Repenser les capteurs afin qu'ils soient mieux adaptés à l'intelligence artificielle. Déjà utilisateurs de l'IA, l'objectif des scientifiques est de se former davantage pour améliorer et optimiser son usage. Une réflexion est en cours sur la transition des modèles basés sur les connaissances vers ceux basés sur les données.

L'utilisation fréquente de ChatGPT pour le codage, la rédaction et la recherche d'informations est en pleine expansion. La création d'un LabCom est prévue pour développer des techniques d'IA appliquées au traitement des données spectrales. L'IA hybride est perçue comme une transition prometteuse, mais des questions subsistent concernant les coûts et la manière de mener cette transition.

4.4. TSCF

Contexte : L'unité est composée d'une soixantaine d'agents (36 permanents) sur 3 sites, d'une ferme expérimentale avec un vivarium de robots. Cinq thèmes de recherche structurent le projet scientifique : Analyse et modélisation de l'environnement, Autonomie et autonomisation des équipements, IoT et observation sobre, Nouvelles technologies : intégration et appropriation, outils et pratiques : agroécologie innovante.

Données : Utilisation des données issues de capteurs pour la prise de décision et le suivi des parcelles agricoles. Traitement d'images pour la détection et l'identification des maladies des plantes. Exploitation des données de capteurs afin d'optimiser et d'apprendre les lois de contrôle des robots. Analyse et modélisation de l'environnement à partir de données multi-capteurs (laser, caméra, radar, etc.).

Méthodes : Algorithmes de traitement d'image, utilisation de bibliothèques de deep learning (DL), avec un peu de reinforcement learning (RL) appliqué à des données de trajectoires de robots en simulation. Réseaux de neurones pour l'optimisation des paramètres des lois de contrôle et pour la reconnaissance multi-capteurs.

Initiation à la planification de tâches par IA. Perte de compétences en raisonnement automatique (Description Logic et ontologies) avec le départ de Catherine Roussey, à MISTEA.

Pas d'usage des modèles de langues. Méthodes d'optimisation heuristique appliquées aux réseaux de capteurs. Utilisation des outils IA sans développement approfondi, combinée avec des approches plus classiques.

L'unité se décrit comme principalement utilisatrice d'IA, avec peu de compétences internes en la matière. Cependant, 3 à 4 personnes sont directement impliquées dans l'utilisation de bibliothèques IA et une dizaine de personnes les exploitent.

Calcul : Centre de calcul local, pas encore besoin des ressources nationales (Jean Zay), collaboration avec le CEA (FactorAI). Expertise d'exploitation parallèle d'une trentaine de GPU via une doctorante.

Bibliographie : En robotique, la vague de l'IA est omniprésente : 70 à 80% des papiers de robotique avec des "bouts" d'IA (RN, ML, deep, etc.).

A l'inverse, la vague est moins évidente dans les revue web sémantique.

Perspective : L'unité a un positionnement bien défini, mais elle doit veiller à ne pas se laisser distancer. L'intelligence artificielle (IA) vient compléter l'expertise historique des chercheurs sans en prendre la place.

Il y a un rapprochement en cours entre les deux axes principaux : Données (IA symbolique) et Robotique (contrôle et IA), grâce à une restructuration des équipes en thèmes.

Des défis se posent autour des volumes de données, notamment avec les caméras 3D (2 minutes de capture représentant 7 Go de données). Le LabCom AgriVIA (en partenariat avec la société Inogura) explore également l'adaptation des robots à leur environnement en se basant sur l'analyse des données multi-capteurs.

Il est nécessaire de recruter des Chargés de Recherche (CR) pour approfondir les aspects pratiques, sans quoi les travaux des doctorants ne pourront pas être réutilisés efficacement. Deux CR sont attendus pour faciliter la transition vers l'utilisation de l'IA générative.

4.5. MISTEA

Contexte : L'unité compte environ 60 personnes (hors stagiaires), dont une trentaine de titulaires. Elle se consacre à un thème transversal : la compréhension et l'aide à la décision sur des enjeux liés à l'agriculture ainsi qu'à la gestion et la protection des ressources naturelles.

Les disciplines représentées incluent la statistique, les probabilités, l'apprentissage automatique, les systèmes dynamiques, l'automatique, l'informatique, l'intelligence artificielle symbolique, les ontologies, l'intelligence artificielle par contraintes et la représentation des connaissances. L'unité développe et applique des méthodes en intelligence artificielle.

L'unité est organisée en 3 axes (Systèmes Dynamiques, Probabilités Statistiques et Informatique) en plus d'une cellule d'ingénierie OpenSilex.

Données : Les capteurs et l'imagerie sont au cœur du laboratoire. Il y a une certaine utilisation du spectre ainsi que des données annotées par des experts. Les données sont à la fois temporelles et spatiales, avec une forte composante temporelle ou fonctionnelle. Des indicateurs sont extraits des images pour permettre leur analyse dans un référentiel explicable. En revanche, il y a moins de données disponibles pour la partie contrôle.

Les métadonnées, très structurées, jouent également un rôle important mais le laboratoire ne travaille pas sur le texte.

Méthodes : L'utilisation du deep learning (DL) permet de calculer des indicateurs à partir d'images, comme par exemple le nombre de feuilles.

Concernant la modélisation de systèmes dynamiques, le machine learning (ML) est utilisé pour une partie du modèle. Des travaux débutent également avec le deep learning et les réseaux de neurones (DL/RN) pour la prédiction, notamment à travers les réseaux de neurones pour la physique (PINNs). À ce stade, il n'y a pas encore d'application de l'apprentissage par renforcement (RL), mais un peu d'assimilation de données. Des tests sont en cours avec les GAN pour l'augmentation de données.

En ce qui concerne la statistique, plusieurs techniques sont explorées : deep learning (DL), méthodes ABC (Approximate Bayesian Computation), mélanges, random forests, apprentissage non supervisé, clustering bayésien à grande échelle, stochastic bloc models, bandits et apprentissage séquentiel.

Modèles hybrides pour la gestion des connaissances (modélisation RDF et apprentissage de représentation).

Calculs : Utilisation des langages de programmation R, Java, Ruby, Python, un peu de Julia, quasiment plus de Matlab. Usage de Copilot pour le développement.

Très bonnes ressources régionales pour le calcul : meso@LR cluster régional, le Cines aussi qui pourrait être sollicité, mais pas encore de fortes demandes pour du ML à très grosses échelles. Pas d'utilisation de Jean Zay.

Bibliographie : Les publications sont possibles dans des conférences prestigieuses du machine-learning telles que NeurIPS, ICML, COLT, ou dans des journaux classiques de statistiques, ainsi que dans des revues spécialisées dans les domaines d'application.

Pour les travaux liés au web sémantique, il est possible de viser des supports spécialisés. Cependant, il est plus difficile de publier dans des journaux généralistes situés à l'interface, comme *Computers in Agriculture*. Les articles de conférences sur l'intelligence artificielle sont généralement orientés vers les domaines d'application spécifiques.

Perspective : Suite à de bons recrutements, l'unité est bien positionnée dans le domaine du deep-learning, avec une dynamique très réactive et attentive aux évolutions. Les compétences se sont bien diffusées, le positionnement actuel facilite le recrutement des stagiaires et des permanents. La transition est déjà bien enclenchée et l'ambiance sereine.

Il reste des besoins en formation (FIDLE), en recrutement et à conserver une veille technologique sur les outils mais l'unité a du recul et de la maturité sur les techniques d'IA moderne. La présence d'acteurs locaux tels que l'Université de Montpellier et l'Inria Zenith (PlantNet) constitue un atout majeur.

4.6. MIAT

Contexte : L'unité compte deux équipes de recherche SCIDYN (Simulation, Contrôle et Inférence de Dynamiques Agro-environnementales et Biologiques) et SaAB (Statistique et Algorithmique pour la Biologie) et deux plateformes Genotoul Bioinfo et RECORD.

Données : Pour l'équipe SaAB, les données sont essentiellement de type omique, incluant les séquences, les données RNAseq, les comptages, ainsi que d'autres types de données à très grande dimension. Ces données peuvent aussi prendre la forme de graphes, avec des mesures à différents niveaux, telles que les protéines ou les métabolites. Des données spectrales et, dans une moindre mesure, des images sont également exploitées. En revanche, l'équipe SCIDYN travaille avec peu de données, quelques trajectoires GPS d'animaux ou d'état des patients. Son travail repose principalement sur des données issues de simulations, ainsi que sur quelques enquêtes fournissant des tableaux de petite dimension.

Méthodes : Les méthodes utilisées incluent l'apprentissage automatique (ML) et l'apprentissage profond, avec un rôle à la fois d'utilisateur et de développeur. L'accent est moins mis

sur les preuves de convergence. Le travail se situe entre le développement de nouvelles méthodes et la création d'architectures d'apprentissage profond spécifiques (notamment pour les protéines et le génome). Des techniques issues des LLM sont testées sur les données génomiques. AlphaFold est utilisé pour la modélisation 3D des protéines.

Par ailleurs, des simulations multi-agents, de l'apprentissage par renforcement (RL et DeepRL) et des approches issues de la théorie des jeux sont également exploitées, à l'interface entre la recherche opérationnelle (RO) et l'intelligence artificielle (IA). Des tests sont en cours sur les modèles de diffusion pour la génération de protéines (en alternative à des approches combinatoires).

Calculs : Développements en R, Python, scikit-learn, TensorFlow, PyTorch pour le ML/DL, C C++, de moins en moins de Matlab. L'unité s'est équipée de GPU puissants (H100).

Bibliographie : Les travaux en apprentissage automatique (ML) sont publiés dans des conférences du domaine (comme ICML). Les contributions en biologie computationnelle et, dans une moindre mesure, en statistique computationnelle, sont souvent publiées dans des revues spécialisées. La participation à des conférences nationales en statistique permet de rester en contact avec la communauté. En apprentissage profond (DL) et ML, les conférences sont souvent suivies en visio. Des publications sont également réalisées dans des revues telles que NAR Genomics and Bioinformatics et lors de conférences comme l'ECCB, à la croisée entre méthodologies et applications. L'accent est également mis sur les grandes conférences en intelligence artificielle (IA), telles que IJCAI, AAAI, et UAI, qui continuent de traiter de sujets comme la programmation par contraintes (CP), le raisonnement logique et l'optimisation. Enfin, plusieurs revues se situent à l'interface entre écologie et modélisation, un domaine où la part de l'apprentissage automatique (ML) et de l'apprentissage profond (DL) a considérablement augmenté.

Perspective : L'IA est encore là, mais tout est très marqué par le DL, l'unité se positionne plutôt en utilisatrice, à l'exception de quelques individus. La vitesse de changement est étonnante, c'est difficile à suivre et la pression est grande.

Certains choisissent de se focaliser sur les approches frugales, à l'opposé du deep learning.

Les compétences en deep-learning sont bien présentes, les formations (FIDLE) suivies.

4.7. MAIAGE

Contexte : L'unité de recherche MaIAGE regroupe des mathématiciens appliqués, informaticiens et biologistes pour étudier des questions de biologie et d'agro-écologie à différentes échelles. Elle est composée de cinq équipes : Dynenvie (Modélisation dynamique et Statistique pour les écosystèmes, l'épidémiologie et l'agronomie), Bibliome (Acquisition et formalisation de connaissances à partir de textes), BioSys (biologie des systèmes), StatInfOmics (Bioinformatique et Statistique pour les données omiques) et Migale (plateforme de bioinformatique).

Données : Aux petites échelles, ces données incluent des informations omiques et méta-omiques issues du séquençage (génome, transcriptome), mais également métabolome ou

traits phénotypiques permettant d'étudier les écosystèmes microbiens, et le microbiote et les plantes. Elles peuvent être statiques ou résulter d'un suivi temporel. Aux échelles plus larges, elles se concentrent également sur les aspects temporels, tels que l'épidémiologie (végétale et animale) et l'évolution des parcelles agricoles, les interactions entre plantes et l'évolution des phénotypes.

Un contexte fréquent des travaux de l'unité est l'analyse d'un faible nombre d'échantillons, mais avec un grand nombre de covariables.

L'unité travaille sur des études multi-échelles, par exemple sur la chaîne alimentaire et les liens entre différents compartiments comme l'herbe, le lait et le fromage. Enfin, l'unité traite également des données textuelles et conduit des recherches sur le traitement du langage naturel (NLP), avec une expertise sur l'ensemble de la chaîne, depuis la collecte des données et la gestion de grandes bases documentaires à leur représentation formalisée à travers des ontologies et des bases de connaissances en se focalisant sur le développement d'outils d'extraction d'information.

Méthodes : L'unité a une expertise dans la modélisation statistique, par exemple sur les modèles à variables latentes. Récemment, les chercheurs se sont tournés notamment vers les Variational Autoencoders (VAE) et des transformers, tout en exploitant l'auto-différentiation pour optimiser les modèles. Elle utilise également AlphaFold pour la prédiction 3D des structures protéiques. L'unité de recherche développe depuis peu des outils d'analyse des données omiques basés sur le traitement du langage naturel (NLP).

De façon générale et sans se focaliser sur les domaines d'application, les chercheurs développent des modèles dynamiques très variés, qu'ils soient déterministes ou stochastiques, continus ou discrets en temps ou en état. Ces approches génériques de modélisation impliquent classiquement l'inférence de paramètres, l'analyse de sensibilité ou la quantification d'incertitude. Différentes collaborations avec d'autres unités permettent d'aborder des problématiques spécifiques complexes. Les modèles stochastiques en épidémiologie font également partie des outils de l'unité.

Dans le domaine de l'extraction d'information, l'unité développe des outils de NLP originaux (en reconnaissance d'entités nommées, détection de relations et identification/alignement). Elle combine des approches historiques du NLP avec des contributions LLM, incluant le prompting. En collaboration avec Migale, les modèles sont valorisés et mis en production dans des chaînes de traitement.

Calcul : L'unité utilise Python, JAX et PyTorch pour ses calculs en plus des outils classiques (Python, R, C/C++).

Les ressources locales intègrent la plateforme bioinformatique Migale. L'unité utilise également le mesocentre LabIA, bénéficiant de l'accompagnement des ingénieurs de LabIA, ce qui constitue un atout majeur. Enfin, Jean Zay est utilisé régulièrement, bien que l'accès soit limité pour les personnels étrangers, notamment chinois et russes (ZRR).

Notons aussi les réflexions de l'équipe Migale autour du déploiement des GPU en production qui pourraient servir à d'autres unités.

Bibliographie : L'unité publie sur les données omiques dans différentes conférences de machine learning et journaux de

statistiques (e.g. ICML).

La bibliographie de MAIAGE, en sources comme en production, est fortement influencée par le deep learning : il s'agit d'une tendance de fond pour les chercheurs spécialisés en traitement de données comme pour ceux qui sont plus centrés sur des domaines applicatifs. Dans cette transition, les chercheurs remarquent fréquemment des papiers qualifiés de "paresseux", qui se contentent d'utiliser les derniers modèles en vogue (de BERT à GPT sans proposer d'analyse convaincante).

En NLP, des plateformes comme GitHub et Paperwithcode sont devenues des ressources bibliographiques importantes pour l'équipe Bibliome ; les débouchés sont divers pour la publication des contributions de l'équipe avec notamment des revues comme PLOS ONE et BMC Bioinformatics.

Perspective : Le deep-learning intéresse la majorité de l'unité : ceux qui sont déjà largement dedans, ceux qui utilisent ponctuellement des outils et ceux qui ressentent le besoin de se former à ces nouvelles techniques. Tous soulignent la difficulté à suivre le rythme rapide des avancées dans ce domaine et la pression toujours plus grande dans une communauté scientifique en expansion. Il est parfois nécessaire de trouver des niches spécifiques à exploiter afin de rester compétitifs et visibles dans ce contexte.

Le laboratoire se (ré)organise pour combler les lacunes en épidémiologie animale suite au décès d'une collègue. Plusieurs collègues soulignent leur intérêt pour les Physics Informed Neural Networks (PINNs) à l'interface entre les méthodes mécanistes et les architectures neuronales. Cet axe est renforcé par une collaboration avec l'équipe MIA-PS. Un recrutement en statistique via une Chaire de Professeur Junior (CPJ) est en cours pour renforcer l'expertise de l'unité en apprentissage statistique et en deep-learning.

Enfin, l'unité s'interroge sur les moyens de rivaliser avec les équipes géantes chinoises ou américaines et sur la manière de construire une masse critique pour rester à la pointe à la fois en recherche et en ingénierie.

4.8. MIA-PS

Contexte : L'UMR MIA Paris-Saclay, associée aux tutelles AgroParisTech, INRAE et Université Paris-Saclay est structurée en deux équipes (respectivement centrées sur les statistiques et l'informatique) : Statistical mOdelling and Learning for environnemenT and lIfE Sciences (SOLsTIS) Expert Knowledge, INteractive modellINg and learnINg for understandINg and decisiOn makINg in dyNamic Complex Systems (EkI-Nocs).

Données : Les données sont très variées, à la fois du côté statistique et du côté informatique. Informations génétiques et biologiques, SNIPs et les marquages, relatives aux plantes et aux animaux. Données provenant de capteurs et de l'observation de phénomènes évolutifs. Images satellites (rendement, végétation) et suivi des territoires. Trajectoires de bateaux et d'animaux, ainsi que des cinétiques. Données d'interaction diverses pour la construction de profils. Données issues de caméras et de drones. Forte dimension multi-modale avec un objectif prédictif marqué. Analyse de séries temporelles. Données textuelles et connaissances, telles que les ontologies. etc...

Méthodes : Modélisation statistique et optimisation (e.g. PLN *Poisson Log-Normal*). Beaucoup d'expérience autour des modèles hybrides : Exemples de PINNs, intégration hybride entre statistiques, systèmes d'équations différentielles et données pharmacocinétiques. Hybridation entre deep learning et méthodes statistiques autour des VAE (Variational Auto-Encoder).

Du côté informatique, des approches très éclectiques : extraction d'information, génération d'explications ou de recommandations. Utilisation de Random Forest, méthodes d'ensemble, ... Et des approches hybrides avec de la modélisation par agents et des algorithmes génétiques pour l'optimisation de systèmes complexes.

Approche large, allant de la preuve de convergence à l'application de méthodes améliorées, telles que les GAN ou plus récemment les normalized flows. Des développeurs en deep learning, mais aussi des utilisateurs simples (comme pour l'analyse d'images). Côté statistiques, chacun doit comprendre les architectures, avec quelques titulaires plus qualifiés pour approfondir ces aspects et diffuser les connaissances dans les équipes.

Calcul : Les équipes utilisent principalement R et python, du C et JAVA pour les développements spécifiques. Les moyens de calcul sont locaux (des ordinateurs avec GPU) et Jean Zay.

Bibliographie : Les publications couvrent un large éventail, incluant des articles dans des revues et des conférences en informatique et en statistiques. En statistique, on observe une tendance croissante vers l'intelligence artificielle (IA) et le machine learning (ML). Des publications sont également présentes dans des domaines d'application spécifiques, tels que la biologie et l'écologie.

Perspective : Les équipes disposent de compétences avancées et visent des objectifs ambitieux en matière de formation et de recherche. Des idées et projets autour de l'IA générative, déjà bien développés à l'interface entre la statistique et le machine learning, montrent un positionnement novateur. La stratégie de l'unité est claire et le site bénéficie d'un environnement très favorable, offrant de nombreuses opportunités d'interactions avec des institutions prestigieuses comme Télécom, Centrale, Polytechnique, entre autres.

4.9. OPAALE

Contexte : Une unité composée de 49 permanents, répartis en quatre équipes. L'équipe ACTA travaille sur des techniques de contrôle des écoulements d'air et des transferts thermiques. L'équipe IRM-Food développe des méthodes pour mieux comprendre les transformations des bio-matrices. L'équipe PANDOR se concentre sur la caractérisation et la valorisation des déchets organiques agricoles et urbains. Enfin, l'équipe SAFIR évalue et améliore les procédés de valorisation des biomasses résiduelles, notamment via la méthanisation et le compostage.

Données : En mécanique des fluides, les données sont collectées en laboratoire, dans des environnements contrôlés ou sur le terrain, à l'aide de capteurs d'images et de vidéos. Le défi est de travailler avec des données non uniformes souvent en très grande dimension.

Les autres équipes disposent également d'une grande quantité de données, incluant des images IRM, des signaux RMN,

ainsi que des données captées par divers capteurs, images, et des informations issues d'enquêtes sous forme textuelle ou tabulaire.

Méthodes : Modèles de mécanique des fluides, modélisation numérique, et le couplage entre modèles et données, ainsi que la modélisation stochastique en dynamique des fluides. L'équipe est déjà compétente sur des modélisations hybrides.

Les scientifiques exploitent les équations différentielles ordinaires (EDO) et partielles (EDP), avec des contributions en mathématiques appliquées, analyse numérique et statistiques. Une thèse est encadrée sur l'intelligence artificielle générative et la modélisation stochastique. Les modèles de bruits sont définis par des principes physiques.

L'utilisation de l'apprentissage automatique et profond est récente, notamment pour le couplage avec les modèles physiques, mais il n'y a pas encore de recours aux réseaux neuronaux pour la validation des modèles ou les problèmes inverses. Des méthodes d'analyse de cycle de vie sont également étudiées.

Calcul : Utilisation des langages Matlab, scilab, comsol, codes Navier-Stokes. Python Pytorch. Utilisation de diverses ressources pour le calcul : GPU locaux + Jean Zay (partie CPU), Meso LR, TGCC CEA.

Bibliographie : Publications en statistiques, en physique et dans les domaines applicatifs (respectivement sur le retraitement et la mécanique des fluides).

En mécanique des fluides, beaucoup de travaux à l'interface avec l'IA : combinaisons de modèles (local/dynamique). Un certain nombre de ponts existent avec le domaine de la vision (ECCV, ICCV, CVPR...)

Perspective : Dans les procédés biologiques, la vague IA reste mesurée : des formations et recrutements ciblés sont envisagés pour identifier et exploiter les pistes prometteuses.

Côté ACTA, l'équipe sait bien exploiter les techniques d'IA, en partie suite au recrutement de compétences dans le domaine et grâce à des collaborations avec l'INRIA. Les PINNs représentent à la fois une grande opportunité de faire progresser la recherche... Mais ce sont aussi de très bons attracteurs de financements (hybridation, PINNs, ...) qui augmentent la visibilité de l'équipe et facilitent les recrutements.

La formation repose sur les MOOC et le réseau eNeuron. Il existe des interrogations autour d'une possible transition vers le langage JULIA, bien que Python soit actuellement utilisé.

Concernant la gestion des déchets, le sujet suscite un intérêt croissant, bien que parfois perçu comme surestimé, avec un risque de léger recul à moyen terme.

Différents usages en expansion rapide pour chatGPT (écriture, développement, ...). Des questions en suspens sur les bonnes pratiques.

4.10. TETIS

Contexte : L'unité regroupe 120 agents au total, spécialisés dans le domaine de l'information géospatiale au sens large, avec un focus sur la caractérisation des états des systèmes territoriaux et leurs dynamiques. Elle est structurée en trois équipes scientifiques et héberge aussi l'EP Inria Evergreen, qui intègre des technologies d'intelligence artificielle en collaboration avec les équipes ATOS et Misca.

Données : Les données proviennent de diverses sources, incluant l'imagerie (satellites, drones) et des capteurs terrestres (optique, lidar, radar, hyperspectral), ainsi que des bases de données, des textes, des relevés de terrain et des contributions participatives. Les séries temporelles et la très haute résolution sont des spécialités de l'unité.

En complément, des données tabulaires sont utilisées pour l'analyse des réseaux. Ces données, très variées, sont de nature multi-sources, hétérogènes et temporelles.

Méthodes : la représentation des connaissances et le text mining, le web sémantique, ainsi que l'apprentissage automatique appliqué aux données d'imagerie.

Elles intègrent également le traitement du signal, l'optimisation, la vision par ordinateur et les réseaux neuronaux. Les techniques de modélisation multi-agents, les graphes et l'analyse des réseaux sont aussi utilisés.

Enfin, des méthodes statistiques et des algorithmes de machine learning, comme les SVM, sont régulièrement utilisés. Certains membres se concentrent sur le développement de nouvelles méthodes, tandis que d'autres se spécialisent davantage dans leurs données, avec des compétences méthodologiques variées.

Calculs : Les calculs se font principalement en local, avec une vingtaine de GPU disponibles au sein du laboratoire. Des ressources supplémentaires sont utilisées via un serveur Inria, ainsi qu'un accès (limité pour l'instant) à Jean Zay.

Bibliographie : Les sources bibliographiques sont très variées sur l'ensemble de l'unité. Elles incluent des revues spécialisées en télédétection (remote sensing), en science des données, ainsi qu'à l'interface de ces domaines. Elles couvrent également des publications sur les systèmes experts et les systèmes complexes, entre autres.

Perspective : L'unité est déjà bien avancée dans la transition IA dans deux de ses trois équipes, notamment sur les domaines de l'image et du texte. Le partenariat avec Inria apporte une visibilité accrue sur les thématiques de l'apprentissage automatique et du deep learning. L'équipe se considère actuellement en position de force dans le domaine à l'interface de l'intelligence artificielle (IA) avec diverses modalités de données (connaissances géographique, image, texte, réseaux).

L'équipe estime que la place du chercheur en informatique s'est renforcée ces dernières années par rapport au périmètre et aux outils de l'IA.

ChatGPT est largement adopté pour la rédaction et le codage, bien que des réflexions sur les bonnes pratiques d'utilisation soient encore en cours.

5. Conclusion

A l'issue de cette enquête menée auprès des unités du département MathNum, nous disposons d'une vision assez complète de l'IA et de sa place dans les recherches du département, en termes de forces quantitatives, de dynamiques mais aussi de ressentis au sein des unités et des équipes. Nous présentons en annexe A certaines de ces données sous forme de tableaux ou de figures.

Si nous avons souhaité considérer dans l'enquête une définition la plus large possible de l'IA, il est certain que par la force des choses, nos entretiens et nos questionnaires ont abordé

plus en profondeur et en détail l'impact de la vague actuelle autour de l'apprentissage automatique, de l'apprentissage profond et des modèles de langues. Néanmoins, nous pensons que l'ensemble du rapport et de ses résultats pourront servir de base pour l'élaboration du prochain schéma stratégique de département, dans toutes ses dimensions relevant de l'IA, en lien donc avec la statistique, l'informatique, et les autres disciplines de MathNum.

Sans anticiper les éléments de ce futur SSD, nous pouvons alors simplement lister quelques points clés qui nous semblent ressortir d'une première analyse :

- les deux premières composantes disciplinaires du département restent, en lien avec l'IA, la modélisation statistique (ML1) et la modélisation informatique (SIM2);
- les composantes relatives à l'IA « historique » symbolique et combinatoire (IAS, KB, RO, GR, SIM2) ont encore aujourd'hui un poids en équivalent temps-plein (etp) supérieur aux composantes apprentissage statistique et apprentissage machine (ML2-3-4);
- les etp en apprentissage profond ou génératif (DL1-5) sont déjà supérieurs à ceux en apprentissage statistique et machine, et sont presque au niveau de l'IA symbolique et combinatoire;
- les forces en décision séquentielle (RL1-3), contrôle (CTR), processus (Seq), etc. sont bien en deçà (en etp);
- 40% de nos travaux en IA se positionnent sur une posture de développement, 30% en théorique et idem en utilisateur. Mais le théorique remonte fortement pour la modélisation statistique ou informatique, pour l'IA combinatoire et symbolique, et diminue pour l'apprentissage profond;
- de nombreuses équipes et unités de MathNum croisent plusieurs des composantes disciplinaires de l'IA;
- l'étude des données spatio-temporelles (SpT, ST) et omiques (OmQ) est très majoritaire dans le département. Les données tabulaires et textuelles (Tab, TAL), puis images (Vis) complètent l'ensemble. Le département affiche une posture théorique plutôt sur les données tabulaires, temporelles, spatio-temporelles et textuelles.
- les types de données et de méthodes ne sont pas orthogonaux : on voit bien que les spécialisations des équipes sur des données vont souvent de pair avec des approches spécifiques en IA.

Le point relatif au croisement entre disciplines (IA et statistique, informatique, automatique) au sein des unités du département nous semble particulièrement important à souligner. Il faut noter également la richesse des types de données abordées au sein des équipes. De par son histoire, le département apparaît ainsi comme bien positionné pour aborder sereinement la vague de l'IA actuelle où l'hybridation des méthodes est centrale.

A. Premiers résultats

La figure 2 présente la synthèse des réponses des équipes en termes d'etp distribués sur les disciplines et les types de données. Il est à noter ici que ces nombres d'etp ne correspondent pas aux seuls agent Inrae du département MathNum, mais plus largement à l'ensemble des personnels scientifiques au sein de ces unités, souvent pluri-instituts, qui s'inscrivent sur les missions du département.

Sur les figures 3 et 4 sont représentées les projections de ce tableau sur chacune des deux dimensions de l'enquête.

Si l'on prend par exemple les deux regroupements ML (statistique et apprentissage automatique) et DL (apprentissage profond), on peut alors comparer les répartitions des forces en fonction des types de donnée (figures 5 et 6).

La même chose peut être faite sur le regroupement ingénierie des connaissances et ontologies (figures 7).

En représentant sur un même graphique les types de données et les méthodes utilisées sur ces données, nous voyons bien que les deux sont souvent liés (Figure 8).

B. Guide d'entretien

Ce questionnaire a servi de guide général : toutes les questions n'ont pas été systématiquement posées et rarement dans cet ordre. L'idée de cette enquête était de construire un référentiel, de positionner les acteurs MathNum dans ce référentiel et surtout de comprendre les directions résultant des chocs successifs du deep-learning et de l'IA générative, y compris sur des thématiques a priori assez lointaines.

B.1. Applications & Modalités

Après un tour de table et une introduction indiquant explicitement une acception très large du terme *Intelligence Artificielle*.

- Quelles sont les applications marquantes développées ces dernières années dans votre équipe, centrées sur l'IA ou utilisant de l'IA ?
- Quel(s) type(s) de données sont en jeu ?
 - Données textuelles/images/séries temporelles/omics/données tabulaires/télémétrie ?
 - Données d'expert, processus, enquêtes de terrain ?
 - Bases de connaissances ?
- Quelle volumétrie de données est associée à vos applications ?
 - Avez-vous des problématiques de très petits jeux de données, des problèmes d'ingénierie avec de très gros volumes ?
- Quel est le cœur de compétence de votre équipe en terme d'application et de types de données ?
- Combien de personnes impliquées ? Combien de thèse en ce moment ?

B.2. De la théorie à la pratique

Idéalement en reprenant les applications phares mentionnées dans la section précédente

- Entre le métier et la science des données, qui fait quoi ?
 - Une personne multi-compétente / une équipe avec une répartition des rôles ? / une équipe métier et un appui technique externe pour l'analyse des données ?

- Pensez-vous être utilisateur des outils d'IA ou développeur de ces outils ?
 - Voyez-vous une frontière distincte entre utilisation et développement ?
- Quelle capacité de généralisation dans les outils développés ?
 - Quelle force de développement pour adapter les outils à de nouveaux cas de figure ?
- Quelle balance entre statistiques et informatique ? Comment situer les personnes impliquées dans ces projets sur cet axe ?
 - Quels objectifs pour les chercheurs ? Démontrer une capacité opérationnelle ? Obtenir des garanties théoriques sur les modèles, sur la convergence de l'optimisation ?
- Quelle interface entre la recherche et l'ingénierie ?
 - Dans le cadre des projets en IA, êtes-vous utilisateur des outils INRAe PEPI (Partage d'Expériences et de Pratiques en Informatique) / CATI (Centre Automatisé de Traitement de l'Information) ?

B.3. Types d'IA, cadre de développement

Les questions suivantes sont très largement liées à des questions précédentes. L'idée est de re-questionner les applications abordées sous l'angle du cadre de développement des chercheurs.

- Quelles familles de modèles utilisez-vous ?
 - Analyse de données : Machine-Learning, Deep-Learning, Apprentissage par renforcement, Modélisation statistique (choix ou construction de lois de probabilité)
 - Ingénierie des Connaissances : logique & représentation des connaissances (graphes, lexiques, ontologies)
 - Optimisation : recherche opérationnelle, développement heuristique, algorithmes génétiques
 - Modèle mécaniste : modélisation mathématique (e.g. EDP), simulation multi-agent
 - Des approches hybrides mêlant plusieurs propositions (e.g. jumeaux numériques, PINNs, ...)
- Les outils utilisés par les uns et les autres vous semblent-ils discriminants pour distinguer/structurer différents types d'IA ?
- Quels outils sont à la base de vos travaux ?
 - Quels langages ? R ? Python ? C, JAVA ? SQL, SparQL ?
 - Quelles bibliothèques générales ? Cplex, Scikit-Learn, TensorFlow, pytorch, Maelia, GAMA, bibliothèque omics...
 - Quelles bibliothèques plus spécialisées ? (Huggingface, séries temporelles, télémétrie, AlphaFold...)
 - Papier/stylo ? Pour les théorèmes et les garanties théoriques
- Quelle manière de développer ? Exploitation d'outils (scripts python, shell, ...), construction de modèles à partir d'outils (architecture de machine learning pré-entraînée), construction à partir de zéro (SMA, modèles mécanistes, ...)

		Machine learning												Deep-Learning					Reinforcement			Optimisation		Simulation	
		IAS	KB	RI	ML1	ML2	ML3	ML4	DL1	DL2	DL3	DL4	DL5	CTR	RL1	RL2	RL3	Seq	OPT	RO	GR	SIM1	SIM2		
		Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff	Eff		
(Tab) Données tabulaires	U	1	2	0	0.6	0.45	0.45	0.7	0	0.2	0	0	0	0	0	0	0	0.2	0.2	0.05	0.25	0	0		
	D	0	0	0	0.5	0.3	0.55	0.35	0.4	0.1	0	0.1	0	0	0.2	0	0	0	0	0.85	0.3	0.3			
(ST) Séries temp., systèmes dyn., capteurs	T	0.1	0	0	1.3	1	0	0.4	0	0	0	0	0	0	0.2	0.2	0	0.2	0.2	0.4	0.2	0.3			
	U	0	1.25	0	0.4	0.65	0.25	0.25	0	0.2	0.1	0.2	0.3	0	0	0	0	0.2	0.8	0.2	0.1	0.45			
(SpT) Données spatiales et séries spatio-temporelles	D	0	0.5	0	0.65	0.2	0.7	0.2	0	0.25	0	0	1.15	0.6	0	0	0	0.2	0.35	0.55	0	0.65			
	T	0	0	0	1.6	0.3	0.25	0.2	0	0.2	0	0.25	0	0.35	0.3	0.2	0	0.2	0.5	0.75	1.2	1.6			
(PM) Paroles et musique	U	0	1	0	0.7	0.3	0.6	0.85	1.3	1	0.3	0.3	0	0.25	0	0	0.3	0	0.1	0.25	0.25	0			
	D	0	0	0	3.65	0.675	0.625	0.325	0.125	0.25	0.2	0.45	1.5	0	0	0	0	0.2	1	0	0	2.2			
(Vis) Image/vision	T	0.1	0	0	2.15	0	0	0	0.3	0.2	0.2	0.75	0	0	0	0	0	0	0	0	0	2.2			
	U	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
(Vid) Vidéo	D	0	0	0	0	0	0	0	0	0	0	0	0	0.25	0	0	0	0	0	0	0	0			
	T	0	0	0.2	0	0	0	0	0	0	0	0	0	0.25	0	0	0	0	0	0	0	0			
(Tel) Télémétrie	U	0	0	0	0	0	0	0	0	0	0	0	0	0	0.25	0	0	0	0	0	0	0			
	D	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
(Spm) Spectrométrie	T	0	0	0	0.85	0.85	0.2	0.2	0.25	0.25	0.25	0.25	0.25	0	0	0	0	0	0	0	0	0.2			
	U	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
(TAL) Texte	D	0.1	0.6	0	0	0	0	0	0.2	0.2	0.1	0.6	0.4	0	0	0	0	0.1	0	0	0	0			
	T	0.1	0.1	0	0.1	0	0	0	0.2	0.2	0.2	1.15	0.3	0	0	0	0	0	0	0	0	0			
(Exp) Données d'expertise, ontologies	U	0	2.5	0	0	0	0	0	0	0	0.1	0	0	0	0	0	0	0	0	0.25	0.25	0			
	D	0.3	0.2	0	0	0	0	0	0.1	0	0	0	0	0	0	0	0	0	0	0	0	0			
(Omng) Données omniques	T	0	0.8	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
	U	0	0	0	0.7	0.6	0.2	1.1	0	0	0.35	0	0.2	0	0	0	0	0.1	0.1	0.25	0	0			
(Rec) Données d'interaction & recommandation	D	0	1.1	0	1.65	1.75	0.65	1.1	1.2	0	0.3	0.6	0	0	0	0	0	0	0	1.1	0.2	0.2			
	T	0.2	0.6	0	0.9	0.2	0	0.2	0	0	0.1	0.25	0	0	0	0	0	0.2	0	0	0.1	0.1			
(Eng) Données d'enquête	U	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			
	D	0	0	0.2	0	0	0	0	0	0	0	0	0	0.25	0	0	0	0	0	0	0	0.45			
(Par) Données de paramétrage IA	T	0	0	0	0	0	0	0	0	0	0	0	0	0.25	0	0	0	0	0	0	0	0			
	U	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0			

FIGURE 2. Résultats de l'enquête menée auprès de toutes les équipes du département (voir figure 1).

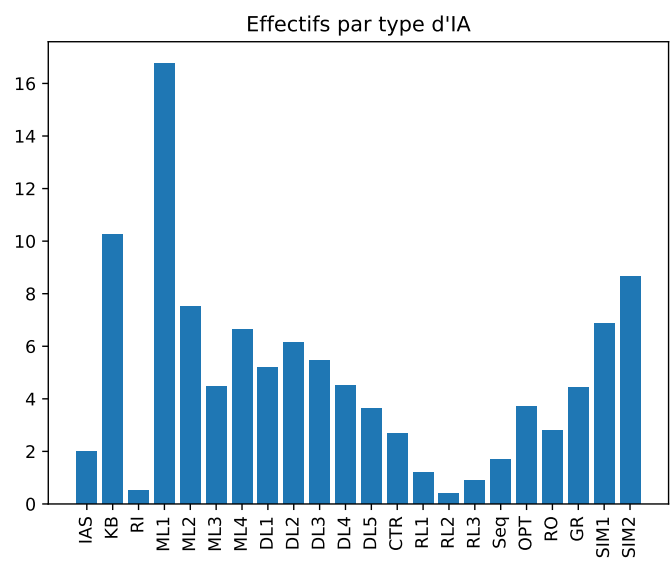


FIGURE 3. Effectifs en IA par type de méthodes (en etp).

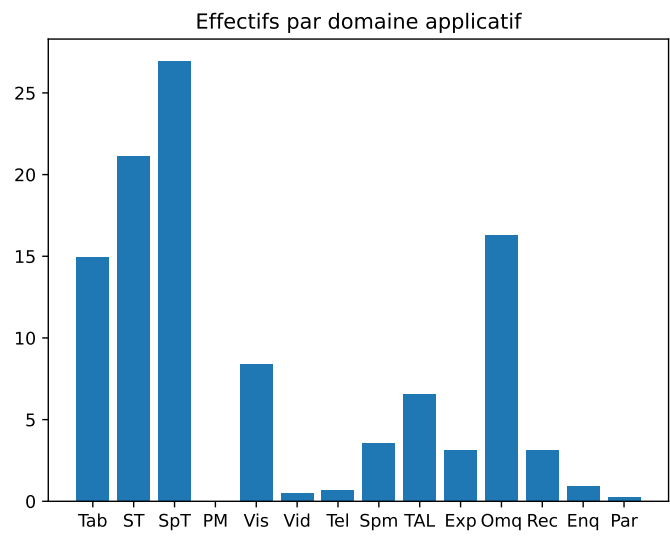


FIGURE 4. Effectifs en IA par type de données (en etp).

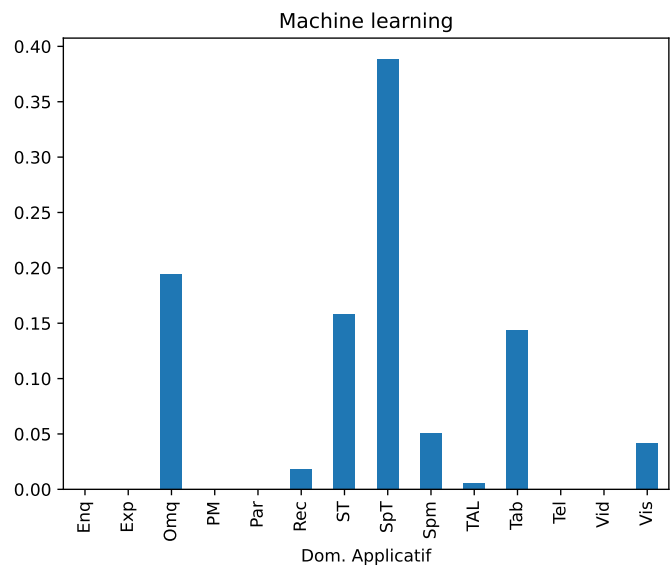


FIGURE 5. Effectifs en ML par type de données (en %).

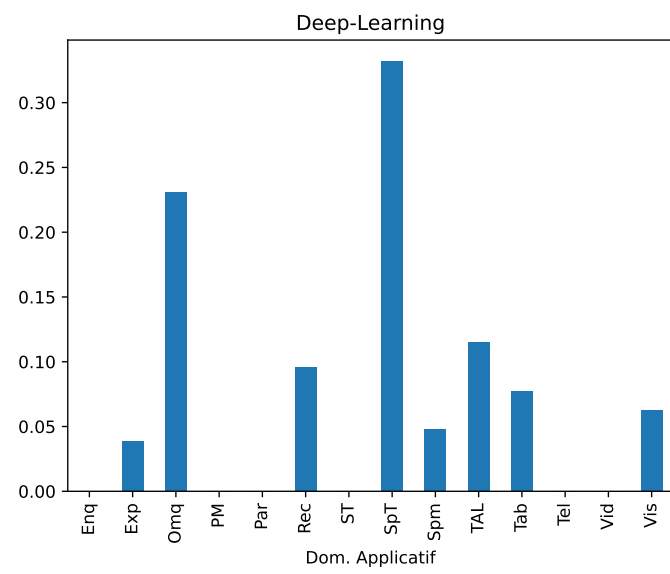


FIGURE 6. Effectifs en DL par type de données (en %).

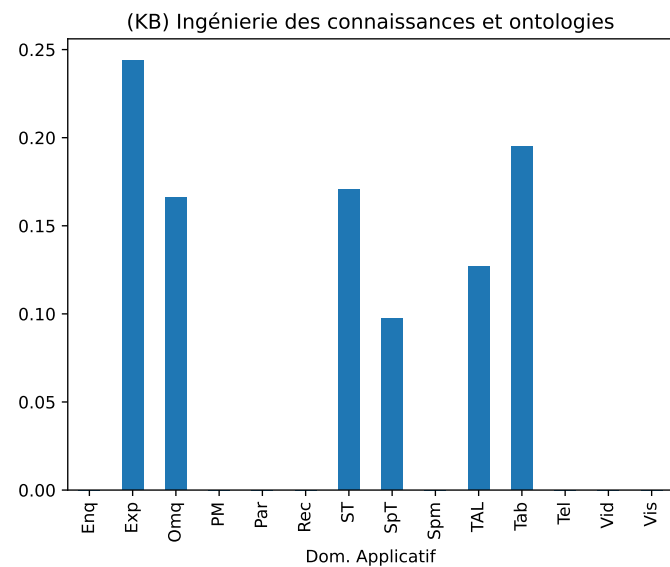


FIGURE 7. Effectifs en KB par type de données (en %).

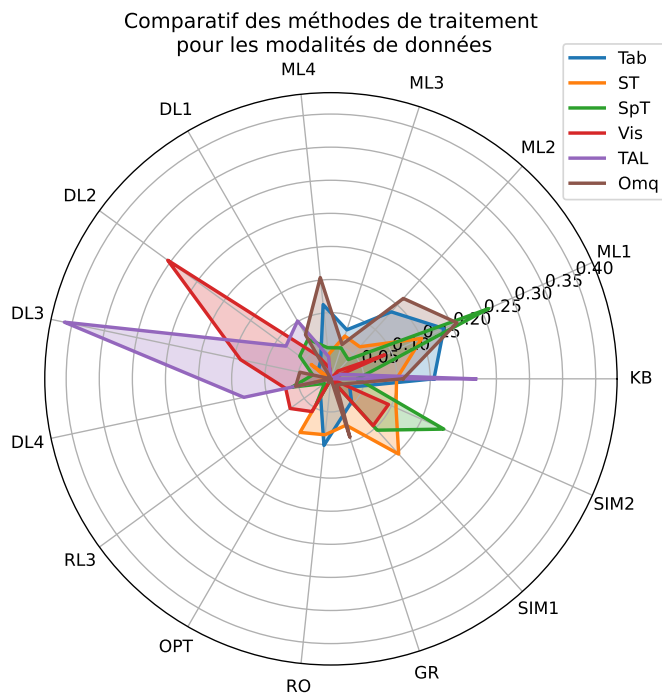


FIGURE 8. Répartition des méthodes d'IA utilisées pour les 6 types de données majoritaires (en %).

- Quelle infrastructure matérielle ? simple ordinateur, HPC, GPU, serveur
 - Quelle infrastructure particulière ? Serveur local, régional, national (GENCI/Jean Zay)
 - Qui s'occupe du maintien ?

B.4. Bibliographie

L'enjeu de cette section est de catégoriser les lieux de publication, de comprendre les enjeux scientifiques et de valorisation. Les sources bibliographiques elles-mêmes sont impactées par les vagues actuelles d'IA et d'IA générative : nous voulons savoir comment les équipes se positionnent par rapport à ces changements. L'enjeu de ce rapport n'est en aucun cas quantitatif : il ne s'agit pas de refaire ou de compléter l'HCERES.

- Distinguez-vous à l'intérieur des projets des publications plus orientées *application* et d'autres plus orientées *techniques de l'IA* ?
 - Quelles conférences/revues citeriez-vous dans chaque domaine ?
- Est ce que vous ciblez les mêmes sources pour vos publications et vos lectures ?
 - Que penser du niveau de sélectivité dans les grandes sources de machine-learning ?
- Comment vos sources ont-elles été impactées par les dernières vagues d'IA (en particulier par l'IA Générative) ?
 - Avez-vous changé de stratégie de publication ?
 - Quel est votre ressenti par rapport à ces changements ?

B.5. Prospective

- Globalement, comment voyez-vous le futur vis-à-vis des propositions actuelles en IA/IA générative ? Pas d'impact,

quelques nouveaux outils, de nouvelles technologies in-contournables à prendre en main

- Quels sont les enjeux généraux de l'IA pour vous dans les années à venir ?
 - Au niveau des données métier : mieux exploiter les données actuelles ?
 - Nouveaux types d'IA sur les données actuelles (générative AI, renforcement, ...)
 - Une opportunité pour exploiter de nouveaux *gismes* ?
 - Y a-t-il des enjeux spécifiques à l'IA générative ?
 - L'hybridation des méthodes est-elle une option dans votre domaine ?
- Comment prendre le virage ? [s'il y en a besoin] Quels besoins ?
- RH : une transition passe-t-elle, pour vous, par des recrutements ? Quelles attentes et quelles interactions avec les recrues ?
 - Quelles difficultés actuelles pour recruter ? Quels risques pour le futur ?
 - Quels besoins en formation ?
 - Quels besoins en recrutement ? Ingénieur, chercheur ?
- chatGPT : qui s'en sert ? Pour quoi faire ? Un rapport à l'outil générationnel ?
 - Est-ce bien ? mal ? comment encadrer les pratiques ?
- Les infrastructures matérielles auxquelles vous avez accès sont-elles suffisantes ?
- Quels projets concrets envisagez-vous dans les 5 ans ? Comment les positionner dans le référentiel précédent : type d'IA, type de publication, type d'environnement logiciel ?

■ Références

- [1] M. GHALLAB, D. NAU, P. TRAVERSO et M. MILANO, « Acting, Planning, and Learning, » *American history*, t. 1861, n° 1900, 1945.
- [2] L. RABINER et B.-H. JUANG, *Fundamentals of speech recognition*. Prentice-Hall, Inc., 1993.
- [3] D. JURAFSKY, *Speech & language processing*. Pearson Education India, 2000.
- [4] G. SCHREIBER, *Knowledge engineering and management : the CommonKADS methodology*. MIT press, 2000.
- [5] F. S. HILLIER, *Introduction to operations research*. McGrawHill, 2005.
- [6] R. O. DUDA, P. E. HART et al., *Pattern classification*. John Wiley & Sons, 2006.
- [7] C. D. MANNING, *Introduction to information retrieval*. Syngress Publishing, 2008.
- [8] T. HASTIE, R. TIBSHIRANI, J. FRIEDMAN et al., *The elements of statistical learning*, 2009.
- [9] D. KOLLER et N. FRIEDMAN, *Probabilistic Graphical Models : Principles and Techniques*. The MIT Press, 2009, p. 1231, ISBN : 978-0262013192.

- [10] J. PEARL, *Causality : Models, Reasoning and Inference, 2nd Edition*. Cambridge University Press, 2009, p. 486, ISBN : 978-0521895606.
- [11] A. M. UHRMACHER et D. WEYNS, *Multi-Agent Systems : Simulation and Applications*. CRC Press, 2009, p. 582, ISBN : 978-1420070231.
- [12] F. RICCI, L. ROKACH et B. SHAPIRA, « Introduction to recommender systems handbook, » in *Recommender systems handbook*, Springer, 2010, p. 1-35.
- [13] J. B. CAMPBELL et R. H. WYNNE, *Introduction to remote sensing*. Guilford press, 2011.
- [14] R. DECHTER, *Constraint Processing*. Morgan Kaufmann, 2013, p. 502, ISBN : 978-0123996718.
- [15] O. SIGAUD et O. BUFFET, *Markov Decision Processes in Artificial Intelligence*. Wiley-ISTE, 2013, p. 621, ISBN : 978-1118619940.
- [16] G. E. BOX, G. M. JENKINS, G. C. REINSEL et G. M. LJUNG, *Time series analysis : forecasting and control*. John Wiley & Sons, 2015.
- [17] M. BELAISSAOUI, *La Logique et l'Intelligence Artificielle : Raisonnements & Algorithmes*. Éditions Universitaires Européennes, 2016, p. 148, ISBN : 978-3639509434.
- [18] I. GOODFELLOW, Y. BENGIO, A. COURVILLE et Y. BENGIO, *Deep learning*. MIT press Cambridge, 2016, t. 1.
- [19] Y. HASIN, M. SELDIN et A. LUSIS, « Multi-omics approaches to disease, » *Genome biology*, t. 18, n° 1, p. 83, 2017.
- [20] R. S. SUTTON et A. G. BARTO, *Reinforcement Learning, Second Edition : An Introduction*. Bradford Books, 2018, p. 552, ISBN : 978-0262039246.
- [21] B. P. ZEIGLER, A. MUZY et E. KOFMAN, *Theory of Modeling and Simulation : Discrete Event & Iterative System, Computational Foundations*. Academic Press, 2018, p. 692, ISBN : 978-0128133705.
- [22] Y. L. CUN, *Quand la machine apprend : La révolution des neurones artificiels et de l'apprentissage profond*. Odile Jacob, 2019, p. 394, ISBN : 2738149316.
- [23] J. HENDLER, F. GANDON et D. ALLEMANG, *Semantic Web for the Working Ontologist : Effective Modeling for Linked Data, RDFS, and OWL*. ACM Books, 2020, p. 498, ISBN : 978-1450376167.
- [24] S. RUSSELL et P. NORVIG, *Intelligence artificielle - Une approche moderne, 4e édition*. Pearson France, 2021, p. 976, ISBN : 2744074551.
- [25] R. SZELISKI, *Computer vision : algorithms and applications*. Springer Nature, 2022.
- [26] M. LAPAN, *Deep Reinforcement Learning Hands-On : A practical and easy-to-follow guide to RL from Q-learning and DQNs to PPO and RLHF*. Packt Publishing, 2024, p. 716, ISBN : 978-1835882702.
- [27] O. CAPPÉ et C. MARC, *Tout comprendre (ou presque) sur l'intelligence artificielle*. CNRS Éditions, 2025, p. 140.
- [28] S. KONIECZNY et H. PRADE, *Artificial Intelligence-What is it, exactly? Second Edition*. College publication, 2025, ISBN : 978-1-84890-338-8.