

Intelligence Artificielle

Les défis actuels et l'action d'Inria



Les livres blancs d'Inria examinent les grands défis actuels du numérique et présentent les actions menées par nos équipes-projets pour résoudre ces défis.

Ce document est le premier produit par la cellule veille et prospective d'Inria. Cette unité, par l'attention qu'elle porte aux évolutions scientifiques et technologiques, doit jouer un rôle majeur dans la détermination des orientations stratégiques et scientifiques d'Inria. Elle doit également permettre à l'Institut d'anticiper l'impact des sciences du numérique dans tous les domaines sociaux et économiques.

Ce livre blanc a été coordonné par Bertrand Braunschweig avec des contributions de 45 chercheurs d'Inria et de ses partenaires.

Un grand merci à Peter Sturm pour sa relecture précise et complète. Merci également au service STIP du centre de Saclay – Île-de-France pour la correction finale de la version française.



Intelligence artificielle :

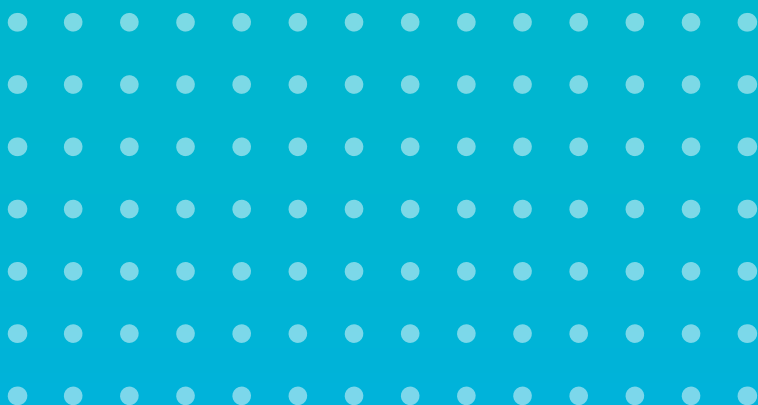
une thématique de recherche majeure pour Inria

1. Samuel et Toi.Net	5
2. 2016, l'année de l'intelligence artificielle (IA) ?	9
3. Les débats sur l'IA	14
4. Les défis de l'IA et les contributions d'Inria	19
4.1 Les défis génériques de l'intelligence artificielle	23
4.2 Dans l'apprentissage automatique	26
4.3 Dans l'analyse des signaux : vision, parole	37
4.4 Dans les connaissances et le web sémantique	46
4.5 En robotique et véhicules autonomes	53
4.6 En neurosciences et sciences cognitives	59
4.7 Dans le traitement du langage	63
4.8 En programmation par contrainte pour l'aide à la décision	65
4.9 En musique et environnements intelligents	68
5. Références Inria : chiffres	73
6. Références supplémentaires	76
Remerciements	81





Samuel et Toi.Net¹





7h15, Sam se réveille et se prépare pour une journée de travail ordinaire. Il prend sa douche et s'installe dans la cuisine pour prendre son petit déjeuner. Toi.Net², sa servante-robot, lui sert un café chaud et une assiette de fruits. « *Toi.Net, passe-moi le sucre s'il te plaît* », demande Sam. La servante-robot prend le sucrier qui se trouve à l'autre bout de la table et le lui apporte. Pourtant, il y a une boîte de sucre dans le placard de la cuisine mais Toi.Net sait qu'il est beaucoup plus simple d'utiliser le sucrier. « *Des nouvelles intéressantes ?* », demande Sam. La servante-robot comprend qu'elle doit trouver l'actualité qui correspond aux centres d'intérêt de Sam. Elle commence par le football.

Toi.Net: « *Monaco a battu Marseille 3-1 à domicile. C'est la première fois depuis douze ans que Monaco marque trois buts contre Marseille. Un coup du chapeau de Diego Suarez* ».

Toi.Net : « *L'élection de Miss France s'est déroulée au Palais des congrès. Miss Martinique est l'heureuse élue* ».

Sam : « *Ne m'embête plus avec ce type d'actualité, les élections de Miss ne m'intéressent pas* ».

Toi.Net : « *D'accord, je ne le ferai plus* ».

Toi.Net : « *Pour Paris, la météo prévoit du soleil dans la matinée et de fortes pluies à partir de 13 heures et toute l'après-midi* ».

1 Le titre de cette section fait référence à l'écrivain britannique du 19^e siècle Samuel Butler, auteur de *Erewhon*, l'un des premiers livres à formuler des hypothèses sur la possibilité d'une intelligence artificielle s'appuyant sur la théorie de la sélection naturelle darwinienne et la reproduction entre machines.

2 Personnage de la servante dans la comédie de Molière "*Le Malade imaginaire*".

Toi.Net : « *M. Lamaison, candidat à la présidence de la région sud-ouest, déclare que le taux de chômage a atteint les 3,2 millions ; son plus haut niveau depuis 2004* ».

Sam : « *Peux-tu vérifier cette information ? J'ai l'impression que ce nombre était plus élevé au milieu des années 2010* ».

Deux secondes après, Toi.Net répond : « *Vous avez raison, le chômage s'élevait à 3,4 millions en 2015, selon les statistiques de l'INSEE* ».

À la fin du petit déjeuner, Sam ne se sent pas bien. Son bracelet connecté indique une pression artérielle anormale et Toi.Net reçoit la notification. « *Où avez-vous mis vos comprimés ?* » demande-t-elle à Sam. « *Sur la table de chevet ou peut-être dans la salle de bain* ». Toi.Net lui apporte la boîte de médicaments et Sam reprend des forces.

Toi.Net : « *C'est l'heure d'aller au travail. Puisqu'il va probablement pleuvoir lorsque vous irez faire de la marche au parc après le déjeuner, je vous ai apporté vos bottines* ».

Une voiture autonome attend devant la maison. Sam entre dans le véhicule qui lui annonce : « *Ce matin, je ferai un détour par l'A-4 en raison d'un accident sur votre trajet habituel et d'un temps d'attente de quarante-cinq minutes à cause d'un embouteillage* ».

Toi.Net est une servante-robot bien éduquée. Elle connaît beaucoup de choses sur Sam, comprend ses demandes, se souvient de ses préférences, est capable de trouver des objets et de les utiliser à bon escient, se connecte à Internet et extrait des informations pertinentes, apprend des situations nouvelles, etc. Cela a été rendu possible grâce aux immenses progrès réalisés dans le domaine de l'intelligence artificielle : le traitement et la compréhension de la parole (comprendre les demandes de Sam) ; la reconnaissance visuelle et la reconnaissance d'objets (localiser le sucrier sur la table) ; la planification automatisée (définir la série d'actions débouchant sur une situation donnée comme chercher une boîte de comprimés dans la salle de bain) ; la représentation des connaissances (pouvoir identifier un coup du chapeau comme une série de trois buts marqués par un même joueur au cours d'un même match de football) ; le raisonnement (choisir de prendre le sucrier présent sur la table plutôt que d'aller chercher la boîte de sucre dans le placard de la cuisine ou recourir aux prévisions météorologiques pour décider quelle paire de chaussures Sam devrait mettre) ; la fouille de données (extraire les informations pertinentes à partir du web, y compris vérifier les faits dans le cas de la déclaration politique) ; l'algorithme incrémental d'apprentissage automatique (lui permettant de se souvenir de ne plus mentionner les concours de Miss à l'avenir). La servante-robot adapte constamment

son interaction avec Sam en dressant le profil de son propriétaire et en détectant ses émotions.

Inria, avec plus de cent-soixante équipes-projets réparties au sein de huit centres de recherche, est très actif sur toutes les facettes de l'IA. Le présent document expose ses points de vue sur les grandes tendances et les principaux défis de l'intelligence artificielle et décrit la manière dont ses équipes conduisent des recherches scientifiques, développent des logiciels et œuvrent pour le transfert technologique afin de relever ces défis.



2016, l'année
de l'intelligence
artificielle (IA) ?



C'est ce qu'ont déclaré récemment plusieurs éminents scientifiques de Microsoft. **L'IA est devenue un sujet en vogue dans les médias et magazines scientifiques** en raison de nombreuses réalisations, dont beaucoup sont le fruit des progrès accomplis dans le domaine de l'apprentissage automatique. De grandes entreprises dont Google, Facebook, IBM, Microsoft mais aussi des constructeurs automobiles à l'instar de Toyota, Volvo et Renault, sont très actifs dans la recherche en IA et prévoient d'y investir davantage encore dans le futur. Plusieurs scientifiques spécialisés dans l'IA dirigent désormais les laboratoires de recherche de ces grandes entreprises et de nombreuses autres.

.....
Les progrès accomplis en matière de robotique, de véhicules autonomes, de traitement de la parole et de compréhension du langage naturel sont tout aussi impressionnants.
.....

La recherche en IA a permis de réaliser d'importants progrès dans la dernière décennie, et ce dans différents secteurs. Les avancées les plus connues sont celles réalisées dans l'apprentissage automatique, grâce notamment au développement d'architectures d'apprentissage profond, des réseaux de neurones convolutifs multicouche dont l'apprentissage s'opère à partir de gros volumes de données sur des architectures de calcul intensif. Parmi les

réalisations de l'apprentissage automatique, il convient de citer la résolution de jeux Atari (Bricks, Space invaders, etc.) par Google DeepMind, utilisant les pixels images affichés à l'écran comme données d'entrée afin de décider quelle action adopter pour atteindre le plus haut score possible à la fin de la partie. Il convient également de mentionner le dernier résultat obtenu par la même équipe dans le jeu de Go, l'emportant sur le meilleur joueur humain en une série de cinq matches, en recourant à une combinaison de recherche Monte Carlo arborescente d'apprentissage profond à partir de vrais matches et d'apprentissage par renforcement en jouant contre son propre clone.

Voici d'autres exemples remarquables :

- la description automatique du contenu d'une image ("Une image est plus parlante qu'un millier de mots"), également par Google (<http://googleresearch.blogspot.fr/2014/11/a-picture-is-worth-thousand-coherent.html>) ;
- les résultats du *Large Scale Visualisation Challenge 2012* d'Imagenet, remporté par un très grand réseau neuronal convolutif développé par l'Université de Toronto, et entraîné avec des GPU de Nvidia (<http://imagenet.org/challenges/LSVRC/2012/results.html>) ;
- la qualité des systèmes de reconnaissance faciale tels que ceux de

Facebook (<https://www.newscientist.com/article/dn27761-facebook-can-recognise-you-in-photos-even-if-youre-not-looking#.VYkVxFzjZ5g>) ;
- etc.

Mais il ne s'agit là que d'une très petite partie des résultats obtenus grâce à l'IA. Les progrès accomplis en matière de robotique, de véhicules autonomes, de traitement de la parole et de compréhension du langage naturel sont tout aussi impressionnants.

Comme exemples, on peut citer :

- le niveau de compétence atteint par les robots au Robotic Challenge de Darpa, remporté par KAIST en 2015. Dans cette compétition, les robots doivent conduire un véhicule utilitaire, ouvrir une porte et entrer dans un bâtiment, fermer une vanne, utiliser un outil pour percer un mur, se déplacer à travers des décombres ou retirer les débris bloquant une porte d'entrée et monter une échelle (https://en.wikipedia.org/wiki/DARPA_Robotics_Challenge) ;
- la compréhension de la parole. Elle est désormais considérée comme une fonctionnalité standard des smartphones et tablettes dotés de compagnons artificiels tels que Siri d'Apple, Cortana de Microsoft ou M de Facebook ;
- la traduction instantanée. Microsoft Skype Translator traduit des conversations dans différentes langues en temps réel ;
- les véhicules sans chauffeur. Des voitures autonomes ont parcouru des milliers de kilomètres sans incidents majeurs.

The screenshot shows a Google search result for 'Mondol Kiri'. On the left, there's a 'Fermé' (Closed) status for Monday, followed by links to TripAdvisor and Yelp reviews. The TripAdvisor review mentions a 4.5 rating from 158 reviews. The Yelp review mentions a 4.5 rating from 54 reviews. On the right, there's a map showing the location at 159 Avenue de Choisy, 75013 Paris. Below the map, there's a detailed description of the restaurant: 'Sièges noirs et murs rouges ou en pierre ornés de tableaux créent un cadre cosy pour des plats cambodgiens.' It also lists the address, phone number (01 53 79 75 96), and hours (Fermé aujourd'hui).

Figure 2 : Informations sémantiques ajoutées aux résultats du moteur de recherche Google

Tout aussi fondamentaux sont les résultats obtenus dans des domaines tels que la représentation des connaissances et le raisonnement, les ontologies et d'autres technologies pour le web sémantique et le web de données :

- **Google Knowledge Graph** améliore les résultats de recherche en affichant des données structurées sur les termes ou expressions de recherche demandés ;

- **Schema.org** contient des millions de triplets RDF décrivant des faits connus : les moteurs de recherche peuvent utiliser ces données pour fournir des informations structurées sur demande ;

- **Facebook utilise le protocole OpenGraph** - qui s'appuie sur RDFa - pour permettre à n'importe quelle page web de devenir un objet enrichi dans un graphe social.

Enfin, une autre tendance importante est la récente ouverture de plusieurs technologies auparavant propriétaires afin que la communauté des chercheurs en IA puisse non seulement en bénéficier, mais aussi contribuer à les enrichir par de nouvelles fonctionnalités. **Comme exemples, on peut citer :**

- **les services cognitifs d'IBM Watson**, disponibles sur des interfaces programmatiques (API), qui offrent jusqu'à vingt technologies différentes telles que les services *Speech To Text* et *Text To Speech*, l'identification des concepts et des liens entre concepts, la reconnaissance visuelle et bien d'autres : <http://www.ibm.com/cloud-computing/bluemix/solutions/watson/> ;

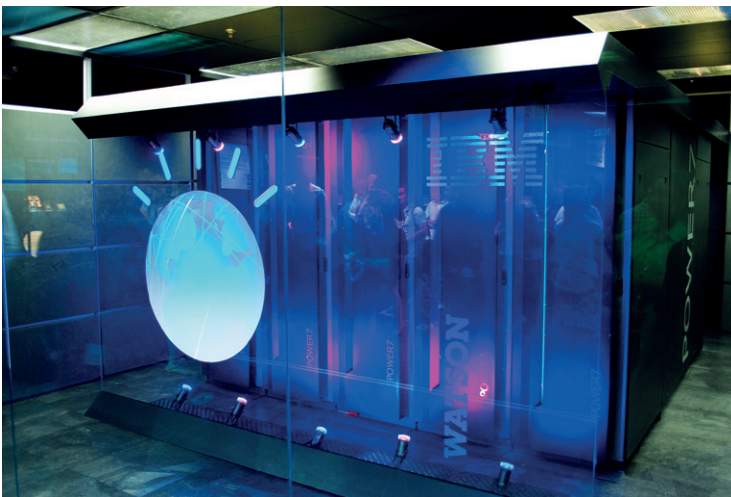


Figure 3 : Ordinateur IBM Watson - © Clockready (CC BY-SA 3.0, via Wikimedia Commons)

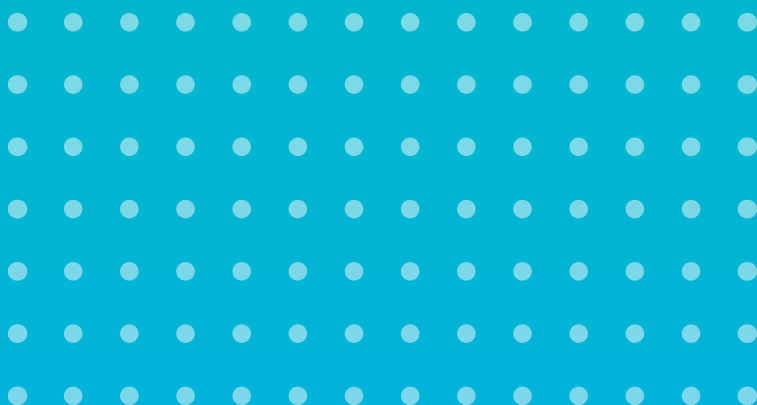
- **la bibliothèque logicielle TensorFlow de Google** utilisée pour l'apprentissage automatique, qui a été mise en *open source* ; <https://www.tensorflow.org/> ;

- la mise en *open source* par Facebook du design de son serveur **Big Sur** pour exécuter de grands réseaux neuronaux d'apprentissage profond sur des GPU : https://code.facebook.com/posts/1687861518126048?_mref=message_bubble.

L'enthousiasme suscité par toutes ces réalisations positives a été cependant tempéré par les inquiétudes exprimées ici et là par des scientifiques de renom. Ce sera le sujet de la section suivante.



Les débats sur l'IA



Les débats sur l'IA ont commencé au 20^e siècle - cf. **les lois de la robotique d'Isaac Asimov** - mais s'intensifient aujourd'hui en raison des récentes avancées, décrites précédemment, dans le domaine.

Selon la théorie de la singularité technologique, une ère de domination des machines sur l'Homme verra le jour lorsque les systèmes d'intelligence artificielle deviendront super-intelligents :

"La singularité technologique est un événement hypothétique lié à l'avènement d'une véritable intelligence artificielle. Ainsi un ordinateur, un réseau informatique ou un robot seraient théoriquement capables d'une auto-amélioration récursive (processus de perfectionnement auto-généré) ou de concevoir et fabriquer des ordinateurs ou robots plus intelligents. Des répétitions de ce cycle pourraient aboutir à un effet d'accélération - une explosion de l'intelligence - c'est-à-dire à des machines intelligentes capables de concevoir des générations successives de machines de plus en plus puissantes, créant une intelligence largement supérieure aux capacités intellectuelles humaines, d'où un risque de perte de contrôle. L'Homme ne pouvant appréhender les capacités d'une telle super-intelligence, la singularité technologique est le point au-delà duquel l'intelligence humaine ne pourrait ni prédire ni même imaginer les événements" (traduction de Wikipédia anglais).

Les défenseurs de la singularité technologique sont proches du **mouvement transhumaniste** qui vise à améliorer les capacités physiques et intellectuelles des humains grâce aux nouvelles technologies. La singularité correspondrait au moment où la nature des êtres humains subirait un changement fondamental, lequel serait perçu soit comme un événement souhaitable (par les transhumanistes), soit comme un danger pour l'humanité (par leurs opposants).

.....
*Une ère de
 domination des
 machines sur
 l'Homme verra
 le jour lorsque
 les systèmes d'IA
 deviendront super-
 intelligents*

Le débat sur les dangers de l'IA a franchi un nouveau palier avec la récente polémique sur **les armes autonomes et les robots tueurs**, suscitée par une lettre ouverte publiée à l'ouverture de la conférence IJCAI en 2015³. La lettre, qui demande l'interdiction des armes capables de fonctionner sans intervention humaine, a été signée par des milliers de personnes, dont Stephen Hawking, Elon Musk, Steve Wozniak et nombre de chercheurs de premier plan en IA ; on retrouve

³ Voir <http://futureoflife.org/open-letter-autonomous-weapons/>

d'ailleurs parmi les signataires certains des chercheurs d'Inria ayant contribué à l'élaboration du présent document.

Parmi les autres dangers et menaces ayant fait l'objet de débats au sein de la communauté, on peut citer : les conséquences financières sur les marchés boursiers du *trading* haute fréquence, lequel représente désormais la grande majorité des ordres passés - dans le *trading* haute fréquence, des logiciels soi-disant intelligents exécutent à grande vitesse des transactions financières, pouvant conduire à des crashes boursiers comme le Flash Crash de 2010 ; les conséquences de la fouille de données massives sur le respect de la vie privée, avec des systèmes de fouille capables de divulguer les caractéristiques personnelles des individus en établissant des liens entre leurs opérations en ligne ou leurs enregistrements dans des bases de données ; et bien entendu le chômage potentiel généré par le remplacement progressif de la main d'œuvre par des machines.

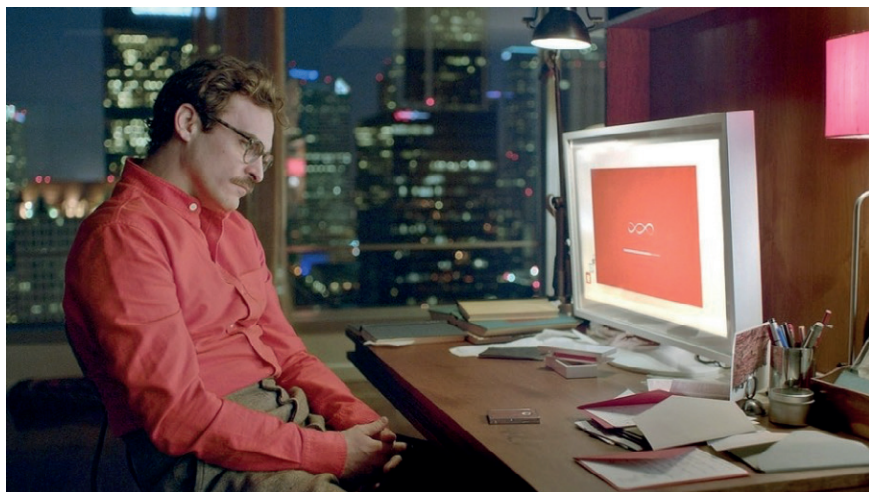


Figure 4 : Dans le film *Her* de Spike Jonze (Annapurna Pictures / Warner Bros, 2013), un homme tombe amoureux de son système d'exploitation intelligent

Plus nous développons l'intelligence artificielle, plus le risque est grand d'étendre uniquement certaines capacités intelligentes (par exemple optimisation et fouille par l'apprentissage) au détriment d'autres, peu susceptibles de générer un retour sur investissement immédiat ou de présenter un quelconque intérêt pour le créateur de l'agent (par exemple : morale, respect, éthique, etc.). Utilisée à grande échelle,

l'intelligence artificielle peut comporter de nombreux risques et constituer quantité de défis pour les humains, en particulier si les intelligences artificielles ne sont pas conçues et encadrées de façon à respecter et protéger les humains. Si, par exemple, l'optimisation et les performances sont le seul objectif de leur intelligence, cela peut conduire à des catastrophes à grande échelle où les utilisateurs sont instrumentalisés, abusés, manipulés, etc. par des agents artificiels inépuisables et sans vergogne. La recherche en IA doit être globale et inclure tout ce qui rend les comportements intelligents, pas uniquement les "aspects les plus raisonnables".

Dietterich et Horvitz ont publié dernièrement une réponse intéressante à certaines de ces questions⁴. Dans leur bref article, les auteurs reconnaissent que la communauté des chercheurs en IA ne devrait pas se focaliser sur le risque de perte de contrôle par les humains car celui-ci ne présente pas de caractère critique dans un avenir prévisible, et recommandent plutôt d'accorder plus d'attention aux cinq risques à court terme auxquels sont exposés les systèmes basés sur l'IA, à savoir :

- les bugs dans les logiciels ;
- les cyberattaques ;
- la tentation de jouer à "l'Apprenti Sorcier", c'est à dire donner la capacité aux systèmes d'IA de comprendre ce que veulent les utilisateurs au lieu d'interpréter littéralement leurs ordres ;
- "l'autonomie partagée", à savoir la coopération fluide des systèmes d'IA avec les utilisateurs de façon que les utilisateurs puissent toujours reprendre le contrôle en cas de besoin ;
- et les impacts socio-économiques de l'IA : en d'autres termes l'IA doit être bénéfique pour l'ensemble de la société et pas seulement pour quelques privilégiés.

Inria n'ignore rien de ces débats et, en tant qu'institut de recherche consacré aux sciences du numérique et au transfert technologique, œuvre pour le bien-être de tous, pleinement conscient de ses responsabilités envers la société. Informer la société et les instances dirigeantes des potentialités et des risques des sciences et technologies du numérique fait partie des missions d'Inria.

Dans cette optique, ces dernières années, Inria a :

- **lancé une réflexion sur l'éthique** bien avant que les menaces de l'IA ne suscitent des débats au sein de la communauté scientifique ;

⁴ Dietterich, Thomas G. and Horvitz, Eric J., Rise of Concerns about AI: Reflections and Directions, Communications of the ACM, October 2015 Vol. 58 no. 10, pp. 38-40

- **contribué à la création de la CERNA⁵ d'Allistene**, une commission de réflexion qui travaille sur les problèmes éthiques soulevés par la recherche en sciences et technologies du numérique ; son premier rapport de recommandations est consacré à la recherche en robotique ;

- **mis en place un nouvel organe chargé d'évaluer au cas par cas les enjeux légaux ou éthiques d'une recherche** : le Comité Opérationnel d'Évaluation des Risques Légaux et Éthiques (COERLE), composé de scientifiques d'Inria et de contributeurs externes. Le COERLE a pour mission d'aider à identifier les risques et à déterminer si l'encadrement d'un projet de recherche est nécessaire.

En outre, Inria encourage ses chercheurs à prendre part aux débats sociétaux lorsque des médias les sollicitent pour s'exprimer sur des questions éthiques telles que celles soulevées par **la robotique, l'apprentissage profond, la fouille de données et les systèmes autonomes**.

Les enjeux scientifiques et technologiques mis au jour par la recherche en IA amènent Inria à développer des stratégies pour relever les multiples défis soulevés. La prochaine section sera consacrée à son positionnement et à la typologie de ces défis.

⁵ Commission de réflexion sur l'Éthique de la Recherche en sciences et technologies du Numérique de l'Alliance des Sciences et Technologies du Numérique : <https://www.allistene.fr/cerna-2/>



Les défis de l'IA et les contributions d'Inria



L'IA est un vaste domaine ; toute tentative de le structurer en sous-domaines peut donner lieu à débat. Toutes les typologies se valant, il a donc été décidé d'opter pour la hiérarchie de mots-clés proposée récemment par la communauté des responsables d'équipes-projets Inria, afin de mieux identifier leurs contributions aux sciences du numérique en général. Dans cette hiérarchie, l'intelligence artificielle est un mot-clé de premier niveau avec sept sous-domaines, certains spécifiques, d'autres renvoyant à différentes sections de la hiérarchie (voir le tableau ci-dessous).

■ CONNAISSANCES

- _Bases de connaissances
- _Extraction & nettoyage de connaissances
- _Inférence
- _Web sémantique
- _Ontologies

■ APPRENTISSAGE AUTOMATIQUE

- _Apprentissage supervisé
- _Apprentissage (partiellement) non-supervisé
- _Apprentissage séquentiel et par renforcement
- _Optimisation pour l'apprentissage
- _Méthodes bayésiennes
- _Réseaux de neurones ou neuronaux
- _Méthodes à noyau
- _Apprentissage profond
- _Fouille de données
- _Analyse de données massives

■ TRAITEMENT DU LANGAGE NATUREL

■ TRAITEMENT DES SIGNAUX

- _Parole
- _Vision
 - Reconnaissance d'objets
 - Reconnaissance d'activités
 - Recherche dans des banques d'images et de vidéos
 - Reconstruction 3D et spatio-temporelle
 - Suivi d'objets et analyse des mouvements
 - Localisation d'objets
 - Asservissement visuel

■ ROBOTIQUE (y compris les véhicules autonomes)

- _Conception
- _Perception
- _Décision
- _Action

- _ Interactions avec les robots (environnement/humains/robots)
- _ Flottes de robots
- _ Apprentissage des robots
- _ Cognition pour la robotique et les systèmes

■ **NEUROSCIENCES, SCIENCES COGNITIVES**

- _ Compréhension et stimulation du cerveau et du système nerveux
- _ Sciences cognitives

■ **ALGORITHMIQUE DE L'IA**

- _ Programmation logique et ASP
- _ Dédution, preuve
- _ Théories SAT
- _ Raisonnement causal, temporel, incertain
- _ Programmation par contraintes
- _ Recherche heuristique
- _ Planification et ordonnancement

■ **AIDE À LA DÉCISION**

Tableau 1: Hiérarchie de mots-clés pour l'IA établie par Inria

Les contributions d'Inria sont présentées ci-après par équipe-projet.

Inria s'appuie sur un modèle de recherche original qui repose sur une entité de base : l'équipe-projet. Les équipes-projets réunissent, autour d'une personnalité scientifique, un groupe de chercheurs, d'enseignants chercheurs, de doctorants et d'ingénieurs. Elles ont toutes un objectif commun : relever un défi scientifique et technologique dans l'un des domaines de recherche prioritaires de l'institut définis dans le plan stratégique. L'équipe-projet a une durée maximale de vie de douze ans et une durée moyenne de huit ans. (Les noms des équipes-projets seront écrits en **CAPITALES ORANGE** afin de les distinguer des autres noms.)

PLAN STRATÉGIQUE "OBJECTIF INRIA 2020"

La stratégie scientifique d'Inria est construite autour de deux axes complémentaires, sur lesquels s'articulera la contribution de l'institut :

• **Les sciences et technologies du numérique utiles à l'humain, à la société et la connaissance.**

Les chantiers majeurs sur lesquels Inria s'engage mettent l'humain au cœur des problématiques du numérique :

- l'humain en tant que tel : santé et bien-être ;

- l'humain et ses environnements : de l'individu à la société, de l'habitat à la planète ;
- l'humain et la connaissance : émergence, médiation et éducation.

• **Les développements scientifiques prioritaires au cœur de nos sciences :**

- calculer le futur : modèles, logiciels et systèmes numériques ;
- maîtriser la complexité : données, réseaux et flux ;
- interagir avec les mondes réel et numérique : interactions, usages et apprentissage.

Nous abordons des questions clés posées par d'autres sciences ou par de grands domaines d'application, auxquels les sciences informatiques et mathématiques devront contribuer :

- **la santé et le bien-être ;**
- **l'énergie et les ressources naturelles ;**
- **l'environnement et le développement durable ;**
- **la société et l'éducation.**

Après un premier paragraphe consacré aux défis génériques, seront abordés des défis plus spécifiques par ordre décroissant des investissements d'Inria : cela signifie que l'inventaire commencera par les domaines dans lesquels Inria compte le plus grand nombre d'équipes et de chercheurs actifs. Cela ne signifie pas que les équipes contribuant aux défis présentés en dernier sont de moindre qualité, seulement qu'elles sont moins nombreuses.

Les définitions de l'IA et des sous-domaines ne seront pas reprises dans cet exposé car elles font l'objet d'une abondante littérature. De bonnes définitions figurent également sur Wikipédia, par exemple :

https://en.wikipedia.org/wiki/Artificial_intelligence

https://en.wikipedia.org/wiki/Machine_learning

<https://en.wikipedia.org/wiki/Robotics>

https://en.wikipedia.org/wiki/Natural_language_processing

https://en.wikipedia.org/wiki/Semantic_Web

https://en.wikipedia.org/wiki/Knowledge_representation_and_reasoning

etc.

4.1 Les défis génériques de l'intelligence artificielle

Les principaux défis génériques identifiés par Inria sont les suivants : (i) IA située ; (ii) processus d'apprentissage avec intervention humaine ; (iii) ouverture à d'autres disciplines ; (iv) multitâches ; (v) validation et certification des systèmes d'IA.

IA située

Les systèmes d'IA doivent opérer dans le monde réel et interagir avec leur environnement : recevoir des données de capteurs, déterminer le contexte dans lequel ils opèrent, agir sur le monde réel ; ils doivent par ailleurs se comporter de façon autonome et conserver leur intégrité dans diverses conditions. Pour répondre à ces exigences, les systèmes d'IA doivent gérer des données non structurées ainsi que des données sémantiques.

Processus d'apprentissage avec intervention humaine

Les systèmes d'IA ont vocation à interagir avec des utilisateurs humains : ils doivent donc être capables d'expliquer leur comportement, de justifier d'une certaine manière les décisions qu'ils prennent afin que les utilisateurs humains puissent comprendre leurs actions et leurs motivations. Si cette compréhension n'est pas au rendez-vous, les utilisateurs humains n'auront qu'une confiance limitée, voire inexistante, dans les systèmes, qui ne seront donc pas acceptés. En outre, les systèmes d'IA ont besoin d'une certaine flexibilité et de capacités d'adaptation afin de pouvoir gérer différents utilisateurs et différentes attentes. Il est important de développer des mécanismes d'interaction qui favorisent une bonne communication et interopérabilité entre les humains et les systèmes d'IA. Certaines équipes d'Inria travaillent sur la collaboration entre l'IA et l'interface Homme-machine car il existe de nombreuses attentes dans ce domaine.

Ouverture à d'autres disciplines

Une IA sera souvent intégrée dans un système élargi composé de nombreux éléments. L'ouverture signifie donc que des scientifiques et développeurs en IA devront collaborer avec des spécialistes d'autres disciplines des sciences informatiques (par exemple la modélisation,

la vérification et la validation, les réseaux, la visualisation, l'interaction Homme-machine, etc.) composant le système élargi, ainsi qu'avec des scientifiques d'autres champs de compétences contribuant à l'IA comme les psychologues, les biologistes (en biomimétique, notamment), les mathématiciens, etc.

Le deuxième aspect à prendre en compte est l'impact des systèmes de l'IA sur plusieurs facettes de notre vie, de notre économie et de notre société et donc la nécessaire collaboration avec des spécialistes d'autres domaines. Il serait trop long de tous les mentionner mais à titre d'exemples on peut citer les économistes, les environnementalistes, les juristes.

Multitâches

De nombreux systèmes d'IA excellent dans un domaine spécifique, mais se révèlent incompetents en-dehors de celui-ci. Pourtant, les systèmes évoluant dans un environnement réel, tels que les robots, doivent être capables de réaliser plusieurs actions en parallèle, comme mémoriser des faits, assimiler de nouveaux concepts, agir sur le monde réel et interagir avec les humains.

Validation et certification

Composantes incontournables des systèmes critiques, la certification des systèmes d'IA ou leur validation par des moyens appropriés, constituent de véritables défis, en particulier si ces systèmes répondent aux attentes citées ci-dessus (adaptation, multitâches, processus d'apprentissage avec intervention humaine). La vérification, la validation et la certification des systèmes classiques (qui ne relèvent donc pas de l'IA) constituent déjà des tâches difficiles - même s'il existe déjà des technologies exploitables, dont certaines sont développées par des équipes-projets d'Inria. L'application de ces outils aux systèmes d'IA complexes est une tâche titanesque à laquelle il convient de s'attaquer pour être en capacité d'utiliser ces systèmes dans des environnements tels que les avions, les centrales nucléaires, les hôpitaux, etc.

Autres défis génériques

Outre les précédents défis, les exigences ci-dessous concernant les systèmes d'IA devraient donner lieu à de nouvelles activités de recherche : certaines sont extrêmement complexes et ne peuvent pas être satisfaites à court terme, mais méritent qu'on s'y intéresse.

Inculquer des normes et des valeurs aux systèmes d'IA va bien au-delà des sciences et technologies existantes : par exemple, un robot qui va acheter du lait pour son propriétaire doit-il s'arrêter en chemin pour aider une personne dont la vie est en danger ? Une technologie d'IA puissante pourrait-elle être utilisée par des terroristes artificiels ? À l'heure actuelle, la recherche en IA est bien loin de pouvoir répondre à ces exigences.

L'exigence de protection de la vie privée est particulièrement importante pour les systèmes d'IA confrontés aux données personnelles, tels que les assistants/compagnons intelligents ou les systèmes de fouille de données. Cette exigence valait déjà pour les systèmes classiques, mais les systèmes d'IA ont ceci de particulier qu'ils généreront de nouvelles connaissances à partir des données privées et les rendront probablement publiques à défaut de moyens techniques capables d'imposer des restrictions.

Enfin, un dernier défi concerne le passage à échelle. **Les systèmes d'IA doivent être capables de gérer de vastes quantités de données et de situations.** On pense aux algorithmes d'apprentissage absorbant des millions de points de données (signaux, images, vidéos, etc.) et aux systèmes de raisonnement à grande échelle, tels que le système Watson d'IBM, utilisant des connaissances encyclopédiques. Cependant, la question du passage à échelle pour les nombreux V (variété, volume, vitesse, vocabulaires, etc.) reste entière.

Applications de l'IA

Il ne s'agit pas à proprement parler d'un défi pour l'IA, mais il est important d'insister sur le fait que **les systèmes d'IA contribuent à résoudre des problèmes sociétaux :** les applications d'IA couvrent tout le spectre des activités humaines, telles que l'environnement et l'énergie, la santé et l'assistance à l'autonomie à domicile, le transport et les villes intelligentes, etc. Ils peuvent être bénéfiques pour l'humanité et l'économie, mais ils peuvent également représenter des menaces s'ils ne sont pas contrôlés (voir la section 3).

4.2 Les défis génériques dans l'apprentissage automatique

Les algorithmes et systèmes d'apprentissage automatique ont connu d'importantes avancées ces dernières années grâce à la disponibilité de grands volumes de données et du calcul intensif, sans oublier les avancées intéressantes en optimisation. Une caractéristique majeure de l'apprentissage profond (*deep learning*) est sa capacité à apprendre les descripteurs tout en effectuant le classement (*clustering*).

Il subsiste toutefois nombre de limites et défis que nous avons classés comme suit : i) sources de données ; ii) représentations symboliques vs représentations continues ; iii) apprentissage continu et sans fin ; iv) apprentissage sous contraintes ; v) architectures de calcul ; vi) apprentissage non-supervisé ; vii) processus d'apprentissage avec intervention humaine, explications.

Sources de données

Les défis dans ce domaine sont nombreux : apprendre à partir de données hétérogènes, disponibles sur une multiplicité de canaux ; gérer des informations incertaines ; identifier et traiter des événements rares au-delà des approches purement statistiques ; travailler en combinant des sources de connaissances et des sources de données ; intégrer des modèles et des ontologies dans le processus d'apprentissage ; et enfin obtenir de bonnes performances d'apprentissage avec peu de données, lorsque des sources de données massives ne sont pas disponibles.

ORPAILLEUR est une équipe-projet d'Inria et du Loria, créée début 2008. Il s'agit d'une équipe pluridisciplinaire qui comprend des scientifiques en informatique, mais également un biologiste, des chimistes et un physicien. Les sciences de la vie, la chimie et la médecine sont des domaines d'application de première importance et l'équipe développe des systèmes de travail à leur intention.

L'extraction de connaissances à partir de bases de données (*Knowledge Discovery in Databases* - KDD - en anglais), consiste à traiter un grand volume de données afin d'extraire des unités de connaissance significatives et réutilisables. Si on assimile les unités de connaissance à des pépites d'or et les bases de données à des terres ou des rivières à explorer, le processus KDD peut être comparé à l'orpaillage¹.

Le processus KDD est itératif, interactif et généralement contrôlé par un expert du domaine de données, appelé l'analyste. L'analyste sélectionne et interprète un sous-ensemble d'unités extraites pour obtenir des unités de connaissance présentant une certaine validité vis-à-vis des données analysées. À l'instar d'une personne cherchant de l'or et dotée d'une certaine expérience de la tâche et du lieu, l'analyste peut utiliser ses propres connaissances mais également les connaissances du domaine de données pour améliorer le processus KDD.

Le processus KDD exploite les connaissances du domaine notamment par une mise en rapport avec les ontologies relatives au domaine de données, pour faire un pas en direction de la notion d'extraction de connaissances guidée par les connaissances du domaine ou KDDK (pour *Knowledge Discovery guided by Domain Knowledge*). Dans le processus KDDK, les unités de connaissance extraites ont encore "une vie" après la phase d'interprétation : elles sont figurées par un formalisme de représentation des connaissances à intégrer dans une ontologie et réutilisées pour les besoins de résolution de problèmes. Ainsi, l'extraction des connaissances sert à étendre et à actualiser les ontologies existantes, en montrant que l'extraction des connaissances et leur représentation sont des tâches complémentaires permettant de donner corps à la notion de KDDK.

¹ Ceci explique le nom de l'équipe de recherche puisque le terme "orpailleur" désigne une personne qui cherche de l'or dans les rivières ou les montagnes.

Représentations symboliques vs représentations continues

Les représentations continues permettent à l'algorithme d'apprentissage automatique d'approcher des fonctions complexes, tandis que les représentations symboliques sont utilisées pour apprendre des règles et des modèles symboliques. Les avancées récentes les plus significatives concernent les représentations continues. Celles-ci laissent néanmoins de côté le raisonnement alors qu'il serait souhaitable de l'intégrer dans la représentation continue pour pouvoir faire des inférences sur les données numériques. Par ailleurs, afin d'exploiter la puissance de l'apprentissage profond, il peut s'avérer utile de définir des représentations continues de données symboliques, comme cela a été fait par exemple pour le texte avec word2vec et text2vec.

Les travaux de **LINKMEDIA** portent sur l'interprétation automatique de contenus multimédia professionnels et sociaux dans toutes leurs modalités. Dans ce contexte, l'intelligence artificielle s'appuie à la fois sur la conception de modèles de contenus et sur les algorithmes d'apprentissage associés pour extraire, décrire et interpréter des messages édités pour les humains. **LINKMEDIA** développe des algorithmes d'apprentissage automatique principalement basés sur des modèles statistiques et neuronaux pour extraire des structures, des connaissances, des entités ou des faits à partir de documents et collections multimédia.

La multimodalité et la modalité transversale en vue d'établir un lien entre les représentations symboliques (les mots ou des concepts dans un texte par exemple) et les observations continues (images continues ou attributs de signaux par exemple) sont les deux principaux défis de **LINKMEDIA**, pour lesquels les réseaux neuronaux sont prometteurs. La détection de canulars dans les réseaux sociaux, par une association du traitement de l'image et du traitement naturel du langage, la création d'hyperliens dans des collections de vidéos en exploitant simultanément le contenu parlé et visuel, et les analyses d'actualités interactives par des graphes de proximité sont parmi les principaux sujets sur lesquels travaille l'équipe.

Les analyses de type "processus d'apprentissage avec intervention humaine" pour lesquelles l'intelligence artificielle est au service d'un utilisateur, sont également au cœur du travail de l'équipe et posent des défis quant à l'interprétation du contenu multimédia basé sur des machines supervisées par l'Homme : l'Homme a besoin de comprendre les décisions prises par les machines et d'évaluer leur fiabilité, deux problématiques délicates étant données les approches actuelles guidées par les données. La connaissance et l'apprentissage automatique sont fortement intriqués dans ce scénario, d'où la nécessité de mécanismes permettant aux experts humains d'injecter des connaissances dans les algorithmes d'interprétation des données. Des utilisateurs malveillants manipuleront inévitablement les données pour influencer l'interprétation automatique en leur faveur, situation que l'apprentissage automatique actuel a du mal à gérer.

Enfin et surtout, plus que les mesures objectives sur les données, l'évaluation des algorithmes s'oriente vers des paradigmes de conception centrée utilisateur pour lesquels il n'existe de pas de fonction objectif à utiliser.

Apprentissage continu et sans fin

On attend de certains systèmes d'IA qu'ils soient résilients, c'est-à-dire capables de fonctionner 24/7 sans interruption. Des avancées intéressantes ont été réalisées dans les systèmes d'apprentissage "tout au long de la vie" qui engrangeront continuellement de nouvelles connaissances. La difficulté réside ici dans la capacité des systèmes d'IA à opérer en ligne en temps réel, tout en étant capables de remettre en question des croyances antérieures, et ce de façon autonome. L'auto-amorçage constitue une option pour ces systèmes car il permet d'utiliser les connaissances élémentaires acquises en début d'exploitation pour orienter les futures tâches d'apprentissage, comme dans le système NELL (*Never-Ending Language Learning*) développé à l'Université Carnegie-Mellon (<http://rtw.ml.cmu.edu/rtw/>).

Apprentissage sous contraintes

La protection de la vie privée est sans doute la contrainte la plus importante à prendre en compte. Les chercheurs spécialisés dans l'apprentissage automatique ont reconnu récemment la nécessité de protéger la vie privée tout en poursuivant l'apprentissage à partir de données personnelles. Une théorie de l'apprentissage automatique respectueux de la vie privée est actuellement développée par des chercheurs tels que Michael Jordan (<http://arxiv.org/abs/1210.2085>). Plusieurs équipes-projets d'Inria travaillent sur la protection de la vie privée : en particulier **ORPAILLEUR** dans le domaine de l'apprentissage automatique, mais également des équipes-projets d'autres domaines telles que **PRIVATICS** (sur les algorithmes de protection de la vie privée) et **SMIS** (sur la protection de la vie privée dans les bases de données). De façon plus générale, l'apprentissage automatique peut être amené à prendre en compte d'autres contraintes externes telles que des données décentralisées ou des limites énergétiques. Des recherches sur la problématique générale de l'apprentissage automatique sous contraintes externes sont donc nécessaires.

MAGNET

L'apprentissage automatique, dans une de ses dimensions, a pour objet d'identifier des régularités, des motifs, dans les données et de les représenter dans un modèle. La notion de modèle est variée, statistique, probabiliste ou formelle comme une machine à états finis, etc. L'objectif de l'apprentissage est de pouvoir utiliser ensuite ce modèle pour prédire des valeurs associées à de nouvelles données. L'approche se distingue

en cela de la tentative de raisonner, de représenter et traiter des connaissances, autres branches classiques de l'IA. Dans une formalisation mathématique traditionnelle, toutefois, l'objectif de l'apprentissage automatique est de réussir à bien approximer une fonction inconnue qui prend en entrée des données, représentant la description d'un objet, une entité, et produit un résultat comme une valeur discrète ou une valeur réelle. C'est par exemple attribuer une valeur positive négative ou neutre à la représentation d'un texte court dans le but d'analyser les sentiments dans les tweets ou estimer l'intérêt que peut porter une personne à un nouveau film de cinéma. Retrouver cette fonction, ces motifs, nécessite de disposer de suffisamment de données pour les identifier. Plus le modèle se complexifie, par exemple repose sur des relations de dépendance entre les données comme l'influence que peuvent avoir des amis sur des goûts cinéphiles, plus la tâche devient difficile. C'est aussi le cas lorsque la fonction à apprendre ne calcule plus une valeur scalaire mais une structure à part entière comme un arbre d'analyse d'une phrase par exemple. L'équipe **MAGNET** s'intéresse particulièrement à ce défi de prendre en compte les relations structurelles dans les données, que ce soit en entrée quand l'ensemble des données est représenté sous forme de graphe, ou lorsqu'il faut prédire une structure complexe.

Ce qui n'était pas possible il y a vingt ans le devient d'un point de vue pratique aujourd'hui. L'apprentissage automatique est ancien, mais redevenu populaire à l'ère de l'explosion à la fois des capacités de calcul et de la collecte massive de données. On observe alors des progrès spectaculaires. Toutefois, ce qui est devenu possible du point de vue statistique, l'abondance de données fournissant des statistiques suffisantes pour estimer des fonctions plus complexes, soulève des questions informatiques de passage à l'échelle, d'efficacité des algorithmes. Ces questions fondamentales sont au cœur des problématiques de **MAGNET**. Une représentation sous forme de graphe, qu'elle soit donnée ou calculée de façon adaptative face à une tâche donnée, permet d'explicitier et d'exploiter une approximation de ces dépendances. La question de ces dépendances est également très pertinente dans le cas du traitement de la langue. Par exemple lors de la résolution du problème de coréférence, lorsqu'on cherche à identifier si deux parties de texte font référence à une même entité, une même personne.

Les défis à venir semblent posés par les systèmes d'information modernes imposant de nouvelles contraintes : les données massives sont réparties, les systèmes sont décentralisés, en compétition ou coopératifs. Un enjeu majeur réside dans la capacité d'exploiter ces masses

de données non plus sous l'angle de la centralisation, mais comme un ensemble de tâches d'apprentissage autonomes et personnalisées sur des données personnelles. Le caractère personnel et privatif des données ajoute une contrainte supplémentaire limitant la communication. Des contraintes énergétiques apparaissent évidemment. Apprendre efficacement sous de telles contraintes au sein d'un réseau de données et de machines d'apprentissage semble un défi fondamental et alternatif à l'hypercentralisation observée aujourd'hui.

Architectures de calcul

Les systèmes d'apprentissage automatique modernes nécessitent du calcul intensif et un stockage de données efficace afin de s'adapter à la taille des données et aux dimensions des problèmes. Des algorithmes seront exécutés sur des GPU et d'autres architectures puissantes ; données et processus doivent être distribués sur plusieurs processeurs. De nouvelles recherches doivent s'attacher à améliorer les algorithmes d'apprentissage automatique et les formulations des problèmes afin de tirer le meilleur parti de ces architectures de calcul.

SIERRA s'intéresse avant tout aux problèmes d'apprentissage automatique, le principal objectif étant de faire le lien entre la théorie et les algorithmes, ainsi qu'entre les algorithmes et les applications à fort impact dans différents domaines techniques et scientifiques, en particulier la vision artificielle, la bio-informatique, le traitement du signal audio, le traitement de texte et la neuro-imagerie. Les dernières réalisations de cette équipe de recherche comprennent des travaux théoriques et algorithmiques pour une optimisation convexe à grande échelle, débouchant sur des algorithmes qui effectuent quelques passages sur les données tout en assurant une performance prédictive optimale dans un grand nombre de situations d'apprentissage supervisé.

Les défis futurs comprennent le développement de nouvelles méthodes d'apprentissage non-supervisé et l'élaboration d'algorithmes d'apprentissage pour les architectures informatiques parallèles et distribuées.

Apprentissage non-supervisé

Les résultats les plus remarquables obtenus dans le domaine de l'apprentissage automatique sont basés sur l'apprentissage supervisé,

c'est-à-dire l'apprentissage à partir d'exemples dans lesquels le résultat attendu est fourni avec les données d'entrée. Cela implique d'étiqueter les données avec les résultats attendus correspondants, un processus qui nécessite des données à grande échelle. Le Turc mécanique d'Amazon (www.mturk.com) est un parfait exemple de la manière dont les grandes entreprises mobilisent des ressources humaines pour annoter des données. Mais la vaste majorité des données existe sans résultat attendu, c'est-à-dire sans annotation désirée ou nom de classe. Il convient donc de développer des algorithmes d'apprentissage non-supervisé afin de gérer cette énorme quantité de données non étiquetées. Dans certains cas, un apport minime de supervision humaine peut être utilisé pour guider l'algorithme non-supervisé.

SEQUEL

Cette équipe-projet s'appelle **SEQUEL** pour *sequential learning* (apprentissage séquentiel). En effet, **SEQUEL** se concentre sur **l'apprentissage dans les systèmes artificiels** (matériels ou logiciels) qui accumulent des informations au fil du temps. Ces systèmes sont désignés ci-après comme agents apprenants ou machines apprenantes. Les données recueillies peuvent être utilisées pour faire une estimation de certains paramètres d'un modèle qui, à son tour, peut être utilisé pour sélectionner des actions en vue d'exécuter certaines tâches d'optimisation à long terme.

Ces données peuvent être obtenues par un agent qui observe son environnement : elles représentent ainsi une perception. C'est notamment le cas lorsque l'agent prend des décisions - en vue d'atteindre un certain objectif - qui affectent l'environnement, et par conséquent, le processus d'observation lui-même.

Dans **SEQUEL**, le terme "séquentiel" se réfère à deux aspects :

- l'acquisition séquentielle de données, à partir de laquelle un modèle est appris (apprentissage supervisé et non-supervisé) ;
- la tâche de prise de décision séquentielle, fondée sur le modèle appris (apprentissage par renforcement).

Les problèmes relatifs à l'apprentissage séquentiel incluent notamment l'apprentissage supervisé, l'apprentissage non-supervisé et l'apprentissage par renforcement. Dans tous ces cas, il est posé comme hypothèse que le processus peut être considéré comme stationnaire, au moins pendant un certain temps, puisqu'il évolue lentement.

Il est souhaitable d'avoir des algorithmes "*anytime*", c'est-à-dire qu'à tout moment, une prédiction peut être requise ou une action peut être sé-

lectionnée en tirant pleinement parti, et si possible le meilleur parti, de l'expérience déjà acquise par l'agent apprenant.

La perception de l'environnement par l'agent (à l'aide de ses capteurs) n'est généralement pas la meilleure perception pour faire une prédiction ou pour prendre une décision (processus de décision markovien partiellement observable - POMDP). Par conséquent, la perception doit être projetée sur un espace d'état (ou entrée) plus efficace et pertinent.

Enfin, une problématique importante relative à la prédiction concerne son évaluation : jusqu'à quel point peut-on se tromper lorsque l'on fait une prédiction ? Pour pouvoir contrôler les systèmes réels, il est essentiel de répondre à cette question.

En résumé, les principales problématiques traitées par l'équipe **SEQUEL** concernent :

- l'apprentissage de modèles, et plus spécifiquement de modèles qui projettent un espace d'entrée RP sur R ;
- la projection de l'observation sur l'espace d'états ;
- le choix de l'action à exécuter (dans le cas du problème de prise de décision séquentielle) ;
- les garanties d'exécution ;
- la mise en place d'algorithmes utilisables ;
- le tout compris dans un cadre séquentiel.

Les principales applications concernent les systèmes de recommandation, mais également la prédiction, l'apprentissage supervisé et la classification non-supervisée de données (*clustering*). De nombreuses méthodes sont utilisées, depuis les méthodes à noyaux jusqu'à l'apprentissage profond. Les méthodes non-paramétriques sont souvent privilégiées dans les travaux de **SEQUEL**.

Processus d'apprentissage avec intervention humaine, explications

Les défis portent sur la mise en place d'une collaboration naturelle entre les algorithmes d'apprentissage automatique et les utilisateurs, afin d'améliorer le processus d'apprentissage. Pour ce faire, les systèmes d'apprentissage automatique doivent être en mesure de montrer leur état sous une forme compréhensible pour l'Homme. De plus, l'utilisateur humain doit pouvoir obtenir des explications de la part du système sur n'importe quel résultat obtenu. Ces explications seraient fournies au cours de l'apprentissage et pourraient être liées à des données d'entrée ou à des représentations intermédiaires. Elles pourraient également indiquer des niveaux de confiance, selon le cas.

LACODAM

Les techniques relatives à la science des données ont déjà démontré leur immense potentiel. Toutefois, elles exigent de posséder une solide expertise non seulement sur les données d'intérêt, mais également sur les techniques d'extraction de connaissances à partir de ces données. Seuls quelques experts hautement qualifiés possèdent une telle expertise. Une méthode de recherche prometteuse consiste donc à automatiser (en partie) le processus lié à la science des données et à le rendre accessible à un plus grand nombre d'utilisateurs.

Un des travaux novateurs dans ce domaine a été le projet "**Data Science Machine**" du MIT¹, qui permet d'automatiser le processus délicat du *feature engineering*², afin d'identifier les caractéristiques les plus pertinentes parmi des milliers de caractéristiques possibles et d'exploiter celles-ci en vue d'exécuter des tâches d'apprentissage automatique. Une autre approche intéressante est celle du projet "*Automatic Statistician*" de l'université de Cambridge³, qui découvre des régressions complexes par exploration d'un espace de recherche de combinaisons de noyaux de régression simples et produit des rapports en langage naturel.

Ces deux approches sont conçues pour des tâches supervisées : elles ont accès à des mesures d'évaluation leur permettant de mesurer automatiquement leur performance et de l'améliorer. La nouvelle équipe **LACODAM** s'intéresse à l'automatisation pour mettre en place une approche plus exploratoire, dont le but est de découvrir des connaissances "intéressantes" à partir de données, avec des informations préalables définissant ce qui est intéressant. Premièrement, cela nécessite l'utilisation de méthodes différentes de celles mentionnées ci-dessus, comme l'extraction de motifs ou le regroupement non-supervisé de données (*clustering*). Deuxièmement, cela exige d'interagir avec l'utilisateur, car il est le seul à savoir ce qui l'intéresse. Un problème toujours présent est de trouver la meilleure méthode pour permettre au système et à l'utilisateur d'explorer conjointement les données. La solution exige une collaboration entre l'exploration de données et l'apprentissage automatique ainsi que l'intelligence artificielle ou encore les interfaces homme-machine et la visualisation.

1 James Max Kanter, Kalyan Veeramachaneni: Deep feature synthesis: Towards automating data science endeavors. DSAA 2015: 1-10

2 Feature engineering : construction de nouvelles variables à partir de celles dont on dispose déjà.

3 David K. Duvenaud, James Robert Lloyd, Roger B. Grosse, Joshua B. Tenenbaum, Zoubin Ghahramani: Structure Discovery in Nonparametric Regression through Compositional Kernel Search. ICML (3) 2013: 1166-1174

Apprentissage par transfert

L'apprentissage par transfert est utile lorsque peu de données sont disponibles pour l'apprentissage d'une tâche. Il consiste à utiliser pour une nouvelle tâche les connaissances ayant été acquises à partir d'une autre tâche et pour laquelle un plus grand nombre de données est disponible. Il s'agit d'une idée plutôt ancienne (1993), dont les résultats restent modestes car elle est difficile à mettre en œuvre. En effet, elle implique de pouvoir extraire les connaissances que le système a acquises en premier lieu, mais il n'existe aucune solution générale à ce problème (de quelle manière, comment les réutiliser ? ...).

Une autre approche de l'apprentissage par transfert est le "*shaping*". Il s'agit d'apprendre une tâche simple, puis de la complexifier progressivement, jusqu'à atteindre la tâche cible. Il existe quelques exemples de cette procédure dans la littérature, mais aucune théorie générale.

Apprentissage automatique et aide à la décision : retours d'informations, modèles de causalité, représentations

Il existe une "exubérance irrationnelle" sur les mégadonnées (ou Big data). Cette exubérance et le niveau élevé des attentes de certains utilisateurs (entrepreneurs, journalistes, politiques) pourraient avoir un effet contre-productif et exposer la communauté scientifique à un hiver des mégadonnées, semblable à l'hiver de l'IA (longue période d'espoirs déçus à partir des années 1980, dont l'IA n'est sortie qu'au début de ce siècle).

TAO

I. Programmes informatiques sous-spécifiés : la programmation par *feedback*

Les algorithmes comprennent des fonctionnalités de plus en plus complexes, et l'utilisateur est de moins en moins disposé à lire le manuel d'utilisation. Pour réaliser des logiciels (adaptés à un environnement ouvert et à des utilisateurs dont les préférences évoluent), la programmation par *feedback* investiguée par TAO s'appuie sur l'interaction avec l'utilisateur : itérativement, l'algorithme propose un nouveau comportement et l'utilisateur indique si ce comportement représente ou non une amélioration (par rapport au meilleur comportement antérieur).

Cette approche repose sur deux piliers fondamentaux :

- l'apprentissage automatique, pour modéliser les préférences de l'utili-

sateur en fonction de ses retours ;

- l'optimisation, pour choisir les comportements proposés les plus informatifs et minimiser le temps requis pour atteindre un comportement satisfaisant.

L'une des questions ouvertes concerne la communication avec l'utilisateur, permettant à l'humain dans la boucle de donner du sens au processus d'apprentissage et d'interaction. Les approches envisagées se fondent sur la visualisation d'une part, et sur la définition de tâches ou d'environnements graduellement plus complexes d'autre part.

II. Décision optimale et apprentissage de modèles causaux

Parmi les objectifs de l'exploitation des mégadonnées figure la décision optimale. Or, les modèles permettant de prendre des décisions appropriées sont nécessairement des modèles causaux (par opposition aux modèles prédictifs, qui peuvent s'appuyer sur des corrélations). Par exemple, les notes des élèves sont corrélées à la présence plus ou moins importante de livres à leur domicile. Cependant, envoyer des livres dans chaque foyer n'améliorera **pas** les notes des élèves, toutes choses égales par ailleurs.

La recherche de modèles causaux conduit à revisiter les critères de l'apprentissage statistique. Une approche prometteuse consiste à apprendre un indice de causalité, en partant de l'ensemble des problèmes pour lesquels la causalité est connue (par exemple, l'altitude d'une ville "cause" sa température moyenne).

III. Changements de représentation / apprentissage profond

Les avancées significatives apportées par l'apprentissage de réseaux neuronaux profonds (par exemple dans les domaines de la vision par ordinateur ou du traitement automatique du langage naturel) sont liées à leur capacité à apprendre automatiquement une représentation efficace de domaines complexes à partir d'une grande quantité de données (ou d'interactions intensives avec l'environnement dans le contexte de l'apprentissage par renforcement).

Parmi les perspectives de recherche figurent :

- la recherche d'approches d'optimisation non convexe efficaces d'un point de vue mathématique (e.g. en tenant compte de la géométrie de l'espace) et informatique (e.g. du point de vue de la complexité algorithmique en temps et en mémoire, tirant parti des architectures spécialisées) ;

- l'extension de ces méthodes au cas de petites données, lorsqu'on dispose de connaissances *a priori* sur le domaine d'application (permettant par exemple l'augmentation des données par application d'opérateurs d'invariance) ;
- l'exploitation de sources de données complémentaires (par exemple concernant des données réelles et des données simulées) pour construire des représentations à la fois précises et robustes.

4.3 Les défis génériques dans l'analyse des signaux : vision, parole

L'analyse des signaux, en particulier la vision et la parole, a bénéficié des récents progrès de l'IA. Les systèmes d'apprentissage profond ont gagné de nombreuses compétitions dans le domaine de la reconnaissance des formes et de la reconnaissance visuelle. Ainsi, le système de vision MobilEye renforce les capacités de conduite autonome des voitures Tesla, tandis que des assistants vocaux tels que Siri, Cortana ou Amazon Echo, sont utilisés chaque jour par des millions de personnes.

Les différents défis relatifs à l'analyse des signaux visuels sont les suivants : (i) le passage à l'échelle ; (ii) la transition entre images fixes et vidéo ; (iii) la multimodalité ; (iv) l'introduction de connaissances *a priori*.

THOTH

La quantité d'images et de vidéos numériques disponibles en ligne ne cesse d'augmenter de manière exponentielle : les particuliers

publient leurs films sur YouTube et leurs images sur Flickr, les journalistes et les scientifiques créent des pages web pour diffuser des informations et des résultats de recherche, et les archives audiovisuelles des émissions de télévision sont accessibles au public. En 2018, il est prévu que près de 80 % du trafic sur Internet seront générés par les vidéos et qu'une personne mettrait plus de cinq millions d'années pour visionner la quantité de vidéos qui sera alors publiée chaque mois sur les réseaux IP mondiaux. Par conséquent, il est de plus en plus nécessaire, d'annoter et d'indexer ce contenu visuel pour les utilisateurs particuliers et professionnels. Les métadonnées textuelles et auditives disponibles ne sont généralement pas suffisantes pour répondre à la plupart des requêtes et les données visuelles doivent alors entrer en jeu. Par ailleurs, il n'est pas concevable de définir les modèles des contenus visuels nécessaires

pour répondre à ces requêtes en annotant manuellement et rigoureusement chaque catégorie de concept, d'objet, de scène ou d'action pertinente dans un échantillon représentatif de situations quotidiennes, ne serait-ce que parce qu'il serait difficile, voire impossible, de déterminer *a priori* les catégories pertinentes et le niveau de granularité approprié.

L'idée principale de la proposition de l'équipe-projet **THOTH** est de **développer un nouveau cadre pour apprendre automatiquement la structure et les paramètres de modèles visuels**, en explorant activement de grandes sources d'images et de vidéos numériques (des archives hors ligne, mais également la quantité croissante de contenus disponibles en ligne, avec des centaines de milliers d'images), tout en exploitant le faible signal de supervision émis par les métadonnées qui les accompagnent. Cet énorme volume de données d'apprentissage visuelles permettra à l'apprentissage de modèles complexes non-linéaires avec un grand nombre de paramètres, tels que les réseaux convolutifs profonds et les modèles graphiques d'ordre supérieur.

L'objectif principal de l'équipe **THOTH** est d'explorer automatiquement de grandes collections de données, de sélectionner les informations pertinentes et d'apprendre la structure et les paramètres de modèles visuels. **Il existe trois défis majeurs** : (1) la conception et l'apprentissage de modèles structurés capables de représenter des informations visuelles complexes ; (2) l'apprentissage en ligne de modèles visuels à partir d'annotations textuelles, de sons, d'images et de vidéos ; et (3) l'apprentissage et l'optimisation à grande échelle. Un autre point important est (4) la collecte et l'évaluation des données.

Passage à l'échelle

Les systèmes de vision modernes doivent être capables de traiter des données volumineuses et à haute fréquence. C'est par exemple le cas pour les systèmes de surveillance dans les lieux publics, les robots évoluant dans des environnements inconnus et les moteurs de recherche d'images sur le web qui doivent tous traiter d'énormes quantités de données. Les systèmes de vision doivent non seulement traiter ces données très rapidement mais également le faire en garantissant des niveaux de précision élevés pour éviter toute nécessité de vérification des résultats ou tout post-traitement de la part d'opérateurs humains. Même les taux de précision atteignant 99,9 % pour la classification d'images pour des systèmes critiques ne sont pas suffisants lors du traitement de millions

d'images, puisque les 0,1 % restants nécessiteront des heures de traitement par l'Homme.

Transition entre images fixes et vidéo

La reconnaissance d'objets dans des images reste un véritable défi au sens général du terme, même si des résultats très prometteurs sont disponibles. La reconnaissance d'actions dans des vidéos - et plus généralement la compréhension de scènes - est un domaine en pleine expansion qui a déclenché de nouvelles recherches, comme celles de l'équipe-projet **STARS** présentée ci-après.

STARS

Au cours de ces dernières années, de nombreuses études avancées ont été réalisées dans le domaine de l'analyse des signaux, et en particulier dans celui de la compréhension de scènes. La compréhension de scènes est un processus, souvent réalisé en temps réel, de perception, d'analyse et d'élaboration d'une interprétation d'une scène dynamique en 3D observée par l'intermédiaire d'un réseau de capteurs (caméras vidéo, par exemple). Ce processus consiste essentiellement à associer les informations des signaux issues des capteurs observant la scène avec les modèles développés pour comprendre la scène. Par conséquent, la compréhension de scènes ajoute et extrait de la sémantique à partir des données des capteurs qui caractérisent une scène. Cette scène peut contenir un certain nombre d'objets physiques différents (personnes, véhicules) qui interagissent entre eux ou avec un environnement (équipements, par exemple) plus ou moins structuré. La scène peut durer quelques instants (chute d'une personne) ou quelques mois (dépression d'une personne) et peut être limitée à une lame de verre de laboratoire observée à travers un microscope ou dépasser la taille d'une ville. Les capteurs sont généralement des caméras (omnidirectionnelles, infrarouges, de profondeur), mais il peut également s'agir de microphones ou d'autres types de capteurs (tels que les cellules optiques, les capteurs de contact, les capteurs physiologiques, les accéléromètres, les radars, les détecteurs de fumée, les smartphones, etc.).

La compréhension de scènes s'inspire de la vision cognitive et nécessite l'association d'au moins trois domaines : **la vision par ordinateur, les sciences cognitives et le génie logiciel.**



Figure 5 : Reconnaissance d'activité dans un hôpital - © Inria / Photo H. Raguet

La compréhension de scènes peut atteindre cinq niveaux de fonctionnalités génériques de vision par ordinateur : la détection, la localisation, le suivi, la reconnaissance et la compréhension. Toutefois, les systèmes de compréhension de scènes vont au-delà de la détection de caractéristiques visuelles, telles que les angles, les bords ou les zones en mouvement, pour extraire les informations liées au monde physique qui sont importantes pour les opérateurs humains. La compréhension de scènes exige également des fonctionnalités de vision par ordinateur plus robustes, plus solides et plus flexibles, en leur donnant une compétence cognitive : la capacité d'apprendre, de s'adapter, d'évaluer les solutions alternatives et de développer de nouvelles stratégies pour l'analyse et l'interprétation.

Concernant la compréhension de scènes, l'équipe **STARS** a développé des systèmes automatisés originaux afin de comprendre les comportements humains dans un grand nombre d'environnements et pour différentes applications :

- dans les stations de métro, dans les rues et à bord des trains : violences, bagages abandonnés, graffitis, fraudes, comportement de la foule ;
- dans les zones aéroportuaires : arrivée des avions, ravitaillement des avions, chargement/déchargement des bagages, guidage au sol des avions ;
- dans les banques : braquages, contrôles d'accès dans les bâtiments,

utilisation de distributeurs automatiques de billets ;

- applications de soins à domicile pour le suivi des activités des personnes âgées : cuisine, repos, préparation d'un café, visionnage de programmes télévisés, préparation de la boîte à pilules, chute de personnes ;
- maison intelligente, suivi du comportement au bureau pour l'intelligence ambiante : lecture, consommation de boissons ;
- surveillance dans les supermarchés: arrêts devant un produit, files d'attente, retraits d'articles ;
- applications biologiques : étude du comportement des guêpes.

Pour élaborer ces systèmes, l'équipe **STARS** a conçu de nouvelles technologies pour la reconnaissance des activités humaines notamment à l'aide de caméras vidéo 2D et 3D. Plus précisément, les chercheurs ont combiné **trois catégories d'algorithmes pour reconnaître les activités humaines** :

- des moteurs de reconnaissance qui utilisent des ontologies exprimées par des règles représentant les connaissances de l'expert. Ces moteurs de reconnaissance d'activités peuvent être facilement étendus et permettent d'intégrer ultérieurement des informations supplémentaires issues des capteurs lorsqu'elles sont disponibles [König 2015] ;
- des méthodes d'apprentissage supervisé qui s'appuient sur des échantillons positifs/négatifs représentatifs des activités cibles spécifiées par les utilisateurs. Ces méthodes s'appuient généralement sur le modèle de "sac de mots" (*bag-of-words*) qui calcule une grande variété de descripteurs spatio-temporels [Bilinski 2015].
- des méthodes d'apprentissage non-supervisé (entièrement automatisées) qui découvrent des classes d'activités dans des grandes bases de données [Negin 2015].

Bibliographie :

- A. König, C. Crispim, A. Covella, F. Bremond, A. Derreumaux, G. Bensadoun, R. David, F. Verhey, P. Aalten et P.H. Robert. Ecological Assessment of Autonomy in Instrumental Activities of Daily Living in Dementia Patients by the means of an Automatic Video Monitoring System, *Frontiers in Aging Neuroscience* - publication en libre accès - <http://dx.doi.org/10.3389/fnagi.2015.00098>, 02 June 2015.

- P. Bilinski et F. Bremond. Video Covariance Matrix Logarithm for Human Action Recognition in Videos. *The International Joint Conference on Artificial Intelligence, IJCAI 2015*, Buenos Aires, Argentine du 25 au 31 juillet 2015.

- F. Negin, S. Cosar, M. Koperski, et F. Bremond. Generating Unsupervised Models for Online Long-Term Daily Living Activity Recognition. *The 3rd IAPR Asian Conference on Pattern Recognition, ACPR2015*, Kuala Lumpur, Malaisie, 4-6 novembre 2015.

Multimodalité

La compréhension des données visuelles peut être améliorée de différentes manières : sur le web, les métadonnées qui accompagnent les images et les vidéos peuvent être utilisées pour éliminer certaines hypothèses et guider le système vers la reconnaissance d'objets, d'événements et de situations spécifiques. Une autre option consiste à utiliser la multimodalité, c'est-à-dire des signaux émanant de différents canaux (infrarouges, lasers, signaux magnétiques, etc.). Il est également souhaitable d'utiliser une combinaison de signaux sonores en complément de la vision (images ou vidéos) le cas échéant. Par exemple, un défi a été proposé par l'Agence Nationale de la Recherche et la DGA en 2010 (<http://www.defi-repere.fr/>) : l'identification multimodale de personnes dans des émissions d'information, en utilisant l'image dans laquelle les personnes sont visibles, les textes en incrustation dans lesquels le nom des personnes apparaît, la bande-son dans laquelle la voix des locuteurs est reconnaissable et le contenu du signal de parole dans lequel le nom des personnes est prononcé.

Le programme de recherche de l'équipe-projet **PERCEPTION** repose sur l'étude et la mise en place de modèles informatiques, en vue d'établir un lien entre images et sons d'une part et significations et actions d'autre part. Les membres de l'équipe-projet **PERCEPTION** relèvent ce défi grâce à une approche interdisciplinaire allant de la vision par ordinateur au traitement de signaux sonores, en passant par l'analyse de scènes, l'apprentissage automatique et la robotique. Plus précisément, **PERCEPTION** développe des méthodes pour la représentation et la reconnaissance d'objets et d'événements visuels et sonores, la fusion audiovisuelle, la reconnaissance des actions, des gestes et de la parole de l'Homme, l'audition spatiale et l'interaction Homme-robot.

Les sujets de recherche de cette équipe-projet sont les suivants :

- **la vision par ordinateur** : représentation spatio-temporelle d'informations visuelles en 2D et 3D, reconnaissance des gestes et des actions, capteurs 3D, vision binoculaire, systèmes avec plusieurs caméras ;
- **l'analyse de scènes sonores** : audition binaurale, localisation et séparation de sources sonores, communication verbale, classification d'événements sonores ;
- **l'apprentissage automatique** : modèles de mélange, réduction de dimension linéaire et non-linéaire, apprentissage de variétés, modèles graphiques ;

- **la robotique** : vision robotique, audition robotique, interaction Homme-robot, fusion de données.

L'équipe-projet **WILLOW** s'intéresse à des problèmes fondamentaux de la vision par ordinateur, notamment la perception en trois dimensions, la photographie numérique ou encore la compréhension d'images et de vidéos.

Elle étudie de nouveaux modèles de contenus d'images (quels sont les critères d'un vocabulaire visuel adapté ?) et de processus d'interprétation (qu'est-ce qu'une architecture de reconnaissance efficace ?).

Ses dernières réalisations incluent notamment un travail théorique sur les fondements géométriques de la vision par ordinateur, de nouvelles avancées en matière de restauration d'images (retrait d'effet flou, réduction du bruit, suréchantillonnage), ou encore des méthodes faiblement supervisées pour la localisation temporelle d'actions dans des vidéos. Les membres de l'équipe **WILLOW** travaillent en étroite collaboration avec les équipes **SIERRA** et **THOTH** d'Inria, ainsi qu'avec des chercheurs de l'université Carnegie-Mellon, de l'université de Californie à Berkeley ou encore du laboratoire d'IA de Facebook, collaborations qui mettent en avant la forte synergie qui unit l'apprentissage automatique et la vision par ordinateur, avec de nouvelles opportunités qui émergent également dans des domaines comme l'archéologie ou la robotique.

Les défis futurs incluent notamment le développement de modèles très faiblement supervisés de reconnaissance visuelle dans des ensembles de données d'images et de vidéos à grande échelle, à l'aide de métadonnées disponibles sous forme de texte ou de parole, par exemple.

Introduction de connaissances *a priori*

Une autre option visant à améliorer les applications de vision consiste à introduire des connaissances *a priori* dans le moteur de reconnaissance. L'exemple ci-dessous, issu de l'équipe-projet **ASCLEPIOS**, consiste à ajouter des informations sur l'anatomie et la pathologie d'un patient pour garantir de meilleures analyses des images biomédicales. Dans d'autres domaines, il est possible d'utiliser des informations contextuelles, des informations sur une situation ou une tâche, ou encore des données de localisation, pour lever toute ambiguïté sur les inter-

prétations possibles. Toutefois, la question de savoir comment fournir ces connaissances *a priori* n'est pas résolue de manière générale : des méthodes et des représentations de connaissances spécifiques doivent être établies pour traiter une application cible en vision artificielle.

L'équipe-projet **ASCLEPIOS** a trois axes de recherches principaux :

- **l'analyse d'images biomédicales** à partir de modèles géométriques, statistiques, physiques et fonctionnels avancés ;
- **la simulation de systèmes physiologiques** à partir de modèles informatiques s'appuyant sur des images biomédicales et d'autres signaux ;
- **l'application de ces outils à la médecine et à la biologie** pour assister la prévention, le diagnostic et la thérapie.



Figure 6 : Modélisation biophysique à partir d'analyse d'images médicales

© Inria / Photo H. Raguet

Le point sur l'analyse d'images médicales : la qualité des images biomédicales tend à s'améliorer de façon constante (meilleure résolution spatiale et temporelle, meilleur rapport signal/bruit). Non seulement les images sont multidimensionnelles (trois coordonnées spatiales et éventuellement une dimension temporelle), mais les protocoles médicaux ont aussi tendance à inclure des images multiséquences (ou multiparamétriques) et multimodales pour chaque patient. Malgré les efforts et les avancées considérables réalisés au cours de ces vingt dernières années, les problèmes les plus importants relatifs à la segmentation et

au recalage d'image n'ont pas été résolus de manière générale.

Le premier objectif de cette équipe-projet, à court terme, est de travailler sur des versions spécifiques de ces problèmes, en prenant en compte autant d'informations *a priori* que possible sur l'anatomie et la pathologie du patient. Leur second objectif vise à inclure un plus grand nombre de connaissances sur les propriétés physiques de l'acquisition d'images et les tissus observés ainsi que sur les processus biologiques mis en œuvre. Les autres thèmes de recherche de l'équipe-projet **ASCLEPIOS** sont l'analyse d'images biologiques, l'anatomie algorithmique et la physiologie algorithmique.

Les défis de l'analyse des signaux vocaux et sonores partagent beaucoup de points communs avec la liste précédente : le passage à l'échelle, la multimodalité et l'introduction de connaissances *a priori* sont également pertinents pour les applications audio. Les applications cibles sont l'identification du locuteur, la compréhension de la parole, le dialogue (y compris pour les robots), ainsi que la traduction automatique en temps réel. Dans le cas de signaux audio, il est également nécessaire de développer ou d'avoir accès à des données volumineuses pour l'apprentissage automatique. L'apprentissage incrémental en ligne pourrait être nécessaire pour le traitement de la parole en temps réel.

MULTISPEECH est une équipe de recherche commune de l'Université de Lorraine, d'Inria et du CNRS. Elle fait partie du département D4 "Traitement automatique des langues et des connaissances" du Loria.

Outre la vision, l'audition est un sens porteur de précieuses informations pour les interactions Homme-machine (compréhension de la parole, reconnaissance du locuteur et des émotions) et la perception de l'environnement (informations sonores sur des activités, un danger, etc.). La technologie d'apprentissage profond fait aujourd'hui partie de l'état de l'art et a entraîné ces cinq dernières années la diffusion d'interfaces vocales avec microphone rapproché. Les grands défis qui subsistent sont notamment l'utilisation de microphones distants affectés par la réverbération et le bruit acoustique, l'exploitation de grandes quantités de données vocales et audio non étiquetées, l'amélioration de la robustesse face à la variabilité des locuteurs et des événements sonores, l'exploitation d'informations contextuelles et l'intégration robuste d'autres techniques de détection.

L'équipe **MULTISPEECH** s'emploie actuellement à relever ces défis en développant des architectures d'apprentissage profond adaptées à différents sous-problèmes, en les regroupant dans des systèmes plus larges et en essayant d'évaluer et de transposer le niveau de confiance d'un système à un autre. Outre les interfaces vocales, les applications de ces technologies incluent notamment les dispositifs de surveillance pour l'assistance à domicile ou les villes intelligentes.

4.4 Les défis génériques dans les connaissances et le web sémantique

D'après la définition initiale de Tim Berners-Lee, *"le web sémantique est une extension du web actuel dans laquelle l'information se voit associée à un sens bien défini, améliorant ainsi la capacité des ordinateurs et des humains à travailler en coopération."* La *"semantic tower"* s'appuie sur des URI et sur le format XML, à travers des schémas RDF représentant des triplets de données, jusqu'à des ontologies permettant un raisonnement et un traitement logique. Les équipes d'Inria impliquées dans la représentation, le raisonnement et le traitement des connaissances relèvent les défis du web sémantique de différentes manières : (i) en traitant de grands volumes d'informations issues de sources réparties hétérogènes ; (ii) en établissant des liens entre les données massives stockées dans des bases de données à l'aide de technologies sémantiques ; (iii) en développant des applications s'appuyant sur la sémantique et maîtrisant ces technologies.

Traitement de grands volumes d'informations issus de sources distribuées hétérogènes

Avec l'omniprésence d'Internet, nous sommes désormais face à l'opportunité et au défi de passer de systèmes intelligents artificiels locaux à des sociétés et des intelligences artificielles largement réparties. Concevoir et maintenir en fonctionnement des systèmes fiables et efficaces regroupant des données liées depuis des sources distantes, au travers de flux de traitement et de services répartis, est toujours un problème. Les aspects suivants doivent être traités à grande échelle et de manière continue : la qualité des données et la traçabilité des processus qui les manipulent, le niveau de précision de leur extraction et de leur capture, l'exactitude de leur alignement et de leur intégration ou encore la disponibilité et la qualité des modèles partagés (ontologies, vocabu-

lares) nécessaires pour représenter, échanger ou encore raisonner sur les données.

Le sujet principal de l'équipe-projet **EXMO** est l'hétérogénéité des connaissances et leur représentation.

L'opportunité offerte par le web de partager des connaissances à grande échelle a fait de cette préoccupation centrale un point critique. L'équipe-projet **EXMO** a effectué des travaux importants sur l'alignement d'ontologies dans le contexte du web sémantique. L'alignement d'ontologies consiste à trouver des éléments liés dans des ontologies hétérogènes et à les exprimer comme des alignements.

L'équipe-projet **EXMO** a conçu des langages d'alignement avec une sémantique formelle et un support logiciel (*the Alignment API*). Elle a développé des aligneurs d'ontologies par mesures de similarité. Enfin, l'équipe-projet est à l'origine des campagnes annuelles *Ontology Alignment Evaluation Initiative*, qui évaluent les aligneurs d'ontologies. Des travaux plus récents ont porté sur l'interconnexion des données, c'est-à-dire l'alignement de données plutôt que d'ontologies, avec la notion de clés de liaison (*link keys*).

Un autre sujet de l'équipe est l'évolution des connaissances. **EXMO** a contribué à la révision de réseaux d'ontologies. Toutefois, pour garantir une évolution homogène des connaissances, ses futurs travaux adapteront les techniques d'évolution culturelle à la représentation des connaissances.

Un second aspect du web est qu'il offre non seulement un cadre d'application universel pour Internet, mais également un espace hybride dans lequel les humains et les agents logiciels peuvent interagir à grande échelle et former des communautés mixtes. **Aujourd'hui, des millions d'utilisateurs et d'agents artificiels interagissent quotidiennement** dans des applications en ligne, entraînant ainsi la nécessité d'étudier et de concevoir des systèmes très complexes. Il est nécessaire de disposer de modèles et d'algorithmes qui produisent des justifications et des explications, et qui acceptent les retours d'information pour prendre en charge des interactions avec différents utilisateurs. Il faut pour cela prendre en compte les utilisateurs en tant que composants intelligents des systèmes complexes, qui interagissent avec d'autres composants (intelligence artificielle dans les interfaces, interaction en langage naturel), qui participent

au processus (calcul par l'humain, *crowdsourcing*, machines sociales) et qui peuvent être améliorés par le système (amplification de l'intelligence, augmentation cognitive, intelligence augmentée et cognition distribuée).

WIMMICS

Les applications web (Wikipedia, par exemple) proposent des espaces virtuels dans lesquels les personnes et les logiciels interagissent dans des communautés mixtes qui échangent et utilisent des connaissances formelles (ontologies, bases de connaissances) et des contenus informels (textes, messages, étiquettes).

L'équipe-projet **WIMMICS** étudie des modèles et des méthodes qui mettent en relation la sémantique formelle et la sémantique sociale sur le web. Elle suit une approche multidisciplinaire pour analyser et modéliser ces espaces, leurs communautés d'utilisateurs et leurs interactions. Elle fournit également des algorithmes de calcul pour produire ces modèles à partir des traces sur le web, notamment grâce à l'extraction de connaissances depuis un texte, l'analyse sémantique des réseaux sociaux et grâce à la théorie de l'argumentation.

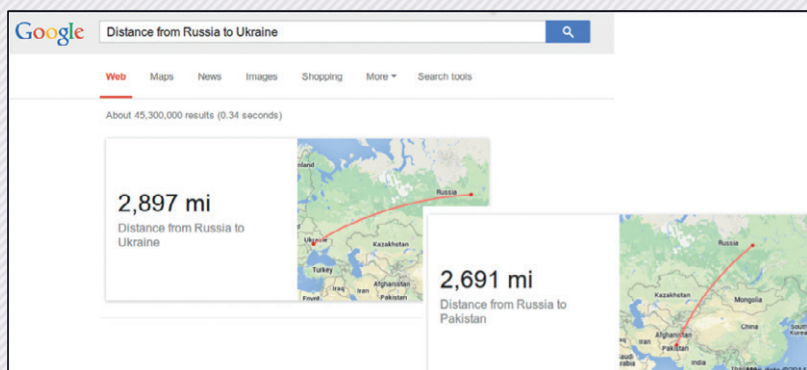


Figure 7 : Sans la sémantique, la Russie apparaît plus près du Pakistan que de l'Ukraine. Extrait de l'article "Why the Data Train Needs Semantic Rails" par Janowicz et al., *AI Magazine*, 2015.

Pour pouvoir formaliser et raisonner sur ces modèles, **WIMMICS** propose des langages et des algorithmes qui utilisent et étendent les approches basées sur les graphes de connaissances pour le web sémantique et le web de données, notamment le *Resource Description Framework* (RDF). Ensemble, ces contributions se concrétisent par des outils d'analyse et des indicateurs qui fournissent de nouvelles fonctionnalités pour les communautés.

Les résultats de ces recherches sont intégrés, évalués et transférés par l'intermédiaire de logiciels génériques (moteur de recherche pour le web sémantique CORESE) et d'applications spécialisées (moteurs de recherche **DiscoveryHub** et **QAKIS**).

Établissement de ponts entre les données massives stockées dans des bases de données à l'aide de technologies sémantiques

Le web sémantique traite l'intégration massive de sources de données très diverses (capteurs des villes intelligentes, connaissances biologiques extraites d'articles scientifiques ou encore descriptions d'événements sur les réseaux sociaux) qui utilisent des vocabulaires très différents (schémas relationnels, thésaurus, ontologies formelles), dans des raisonnements très variés (prises de décision par dérivation logique, enrichissements grâce à l'induction, fouilles de données, etc.). Sur le web, le graphe original des hyperliens est accompagné d'un nombre croissant d'autres graphes et est combiné à des sociogrammes qui capturent la structure des réseaux sociaux, des *workflows* qui définissent les chemins de décision à suivre, des *logs* qui capturent les traces de navigation des utilisateurs, des compositions de services qui spécifient des traitements répartis, ou encore des ensembles de données liées. De plus, ces graphes ne sont pas disponibles dans un entrepôt central unique, mais sont répartis sur de nombreuses sources différentes. Certains sous-graphes sont publics ([dbpedia http://dbpedia.org](http://dbpedia.org), par exemple), tandis que d'autres sont privés (données d'entreprises). Certains sous-graphes sont petits et stockés localement (le profil d'un utilisateur sur un appareil, par exemple), d'autres sont plus grands et hébergés sur des *clusters* (Wikipedia), certains sont largement stables (thésaurus du latin), d'autres changent plusieurs fois par seconde (statuts sur les réseaux sociaux), etc. **Un type de réseau n'est pas isolé, il interagit avec d'autres réseaux** : les réseaux sociaux influencent les flux de messages, ainsi que leurs sujets et leur nature, les liens sémantiques entre les termes interagissent avec les liens entre les sites et vice versa, etc. La mise en place de méthodes permettant de représenter et d'analyser les types de graphes, de les regrouper et de regrouper les processus qui s'appliquent à ces graphes, représente un défi considérable.

CEDAR

Pour donner un sens aux mégadonnées (*Big data*), il est nécessaire de les interpréter à travers le prisme des connaissances sur le contenu,

leur organisation et leur signification. De plus, la connaissance d'un domaine est souvent le langage le plus proche des utilisateurs, qu'ils soient experts d'un domaine ou des utilisateurs novices d'une application manipulant de grands volumes de données. Des outils évolutifs et expressifs de type **OBDA (accès aux données par des ontologies)** représentent par conséquent un facteur clé dans la réussite des applications de mégadonnées.

L'équipe-projet **CEDAR** travaille à l'interfaçage entre des formalismes de représentation des connaissances (notamment logiques de description ou règles existentielles) et des moteurs de bases de données. Elle conçoit des outils de type OBDA très efficaces, en accordant une attention particulière au passage à l'échelle vers des bases de données plus grandes, en intégrant par exemple des capacités de raisonnement aux moteurs des bases de données ou en les déployant dans le *cloud*, pour passer à l'échelle. L'équipe-projet **CEDAR** étudie également de nouvelles méthodes d'interaction avec des bases de connaissances et de données plus grandes et plus complexes, comme celles référencées dans le *Linked Open Data cloud* (<http://lod-cloud.net>). La sémantique est également envisagée comme moyen de donner du sens à des contenus complexes hétérogènes, et de les intégrer dans des banques de données web riches et hétérogènes. Une application particulière concerne la vérification de faits journalistiques (*fact-checking*).

GRAPHIK

Exploiter les différentes données disponibles aujourd'hui passe par une prise en compte de leur sémantique, c'est-à-dire par la connaissance. Cette exigence, largement reconnue, a donné naissance à un nouveau paradigme, l'**OBDA** (voir ci-dessus), qui s'appuie sur les ontologies de domaine pour accéder aux données. En d'autres termes, les bases de données deviennent des bases de connaissances dans lesquelles on a ajouté une couche ontologique sur les données. L'ajout de cette couche présente plusieurs intérêts : il permet de déduire des informations non explicitement codées dans les données, d'adapter le vocabulaire de requête aux besoins spécifiques et d'accéder à des sources de données hétérogènes de façon uniforme.

Au cours des dix dernières années, un nombre significatif de recherches a abouti à toute une série de langages d'ontologies offrant différents degrés de compromis entre expressivité et complexité. Cependant, il reste encore beaucoup de progrès à accomplir avant d'obtenir des systèmes

évolutifs allant au-delà des ontologies de base. De plus, les recherches sur les systèmes **OBDA** se fondent sur l'hypothèse selon laquelle les données suivent des structures relationnelles, alors même que des volumes croissants de sources de données ne sont pas relationnels, comme les bases de données **NOSQL** pour le "*Big data analytics*". La question de savoir si le paradigme **OBDA** peut être adapté aux systèmes de gestion de données non relationnels reste ouverte.

L'une des fonctions importantes des techniques à base de connaissances est leur capacité d'explication, c'est-à-dire leur capacité potentielle à justifier les conclusions obtenues. Être capable d'expliquer, de justifier ou d'argumenter est une caractéristique obligatoire dans de nombreuses applications de l'IA dans lesquelles les utilisateurs ont besoin de comprendre les résultats du système afin de pouvoir aussi bien s'y fier que le maîtriser. En outre, cela devient un problème crucial d'un point de vue éthique dès que des décisions automatisées ont un impact potentiel sur les êtres humains.

Les recherches de **GRAPHIK** portent principalement sur le domaine de la représentation des connaissances et du raisonnement. **L'équipe-projet est particulièrement intéressée par les formalismes d'accès aux données**, comme les logiques de description et les règles existentielles / Datalog+. Les contributions de l'équipe-projet portent à la fois sur l'aspect théorique et sur l'aspect appliqué. **GRAPHIK** s'attaque actuellement à la question de l'**OBDA**, en mettant l'accent sur le développement de formalismes appropriés avec des algorithmes efficaces et sur la question du raisonnement avec des informations imparfaites, en insistant sur l'explication et la justification des conclusions tirées de sources de données incohérentes.

LINKS

L'apparition du web de données a fait naître le besoin de nouvelles technologies de gestion de bases de données pour les collections de données liées. Les défis classiques de recherche sur les bases de données s'appliquent maintenant aux données liées : comment définir des requêtes logiques exactes, comment gérer des mises à jour dynamiques et comment automatiser la recherche des requêtes appropriées. Contrairement aux données ouvertes liées ordinaires, les recherches de l'équipe-projet **LINKS** sont axées sur les collections de données liées sous différents formats, partant de l'hypothèse que les données sont

correctes dans la plupart des dimensions. Les problèmes restent difficiles à résoudre en raison de données incomplètes, de schémas non-informatifs ou hétérogènes, ainsi que de données erronées ou ambiguës. L'équipe-projet développe des algorithmes pour l'évaluation et l'optimisation des requêtes logiques sur les collections de données liées, des algorithmes incrémentaux pouvant superviser les flux du web de données et gérer les mises à jour dynamiques des collections de données liées, ainsi que des algorithmes d'apprentissage symbolique qui peuvent déduire, à partir d'exemples, les requêtes appropriées pour des collections de données liées. **LINKS** développe aussi des langages de programmation répartis pour la gestion dynamique des collections de données.

Ses principaux objectifs de recherche se structurent comme suit :

- **requêtes sur des données liées hétérogènes** : développement de nouveaux types de mises en correspondance de schémas pour les jeux de données semi-structurés en formats hybrides, y compris les collections RDF et les bases de données relationnelles. Ceux-ci impliquent des requêtes récursives sur les collections de données liées avec des algorithmes d'évaluation, une analyse statique et allant vers des applications concrètes ;

- **gestion des données liées dynamiques** : développement des langages de programmation distribués, centrés sur les données, en utilisant des flux et le parallélisme, pour gérer les collections dynamiques de données liées et les flux de traitement. En se basant sur de nouveaux algorithmes de requêtes incrémentales, l'équipe-projet étudie la propagation des mises à jour de données dynamiques par mises en correspondance des schémas, ainsi que les méthodes d'analyse statique pour les *workflows* du web de données ;

- **graphes de liaison** : développement des algorithmes d'apprentissage symbolique pour générer des requêtes et des mises en correspondance (*mappings*) de données liées, en différentes représentations graphiques, à partir d'exemples annotés.

DYLISS

Les sciences expérimentales subissent une révolution des données en raison de la multiplication des capteurs qui permettent de mesurer l'évolution dans le temps de milliers d'éléments interdépendants, physiques ou biologiques. Lorsque les mesures sont suffisamment précises et variées, elles peuvent être intégrées dans un système d'apprentissage machine afin de connaître le rôle et la fonction des éléments dans le système expérimental analysé.

Toutefois, dans de nombreux domaines tels que la biologie moléculaire, la chimie et l'environnement, les mesures ne sont pas suffisantes pour identifier le rôle spécifique de chaque composant.

Dans ce cas, il est essentiel de confronter le résultat des analyses de données au corpus de savoir qui est actuellement en cours de structuration dans le cadre de l'initiative **Linked Open Data** (données ouvertes liées) associée aux technologies du web sémantique. En effet, il existe aujourd'hui plus de mille cinq-cents entrepôts de connaissances sur les sciences expérimentales, stockées au format RDF. L'un des principaux défis posés est donc de permettre l'exploration, le tri et l'extraction de candidats pour les fonctions de composants. C'est pourquoi l'équipe **DYLISS** a pour objectif de développer des techniques de programmation logiques permettant de combiner efficacement classification double symbolique, basée sur l'analyse de concepts formels, avec les technologies du web sémantique.

Autres équipes-projet dans ce domaine : **TYREX**, Grenoble ; **ZENITH**, Montpellier.

4.5 Les défis génériques en robotique et véhicules autonomes

La robotique associe plusieurs sciences et technologies, allant des techniques "bas niveau" - comme la mécanique, la mécatronique, l'électronique, le contrôle-commande - jusqu'aux "haut niveau" comme la perception, les sciences cognitives, la collaboration et le raisonnement.

Dans cette sous-section, même si l'intelligence artificielle en robotique implique l'utilisation de fonctions de "bas niveau" pour certains traitements, nous n'aborderons ici que les aspects directement reliés au champ de l'IA.

Les progrès récemment accomplis en matière de robotique sont impressionnants. Les robots humanoïdes peuvent marcher, courir, se déplacer dans des environnements connus ou inconnus, effectuer des tâches simples telles que saisir des objets ou manipuler certains appareils ; les robots bio-inspirés sont capables de reproduire les comportements d'êtres vivants très différents (insectes, oiseaux, reptiles, rongeurs...) et d'utiliser ces comportements pour résoudre des problèmes complexes efficacement. Ainsi, le robot bipède Atlas, de Boston Dynamics (http://www.bostondynamics.com/robot_Atlas.html) peut se déplacer à l'exté-

rieur, sur terrain difficile et porter des objets lourds, à l'instar du robot quadrupède BigDog de la même société.

.....
*Sur le plan cognitif,
les robots peuvent
jouer de la musique,
accueillir les visiteurs
dans des centres
commerciaux, parler
avec des enfants.*
.....

Sur le plan cognitif, grâce aux progrès réalisés dans le traitement de la parole, l'interprétation des images et des scènes grâce à l'information produite par des capteurs, et grâce aux capacités de raisonnement mises en œuvre, les robots peuvent jouer de la musique, accueillir les visiteurs dans des centres commerciaux, parler avec des enfants. Avec des fonctions de coordination de groupe, les robots sont capables de jouer au football ensemble, mais aucune équipe de robots n'est encore capable de battre une équipe d'hu-

mans même peu habiles. Des véhicules autonomes peuvent rouler sans danger sur de longues durées et certains pays ainsi que des États américains pourraient les autoriser à rouler sur des routes publiques dans un futur proche, bien que beaucoup de questions, notamment éthiques, restent ouvertes.

Les équipes d'Inria chargées de la recherche sur les robots et les véhicules autonomes traitent les problèmes suivants : (i) compréhension de situation par perception multisensorielle ; (ii) raisonnement prenant en compte l'incertitude, la résilience ; (iii) association de plusieurs approches pour la prise de décision.

Compréhension de situation par perception multisensorielle

Pour qu'un robot puisse se déplacer dans des zones inconnues, par exemple pour une voiture autonome dans la circulation ou pour un robot d'assistance personnelle comme Toi.Net (voir section 1), il est essentiel de percevoir son environnement et d'identifier la situation, ce qui est rendu possible grâce à une captation multisensorielle (vision, laser, son, Internet, détection de l'environnement routier dans le cas de véhicules). Les situations peuvent correspondre à de simples symboles, à des ontologies ou à des représentations plus sophistiquées des acteurs et des objets présents dans un environnement. Une bonne analyse de la situation peut permettre au robot de prendre des décisions complexes - allant même dans certains cas jusqu'à enfreindre la loi ou les réglementations pour sauver la vie des passagers d'une voiture.

RITS est un projet multidisciplinaire portant sur la robotique appliquée aux systèmes de transport intelligents. L'équipe-projet cherche notamment à combiner les outils et les techniques mathématiques pour concevoir des systèmes robotiques intelligents permettant une mobilité autonome et durable.

Parmi les thématiques scientifiques couvertes :

- traitement des signaux multicapteurs (traitement de l'image, données télémétriques laser, données GPS, UMI...) et fusion des données ;
- perception évoluée pour la modélisation et l'interprétation de l'environnement ;
- contrôle-commande de véhicule (accélération, freinage, direction) ;
- communications sans fil (véhicule-véhicule, véhicule-infrastructure) ;
- modélisation et simulation de trafic à grande échelle ;
- contrôle et optimisation des systèmes de transport routier ;
- développement et déploiement de véhicules automatisés (cyber-voitures, véhicules particuliers...).

Ces recherches ont pour objectif d'améliorer le transport routier en termes de sécurité, d'efficacité, de confort et d'impact écologique. L'approche technique est centrée sur les aides à la conduite, pouvant aller jusqu'à une automatisation totale.



Figure 8 : Voitures autonomes de RITS. - © Inria / Photo H. Raguet

Un système expérimental de démonstration basé sur des véhicules totalement automatisés a été installé sur le site Inria de Rocquencourt.

L'équipe-projet dispose de moyens importants dont une flotte d'une douzaine de véhicules autonomes, différents capteurs ; elle dispose également de moyens informatiques sophistiqués, dont un outil de simulation.

Raisonnement prenant en compte l'incertitude, la résilience

Les robots sont actifs dans le monde physique et doivent faire face à des défaillances de toutes sortes : coupures réseau, capteurs défectueux, risques électroniques, etc. Certains capteurs fournissent des informations incomplètes ou présentent des marges d'erreur qui induisent une incertitude par rapport aux données transmises. Cependant, un robot mobile autonome doit fonctionner continuellement, sans aucune intervention humaine et sur de longues durées. **L'un des principaux problèmes pour les systèmes robotiques provient des données transmises, qui peuvent être incertaines, insuffisantes ou encore disponibles de manière fractionnée** en raison des différents temps d'acquisition. Des algorithmes *anytime*, qui fournissent un résultat à la demande, peuvent constituer une solution dans les cas où il est nécessaire de prendre une décision rapidement, même si cette décision n'est pas parfaite.

Prise de décision basée sur plusieurs approches

Un robot prend une décision en se fondant sur un certain nombre de données et d'informations disponibles : les données des différents capteurs, des informations sur l'environnement sous forme d'évaluation de la situation, le souvenir des décisions passées, les règles et règlements stockés dans la mémoire du robot. Il est nécessaire de combiner ces faits avec les données et d'obtenir un raisonnement hybride à partir des données numériques, continues ou discrètes et des représentations sémantiques. De plus, comme nous l'avons vu précédemment, ce raisonnement doit aussi prendre en compte le facteur d'incertitude : c'est le défi que doit relever la recherche sur la prise de décision des robots. L'apprentissage non-supervisé et l'apprentissage par renforcement des situations et de leurs interprétations sémantiques peuvent constituer des approches pertinentes.

L'objectif de l'équipe **LARSEN** est de faire sortir les robots des laboratoires de recherche et des usines. En effet, les robots actuels sont loin d'être des machines autonomes, fiables et interactives qui pourraient

coexister avec nous, dans notre société, et fonctionner pendant des jours, des semaines ou des mois. Alors qu'il reste sans aucun doute des progrès à faire au niveau matériel, les plates-formes de robotique évoluent rapidement. Les principaux problèmes restant à résoudre pour atteindre cet objectif se situent, pour cette équipe-projet, du côté logiciel. Ainsi, **LARSEN** espère que son logiciel puisse fonctionner sur des robots peu coûteux qui ne sont pas équipés de capteurs ou d'actionneurs à haute performance (et à coût élevé), de sorte que ses techniques puissent être déployées de façon réaliste et évaluées dans des conditions réelles, comme dans des applications robotiques de service et d'assistance. L'équipe-projet conçoit ces robots comme capables de coopérer non seulement les uns avec les autres, mais aussi avec des espaces ou des appartements intelligents, qui peuvent être vus comme une extension des robots au cadre de vie dans son ensemble.

Développer de nouvelles formes d'interactions physiques entre les robots et les humains

Les membres de l'équipe **LARSEN** organisent leurs recherches autour de deux axes principaux, **l'autonomie permanente et l'interaction naturelle avec les systèmes robotiques** :

- concernant l'autonomie permanente, les chercheurs ont identifié deux problèmes à traiter. Le premier est de parvenir à doter les robots d'une conscience de la situation qui reste stable dans des environnements ouverts et dynamiques. Le second est de leur permettre de se remettre de dommages physiques. L'une des approches proposées par l'équipe-projet est de laisser le robot expérimenter grâce à de nouveaux algorithmes essais-erreurs qui lui permettront d'apprendre avec un nombre d'essais réduit (une douzaine en principe). Une autre approche que l'équipe-projet va étudier consiste à déployer plusieurs robots ou un essaim de robots.
- concernant l'interaction naturelle avec les systèmes robotiques, l'équipe-projet modifie le comportement des robots de façon à améliorer leur degré d'acceptabilité. Elle traite une question plus originale qui est de développer de nouvelles formes d'interactions physiques entre les robots et les humains. Pour ce faire, le robot doit pouvoir réagir et s'adapter en fonction de la réaction de l'humain, en exploitant l'ensemble des signaux verbaux et non-verbaux que l'humain émet naturellement lors des interactions physiques ou sociales. L'un des principaux objectifs est de faire en sorte que le robot soit capable de détecter la présence d'un être humain et de comprendre son activité.

Assistance personnelle : Inria Project Lab PAL

Au cours des cinquante dernières années, les progrès de la médecine et l'amélioration de la qualité de la vie ont permis d'allonger l'espérance de vie dans les sociétés industrielles. L'augmentation du nombre de personnes âgées induit des problématiques nouvelles en santé publique car, même si les seniors vieillissent en bonne santé, leur âge entraîne progressivement et naturellement une fragilisation, notamment sur le plan physique, qui peut aboutir à une perte d'autonomie. Ce constat oblige la société à repenser le modèle actuel de prise en charge des personnes du troisième âge. En raison d'une capacité d'accueil limitée dans les instituts spécialisés et du souhait de la majorité des personnes âgées de rester chez elles le plus longtemps possible, la demande de services spécifiques à domicile est de plus en plus importante.

Les technologies de l'intelligence ambiante et la robotique ont leur rôle à jouer face à cet enjeu de société. Inria a lancé entre 2010 et 2015 un Inria Project Lab (IPL) nommé **PAL**, pour *Personally Assisted Living* (vie avec assistance personnelle). En se basant sur les compétences et les objectifs des équipes impliquées, l'IPL **PAL** a défini et traité **quatre thèmes de recherche** :

- évaluation du degré de fragilité de la personne âgée ;
- mobilité des personnes ;
- rééducation, transfert et assistance à la marche ;
- interaction sociale.

CHROMA

L'objectif global de **CHROMA** est de traiter des questions fondamentales non résolues, situées à l'intersection des domaines de recherche émergents que sont la "robotique basée sur l'humain" et les "systèmes multirobots".

Il s'agit de **concevoir des algorithmes** et de **développer des modèles permettant à des robots mobiles de se déplacer et d'agir dans des environnements dynamiques peuplés d'êtres humains**.

CHROMA s'intéresse aux décisions relatives aux tâches de navigation du robot ou multirobots, y compris la perception et la génération de mouvement. L'approche de l'équipe-projet, sur ce point, consiste à réunir les méthodes probabilistes, les techniques de planification de mouvement et les modèles décisionnels multiagents. Ceci implique de

collaborer avec des spécialistes de plusieurs disciplines, telles que la psychosociologie pour la prise en compte des modèles humains.

Les recherches de **CHROMA** portent essentiellement sur **deux aspects de la navigation robotique** :

- i) perception et conscience de la situation dans un environnement peuplé d'humains, en se concentrant sur la perception Bayésienne et la fusion des données des capteurs ;
- ii) passage à l'échelle de la planification de déplacement pour un ou plusieurs robots, en combinant les modélisations d'incertitude, les modèles décentralisés (essaim de robots) et la prise de décision séquentielle multiagent.

L'équipe-projet **CHROMA** vise des applications et le transfert de ses résultats scientifiques. Elle travaille avec des partenaires industriels et des start-up. Les principaux domaines d'application concernent la conduite de véhicules autonomes (avec Renault et Toyota), les robots aériens pour les tâches de surveillance et la robotique de service.

Autres équipes-projets dans ce domaine : **LAGADIC**, Rennes, **HEPHAISTOS**, Sophia-Antipolis.

4.6 Les défis génériques en neurosciences et sciences cognitives

L'objectif principal de l'équipe **FLOWERS** est la modélisation du développement évolutif chez les humains et ses applications dans les domaines de la robotique, de l'interaction Homme-ordinateur et des technologies éducatives.

L'un des principaux enjeux des sciences cognitives et de l'intelligence artificielle est de **comprendre comment les organismes peuvent acquérir un éventail de compétences sur une durée étendue**. Le processus de développement sensorimoteur, cognitif et social chez l'enfant est organisé en étapes ordonnées et résulte de l'interaction complexe entre le cerveau, le corps et l'environnement physique.

Pour faire avancer la compréhension fondamentale des mécanismes de développement, l'équipe-projet **FLOWERS** a mis au point des modèles

informatiques et robotiques, en étroite collaboration avec des spécialistes de psychologie du développement et de neurosciences. L'équipe-projet a développé des modèles de processus d'exploration guidée qui permettent aux apprenants de collecter des données efficacement dans des espaces de grande dimension multitâches. Cela inclut des mécanismes d'apprentissage par l'action et par recherche d'information (également appelé apprentissage motivé par la curiosité, *curiosity-driven learning*), d'apprentissage par imitation et la libération maturationnelle des degrés de liberté.

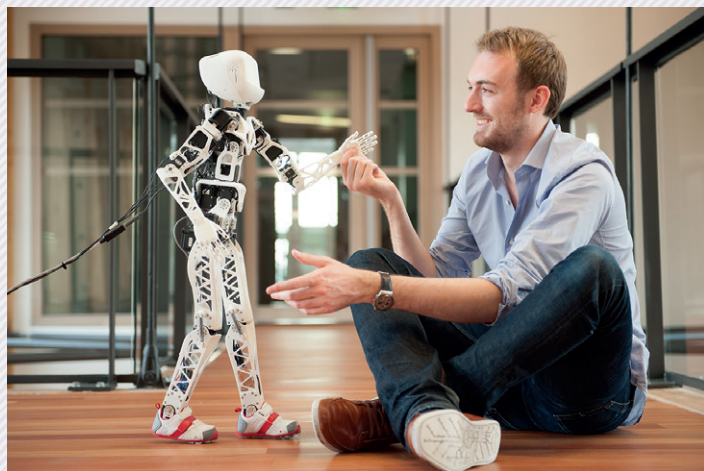


Figure 9 : Le robot humanoïde Poppy, imprimé en 3D et open-source - © Inria / Photo H. Raguet

FLOWERS a non seulement abouti à élaborer de nouvelles théories et de nouveaux paradigmes expérimentaux qui permettent de comprendre le développement humain, mais elle a également exploré l'étendue des possibilités d'application de ces modèles à la robotique, à l'interaction Homme-ordinateur et aux technologies éducatives. En robotique, l'équipe-projet a démontré que la curiosité artificielle combinée à l'apprentissage par imitation peut fournir des blocs fonctionnels essentiels permettant aux robots d'acquérir plusieurs capacités par interaction naturelle avec des utilisateurs profanes humains, par exemple dans le contexte de la robotique d'assistance. L'équipe-projet a également montré que **les modèles d'apprentissage motivé par la curiosité peuvent être transposés dans des algorithmes intelligents permettant aux logiciels éducatifs de s'adapter** de façon incrémentielle et dynamique aux particularités de chaque apprenant humain et proposant des séquences personnalisées d'activités pédagogiques. Dans

le domaine de l'interaction Homme-ordinateur, **FLOWERS** a montré comment les algorithmes incrémentaux d'apprentissage peuvent être utilisés pour supprimer la phase de calibration dans certaines interfaces cerveau-ordinateur.

Information Seeking, Curiosity and Attention : Computational and Neural Mechanisms Gottlieb, J., Oudeyer, P-Y., Lopes, M., Baranes, A. (2013) Trends in Cognitive Science, 17(11), p. 585-596.

How Evolution may work through Curiosity-driven Developmental Process Oudeyer, P-Y. et Smith. L. (sous presse) Topics in Cognitive Science.

Self-Organization in the Evolution of Speech, 2e édition (2016) Oudeyer, P-Y., Oxford University Press

Multi-Armed Bandits for Intelligent Tutoring Systems, Manuel Lopes, Benjamin Clement, Didier Roy, Pierre-Yves Oudeyer. Journal of Educational Data Mining (JEDM), Vol 7, No 2, 2015.

Exploiting task constraints for self-calibrated brain-machine interface control using error-related potentials, I. Iturrate, J. Grizou, J. Omedes, P-Y. Oudeyer, M. Lopes and L. Montesano. PLOS One, 2015.

Inverse Reinforcement Learning in Relational Domains, Thibaut Munzer, Bilal Piot, Mathieu Geist, Olivier Pietquin et Manuel Lopes. International Joint Conference on Artificial Intelligence (IJCAI'15), Buenos Aires, Argentine, 2015.

ARAMIS

Comprendre les fonctions du cerveau et leurs altérations requiert l'intégration de plusieurs niveaux d'organisation, fonctionnant à des échelles spatiales (allant du niveau de la molécule à celui de l'ensemble du cerveau) et temporelles (de la milliseconde à la durée d'une vie entière) différentes et représentant plusieurs types de processus biologiques (processus anatomique, fonctionnel, moléculaire et cellulaire). Plusieurs aspects de ces processus peuvent maintenant être quantifiés chez des sujets humains vivants grâce au développement de diverses technologies dont la neuro-imagerie, l'électrophysiologie, la génomique, la transcriptomique... L'un des principaux objectifs de l'équipe est de développer des approches qui peuvent automatiquement apprendre des modèles adaptés à partir des données multimodales générées par ces techniques.

L'objectif de l'équipe-projet **ARAMIS** est de **développer des approches informatiques et statistiques permettant l'apprentissage à partir de données cérébrales multimodales** (neuro-imagerie, électrophysiologie, génomique...). Une première méthode de recherche consiste à modéliser la structure du cerveau et sa variabilité statistique selon les populations. Cela implique le développement d'une méthode de segmentation pour extraire les structures cérébrales des images et des modèles statistiques de forme (Gori et al., 2013), mise en œuvre dans le progiciel

gratuit Deformetrica⁶. Un second axe de recherche a pour objectif de **modéliser les interactions fonctionnelles entre des zones distantes du cerveau** qui sont à la base des processus cognitifs. Il s'appuie sur des approches pouvant modéliser l'organisation des réseaux cérébraux complexes (De Vico Fallani et al., 2014). Ces modèles s'appliquent à la conception de nouveaux dispositifs, aux interfaces cerveau-ordinateur et au *neurofeedback*, pour la rééducation de patients en neurologie. Enfin, l'équipe-projet développe des méthodes pour étudier des modèles des maladies neurologiques et de leur progression. Ces derniers permettent d'étudier les modèles topographiques (Cuingnet et al., 2013) et spatio-temporels (Schiratti et al., 2015) caractéristiques d'une maladie donnée. Ils pourraient déboucher sur la mise au point de nouveaux outils de personnalisation du diagnostic, pronostic et traitement pour des troubles comme la maladie d'Alzheimer, la maladie de Parkinson, l'épilepsie et les troubles cérébrovasculaires.

⁶ <http://www.deformetrica.org>

Bibliographie :

Cuingnet R, Glaunès JA, Chupin M, Benali H, and Colliot O (2013). Spatial and Anatomical Regularization of SVM: A General Framework for Neuroimaging Data. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 35, 682–696.

Gori P, Colliot O, Worbe Y, Marrakchi-Kacem L, Lecomte S, Poupon C, et al. (2013). Bayesian Atlas Estimation for the Variability Analysis of Shape Complexes. In *Proc. Medical Image Computing and Computer Assisted Intervention*, Nagoya, Japan, pp. 267–274.

Schiratti J-B, Allassonniere S, Colliot O, and Durrleman S (2015). Learning spatiotemporal trajectories from manifold-valued longitudinal data. In *Neural Information Processing Systems*, Montréal, Canada,.

De Vico Fallani F, Richiardi J, Chavez M, and Achard S (2014). Graph analysis of functional brain networks: practical issues in translational neuroscience. *Philosophical Transactions of the Royal Society of London Series B, Biological Sciences* 369.

MNEMOSYNE

Aux frontières des neurosciences intégratives et computationnelles, l'équipe-projet **MNEMOSYNE** propose de modéliser le cerveau comme un système de mémoires actives en synergie et en interaction avec les mondes interne et externe et de le simuler comme un tout et en situation.

Les principales fonctions cognitives et comportementales (ex: attention, reconnaissance, planification, décision) émergent de boucles sensori-motrices impliquant le monde extérieur, le corps et le cerveau. L'équipe étudie, modélise et implémente ces boucles et leurs interactions dans la perspective d'un comportement complètement autonome. Avec cette

approche "systémique", les chercheurs indiquent que de tels systèmes complexes ne peuvent être vraiment appréhendés que comme un tout et dans des situations comportementales naturelles. Pour mettre au point les caractéristiques fonctionnelles et adaptatives de tels modèles au niveau de leur circuiterie neuronale et pour les implémenter dans des systèmes interagissant avec le monde, ils doivent combiner les principes, les méthodes et les outils de différents domaines scientifiques.

- **L'équipe-projet modélise les principales structures cérébrales et les principaux flux d'information du cerveau** (comme en neurosciences intégratives et cognitives) en insistant sur les liens entre le cerveau, le corps et l'environnement (cognition incarnée).
- **Elle utilise des formalismes de calcul distribué lui permettant d'implanter des modèles à différentes échelles de description** (comme en neurosciences computationnelles).
- **Elle déploie ses modèles à large échelle** (calcul à haute performance), les incarne dans des corps en interaction avec l'environnement (robotique autonome) et les simule de façon interactive avec les événements rencontrés par un robot virtuel/réel.

4.7 Les défis génériques dans le traitement du langage

Le domaine du traitement automatique de langage naturel (TALN) est apparu dans les années 1950. Il est toujours d'une importance capitale aujourd'hui, dans la société de l'information. Son objectif est de traiter des textes en langage naturel, qu'il s'agisse d'analyser des textes existants ou d'en générer de nouveaux, afin d'atteindre un traitement de langage semblable à celui de l'Homme pour toute une série de tâches ou d'applications. Ces applications, regroupées sous le terme "**ingénierie du langage**", regroupent la traduction automatique, la réponse aux questions, la recherche d'informations, l'extraction d'informations, l'exploration de données, la lecture et l'écriture assistées, etc. D'un point de vue recherche, la linguistique empirique et les humanités numériques peuvent aussi être considérées comme des domaines d'application du TALN.

Le TALN est un domaine transdisciplinaire qui requiert une expertise en matière de linguistique formelle et descriptive (afin de développer des modèles linguistiques des langages humains), d'informatique et d'algo-

rithmique (pour concevoir et développer des programmes efficaces qui peuvent prendre en charge ces modèles), et de mathématiques appliquées (pour l'acquisition automatique des connaissances linguistiques ou générales). Le traitement des textes en langage naturel représente une tâche difficile, notamment en raison de la forte dose d'ambiguïté présente dans le langage naturel et des spécificités des langages individuels et des dialectes, mais aussi parce que de nombreux utilisateurs ne respectent pas nécessairement les conventions grammaticales et orthographiques lorsque celles-ci existent.

Au cours des premières décennies, le TALN s'est principalement concentré sur des approches symboliques, en offrant des contributions importantes à l'informatique, notamment dans la théorie de la grammaire formelle et des techniques d'analyse syntaxique. Le savoir linguistique a essentiellement été codé sous forme de grammaires et de bases de données lexicales développées manuellement. Au cours des deux dernières décennies, les approches statistiques, basées sur l'apprentissage automatique, ont nettement renouvelé ce domaine, amenant les corpus annotés sur le devant de la scène et améliorant l'état de l'art de façon significative. Cependant, les approches symboliques conservent des avantages spécifiques et les meilleurs résultats sont souvent obtenus en exploitant tous les types de ressources, y compris les connaissances linguistiques dans les systèmes hybrides associant les techniques symboliques et statistiques.

L'équipe-projet **ALPAGE** (Inria Paris & Université Paris-Diderot) est spécialisée dans le TALN et ses applications. Ses recherches s'étendent à tous les niveaux de l'analyse linguistique (morphologie, syntaxe, sémantique, discours), rendant ainsi possible une étude globale multiéchelle des données texte. **ALPAGE a développé de nombreuses ressources linguistiques, surtout pour le français, utilisées au niveau international.** Elle a aussi contribué à améliorer l'état de l'art dans des domaines comme l'analyse syntaxique symbolique, statistique et hybride (pour le français, l'anglais et d'autres langues, y compris des variantes non-standard comme celles trouvées sur le web 2.0), l'analyse sémantique et la modélisation du discours. L'équipe a commencé à explorer une approche émergente qui consiste à coupler l'analyse sémantique à des modèles du monde auquel le texte fait référence (par exemple, l'association du monde artificiel d'un jeu vidéo à la représentation sémantique des conversations en ligne du joueur).

ALPAGE est également fortement impliquée dans la recherche appliquée, de deux manières différentes. D'une part, l'équipe-projet collabore avec le monde industriel pour développer des outils et des techniques TALN, pour des tâches comme l'exploration de données (surtout pour des textes imparfaits), la recherche d'information, l'extraction d'information et la génération automatique de texte. D'autre part, elle est fortement impliquée dans la **recherche en linguistique informatique et empirique, et des humanités numériques**. Cela implique l'étude de la morphologie computationnelle et de la syntaxe empirique, y compris l'étude de l'évolution historique des langages. Les axes de recherche futurs nécessitent une collaboration avec des chercheurs des sciences humaines et sociales (par exemple, des historiens), ainsi que le développement de modèles d'évolution linguistique riche, couvrant les aspects lexicaux, morphologiques et/ou syntaxiques.

Autres équipes-projets dans ce domaine : **SEMAGRAMME**, Nancy.

4.8 Les défis génériques en programmation par contraintes pour l'aide à la décision

La programmation par contraintes apparaît dans les années 1980 et se développe à l'intersection de l'intelligence artificielle et de la recherche opérationnelle, de l'informatique et des mathématiques. De nature multidisciplinaire, cette technique continue à utiliser des connaissances de différents domaines, tels que les mathématiques discrètes, l'informatique théorique (théorie des graphes, combinatoire, algorithmes, complexité), les analyses fonctionnelles et l'optimisation, l'ingénierie informatique et logicielle. La programmation par contraintes a été identifiée comme élément stratégique de l'informatique en 1996 par l'ACM. Au tournant du siècle, la technologie de l'optimisation se développe dans l'industrie (avec notamment Ilog, IBM, Dash et, plus récemment, Microsoft, Google et Dynadec), avec les domaines scientifiques correspondants, au point de convergence de la programmation par contrainte, de la programmation mathématique, de la recherche locale et de l'analyse numérique. La technologie de l'optimisation aide

.....
*La prise de
 décision assistée
 par ordinateur
 et l'optimisation
 deviennent
 des éléments
 fondamentaux
 de l'assistance aux
 activités humaines
 de toutes sortes.*

maintenant le secteur public, les entreprises et, dans une certaine mesure, les personnes, à prendre des décisions utilisant au mieux les ressources et répondant à des exigences spécifiques dans un monde de plus en plus complexe. En effet, la prise de décision assistée par ordinateur et l'optimisation deviennent des éléments fondamentaux de l'assistance aux activités humaines de toutes sortes.

Avec la prééminence de la technologie d'optimisation dans la plupart des secteurs industriels, il apparaît évident que les solutions rapides et sur mesure, souvent utilisées aujourd'hui, ne suffisent pas pour le développement à long terme de la technologie d'optimisation et sa large diffusion. Il est patent qu'**il devrait y avoir une connexion plus directe entre les résultats mathématiques et leur réutilisation systématique dans les principaux domaines de la technologie d'optimisation.**

Problèmes généraux

En dépit de son importance, la décision assistée par ordinateur et l'optimisation présentent un certain nombre de faiblesses fondamentales qui les empêchent de tirer complètement parti de leur potentiel et entravent leurs progrès et leur capacité à gérer des situations de plus en plus complexes. Cela est principalement dû à la diversité des acteurs en cause, qui sont décrits ci-après.

Conception de solutions efficaces

- **l'informatique pour les langages, les outils de modélisation et les bibliothèques** : comment mettre au point, de façon systématique, des méthodes efficaces pour résoudre les problèmes d'optimisation et de décision ;

- **les mathématiques appliquées pour la partie théorique** : trouver des abstractions puissantes qui permettent de traiter une classe de problèmes, les implémenter dans des composants logiciels.

Chaque communauté technologique exploite indépendamment son paradigme de solution comme la programmation par contraintes, la programmation linéaire et par nombres entiers, l'optimisation continue, la recherche locale basée sur la contrainte. Dans une certaine mesure, la plupart de ces techniques exploitent de différentes manières les mêmes résultats mathématiques.

Ainsi, la première difficulté rencontrée pour la programmation par contraintes est la conception de systèmes informatiques mettant en

œuvre de façon transparente des solutions techniques efficaces.

Théoriquement, l'utilisateur doit pouvoir décrire son problème dans un langage de modélisation de haut niveau, sans se soucier des mécanismes de résolution sous-jacents utilisés, indépendamment de tout langage de programmation informatique et de tout moteur de résolution.

Les systèmes doivent également offrir une base de connaissances sur la résolution de problèmes, qui présente des modèles et des heuristiques pour un grand ensemble de problèmes bien identifiés.

Enfin, l'utilisateur doit avoir la possibilité d'interpréter les solutions obtenues, notamment dans le contexte de problèmes surcontraints où il est nécessaire d'atténuer certaines contraintes, et ce de la façon la plus réaliste possible.

Application à grande échelle

Le deuxième défi est d'**augmenter la performance de résolution, surtout dans le contexte de données à grande échelle**. Il est indispensable d'adapter des techniques telles que des algorithmes de consistance, algorithmes de graphes, programmation mathématique, méta-heuristiques, et de les intégrer dans l'application de programmation par contraintes. Cette intégration soulève de nouvelles questions comme la conception d'algorithmes incrémentaux, la décomposition automatique ou la reformulation automatique des problèmes.

Problèmes industriels complexes

Enfin, le troisième problème porte sur **l'utilisation de la programmation par contraintes dans le contexte de problèmes industriels complexes**, surtout pour des systèmes hybrides mêlant discret et continu. Il existe de nombreux facteurs de complexité, comme :

- la combinaison des aspects temporels et spatiaux ou des aspects continus et discrets ;
- le caractère dynamique de certains phénomènes induisant une modification des contraintes et des données dans le temps ;
- la difficulté à exprimer certaines contraintes physiques, par exemple, l'équilibrage de la charge et la stabilité temporelle ;
- la décomposition de grands problèmes en sous-problèmes, ce qui peut pénaliser l'efficacité de résolution.

TASC

La recherche fondamentale de l'équipe **TASC** est guidée par les questions soulevées précédemment : classer et enrichir les modèles, automatiser la reformulation et la résolution, dissocier les connaissances déclaratives et les connaissances procédurales, développer des outils de modélisation et trouver des outils de résolution qui passent à l'échelle.

L'équipe développe une base de connaissances sur la résolution de problèmes combinatoires : le catalogue général des contraintes. Les aspects relatifs à la résolution sont capitalisés dans le système de résolution de contraintes, CHOCO. Enfin, dans le cadre de ses activités de valorisation, d'enseignement et de partenariat de recherche, l'équipe-projet fait appel à la programmation par contraintes pour résoudre différents problèmes concrets.

Le défi est double : augmenter la visibilité des contraintes dans les autres disciplines informatiques, et contribuer à une diffusion plus large de la programmation par contraintes dans l'industrie.

4.9 Les défis génériques en musique et environnements intelligents

Cette partie présente les deux équipes-projets qui développent un travail original, étroitement lié à l'IA mais qui ne peut être aisément rattaché aux parties précédentes.

MUTANT

Musique & intelligence artificielle : des paradigmes pour les enjeux futurs

Les traitements de signaux traditionnels et les premiers modèles probabilistes pour l'apprentissage automatique ont permis d'améliorer considérablement le traitement automatique de la parole et du langage naturel, surtout dans les années 1990, mais n'ont jamais permis d'atteindre ce résultat dans le domaine de la musique, qu'il s'agisse de la recherche d'informations par le contenu, des systèmes de reconnaissance musicale ou des pratiques musicales assistées par ordinateur. Cela s'explique par les schémas temporels complexes des signaux musicaux qui ne peuvent pas être aisément réduits, par la nature hétérogène des informations musicales et par la nature très adaptative de nombreux concepts musicaux intéressants (que signifie vraiment genre de la musique au-

jourd'hui ?). Depuis le début du XXI^e siècle, les scientifiques ont traité les problèmes de recherche d'informations musicales en général et ont formé des communautés actives, comparables en taille, impact industriel et sociétal, à celles existant dans les domaines de la vision artificielle. Les développements se sont concentrés sur les algorithmes hors ligne agissant sur les grandes bases de données, alors que la majorité de l'activité musicale (production, écoute, création) est en ligne par nature et implique un apprentissage par l'action. **Aujourd'hui, il existe des algorithmes de vision artificielle capables d'interpréter des données visuelles à mesure que le système les reçoit de façon incrémentale.** Mais cette révolution doit encore se faire dans le domaine de l'audition artificielle et de la lecture musicale.



Figure 10 : Utilisation d'Antescofo lors d'une représentation en temps réel

© Inria / Photo H. Raguet

Depuis sa création en 2013, l'équipe-projet **MUTANT** se positionne sur les domaines de la reconnaissance des données musicales et audio et de l'interaction en temps réel. **Antescofo, le logiciel d'écoute artificielle en temps réel**, développé par **MUTANT**, est capable de suivre des musiciens et de décoder leurs paramètres musicaux pendant qu'ils jouent. Il est constitué d'un algorithme, extrêmement simple, de traitement des signaux audio pour l'observation des données. Sa solidité, démontrée lors de différentes représentations publiques, réside aussi bien dans le fait que l'apprentissage des paramètres et l'adaptation se font en ligne et en temps réel (apprentissage par l'action) que dans ses capacités inhérentes de fusion des informations, qui permettent

d'obtenir de bons résultats en mixant plusieurs sources d'information faillibles présentant une incertitude élevée. La qualité du résultat final n'est pas due au traitement du signal de meilleure qualité ou de plus grande clarté mais à une gestion intelligente des différentes sources d'information à mesure qu'elles arrivent. **Le traitement en temps réel de l'incertitude non-filtrée est un élément clé de l'intelligence en ligne** pour rendre possibles les applications. Ceci est une analogie avec l'écoute humaine : l'être humain ne décode aucune information musicale de manière parfaite au moment où il l'écoute, qu'il interagisse avec la musique ou qu'il la crée. Il s'agit de l'approche adoptée par l'équipe pour obtenir des algorithmes d'écoute artificielle efficaces et opérationnels.

.....
L'intelligence artificielle, l'écoute artificielle et la reconnaissance des données musicales sont inutiles si elles ne sont pas couplées à des actions.
.....

La question de la génération de musique soulève toute une série de défis pour l'intelligence artificielle. Comment concevoir des machines capables d'improviser de la musique avec ou sans êtres humains (comme pour l'imitation de styles musicaux) avec le même niveau d'excellence ? À ce stade, il convient de prendre explicitement en compte la nature hétérogène de la durée et des informations musicales dans la modélisation. La prise de décision requiert l'application de règles à la fois au niveau local (par

exemple, par rapport aux relations sémantiques dans la progression de la musique) et à long terme (par exemple, pour les phrases et les formes musicales). Elle repose aussi sur l'apprentissage concurrentiel et collaboratif entre différents agents observant des aspects hétérogènes de la même source d'information (hauteur tonale, timbre, rythme, flux d'informations, etc.). Ce problème à lui seul nécessite d'importantes avancées dans le domaine de l'apprentissage automatique et de l'intelligence artificielle, en ligne et hors ligne.

L'équipe-projet **MUTANT** a mis au point des **agents d'apprentissage interactifs avec un apprentissage sous multiples perspectives** (en utilisant les paradigmes de l'apprentissage par renforcement) permettant de générer de la musique dans différents styles, sans connaissance préalable ni codage manuel, et également d'improviser avec des musiciens de jazz en live. Toutefois, les résultats sont loin d'être satisfaisants et les futures recherches devront envisager des architectures combinant différentes approches de découverte d'informations et de prise de décision.

L'intelligence artificielle, l'écoute artificielle et la reconnaissance des données musicales sont inutiles si elles ne sont pas couplées à des actions. L'écoute artificielle avec **Antescofo** permet aux utilisateurs d'exprimer des concepts au travers d'éventuelles expériences perceptives plutôt que par les mathématiques. Cette capacité d'expression avec des concepts humains associés à un langage en temps réel, de haut niveau, leur permet de **fabriquer des machines avec des scénarios temporels complexes** en environnement actif. C'est la conception choisie par **MUTANT** pour la réalisation de systèmes multimédia cyber-physiques, afin que les systèmes experts à base de connaissances soient accessibles aux utilisateurs quotidiens qui peuvent mieux s'exprimer avec des concepts humains qui, "sous le capot", sont exprimés dans des langages temps réel.

L'objectif du projet **PERVASIVE INTERACTION** est de développer les bases scientifiques et technologiques pour des environnements humains capables de percevoir, agir, communiquer et interagir avec les humains afin de leur fournir des services. L'élaboration de ces environnements comporte de nombreux problèmes liés à l'interprétation des informations de capteurs, à l'apprentissage, à la compréhension machine, à la composition dynamique des éléments et à l'interaction Homme-machine. **Le but du projet est de réaliser des progrès sur les bases théoriques de la perception et de la cognition**, ainsi que de développer de nouvelles formes d'interaction Homme-machine, en utilisant des environnements interactifs servant d'exemples pour les problèmes possibles.

Un environnement est une région (un volume) connectée de l'espace 3D. Un environnement est "perceptif" lorsqu'il est capable de reconnaître et de décrire les choses, les personnes et les activités à l'intérieur de ce volume. Il est possible de réaliser des formes simples de perception spécifique à l'aide d'un seul capteur. Cependant, pour être efficace et générique, la perception doit intégrer des informations provenant de plusieurs capteurs et avec plusieurs modalités. Le projet **PERVASIVE INTERACTION** a pour objet de créer et développer des techniques de perception artificielle associant la vision artificielle, la perception acoustique, les capteurs de distance et les capteurs mécaniques afin que les environnements puissent percevoir et comprendre les humains et les activités humaines.

Un environnement est "actif" lorsqu'il est capable de changer son état interne. Parmi les formes courantes de changement d'état, on

trouve le réglage de la température ambiante, du niveau acoustique ou de l'éclairage. Des formes plus innovantes comportent la présentation contextualisée des informations, ainsi que des services de nettoyage, d'organisation matérielle et de logistique. L'utilisation de surfaces à plusieurs écrans associée à la localisation offre la possibilité de **modifier automatiquement l'affichage des informations pour s'adapter à l'activité en cours** des groupes. L'utilisation de la reconnaissance d'activités et les dispositifs d'écoute offrent la possibilité d'enregistrer un journal d'interactions humain-humain, ainsi que de fournir des informations appropriées sans interruption. L'utilisation de projecteurs vidéo orientables (avec détection visuelle intégrée) permet d'utiliser n'importe quelle surface pour les présentations, les interactions et les communications.

Un environnement peut être considéré comme "interactif" lorsqu'il est capable d'interagir avec les humains en utilisant la perception et l'action, de façon étroitement liée. Des formes simples d'interaction peuvent être basées sur l'observation de la manipulation d'objets physiques ou sur la détection visuelle de doigts, mains ou bras. Les formes plus évoluées impliquent la perception et la compréhension des activités et du contexte. **PERVASIVE INTERACTION** a développé une

nouvelle théorie pour la modélisation de situations pour la compréhension machine de l'activité humaine, basée sur les techniques utilisées en psychologie cognitive. L'équipe-projet explore de nombreuses formes d'interactions, dont les *widgets* d'interaction, l'observation de manipulation d'objets, la fusion des informations acoustiques et visuelles, et les systèmes qui modèlent le contexte d'interaction afin de déterminer les actions et services appropriés de l'environnement.

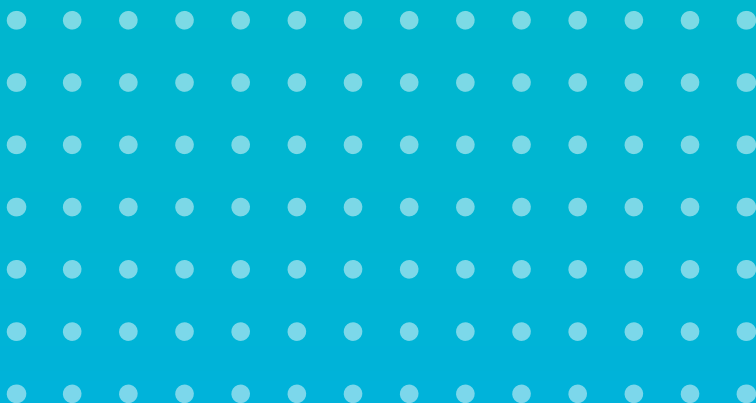
Pour la conception et l'intégration de systèmes de perception des humains et de leurs actions, **PERVASIVE INTERACTION** a développé :

- un environnement à base de modèles de situation pour les services contextuels ;
- des techniques efficaces de vision artificielle ;
- une architecture logicielle distribuée, autonome pour les systèmes de perception multimodale.

Les expériences de **PERVASIVE INTERACTION** ont pour objectif **le développement de services interactifs pour les environnements intelligents**. Les champs d'application sont, entre autres, les services de surveillance pour une vie saine, les objets et services intelligents pour la maison et de nouvelles formes d'interaction Homme-machine basées sur la perception.



Références Inria : chiffres



Sur la période 2005-2015, les chercheurs d'Inria ont publié **plus de 400 articles sur l'IA et plus de 1000 communications à des conférences sur le sujet**. Les références équipes-projets, incluant d'autres supports comme des chapitres d'ouvrages scientifiques et les rapports internes d'Inria sont accessibles sur <https://hal.inria.fr/> et sur la page web présentant le rapport annuel : <http://raweb.inria.fr/rapportsactivite/RA2015/index.html>

Les tableaux ci-dessous indiquent le nombre minimum de publications et travaux de conférences publiés, car certains articles de presse et communications de colloques n'étaient pas indexés dans la base de données d'Inria sur toute la période d'intérêt.

PUBLICATIONS DANS DES REVUES

AI Communications	3
Artificial Intelligence	7
Computer Speech and Langage	13
Computer Vision and Image Understanding	28
Constraints	17
Expert Systems with Applications	17
IEEE Robotics and Automation Magazine	8
IEEE Transactions on Evolutionary Computation	4
IEEE Transactions on Image Processing	45
IEEE Transactions on Intelligent Transportation Systems	14
IEEE Transactions on Knowledge and Data Engineering	15
IEEE Transactions on Pattern Analysis and Machine Intelligence	3
IEEE Transactions on Robotics	48
International Journal of computer vision	68
International Journal of Robotics Research	29
Journal of Artificial Intelligence Research	7
Journal of Data Semantics	6
Journal of Intelligent and Robotic Systems	3
Journal of Intelligent Systems	2
Journal of Langage Technology and Computational Linguistics	1
Journal of machine learning Research	33
Journal of Web Semantics	7
Machine Vision and Applications	9
Pattern Recognition Letters	18
Revue d'Intelligence Artificielle	5
TAL - Traitement Automatique des Langues	20
Trends in Cognitive Sciences	1
Total	431

ACTES DE CONFÉRENCES

Annual Conference of the International Speech Communication Association	43
Annual Conference on Neural Information Processing Systems	76
Conférence Francophone sur l'Apprentissage Automatique	11
Congrès Francophone de Reconnaissance des Formes et Intelligence Artificielle	43
European Conference on Artificial Intelligence	21
European Conference on Machine Learning	5
European Conference on Principles and Practice of Knowledge Discovery in Databases	2
European Semantic Web Conference	5
European Symposium on Artificial Neural Networks	8
Genetic and Evolutionary Computation Conference	71
IEEE Conference on Computer Vision and Pattern Recognition	60
IEEE Conference on Intelligent Transportation Systems	4
IEEE International Conference on Acoustics, Speech and Signal Processing	91
IEEE International Conference on Computer Vision	38
IEEE International Conference on Robotics and Automation	109
IEEE International Conference on Tools with Artificial Intelligence	17
IEEE RSJ International Conference on Intelligent Robots and Systems	110
International Conference on Artificial Evolution	7
International Conference on Autonomous Agents and Multiagent Systems	12
International Conference on Machine Learning	43
International Conference on Practical Applications of Agents and Multi-Agent Systems	5
International Conference on Principles and Practice of Constraint Programming	20
International Conference Parallel Problem Solving from Nature	22
International Joint Conference on Artificial Intelligence	24
International Semantic Web Conference	14
International Workshop on Principles of Diagnosis	5
Journées d'Extraction et Gestion des Connaissances	58
Journées Francophones sur les Systèmes Multi-Agents	7
National Conference on Artificial Intelligence	9
Traitement Automatique du Langage Naturel	18
Workshop on machine learning for System Identification	1
Total	959



Références supplémentaires



Cette partie recense d'autres références, identifiées comme les plus pertinentes pour une lecture approfondie, regroupées en thématiques.

L'IA en général

Cap Digital. **Intelligence Artificielle : technologies, marchés et défis.**

Document # 15-167, www.capdigital.com/publications, 2015

Ernest Davis and Gary Marcus. **Commonsense Reasoning and Commonsense Knowledge in Artificial Intelligence.** Communications Of The ACM Vol. 58 No. 9, 2015

Jonathan Grudin. **AI and HCI: Two Fields Divided by a Common Focus.**

AI magazine, 30(4), 48-57, 2008

Kevin Kelly. **The Three Breakthroughs That Have Finally Unleashed AI On The World.** <http://www.wired.com/2014/10/future-of-artificial-intelligence/>, 2014

Pierre Marquis, Odile Papini, Henri Prade (eds). **Panorama de l'Intelligence Artificielle. ses bases méthodologiques, ses développements.** 3 vols. Cepadue, 2014

Stuart Russell and Peter Norvig. **Artificial Intelligence: A Modern Approach.** <http://aima.cs.berkeley.edu>

Alan Turing. **Intelligent Machinery, a Heretical Theory.** Philosophia Mathematica 4 (3): 256-260, 1996 - Original article from 1951

Terry Winograd. **Shifting viewpoints: Artificial intelligence and human-computer interaction.** Artificial Intelligence 170(18):1256-1258, 2006

Les débats autour de l'IA

Ronald C. Arkin. **The Case for Ethical Autonomy in Unmanned Systems.** Journal of Military Ethics, 9(4), 12/2010

Samuel Butler. **Erewhon.** Free eBooks at Planet eBook.com, 1872

Dominique Cardon. **A quoi rêvent les algorithmes.** Seuil, 2015

Thomas G. Dietterich and Eric J. Horvitz. **Rise of Concerns about AI: Reflections and Directions.** Communications of the ACM | Vol. 58 No. 1, October 2015

Moshe Vardi. **On Lethal Autonomous Weapons.** Communications of the ACM, vol. 58 no. 12, December 2015

L'apprentissage automatique (machine learning)

Martin Abadi et al. **Large-Scale machine learning on Heterogeneous Distributed Systems.** Software available from tensorflow.org, 2015

Christopher Bishop. **Pattern Recognition and Machine Learning**. Springer, 2006

Léon Bottou. **From machine learning to machine reasoning: an essay**, Machine Learning, 94:133-149, January 2014

Yann Le Cun. **The Unreasonable Effectiveness of Deep Learning**. Facebook AI Research & Center for Data Science, NYU. <http://yann.lecun.com>, 2015

Sumit Gulwani, William R. Harris, and Rishabh Singh. **Spreadsheet Data Manipulation Using Examples**. Communications of the ACM, Vol. 55 no. 8, 2012

Michael I. Jordan and Tom M. Mitchell. **Machine learning: trends, perspectives, and prospects**. Science, Vol. 349 Issue 6245, 2015

James Max Kanter and Kalyan Veeramachaneni. **Deep Feature Synthesis: Towards Automating Data Science Endeavors**. International IEEE/ACM Data Science and Advance Analytics Conference, 2015

Tom M. Mitchell. **The Discipline of Machine Learning**. CMU-ML-06-108 School of Computer Science Carnegie Mellon University, 2006

Volodymyr Mnih et al. **Human-level control through deep reinforcement learning**. Nature 518, p. 529-533, 2015

Michèle Sebag. **A tour of Machine Learning: an AI perspective**. AI Communications, IOS Press, 27 (1), p.11-23, 2014

Vision

Peter Auer et al. **A Research Roadmap of Cognitive Vision**.

ECVision : the European Research network for Cognitive Computer Vision Systems. <https://www.researchgate.net>, 2005

Nicholas Ayache. **Des images médicales au patient numérique**, Leçons inaugurales du Collège de France. Collège de France / Fayard, March 2015

Slawomir Bak. **Human Reidentification Through a Video Camera Network**. Computer Vision and Pattern Recognition. PhD, Université Nice Sophia Antipolis, 2012

Svetlana Lazebnik, Cordelia Schmid, Jean Ponce. **Spatial pyramid matching**. Sven J. Dickinson and Alès Leonardis and Bernt Schiele and Michael J. Tarr. Object Categorization: Computer and Human Vision Perspectives, Cambridge University Press, p. 401-415, 2009

Sancho McCann, David G. Lowe. **Efficient Detection for Spatially Local Coding**. Lecture Notes in Computer Science Volume 9008, p. 615-629, 2015

Farhood Negin, Serhan Cosar, Michal Koperski, François Bremond. **Generating Unsupervised Models for Online Long-Term Daily Living Activity Recognition**. Asian conference on pattern recognition (ACPR 2015), 2015

Oriol Vinyals, Alexander Toshev, Samy Bengio & Dumitru Erhan. **Show and Tell: A Neural Image Caption Generator**. <https://arxiv.org/pdf/1502.03044>, 2015

Représentation des connaissances, web sémantique, données

Nathalie Aussenac-Gilles, Fabien Gandon. **From the knowledge acquisition bottleneck to the knowledge acquisition overflow: A brief French history of knowledge acquisition**. International Journal of Human-Computer Studies, Elsevier, vol. 71 (n 2), p. 157-165, 2013

Tim Berners-Lee, James Hendler and Ora Lassila. **The Semantic Web**. Scientific American, May 2001

Meghyn Bienvenu, Balder Ten Cate, Carsten Lutz, Frank Wolter. **Ontology-Based Data Access: A Study through Disjunctive Datalog, CSP, and MMSNP**. ACM Transactions on Database Systems, Vol. 39, No. 4, Article 33, 2014

Fabien Gandon. **The three 'W' of the World Wide Web call for the three 'M' of a Massively Multidisciplinary Methodology**. Valérie Monfort; Karl-Heinz Krempels. 10th International Conference, WEBIST 2014, Barcelona, Spain. Springer International Publishing, 226, Web Information Systems and Technologies, 2014

Janowicz, K.; Hitzler, P.; Hendler, J.; and van Harmelen, F. **Why the Data Train Needs Semantic Rails**. AI Magazine, 36(5-14), 2015

Antonella Poggi et al. **Linking Data to Ontologies**. Journal on data semantics X p. 133-173. Springer-Verlag Berlin, Heidelberg, 2008

Jérôme Euzenat, Pavel Shvaiko. **Ontology matching**. Springer, Heidelberg, 2013

Robotique et véhicules autonomes

J. Christian Gerdes, Sarah M. Thornton, **Implementable Ethics for Autonomous Vehicles**. Autonomes Fahren: Technische, rechtliche und gesellschaftliche Aspekte. Springer, Berlin, 2015

Philippe Morignot, Joshue Perez Rastelli and Fawzi Nashashibi. **Arbitration for Balancing Control between the Driver and ADAS Systems in an Automated Vehicle: Survey and Approach**. 2014 IEEE Intelligent Vehicles Symposium (IV), 2014

Robotique développementale

Antoine Cully, Jeff Clune, Danesh Tarapore & Jean-Baptiste Mouret. **Robots that can adapt like animals**. Nature Vol 521 503-507, 2015

Pierre-Yves Oudeyer. **Developmental Robotics**. Encyclopaedia of the Sciences of Learning, N.M. Seel ed., Springer References Series, Springer, 2012

IA et sciences cognitives

Jacqueline Gottlieb, Pierre-Yves Oudeyer, Manuel Lopes and Adrien Baranes. ***Information-seeking, curiosity, and attention: computational and neural Mechanisms***. Trends in Cognitive Science (2013) 1-9, 2013

Douglas Hofstadter & Emmanuel Sander. ***L'analogie, cœur de la pensée***. Ed. Odile Jacob, 2013

Steels, Luc. ***Self-organisation and selection in cultural language evolution***. In Luc Steels (Ed.), Experiments in Cultural Language Evolution, 1 – 37. Amsterdam : John Benjamins, 2012

Langage naturel, parole, audio, musique

Kenneth Church. ***A Pendulum Swung Too Far***. Linguistic Issues in Language Technology – LiLT. Vol. 2, Issue 4, 2007

Arnaud Dessein, Arshia Cont. ***An information-geometric approach to real-time audio segmentation***. IEEE Signal Processing Letters, Institute of Electrical and Electronics Engineers, 20 (4), p. 331-334, 2013

G. Hinton, L. Deng, D. Yu, G.E. Dahl, A. Mohamed, N. Jaitly, A. Senior, V. Vanhoucke, P. Nguyen, T.N. Sainath, B. Kingsbury. ***Deep neural networks for acoustic modeling in speech recognition: the shared views of four research groups***. IEEE Signal Processing Magazine, 29(6):82-97, 2012

Remerciements

Chercheurs des équipes-projets et des centres Inria qui ont contribué à ce document
(qui ont accordé des entretiens, fourni des textes ou les deux)

Abiteboul Serge, équipe-projet DAHU, Cachan

Ayache Nicholas, responsable de l'équipe-projet ASCLEPIOS, Sophia-Antipolis

Bach Francis, responsable de l'équipe-projet SIERRA, Paris

Beldiceanu Nicolas, responsable de l'équipe-projet TASC, Nantes

Boujemaa Nozha, conseiller Big data auprès du Président d'Inria

Braunschweig Bertrand, directeur du centre de recherche Inria Saclay-Île-de-France

Brémond François, responsable de l'équipe-projet STARS, Sophia-Antipolis

Charpillet François, responsable de l'équipe-projet LARSEN, Nancy

Colliot Olivier, responsable de l'équipe-projet ARAMIS, Paris

Cont Arshia, responsable de l'équipe-projet MUTANT, Paris

Cordier Marie-Odile, équipe-projet LACODAM, Rennes

Crowley James, responsable de l'équipe-projet PERVASIVE INTERACTION, Grenoble

De La Clergerie Eric, équipe-projet ALPAGE, Paris

De Vico Fallani Fabrizio, équipe-projet ARAMIS, Paris

Euzenat Jérôme, responsable de l'équipe-projet EXMO, Grenoble

Gandon Fabien, responsable de l'équipe-projet WIMMICS, Sophia-Antipolis

Giavitto Jean-Louis, équipe-projet MUTANT, Paris

Gilleron Rémi, équipe-projet MAGNET, Lille

Giraudon Gérard, directeur du centre de recherche Sophia-Antipolis Méditerranée

Gravier Guillaume, responsable de l'équipe-projet LINKMEDIA, Rennes

Gros Patrick, directeur du centre de recherche Grenoble-Rhône Alpes

Guittou Pascal, membre de la cellule de veille et prospective

Horaud Radu, responsable de l'équipe-projet PERCEPTION, Grenoble

Manolescu Ioana, responsable de l'équipe-projet CEDAR, Saclay

Moisan Sabine, équipe-projet STARS, Sophia-Antipolis

Mugnier Marie-Laure, responsable de l'équipe-projet GRAPHIK, Montpellier

Nashashibi Fawzi, responsable de l'équipe-projet RITS, Paris

Niehren Joachim, responsable de l'équipe-projet LINKS, Lille

Oudeyer Pierre-Yves, responsable de l'équipe-projet FLOWERS, Bordeaux

Pietquin Olivier, équipe-projet SEQUEL, Lille

Ponce Jean, responsable de l'équipe-projet WILLOW, Paris

Preux Philippe, responsable de l'équipe-projet SEQUEL, Lille

Roussel Nicolas, responsable de l'équipe-projet MJOLNIR, Lille

Sagot Benoit, responsable de l'équipe-projet ALPAGE, Paris

Schmid Cordelia, responsable de l'équipe-projet THOTH, Grenoble

Schoenauer Marc, co-responsable de l'équipe-projet TAO, Saclay

Sebag Michèle, co-responsable de l'équipe-projet TAO, Saclay

Seddah Djame, équipe-projet ALPAGE, Paris

Siegel Anne, responsable de l'équipe-projet DYLISS, Rennes

Sturm Peter, adjoint au directeur scientifique

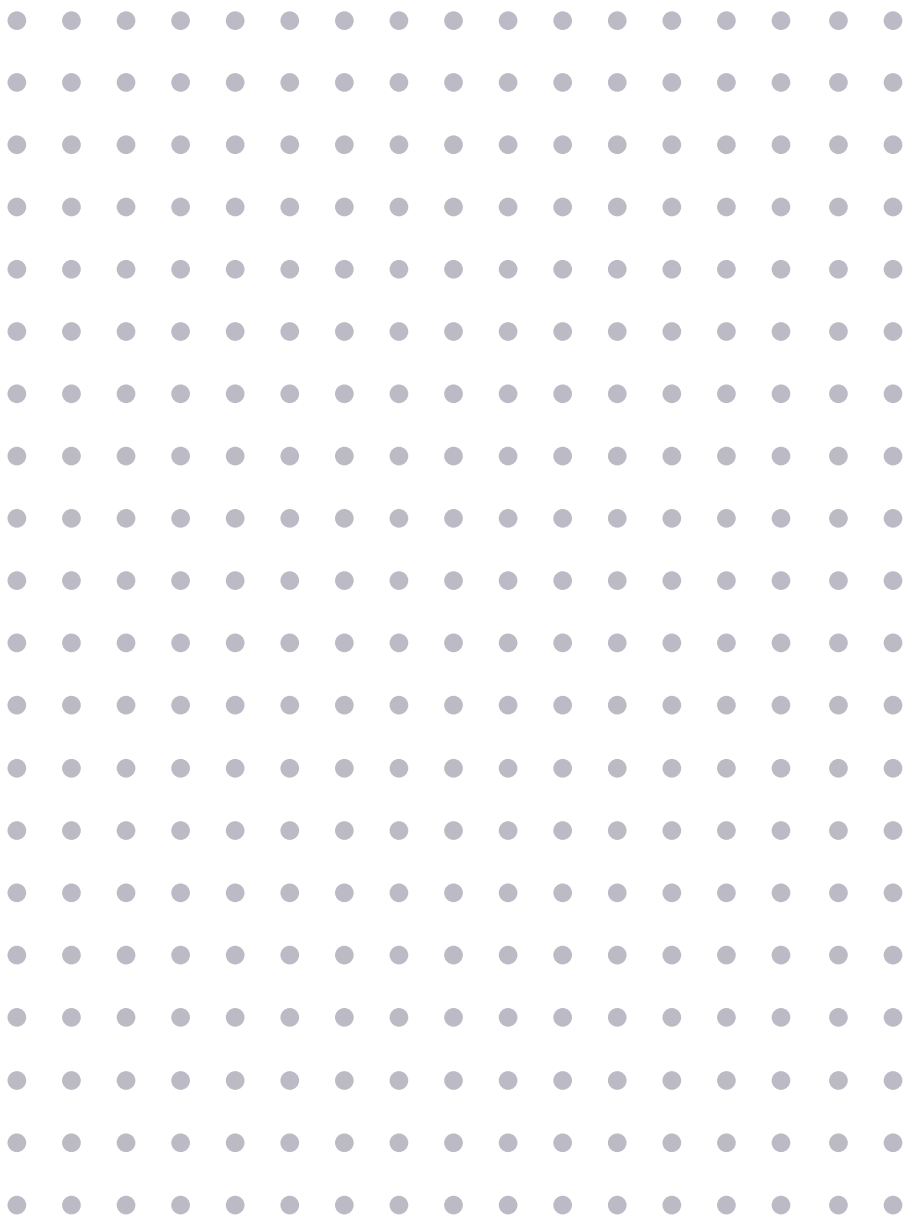
Termier Alexandre, responsable de l'équipe-projet LACODAM, Rennes

Thonnat Monique, directrice du centre de recherche Bordeaux-Sud-Ouest

Tommasi Marc, responsable de l'équipe-projet MAGNET, Lille

Toussaint Yannick, équipe-projet ORPAILLEUR, Nancy

Vincent Emmanuel, équipe-projet MULTISPEECH, Nancy



Domaine de Voluceau, Rocquencourt BP 105
78153 Le Chesnay Cedex, France
Tél. : +33 (0)1 39 63 55 11
www.inria.fr