

English Follow

LUNDI 19 octobre 2026

Lyse Langlois - Garder l'humain dans la boucle : repenser l'agentivité, la décision et la responsabilité dans les systèmes humain-IA

Lyse Langlois est directrice générale de l'Observatoire international sur les impacts sociétaux de l'IA et du numérique (Obvia) depuis 2018 et professeure titulaire au Département des relations industrielles de l'Université Laval. Elle a dirigé pendant huit ans l'Institut d'éthique appliquée (IDÉA) de cette même université et est chercheuse régulière au Centre de recherche interuniversitaire sur la mondialisation du travail (CRIMT) ainsi qu'à l'Institut intelligence et données (IID) de l'Université Laval. En 2025-2026, elle a été professeure invitée à la Korea University avec qui est poursuivi plusieurs projets de collaboration. Ses recherches ont donné lieu à de nombreuses publications scientifiques, ouvrages, articles et rapports sur l'éthique et les transformations du travail liées à la technologie. Sa plus récente contribution porte sur la géopolitique de la science et la diplomatie scientifique, dans un chapitre publié chez Palgrave Macmillan dans *Polycentric Science Diplomacy in an Age of Disruption*, sous la direction de J. Lüder et M. Wählich.

Sarah Brown - Réinterprétations relationnelles de l'IA

La Dre **Sarah M. Brown** est professeure adjointe en informatique à l'Université du Rhode Island. Ses recherches portent sur la manière de rendre les systèmes d'intelligence artificielle meilleurs pour toutes et tous, en étudiant les interactions entre la technologie, les personnes qui la développent et ses impacts sociaux. Ses travaux placent les personnes au centre à toutes les échelles : en développant des interventions qui soutiennent les scientifiques des données dans le développement de systèmes utilisés dans le monde réel, et en développant directement des techniques qui intègrent des valeurs sociales dans les systèmes d'intelligence artificielle. Elle est titulaire d'un baccalauréat, d'une maîtrise et d'un doctorat en génie électrique de la Northeastern University. La Dre Brown est récipiendaire de la Graduate Research Fellowship et du prix CAREER de la National Science Foundation (NSF).

Résumé : De nombreuses valeurs généralement attribuées à l'IA trouvent leur origine dans le langage, et certaines peuvent être façonnées au moyen de stratégies d'alignement post-entraînement. Dans d'autres cas, certaines valeurs sont rattachées à des contraintes techniques. Cependant, souvent, ces valeurs ne sont pas fondamentales à l'objet mathématique ou computationnel lui-même, mais découlent d'une philosophie des sciences dominante. La valeur est imposée dans l'interprétation de l'objet, puis lui reste associée lors de la mise en œuvre de technologies réelles. Ces valeurs sont souvent présentées comme inévitables dans la poursuite du développement technologique, en partie parce que la philosophie des sciences n'occupe pas une place centrale dans la formation dans les domaines computationnels. En réalité, presque

toutes les idées deviennent contestables lorsqu'elles sont examinées en profondeur. Dans cette présentation, je mettrai en évidence quelques exemples de ce phénomène. Je proposerai ensuite un fondement philosophique alternatif : le réalisme relationnel. Je justifierai son utilité et sa pertinence pour l'IA afin d'en présenter les principes fondamentaux, puis je conclurai par des exemples de manières de réinterpréter, depuis cette perspective, les contraintes techniques liées à l'IA, permettant ainsi d'utiliser l'IA de façons qui favorisent des valeurs souvent sous-représentées dans l'IA.

Marc-Antoine Dilhac - Regard sur l'évolution des cadres de gouvernance

Marc-Antoine Dilhac est professeur d'éthique et de philosophie politique à l'Université de Montréal, où ses travaux portent sur l'éthique et la gouvernance de l'intelligence artificielle. En 2017, il a lancé et dirigé le processus de co-construction de la Déclaration de Montréal pour un développement responsable de l'intelligence artificielle. Il est directeur de la gouvernance et de la collaboration internationale à l'OBVIA et membre du Mila, où il a animé une chaire Canada-CIFAR en éthique de l'IA. Il intervient en tant qu'expert de l'UNESCO au sein du réseau « Experts Without Borders » sur l'éthique de l'IA et a dirigé la mise en œuvre de la méthodologie d'évaluation de l'état de préparation (RAM) de l'UNESCO en Ouzbékistan et au Québec. Il a auparavant siégé au Conseil consultatif sur l'IA du gouvernement du Canada et a été l'expert chargé du processus de consultation du Conseil de l'Europe sur la Convention-cadre relative à l'intelligence artificielle, aux droits de l'homme, à la démocratie et à l'État de droit.

Alison Gillwald – Remédier à la répartition inégale des préjudices et des opportunités liés à l'IA en donnant la priorité aux droits de deuxième et de troisième génération et à la valeur de l'« ubuntu »

Alison Gillwald (PhD) est la directrice fondatrice de Research ICT Africa (RIA) et professeure associée à la Nelson Mandela School of Public Governance de l'Université du Cap, où elle a dirigé un programme doctoral panafricain sur les politiques numériques.

Alors que RIA était partenaire de connaissances de la présidence sud-africaine du G20, elle a agi à titre de conseillère technique auprès du Groupe de travail sur l'économie numérique et de la Task Force de haut niveau sur l'intelligence artificielle, la gouvernance des données et l'innovation pour le développement durable. Elle a également coprésidé la Task Force du T20 sur la transformation numérique. Elle est membre du groupe de travail T7 de la présidence française du G20 sur l'IA et les puissances moyennes et siège au sein de l'organe consultatif académique de l'Union internationale des télécommunications sur les technologies avancées. Elle a été active au sein du Partenariat mondial sur l'intelligence artificielle (GPAI), codirigeant les projets sur la transparence algorithmique dans le secteur public et sur la justice des données.

Alison a dirigé l'assistance technique de RIA auprès de la Commission de l'Union africaine, en facilitant les travaux de l'équipe internationale et en rédigeant le Cadre de politique des données

de l'UA et son plan de mise en œuvre. Actuellement, elle est chercheuse principale du projet de la Banque africaine de développement sur les infrastructures publiques numériques pour l'Afrique, qui examine la réduction des risques liés aux investissements bancaires à travers le prisme de l'économie politique.

Ses publications les plus récentes comprennent le document de cadrage pour le Dialogue du Club de Madrid sur la gouvernance des biens publics mondiaux dans un monde fragmenté : biens publics numériques et infrastructures, ainsi que Gillwald, A., et Fuller, J. (2026), Transformation numérique et limites du G20 : enseignements tirés des présidences du G20 des économies émergentes.

MARDI 20 octobre 2026

Alexandra Bensamoun – Titre à venir

Alexandra Bensamoun est professeure de droit à l'Université Paris-Saclay, spécialiste de propriété intellectuelle et de régulation du numérique, directrice du Master 2/LLM Propriété intellectuelle fondamentale et technologies numériques, en codiplômation avec l'Université Laval (Québec) et l'Université Autonome de Madrid (Espagne). Experte pour l'UNESCO et personnalité qualifiée au CSPLA (ministère de la Culture français), elle a siégé en France à la Commission interministérielle de l'intelligence artificielle, qui a rendu un rapport sur la stratégie française au Président de la République, IA : notre ambition pour la France (mars 2024). Ses rapports au ministère de la Culture portent notamment sur la rémunération des contenus culturels utilisés par les IA (à l'origine de la proposition de loi n° 220 votée à l'unanimité au Sénat en avril 2026) ou encore sur le statut des productions de l'IA. Elle est l'auteur de nombreuses recherches individuelles ou collectives, notamment un Traité collectif sur le Droit de l'intelligence artificielle (LGDJ/Lextenso, 3e éd., 2026).

Maurice Jones – Information à venir

Antoine Henry - Opérationnaliser des valeurs dans les systèmes de recommandations culturels : la création d'un Fairness Score

Antoine Henry est maître de conférences en sciences de l'information et de la communication et membre de l'axe 4 de GERIICO dont la thématique est la circulation de l'information et l'organisation des connaissances.

Ses travaux de recherche portent, principalement, sur les questions relatives au numérique (IA, communs numériques, éthique des algorithmes), à la transformation des organisations et à l'intelligence collective.

Il est aussi :

- membre du conseil d'administration du chapitre français de l'association savante ISKO,

- membre du groupe de recherche du Centre Internet et Société (UPR 2000 du CNRS) au sein duquel il co-anime le groupe de travail IA, art et créativité.

Résumé : L'équité algorithmique représente un défi majeur dans le développement des systèmes de recommandation, particulièrement dans le secteur musical. Cette intervention propose une approche novatrice aux cadres actuels d'équité algorithmique en introduisant un nouveau modèle pour évaluer l'équité. Nous développons une méthode permettant de synthétiser neuf dimensions complexes pour influencer la boucle de rétroaction des algorithmes. À travers cette proposition du Fairness Score, nous visons à opérationnaliser l'équité et à offrir aux utilisateurs un contrôle significatif sur leur expérience et leur pratique culturelle.

Christophe Abrassard – Informations à venir

Jocelyn Maclure - L'éthique et la gouvernance de l'IA: Pourquoi l'approche de l'alignement sur les valeurs est mal avisée

Jocelyn Maclure est professeur de philosophie et titulaire de la Chaire Jarislowsky sur la nature humaine et la technologie à l'Université McGill. Ses champs de recherche et d'enseignement sont l'éthique, la philosophie politique, la philosophie de l'IA et l'épistémologie sociale. Ses travaux récents ont été publiés dans des revues comme *Minds & Machines*, *AI & Ethics*, *Digital Society* et *AI & Society*. Il a récemment été professeur invité aux universités de Bonn et de Grenoble, et a présidé la Commission de l'éthique en science et en technologie du Québec de 2017 à 2024.

Le « problème de l'alignement sur valeurs humaines » (PAV) constitue sans doute l'approche la plus influente chez les philosophes et les chercheurs en intelligence artificielle dans les champs de l'éthique et de la gouvernance de l'IA. Le PAV est centré sur la question de savoir comment concevoir des systèmes d'IA avancés dont les objectifs et les résultats sont alignés sur les valeurs humaines. S'il n'y a évidemment rien de problématique à vouloir ancrer le développement des IA de pointe dans des valeurs humaines partagées, cette approche laisse de côté les aspects les plus difficiles de l'éthique pratique et de l'élaboration des politiques publiques. Après avoir examiné les principaux problèmes auxquels se heurte le PAV (les problèmes du pluralisme, de la sous-détermination et de la motivation), je proposerai les grandes lignes d'une approche alternative que j'appelle l'« intégrité normative ».

MERCREDI 21 octobre 2026

Jean-Gabriel Ganascia - IA agentique et/ou garantie humaine?

Ingénieur et philosophe de formation initiale, **Jean-Gabriel Ganascia** s'est très tôt orienté vers l'informatique et l'intelligence artificielle. Il a soutenu en 1983, à l'université Paris-Saclay, une thèse de doctorat sur les systèmes experts puis en 1987, toujours à l'université Paris-Saclay, une

thèse d'État sur l'apprentissage machine. Présentement, professeur émérite à la faculté des sciences de Sorbonne Université, membre honoraire de l'institut universitaire de France, EurAI fellow et chercheur au LIP6, il poursuit des travaux sur l'apprentissage machine, les humanités numériques, la détection d'Infox et l'éthique du numérique. Dans le passé, il a coordonné la recherche française en sciences cognitives en créant puis en animant le Groupement d'Intérêt Scientifique « Sciences de la cognition ». Il est aussi à l'origine du Labex OBVIL (Observatoire de la vie littéraire) qui coordonnait de 2011 à 2021 les efforts de son équipe au sein du LIP6 et des laboratoires de littérature de Sorbonne Université sur le versant littéraire des humanités numériques. Outre ses activités de recherche, il préside le comité d'éthique de France Travail, le comité d'éthique de Docaposte et le comité d'orientation du CHEC (Cycle des Hautes Études de la Culture). Il est aussi membre du COMEST (Commission mondiale d'éthique des connaissances scientifiques et des technologies de l'UNESCO) Enfin, il a présidé le comité d'éthique du CNRS (COMETS) de 2016 à 2021 et l'AFAS (Association Française pour l'Avancement des Sciences) de 2022 à 2026. Au cours de sa carrière, il a publié plus de 600 articles dans des actes de conférences, des livres et des revues scientifiques. Il est aussi l'auteur d'une douzaine d'ouvrages destinés au grand public et il tient depuis 2017 plusieurs chroniques éthiques régulières, l'une dans la revue La Recherche, l'autre dans la revue Sciences et avenir.

Résumé : Suite au boum de l'IA générative, nous entrons, depuis la fin 2025, dans l'âge de l'IA agentique. L'idée est ancienne: un agent capte de l'information, texte, image ou son, l'interprète, prend seul des décisions sur la foi de ses entrées, et agit (d'où le nom d'agent) en vue de réaliser les buts qu'on lui a fixés. Comment ne pas s'en inquiéter: l'IA agentique risque de nous mettre sous tutelle en nous court-circuitant. Dans l'éthique du soin, on a introduit la notion de « garantie humaine » pour s'assurer qu'aucune décision ne se prenne sans qu'un personnel de santé ne l'ait préalablement approuvée. Pourquoi ne pas toujours procéder ainsi? Parce qu'on n'a pas le temps et que l'humain n'est pas une garantie! Errare humanum est: nous nous trompons parfois plus souvent que les automates. L'étude rétrospective d'accidents d'avions ou de centrales nucléaires le montre. Comment résoudre en amont les potentiels conflits d'autorité entre humains et machines et décider de celui qui doit l'emporter? On choisit parfois, sur la base d'arguments rationnels, de laisser la main à l'agent artificiel. Mais, cette délégation de responsabilité ne vaut que pour un temps et doit rester révoable à tout moment. Du moins, l'espère-t-on... Ce sont toutes ces questions qu'il s'agira d'aborder ici afin de se demander ce qui, dans chaque situation, reste admissible et ce qu'il faut condamner.

Maroussia Lévesque - Les droits de la personne, boussole de la gouvernance de l'IA en pratique

Maroussia Lévesque a rejoint la Faculté de droit de l'Université Queen's en tant que professeure adjointe en droit et IA en 2026. Ses recherches portent sur la gouvernance de l'IA, en examinant la conception de la réglementation, les déséquilibres de pouvoir et la concentration du marché dans l'industrie de l'IA, ainsi que les rôles respectifs des acteurs publics et privés dans l'écosystème de gouvernance. Ses projets actuels analysent la manière dont les structures organisationnelles des laboratoires d'IA de pointe façonnent leur

comportement dans la pratique, ainsi que les implications environnementales et les usages judiciaires de l'IA.

Elle est affiliée au Berkman Klein Center for Internet & Society, chercheuse principale au Centre pour l'innovation dans la gouvernance internationale (CIGI) et collaboratrice du projet Abundant Intelligences, dirigé par des Autochtones.

Avant de rejoindre l'Université Queen's, la professeure Lévesque a obtenu son doctorat en sciences juridiques (SJD) à la Faculté de droit de Harvard, où elle a enseigné la gouvernance comparée de l'IA (avec la professeure Martha Minow) et a été chercheuse au Weatherhead Center for International Affairs. Auparavant, elle a dirigé des travaux sur l'IA et les droits de la personne à Affaires mondiales Canada, a été auxiliaire juridique auprès du juge en chef de la Cour d'appel du Québec, a collaboré avec l'Algorithmic Accountability Lab d'Amnistie internationale et a contribué à la mise en place du Partenariat mondial sur l'intelligence artificielle. Elle est membre du Barreau du Québec.

Les travaux de recherche de la professeure Lévesque ont notamment été publiés dans le NYU Journal of Legislation & Public Policy, le University of Illinois Journal of Law, Technology & Policy et les actes de NeurIPS. Elle a écrit pour The Guardian, le Toronto Star, La Presse, The Walrus et Tech Policy Press.

Résumé : Les droits de la personne sont une avenue prometteuse pour guider la gouvernance de l'IA. Leur vocabulaire établi et légitimité internationale sont appuyés par une architecture institutionnelle bien rodée. Par ailleurs, les Principes directeurs des Nations Unies relatifs aux entreprises et aux droits de l'homme étendent ce vocabulaire aux acteurs privés, articulant une responsabilité des entreprises de respecter les droits de la personne distincte de l'obligation contraignante des États de les protéger. Cela dit, demander sur quelles valeurs la gouvernance de l'IA devrait reposer est plus facile que de cerner ce qui rend ces valeurs utiles en pratique.

Cette présentation examinera le potentiel et les limites des droits de la personne comme guide de la gouvernance de l'IA. Au début de l'ascension des modèles d'apprentissage automatique en 2018, la Déclaration de Toronto a traduit l'exigence de diligence raisonnable des Principes directeurs en étapes concrètes, offrant des balises opérationnelles pour cerner les risques, les prévenir et les atténuer, suivre les réponses et offrir des mécanismes de reddition de comptes. Plus récemment, le Conseil de surveillance de Meta (Oversight Board) s'est décrit comme un mécanisme non judiciaire indépendant au sens des Principes directeurs, invoquant les droits de la personne dans son examen de la modération automatisée des contenus. Ces exemples démontrent que les droits de la personne peuvent être mis en œuvre de manière à faire progresser la gouvernance de l'IA.

Par contre, les droits de la personne comportent également des limites. Des engagements non contraignants peuvent conférer un dividende de légitimité inapproprié, les entreprises drapant leurs démarches dans le langage de normes malléables tout en esquivant leur mise en application. Les cadres volontaires permettent également de faire perdurer le statu quo, soit l'absence de réglementation contraignante, normalisant des pratiques problématiques tandis que les attentes sociales s'ajustent. Enfin, les engagements souples cèdent sous la pression

concurrentielle, comme le suggèrent les récents reculs à l'égard d'engagements en matière de droits de la personne.

Du côté des développements en cours, les règles de protection des investisseurs, qui exigent une divulgation rigoureuse des risques, pourraient faire progresser la diligence raisonnable à mesure que les laboratoires d'IA de pointe entrent en bourse. Le projet de traité contraignant sur les entreprises et les droits de la personne obligerait par ailleurs les États à régler, plutôt que de demander aux entreprises de se restreindre elles-mêmes. En définitive, énoncer des valeurs importe moins que la mise en place d'incitatifs et d'obligations contraignantes qui leur donnent prise.

Richard Khoury - L'IA Générative et la guerre contre la diversité linguistique

Richard Khoury a obtenu son baccalauréat et sa maîtrise en Génie Électrique et Génie Informatique de l'Université Laval (Québec, QC) en 2002 et 2004 respectivement, et son doctorat en Génie Électrique et Génie Informatique de l'Université de Waterloo (Waterloo, ON) en 2007. De 2008 à 2016, il a été professeur adjoint, puis professeur agrégé, au département de Génie Logiciel à l'Université Lakehead. En 2016, il s'est joint à l'Université Laval à titre de professeur agrégé, avant d'être nommé professeur titulaire en 2026. De 2021 à 2023 il a également été président de l'Association pour l'intelligence artificielle au Canada. Ses intérêts de recherche sont en forage de données et en traitement du langage naturel, et incluent de plus la gestion des connaissances, l'apprentissage automatique, et l'intelligence artificielle.

Résumé : Les outils d'IA Génératives sont aujourd'hui disponibles dans plus d'une centaine de langues différentes. À première vue, ceci semble être une victoire pour la diversité linguistique : l'intelligence artificielle offre de rejoindre les utilisateurs à travers le monde dans leur langue maternelle. Mais sous cette image positive se cache un problème insidieux. L'intelligence artificielle ne connaît et ne comprend que le dialecte dominant de ces langues, et impose son utilisation aux locuteurs. Cette colonisation linguistique par l'IA menace d'effacer la diversité linguistique mondiale et de supprimer des communautés culturelles minoritaires. Dans cette présentation nous étudierons ce problème, ses origines, ses symptômes, et ses solutions.

Colette Brin – Informations à venir

JEUDI 22 octobre 2026

Emmanuelle Boutin-Gilbert – Information à venir

Pierre Larouche – L'IA, l'ouverture et la flexibilité

Expert en droit de la concurrence et en gouvernance économique ainsi qu'en responsabilité civile, dans les traditions de droit civil et de common law, **Pierre Larouche** est professeur titulaire en droit et innovation à la Faculté de droit de l'Université de Montréal.

Diplômé des facultés de droit des universités McGill, de Bonn et de Maastricht, le Pr Larouche était depuis 2002 professeur titulaire de droit de la concurrence à l'Université de Tilburg, aux Pays-Bas, où il a notamment cofondé le Tilburg Law and Economics Center (TILEC), devenu l'un des plus grands centres de recherche en gouvernance économique au monde. Le Pr Larouche a aussi conçu et lancé un programme de 1er cycle innovateur, le Bachelor Global Law, à l'Université de Tilburg et lancé l'option doctorale « Innovation, science, technologie et droit » à l'Université de Montréal. Lors de son mandat de vice-doyen, il a aussi mené à bien la réforme du programme de droit de 1er cycle. Il a également enseigné au Collège d'Europe, à Bruges, tout en étant professeur invité dans de nombreuses universités en Amérique (Northwestern, Pennsylvanie), en Europe (Sciences Po, Bonn) et en Asie (Singapour).

Le Pr Larouche a publié plus d'une soixantaine de monographies, d'articles et de contributions scientifiques, et ses travaux, cités par la Cour de justice de l'Union européenne et la Cour suprême du Royaume-Uni, ont influencé les politiques de la Commission européenne en matière de communications électroniques, de concurrence et de normalisation.

En plus d'avoir pratiqué le droit en cabinet privé, il a été auxiliaire juridique de l'honorable Charles D. Gonthier à la Cour suprême du Canada.

Résumé : Depuis des décennies, la recherche en sciences sociales, dans un cadre libéral et démocratique, cherche à incorporer les considérations liées à l'information dans le cadre analytique. S'est ensuivie une meilleure compréhension des effets dynamiques, comme en fait foi l'apparition, dans le discours académique et politique, des imperfections et incertitudes informationnelles, à travers des valeurs ou principes comme l'ouverture (au sens du non-déterminisme) et la flexibilité. L'IA, en phase avec l'amélioration de ses performances et sa diffusion dans l'économie et la société, affectera nécessairement ce discours. Cette présentation explore les implications de l'IA sur l'ouverture et la flexibilité dans le contexte plus spécifique du fonctionnement des marchés. Les nombreuses théories économiques du marché peuvent être classifiées selon qu'elles lui attribuent une fin d'efficacité ou de découverte. Selon la perspective théorique, l'impact de l'IA diffère. Si le marché mène à l'efficacité, l'IA pourrait en améliorer le fonctionnement, mais non sans contrepartie. Si par contre, le marché est un outil de découverte, l'IA pourrait en saper les fondements. Dans les deux cas, une intervention juridique pourrait s'avérer nécessaire pour éviter le scénario catastrophe où la société cesse de récolter les bénéfices du marché. Cela ramène à l'avant-scène une autre controverse théorique sur les marchés, qui oppose naturalisme et constructivisme. La présentation se conclut en revenant sur le terrain plus général du discours académique et politique sur l'ouverture et la flexibilité.

Siddarth De Souza - A new theory for the Digital Rule of Law

Siddharth Peter de Souza est professeur adjoint en IA et société au Centre for Interdisciplinary Methodologies de l'Université de Warwick. Ses travaux explorent la manière dont les données sont gouvernées à l'échelle mondiale dans des contextes contestés et pluralistes.

Il est le fondateur de Justice Adda, une entreprise sociale spécialisée en droit et en design en Inde, où il travaille sur des questions liées à l'accès à la justice et à l'autonomisation juridique.

Siddharth a travaillé comme consultant en politiques numériques auprès du Programme des Nations Unies pour le développement et de l'Organisation internationale de droit du développement. Il a auparavant été chercheur postdoctoral au sein du Global Data Justice Project au Tilburg Institute for Law, Technology and Society.

Résumé : Dans cette présentation, j'examinerai pourquoi le concept d'État de droit a besoin d'une nouvelle théorie. Depuis des années, il est critiqué par les chercheurs en droit et développement comme étant un outil servant à justifier le pillage par l'Occident, comme un instrument d'hégémonie juridique dissimulé sous la forme de projets de construction de l'État, et comme un mécanisme de transplantation juridique qui accorde peu d'attention à la nature pluraliste des systèmes juridiques à travers le monde.

Le Secrétaire général des Nations Unies a appelé à une nouvelle vision de l'État de droit qui mette l'accent sur le développement d'une approche centrée sur les personnes quant à la manière dont celui-ci est conceptualisé ainsi qu'institutionnalisé.

Avec la numérisation croissante des services publics dans des domaines tels que la santé, la finance et l'éducation, de nouvelles préoccupations ont émergé pour l'État de droit, notamment en matière de vie privée, de surveillance et de captation des institutions publiques par les entreprises.

Olivier Servais – Informations à venir

VENDREDI 23 octobre 2026

Georges Azzaria - L'effacement de la singularité juridique des individus par les algorithmes

Georges Azzaria est professeur à la Faculté de droit de l'Université Laval à Québec. Ses premières recherches ont porté sur les rapports entre l'art et le droit d'auteur, ainsi que sur le statut socio-économique des artistes. Depuis quelques années, il s'intéresse plus particulièrement à l'intelligence artificielle sous l'angle de la propriété intellectuelle et de la régulation juridique. De 2017 à 2025, il a été le directeur de l'École d'art de son université.

Résumé : La créativité humaine valorisée par la propriété intellectuelle et la consécration de la personnalité juridique ont accordé à l'individu, à partir du 18^e siècle en Occident, une place centrale dans le droit. Or, la capacité de l'intelligence artificielle générative à remplacer l'humain opère aujourd'hui un renversement qui tend à effacer cette singularité et reléguer l'individu au rôle d'un relais informatique.

MONDAY october 19

Lyse Langlois - Keeping the Human in the Loop: Rethinking Agency, Decision-Making and Responsibility in Human–AI Systems

Lyse Langlois has been Executive Director of the International Observatory on the Societal Impacts of AI and Digital Technology (Obvia) since 2018 and is a Full Professor in the Department of Industrial Relations at Université Laval. She directed the Institute of Applied Ethics (IDÉA) at the same university for eight years and is a regular researcher at the Interuniversity Research Centre on Globalization and Work (CRIMT), as well as at Université Laval's Institute Intelligence and Data (IID). In 2025–2026, she was a Visiting Professor at Korea University, with which she is pursuing several collaborative projects.

Her research has resulted in numerous scientific publications, books, articles, and reports on ethics and transformations of work related to technology. Her most recent contribution focuses on the geopolitics of science and science diplomacy, in a chapter published by Palgrave Macmillan in *Polycentric Science Diplomacy in an Age of Disruption*, edited by J. Lüder and M. Wählich.

Sarah Brown - Relational Reinterpretations of AI

Dr. **Sarah M Brown** is Assistant Professor of Computer Science at the University of Rhode Island. Her research focuses on how to make AI systems better for all by studying how the technology, the people who make it, and its social impacts interact. Her work centers people at all scales: developing interventions that support data scientists developing real world systems and directly developing techniques that incorporate social values into AI systems. She holds BS, MS, and PhD degrees in Electrical Engineering from Northeastern University. Dr. Brown is a recipient of the NSF Graduate Research Fellowship and CAREER awards

Abstract : Many values typically attributed to AI are based in the language and some can be shaped with post-training alignment strategies. In other cases, certain values are traced back to technical constraints. However, often these values are not fundamental to the mathematical or computational object itself, but derive to a dominant philosophy of science. The value is imposed in the interpretation of the object and then stick with it through implementation of real technologies. These values are often presented as inescapable in the pursuit of building technology, in part, because philosophy of science is not central in education of computational fields. In reality, nearly all ideas become contestable when examined deeply. In this talk, I will highlight some examples of phenomenon. Then I will pose an alternative philosophical foundation: relational realism. I will motivate its utility and appropriateness for AI to introduce the key tenets and conclude with examples of ways to reinterpret technical constraints relating to AI from this perspective that enable using AI in ways that promote values often under-represented in AI.

Marc-Antoine Dilhac - The Evolution of AI Governance Frameworks: From the Statement of Principles to Their Implementation

Marc-Antoine Dilhac is Professor of Ethics and Political Philosophy at the Université de Montréal, where his work focuses on the ethics and governance of artificial intelligence. In 2017, he initiated and led the co-construction process for the Montréal Declaration for a Responsible Development of Artificial Intelligence. He is Director of Governance and International Collaboration at OBVIA and a member of Mila, where he held a Canada–CIFAR AI Chair in AI Ethics. He serves as a UNESCO expert in the AI Ethics Experts Without Borders network and has led implementations of UNESCO’s Readiness Assessment Methodology (RAM) in Uzbekistan and Québec. He previously served on the Government of Canada’s AI Advisory Council and was the expert responsible for the Council of Europe’s consultation process on the Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law.

Alison Gillwald - Redressing the uneven distribution of AI harms and opportunities through prioritising second and third generation rights and value of ‘ubuntu’

Alison Gillwald (PhD) is the founding director of Research ICT Africa (RIA) and is an Adjunct Professor at the University of Cape Town’s Nelson Mandela School of Public Governance, where she has run a Pan African digital policy doctoral program. With RIA a knowledge partner to South Africa’s G20 Presidency, she served as a technical advisor to the Digital Economy Work Group and the High-Level Task Force on Artificial Intelligence, Data Governance, and Innovation for Sustainable Development. She also co-chaired the T20 Task Force on Digital Transformation. She is a member of the French G20 Presidencies’ T7 working group on AI and Middle Powers and serves on the International Telecommunications Union’s Academic Advisory Body on Enhanced Technologies. She has been active in the Global Partnership on Artificial Intelligence (GPAI), co-leading the Algorithmic Transparency in the Public Sector and Data Justice projects.

Alison led RIA’s technical assistance to the African Union Commission, facilitating the International Task Team and drafting the AU Data Policy Framework and Implementation Plan. Currently, she is the principal investigator on the African Development Bank project on Digital Public Infrastructures for Africa, which examines derisking bank investments through a political economy lens. Her latest publications include the framing paper for the Club de Madrid Dialogue on Governance of Global Public Goods in a Fragmented World: Digital Public Goods and Infrastructure and Gillwald, A., and Fuller, J. (2026). Digital Transformation and the Limits of the G20: Lessons from the Emerging Economy G20 Presidencies.

TUESDAY October 20

Alexandra Bensamoun - Title to come

Alexandra Bensamoun is a Professor of Law at Paris-Saclay University, specializing in intellectual property and digital regulation, and Director of the Master 2/LLM in Fundamental Intellectual Property and Digital Technologies, a joint degree programme with Université Laval (Quebec) and the Autonomous University of Madrid (Spain).

An expert for UNESCO and a qualified member of the CSPLA (French Ministry of Culture), she served in France on the Interministerial Commission on Artificial Intelligence, which submitted a report on France's strategy to the President of the Republic, AI: Our Ambition for France (March 2024). Her reports to the Ministry of Culture notably address the remuneration of cultural content used by AI systems (which led to Bill No. 220, unanimously adopted by the Senate in April 2026), as well as the status of AI-generated productions.

She is the author of numerous individual and collective research works, including a collective treatise on Artificial Intelligence Law (LGDJ/Lextenso, 3rd ed., 2026).

Maurice Jones - Information to come

Antoine Henry - Operationalizing Values in Cultural Recommendation Systems: The Creation of a Fairness Score

Antoine Henry is an Associate Professor in Information and Communication Sciences and a member of GERIICO's Axis 4, whose theme is the circulation of information and the organization of knowledge.

His research focuses mainly on issues related to digital technologies (AI, digital commons, algorithmic ethics), organizational transformation, and collective intelligence.

He is also:

- a member of the Board of Directors of the French chapter of the scholarly association ISKO;
- a member of the research group of the Centre for Internet and Society (CNRS UPR 2000), within which he co-leads the working group on AI, art, and creativity.

Abstract: Algorithmic fairness represents a major challenge in the development of recommendation systems, particularly in the music sector. This presentation proposes an innovative approach to current algorithmic fairness frameworks by introducing a new model for assessing fairness. We develop a method for synthesizing nine complex dimensions in order to influence the feedback loop of algorithms. Through this proposed Fairness Score, we aim to operationalize fairness and provide users with meaningful control over their experience and cultural practice.

Christophe Abrassard – Information to come

Jocelyn Maclure - Why the Value-Alignment Approach is Misguided

Jocelyn Maclure is Professor of Philosophy and the Jarislowsky Chair in Human Nature and Technology at McGill University. His research and teaching interests include ethics, political philosophy, the philosophy of AI, and social epistemology. His recent work has appeared in journals such as *Minds & Machines*, *AI & Ethics*, *Digital Society*, and *AI & Society*. He has recently been a visiting professor at the Universities of Bonn and Grenoble, and served as Chair of Quebec's Commission de l'éthique en science et en technologie from 2017 to 2024.

The “value-alignment problem” (VAP) is arguably the most influential approach among philosophers and AI researchers in the fields of AI ethics and governance. VAP is centered on the question of how to build advanced AI systems whose goals and outputs are aligned with human values. While there is obviously nothing wrong with grounding the development of frontier AI in shared human values, this approach leaves out the most difficult aspects of practical ethics and policy making. After reviewing the main problems facing VAP (the pluralism, underdetermination and motivation problems), I will sketch out an alternative approach, which I call “Normative Integrity”.

WENESDAY October 21

Jean-Gabriel Ganascia - Agentic AI and/or Human Oversight?

An engineer and philosopher by training, **Jean-Gabriel Ganascia** turned his attention to computer science and artificial intelligence at an early stage in his career. In 1983, he completed a doctoral thesis on expert systems at Paris-Saclay University, followed in 1987 by a thèse d'État on machine learning at the same institution. He is currently Professor Emeritus of Computer Science at Sorbonne University and a EurAI Fellow. He is also an honorary member of the Institut Universitaire de France and a member of LIP6 (Sorbonne University's Laboratory of Computer Science). His current research activities focus on artificial intelligence, computational ethics, computer ethics, the detection of fake news and digital humanities. He has previously contributed to French academic research in cognitive science by creating and leading the scientific interest group 'Sciences de la cognition'. He also initiated the OBVIL Labex (Observatoire de la vie littéraire – Observatory of Literary Life), which coordinated the efforts of his team at LIP6 and the literary teams at Sorbonne University in the field of digital humanities. Alongside his research activities, he chairs the Ethics Committee of France Travail (the French employment agency), the Ethics Committee of Docaposte and the Steering Committee of the CHEC (Cycle des Hautes Études de la Culture). He is also a member of COMEST (the UNESCO World Commission on the Ethics of Scientific Knowledge and Technology). He was President of the CNRS Ethics Committee (COMETS) from 2016 to 2021, and President of the Association Française pour l'Avancement des Sciences (AFAS) from 2022 to 2026. Throughout his career, he has published over 600 papers in conference proceedings, scientific journals, and books. He is also the author of a dozen books for a general audience, and since 2017 he has been writing several regular columns on ethics, one in the journal *La Recherche* and the other in the journal *Sciences et avenir*.

Abstract: Following the boom in generative AI, we have been entering the age of agentic AI since the end of 2025. The idea is an old one: an agent captures information—text, image, or

sound—interprets it, makes decisions on its own on the basis of its inputs, and acts (hence the name “agent”) in order to achieve the goals that have been set for it. How can we not be concerned about this? Agentic AI risks placing us under its control by bypassing us. In the ethics of care, the notion of “human guarantee” has been introduced to ensure that no decision is made without first having been approved by a healthcare professional. Why not always proceed in this way? Because we do not have the time, and because humans are not a guarantee! Errare humanum est: we sometimes make mistakes more often than automated systems do. Retrospective studies of accidents involving airplanes or nuclear power plants show this. How can we resolve, in advance, potential conflicts of authority between humans and machines and decide which one should prevail? Sometimes, on the basis of rational arguments, we choose to leave control to the artificial agent. But this delegation of responsibility is only valid for a limited period of time and must remain revocable at any moment. At least, we hope so... These are all the questions that will be addressed here in order to consider what, in each situation, remains acceptable and what must be condemned.

Maroussia Lévesque - Human Rights: the North Star of AI Governance in practice

Maroussia Lévesque joined Queen’s Law as an Assistant Professor in Law & AI in 2026. Her research focuses on AI governance, examining regulatory design, power imbalances and market concentration in the AI industry, and the respective roles of public and private actors in the governance ecosystem. Her current projects analyse how the corporate structures of frontier AI labs shape their behaviour in practice, along with the environmental implications and judicial uses of AI.

She is an affiliate at the Berkman Klein Center for Internet & Society, a Senior Fellow at the Centre for International Governance Innovation (CIGI), and a collaborator on the Indigenous-led Abundant Intelligences project.

Prior to joining Queen’s, Professor Lévesque completed her SJD at Harvard Law School, where she lectured in Comparative AI Governance (with Professor Martha Minow) and was a Fellow at the Weatherhead Center for International Affairs. She previously led AI and human rights work at Global Affairs Canada, clerked for the Chief Justice at the Quebec Court of Appeal, collaborated with Amnesty International’s Algorithmic Accountability Lab, and assisted in setting up the Global Partnership on AI. She is a member of the Quebec Bar.

Professor Lévesque’s scholarship has appeared in the NYU Journal of Legislation & Public Policy, the University of Illinois Journal of Law, Technology & Policy, and NeurIPS proceedings, among others. She has written for The Guardian, the Toronto Star, La Presse, The Walrus and Tech Policy Press.

Abstract: Human rights are an appealing candidate to guide AI governance, as they provide an established vocabulary and international legitimacy, supported by existing institutional machinery. The UN Guiding Principles on Business and Human Rights extend that vocabulary to private actors, setting out a corporate responsibility to respect human rights alongside the state’s binding duty to protect them. Asking which values AI governance should promote is easier than asking what makes a value operative.

This presentation will unpack the potential and limitations of human rights to guide AI governance. Early on in 2018 as machine learning picked up speed, the Toronto Declaration translated the Guiding Principles' due diligence requirement into concrete steps for machine learning: providing actionable guidance to identify risks, prevent and mitigate them, track responses, and report publicly. More recently, Meta's Oversight Board described itself as an independent non-judicial grievance mechanism under the Guiding Principles, invoking human rights when reviewing automated content moderation. These examples show that human rights can be implemented to move the needle on AI governance.

That said, human rights come with limitations. Non-binding commitments can confer unearned legitimacy dividends to voluntary undertakings, with companies shrouding their actions in the language of pliable norms without the substance of enforcement. Voluntary frameworks may also buy companies time in terms of maintaining the status quo as to the absence of binding regulation, normalizing their problematic practices while societal expectations adjust. Finally, soft commitments give way under competitive pressure, as recent retreats from human rights pledges suggest.

Turning to current developments, investor protection rules requiring thorough risk disclosure may advance due diligence further as frontier AI labs go public. The proposed binding treaty on business and human rights would also require states to regulate rather than ask companies to restrain themselves. Ultimately, articulating values matters less than building the scaffolding of incentives and obligations to implement them.

Richard Khoury - Generative AI and the War Against Linguistic Diversity

Richard Khoury obtained his Bachelor's and Master's degrees in Electrical Engineering and Computer Engineering from Université Laval (Quebec City, QC) in 2002 and 2004 respectively, and his PhD in Electrical Engineering and Computer Engineering from the University of Waterloo (Waterloo, ON) in 2007. From 2008 to 2016, he was an Assistant Professor, then an Associate Professor, in the Department of Software Engineering at Lakehead University. In 2016, he joined Université Laval as an Associate Professor, before being appointed Full Professor in 2026. From 2021 to 2023, he was also President of the Canadian Artificial Intelligence Association. His research interests are in data mining and natural language processing, and also include knowledge management, machine learning, and artificial intelligence.

Abstract: Generative AI tools are now available in more than a hundred different languages. At first glance, this seems to be a victory for linguistic diversity: artificial intelligence makes it possible to reach users around the world in their native language. But beneath this positive image lies an insidious problem. Artificial intelligence only knows and understands the dominant dialect of these languages and imposes its use on speakers. This linguistic colonization by AI threatens to erase global linguistic diversity and suppress minority cultural communities. In this presentation, we will examine this problem, its origins, its symptoms, and its solutions.

THURSDAY October 22

Pierre Larouche - AI, open-endedness and flexibility

As an expert in competition law, economic governance, and civil liability within the traditions of civil and common law, **Pierre Larouche** is a full professor in law and innovation at the Faculty of Law at the Université de Montréal. A graduate from the law faculties of McGill University, the University of Bonn, and Maastricht University, Mr. Larouche has been a full professor in competition law at Tilburg University, in the Netherlands, since 2002, where he notably co-founded the Tilburg Law and Economics Center (TILEC), a world-leading research centre in economic governance. Over the course of his academic career in Europe, Mr. Larouche also developed and established an innovative undergraduate program, the Bachelor in Global Law, at Tilburg University. Furthermore, he taught at the College of Europe in Bruges, and was a visiting professor at several universities in America (Northwestern, Pennsylvania), Europe (Sciences Po, Bonn), and Asia (Singapore).

Prof. Larouche has published some sixty monographs, articles, and scientific contributions, and his work, which has been cited by the Court of Justice of the European Union and the Supreme Court of the United Kingdom, informed the electronic communications and competition policies adopted by the European Commission.

In addition to his private law practice, he served as law clerk to the Honourable Charles D. Gonthier of the Supreme Court of Canada. Prof. Larouche also participates in the activities of the Public Law Research Centre (CRDP) and the Centre of the Law of Business and International Trade (CDACI). Additionally, he contributed greatly to the “Innovation, Science, Technology, and Law” doctorate option launched by the Faculty in the fall of 2017.

Abstract: A key plank of the liberal democratic social science research agenda of the last decades has been to incorporate insights from information research. This has typically led to a deeper understanding and appreciation for dynamic effects, as reflected in the embrace, in academic and policy discourse, of informational imperfections and uncertainty through values/principles such as open-endedness (dynamic version of openness) and flexibility. As its performance improves and it spreads throughout the economy and society, AI is likely to affect this area of research and policy. This presentation explores the implications of AI on open-endedness and flexibility in the more specific context of markets. Competing economic theories of the market can be roughly sorted out according to the purpose they ascribe to markets (efficiency or discovery). Depending on the theoretical perspective, the anticipated impact of AI on markets will differ. If markets deliver efficiency, then AI could improve on their performance, but at a price. If however markets are a discovery tool, AI could undermine the very foundation of markets. In both cases, some legal intervention might therefore be necessary to avoid worse-case scenarios where socially-beneficial properties of markets would be lost. This sheds a new light on a different theoretical controversy about markets, namely the extent

to which they are a natural or a constructed phenomenon. The presentation ends by a tentative extrapolation from the case of markets to the broader academic and policy discussion on open-endedness and flexibility.

Siddharth De Souza - A new theory for the Digital Rule of Law

Siddharth Peter de Souza is an Assistant Professor in AI & Society at the Centre for Interdisciplinary Methodologies, University of Warwick. His work explores how data is governed globally in contested and plural settings. He is the founder of Justice Adda, a law and design social venture in India where he works on matters related to access to justice, and legal empowerment. Siddharth has worked as a digital policy consultant with United Nations Development Programme and the International Development Law Organization and was previously a post-doctoral researcher at the Global Data Justice Project at Tilburg Institute for Law, Technology and Society.

Abstract: In this talk, I will explore why the concept of the rule of law needs a new theory. For years, it has been critiqued by law and development scholars for being a tool to justify plunder by the west, as an instrument for legal hegemony masked in the form of state building projects, and as a mechanism of legal transplantation which pays little attention to the plural nature of legal systems across the world. The Secretary General of the United Nations has called for a new vision for the rule of law which focusses on developing a people centric approach to how it is conceptualised as well as institutionalised. With the increasing digitalisation of public services in areas such as health, finance and education, new concerns have emerged for the rule of law including around privacy, surveillance, and the corporate capture of public institutions.

Emmanuelle Boutin-Gilbert – Information to come

Olivier Servais – Information to come

FRIDAY ocotber 23

Georges Azzaria - The Erasure of Individuals' Legal Singularity by Algorithms

Georges Azzaria is a Professor at the Faculty of Law of Université Laval in Quebec City. His early research focused on the relationship between art and copyright, as well as on the socio-economic status of artists. In recent years, he has taken a particular interest in artificial intelligence from the perspective of intellectual property and legal regulation. From 2017 to 2025, he was Director of his university's School of Art.

Abstract: Human creativity, valued by intellectual property, and the recognition of legal personality have, since the 18th century in the West, given the individual a central place in law. However, the capacity of generative artificial intelligence to replace humans is now bringing

about a reversal that tends to erase this singularity and relegate the individual to the role of a computer intermediary.