

Sommet IA : la nécessaire sécession européenne à base de sémantique

Un sommet pour quoi ? Suivre les voies tracées par d'autres puissances pour leurs IA opaques purement probabilistes ou soutenir une voie européenne fondée sur la sémantisation des IA, à l'œuvre clandestinement dans les annotations, et qui redonnerait le contrôle du sens de ces IA. L'Europe doit faire sécession pour affirmer ainsi sa défense d'une civilisation du droit et du débat démocratique.

Dominique Boullier

Pour AOC Février 2025

Version longue

L'heure est au choix, à la bifurcation. Et non à la copie du maître colonisateur ou de ses adversaires, ni à la course après un retard fantasmé si l'on comprend que les LLM (Large Language Models) ne sont qu'un moment provisoire dans l'histoire de l'IA. Le coup de force d'Open AI en 2022 avec son IA générative a provoqué une sidération, puis un suivisme puis une agitation en forme de tunnel cognitif comme dans toute tyrannie du retard : il faut suivre à tout prix pour ne pas rater le train (ce que les investisseurs savent très bien faire et même provoquer). Mais en janvier 2025, DeepSeek a peut-être ouvert les yeux de certains : on peut faire autrement, on aurait pu faire autrement et à quoi bon suivre un train qui va dans le mur ? Si retard il y a, cela nous donne au contraire le temps de bifurquer radicalement. Car toute la mystification des IA génératives fondées sur une approche probabiliste massive apparaît désormais pour ce qu'elle est : une opération de levée de fonds et de prédation sur l'innovation en même temps qu'une emprise colonialiste sur la construction d'une architecture de connaissance partagée qui ne respecte aucune des valeurs des démocraties européennes mais s'aligne très bien sur l'oligarchie autoritaire qui s'installe aux USA et dans plusieurs pays du monde.

Un sommet de l'IA ne doit pas viser une harmonisation impossible entre des technologies portées par des régimes dictatoriaux et des firmes irresponsables et prédatrices. Il doit affirmer que, après la dérive des plateformes de réseaux sociaux, après la captation de rentes des plateformes de commerce et d'information, il est temps de faire sécession et de développer nos propres modèles d'IA. Sécession et non seulement souveraineté, qui reste toujours un mot d'ordre pétri de bonnes intentions mais toujours trop diplomatique et finalement impuissant car sans vision. On a bien compris que le multilatéralisme est mort avec Trump. On sait qu'il n'existera aucune possibilité de gouvernance mondiale de l'IA. Elle n'existe déjà pas pour Internet et ses instances de droit californien tenues en laisse par les GAFAM. Elle n'a jamais réussi à voir le jour pour les réseaux sociaux, malgré les tentatives de réglementation pour corriger les défauts constatés. Car ce ne sont pas des défauts en effet mais bien l'esprit même de non-droit qui gouverne toutes ces plateformes libertariennes. Nous pensions avoir atteint le stade suprême du libéralisme

avec Reagan et Thatcher, nous voyons qu'il peut encore dégénérer en libertarianisme autoritaire disruptif avec Trump et Musk.

Et tout cela était déjà écrit dans l'incapacité structurelle de la tech à instituer des espaces de vie commune vivables et à comprendre pourquoi le droit est au cœur même de notre civilisation. Ils ont décidé de larguer les amarres du droit sans craindre d'emprunter la rhétorique fasciste jusqu'à venir déstabiliser nos régimes politiques, nous devons assumer la sécession pour la préservation de notre civilisation démocratique. L'Europe possède ses références ancestrales, ses capacités créatives, sa puissance industrielle, seule lui manque sans doute cette toute-puissance financière spéculative qui corrompt tout, jusqu'aux processus électoraux. Tant mieux alors, appuyons-nous sur d'autres modèles, comme on le fit pour Airbus, pour le GSM, pour le MP3, pour Linux, pour le Web, tous européens et durables. Mais pour cela, il faut sortir de la mentalité de colonisés qui est la nôtre, de cette servitude volontaire que l'on connaît bien pourtant, et soutenir nos approches ancrées dans le droit « by-design » et dans la recherche du bien commun. En matière d'IA, tout est là, désormais les décideurs peuvent voir qu'il existe d'autres voies, qu'il n'y a aucune fatalité mais qu'il ne s'agit pas de changer d'âne comme les mouches en suivant le Deep Seek chinois.

Annoter : un travail sémantique clandestin et pourtant fondamental

Partons du travail d'annotation des données d'apprentissage, de classification qui demeure invisible et qui pourtant reste indispensable à tous ces LLM soi-disant purement statistiques. Le modèle européen à venir doit rompre avec ces croyances au tout statistique et faire entrer la sémantique en jeu que d'autres traditions d'IA d'inspiration symbolique continuent à faire vivre. Et cela n'est pas une affaire géopolitique de compétition industrielle et financière mais de civilisation européenne s'appuyant sur ses valeurs et ses atouts pour construire la technologie cognitive commune à sa manière.

Casilli et tout le projet DIPlab ainsi que Muldon, Graham et Cant (2025) l'ont mis en évidence. Toute l'IA, en général, l'IA générative en particulier, ne serait rien sans le travail d'annotation effectué par des tacherons du clic, partout dans le monde pour des salaires de misère. Et cela en Chine aussi pour Deep Seek avec une forme de distribution du travail à domicile subventionnée par le gouvernement chinois. Enjeu politique et moral majeur, certes, mais aussi, dans la lignée des travaux de STS (Science and Technology Studies) que Callon, Latour et Law ont initiés, occasion de prendre en compte les conditions matérielles de production d'une architecture de connaissances telle que les systèmes d'IA générative. On ne peut pas penser une stratégie alternative si l'on continue à occulter ce maillon essentiel, pour des raisons politiques et morales, mais aussi pour des raisons d'explicitation de ce que font réellement ces IA génératives. On nous a souvent vendu l'émergence de systèmes d'IA non-supervisés : il est clair que cela n'existe pas et que l'entraînement des modèles suppose toujours à un moment donné une supervision, c'est-à-dire ici des annotations pour créer une « ground truth » robuste quels

que soient les mécanismes d'optimisation de l'apprentissage autonome par la suite, par renforcement par exemple.

Pourquoi cela ? Pourquoi faut-il encore des humains dans la boucle ? Parce que les IA génératives qu'on nous vend ne peuvent se passer d'une description du monde, d'une classification du monde, d'une ontologie fondationnelle, malgré tout ce que les firmes prétendent. Mais alors, pourquoi toujours mettre l'accent sur la puissance de calcul, sur le nombre de paramètres, pourquoi occulter tout le travail de pondération de ces paramètres et évacuer la dépendance structurelle à ces annotations ?

Parce qu'il s'agit d'un choix de business qui est aussi culturel et qui est en train de devenir civilisationnel. Les firmes de la tech californiennes ou chinoises, ces croyants technophiles furieux, sont engagées dans une voie, celle d'une IA purement probabiliste qui pourrait se passer d'une référence au monde et qui serait sans cesse optimisée par l'augmentation de la taille des sets de données d'apprentissage, une forme de « brute force » sans aucune intelligence. Cela suppose des puissances de calcul obtenues par des processeurs spécialisés que NVidia vend très cher pour son plus grand profit, et qui consomment toujours plus d'énergie au finale (d'où l'intérêt pour des centrales nucléaires dédiées aux centres de données !). Ce que DeepSeek montre déjà, c'est que tout ce discours d'urgence stratégique des firmes qui dominent l'IA récente est une arnaque financière et une impasse conceptuelle. Si l'on met les efforts du côté d'une frugalité en termes de volume de données et de puissance de calcul, on peut arriver à faire aussi bien et tous les investisseurs qui ont cru et voulu faire croire aux bénéfices du toujours plus gros en sont pour les frais, mais c'est la règle du jeu de la finance forcément spéculative dans ces domaines. Pourtant, ce point reste secondaire dans l'impératif de sécession qui doit devenir le nôtre, car cette course à la taille des modèles indique surtout une préférence culturelle pour une des dimensions de l'informatique, le calcul, basé sur des proxies (ici des tokens) et un mépris pour l'autre face de l'informatique, à savoir la classification, c'est-à-dire la logique qui permet d'espérer traiter la sémantique même.

Le défaut d'origine des LLM : statistique partout, sémantique nulle part !

Toute la vogue de l'IA connexionniste, depuis 2012 et ses performances en reconnaissance d'image, repose sur un dogme qu'avait exprimé Anderson en 2008 : « The End of Theory ». Le seul calcul de corrélations tous azimuts par ajustement progressif des pondérations en fonction des erreurs (la propagation du gradient de l'erreur) suffit à faire émerger des solutions à des problèmes que l'on n'a même plus besoin de poser en théorie. C'est encore ce que les IA génératives vendent en matière de recherche scientifique, en médecine notamment. Pourtant, tout cela repose sur une littérature scientifique, écrite donc, du matériel langagier, langage dont on peut se dispenser de faire la théorie. La seule décomposition en tokens de tous les éléments visibles constitutifs du

langage permet de détecter statistiquement des patterns qui permettront de prédire non seulement le mot mais la lettre suivante.

Cette option radicale est au cœur d'un choix théorique et culturel majeur qui est en train d'affecter toute notre architecture de connaissance comme s'il allait de soi alors qu'il est remis en cause par de nombreux chercheurs en IA et parce que l'on constate que tous les ratés des LLM proviennent d'une absence de capacité sémantique que l'on doit compenser de plusieurs façons. Lorsque Google a publié son fameux article en 2017 « All You Need Is Attention », il ne faisait pas référence aux enjeux cognitifs majeurs de notre époque mais seulement au fait que ses algos de Natural Language Processing devaient désormais étendre leur zone d'intérêt au-delà du mot suivant pour mémoriser des segments entiers d'expressions et valider leur fréquence d'apparition pour prédire le mot ou la phrase suivante. Ces « Transformers » ont constitué une avancée opérationnelle majeure, performance d'ingénierie qui n'a rien de scientifique, précisons-le. Elle dépend d'ailleurs uniquement des ressources de mémoire et de calcul énormes désormais disponibles. La collecte des données d'entraînement devient essentielle et elle doit être massive pour réduire les erreurs en probabilité. Ce qui n'empêche pas les erreurs grossières et les faux délibérés qu'on dénomme « hallucinations » car les machines sont sommées de donner une réponse même lorsque les marges d'erreur sont élevées et « ignorent qu'elles ignorent ». Cet impératif de réponse non explicitée et sans limite de validité produit une machine à désapprendre pour les humains, une captation de croyance qui ruine tout esprit critique dans l'espace public et toute méthode de test et d'examen traditionnel dans l'industrie.

Quelles probabilités ?

On mobilise ici une version fréquentiste de la probabilité comme le disait Ian Hacking, version dogmatique qui oublie totalement qu'il en existe une autre, la probabilité subjective bayésienne (qui ne relève pas seulement des croyances). Tout se passe comme si les IA génératives avaient décidé de ne plus tenir compte de cette dimension subjective ou épistémique qui nécessairement met au premier plan le rôle des experts convoqués pour valider les estimations. Lorsqu'on évacue cette question, on se trouve contraint de faire appel à des experts quasi clandestinement, avec ce travail de classification des travailleurs du clic, travail qui n'est plus secondaire mais décisif et se retrouve pourtant invisibilisé puisque contredisant de fait le dogme fréquentiste. Or, tous les systèmes d'IA y font pourtant appel, car l'informatique repose sur ses deux jambes, le calcul et la classification et, si l'on n'en choisit qu'une, le type de probabilités qu'on mobilise sera différent.

Et toutes les améliorations de ces IA ou leurs alternatives reposent sur une mobilisation d'expertise humaine qui permet de réduire les risques d'erreurs en fournissant des repères à partir de l'expérience du monde possédée par ces experts (mais qui n'est que leur expérience). Les annotateurs font partie de ces experts de fait, mais on ne les désigne

jamais sous ce vocable, plus noble et qui obligerait à les payer correctement. Tout chercheur ou développeur utilisateur de Mechanical Turk (Amazon) devrait ainsi avoir conscience qu'il soutient une méthode coloniale et esclavagiste qui est pourtant indispensable aux IA connexionnistes et génératives, mais doit rester occultée pour les besoins des discours de promotion des miracles de l'IA générative. Les firmes adoptent cependant des méthodes complémentaires variées pour compenser leur faille constitutive en incorporant par exemple des procédures de renforcement et d'autres dites de Chain of Thought (CoT) qui sont une importation de raisonnement logique pour rendre plus robustes leurs prédictions.

On peut et on sait faire autrement qu'avec des LLM et des annotations clandestines

Des alternatives existent pourtant depuis longtemps. Certaines avaient été labellisées « systèmes experts » dans les années 2000 mais depuis le coup de force des connexionnistes, s'étaient retrouvées disqualifiées alors que leur rôle dans la boucle des inférences reste essentiel. Ainsi Zyed Zalila a développé avec sa société Intellitech son système [Xtractis](#) que j'ai pu utiliser dès 2007 avec Audrey Lohard pour produire un système évaluant la qualité des jeux vidéo ([Boullier, Lohard, 2010](#)). [Boullier, Lohard, 2010](#)). « L'IA raisonnante de confiance » qu'est Xtractis s'appuie sur des données structurées, donc classifiées, pour générer automatiquement en grande quantité des modèles explicites à base de règles que l'expert peut valider ou relancer. Il produit ainsi une « boîte blanche », dans laquelle le rôle des experts n'est pas opaque, n'est pas non plus sous-traité à Amazon Mechanical Turk et dont l'évaluation rétroactive fait partie du contrôle de qualité. De même, Paul-Olivier Dehaye et sa société hestia.ia ont permis avec [Personaldata.io](#) de récupérer les données que collecte Uber de façon opaque et à des fins de surveillance, pour les restituer aux chauffeurs, qui deviennent ainsi des experts de leurs propres activités et peuvent optimiser leur propre activité dans leur lutte contre le donneur d'ordres Uber. De même, Lê Nguyễn Hoang et El Mahdi El Mhamdi, fondateurs de Calicarpa, ont créé [Tournesol](#) pour éviter les effets prescriptifs opaques et marqués commercialement et idéologiquement des moteurs de recommandation des plateformes (Youtube notamment). L'opacité des prescriptions de Youtube fondée désormais sur des LLM peut être combattue si l'on met dans la boucle les utilisateurs des contenus vidéos pour leur demander de valider les contenus et les expériences qu'ils ont vécues avec ces vidéos. C'est aussi ce que fait Wikipédia d'une certaine façon avec ses armées de contributeurs à cette intelligence collective.

De fait, c'est cela que mobilise statistiquement les firmes de l'IA générative : elles prennent pour validation la fréquence d'apparition de certains patterns dans les documents collectés. Mais cet ersatz prétend fournir une approximation de vérité alors qu'il ne fournit qu'une moyenne et qu'un constat de fréquence d'usage tiré de ses corpus d'entraînement. Dès lors la qualité du matériel langagier utilisé en entrée devient décisive et les biais multiples que l'on constate ([K. Crawford, Contre-Atlas de l'IA, 2022](#)) s'expliquent largement ainsi. Le problème n'est pas de faire de la rectification après coup

selon des choix éditoriaux plus ou moins politiquement corrects (ce que font aussi bien Open AI que Deep Seek, pour des motifs différents, en refusant de répondre ou en faisant du relativisme extrême).

Il est temps de reconnaître la place essentielle de l'expertise humaine pour faire entrer de la sémantique dans ces modèles, comme le faisait l'IA symbolique, et d'abandonner la vision fréquentiste qui est la leur. Cette expertise vaut avant tout pour sa capacité à valider la pertinence des descriptions du monde, c'est ce que l'on ne peut déléguer à une quelconque machine qui n'a ni l'histoire, ni le corps, ni l'expérience vécue du monde. Or, tous les modèles actuels d'IA générative s'interdisent d'avoir une quelconque compréhension de la validité des énoncés ou même des paramètres qu'ils utilisent. Ce que Yann LeCun, qui a pourtant développé le LLM Llama, a déjà pointé en indiquant que la prochaine étape des IA sera de s'appuyer sur des données d'expérience du monde, et non plus seulement sur des probabilités issues d'une masse de textes sans que leur référence au monde, et donc à une sémantique, soit utilisée.

L'incontournable dimension sémantique

(Ré)-introduire cette sémantique constitue la nouvelle frontière et elle devrait guider tous les développeurs et les décideurs, non seulement pour corriger les défauts des IA génératives mais pour fonder de nouveaux modèles. C'est d'ailleurs pourquoi les modèles plus petits contrôlables par des experts d'un domaine sont désormais mis en avant. Lorsque les organisations développent des [RAG \(Retrieved Augmented Generation\)](#), elles le font pour contrôler la qualité des sources, qu'elles fournissent elles-mêmes en connaissance de cause, en lieu et place de sets de données d'apprentissage entraînés on ne sait où et en toute généralité. La validité locale d'un savoir doit désormais être reconnue comme une valeur ajoutée importante et intégrée à tous les modèles. Cette expertise locale n'a pas la prétention à la validité statistique ni même épistémique mais elle est ancrée dans le monde, et explicitable, ce qui est une vraie valeur ajoutée pour les organisations qui fonctionnent largement au tacite. Mais entre le tacite d'un côté et le général statistiquement valide de l'autre, il existe bien d'autres choix pour organiser les connaissances et la résolution des problèmes. Il s'agit dans tous les cas de réencastrer le calcul dans la sémantique et dans le « design organisationnel » propre à chaque monde social. En effet la tendance au désencastrement du calcul, expression dérivée de Polanyi, est un risque que l'on mesure de plus en plus précisément avec ses systèmes de probabilité fréquentiste qui prétendent à la domination.

Mais une approche « encadrée by design » dans les univers sociaux oblige à prendre en compte la diversité des groupes sociaux et des cultures. Voilà un enjeu politique majeur qui s'est infiltré sans le dire dans toutes les prétentions des IA génératives. Puisqu'il n'existe que des validations statistiques de réponses qui n'ont aucune connexion avec le monde, il faut et il suffit d'augmenter la taille des sets de données d'apprentissage pour

prétendre couvrir toutes les variétés de connaissance et en faire une moyenne pour présenter la plus probable. Evidemment, il se trouve que cela marche pour l'utilisateur pour un certain type de requêtes, elles-mêmes moyennes. Mais le prix politique et culturel à payer consiste à nier la diversité des points de vue (standpoints) sur le monde (puisque pour les IA génératives, le calcul n'a aucun rapport avec le monde, ce n'est pas leur problème). Les biais que nous avons déjà mentionnés sont inexpugnables d'un tel système puisqu'il ignore par construction les relations avec le monde et la sémantique et donc tous les enjeux de classification. La capacité de traduction -certes remarquable- des IA génératives provient pourtant de cette opération purement statistique, mais il faut noter qu'elle s'appuie sur une langue-pivot, l'anglais toujours, et cela se mesure dans le déclin de la qualité de traitement pour les langues plus rares. De plus, cela ne permet pas de conserver la sémantique propre à une langue, qui n'est désormais qu'une variante statistique si la fréquence des contenus accumulée est suffisante.

Quand la théorie du langage fait défaut

C'est ici que l'on mesure que l'absence de théorie devient dangereuse politiquement et culturellement. La définition spécifique de « langage » en anglais est d'ailleurs très lâche puisque cela peut désigner soit un processus formel de représentation, soit une langue (idiom, très peu utilisé), soit un code spécifique en informatique. Or, le langage se décline socialement en plusieurs langues et aucune langue n'est LA langue humaine, qui n'existe pas car, du fait de son principe d'organisation sociale, l'humanité se distingue sans cesse en groupes différents, elle est donc toujours divisée et travaille politiquement à se recomposer sans cesse, comme l'a théorisé [Jean Gagnepain](#). La mystification qui consiste à rendre tous les matériaux linguistiques équivalents une fois ramenés à leur version anglaise et sans référence au monde spécifique que les langues sont supposées représenter, constitue en fait une opération coloniale et impérialiste de grande ampleur. Alors que les capacités de traduction nouvelles devraient permettre à l'humanité de préserver sa diversité et de faciliter les échanges, on se retrouve en fait à adopter en contrebande une ontologie du monde forgée statistiquement à partir des seules références de la langue dominante, l'anglais.

Or, cette approche est impossible en Europe où la diversité des langues est constitutive et ne peut jamais se réduire à une lingua franca que serait l'anglais dégénéré en globish. Cette expérience de la diversité linguistique est certes un handicap en termes de marché, mais elle doit être valorisée comme fondement même des exigences des IA. Car elle indique qu'il n'existe pas un matériau linguistique purement statistique mais des visions du monde ancrées dans DES langues, qui font pourtant ensemble UNE civilisation forte de ce pluralisme et capable de faire l'effort pacifique de s'harmoniser sans se dénaturer. Tout l'inverse de tout processus colonial, celui des étatsuniens avec les peuples indiens mais aussi celui de toutes les puissances coloniales européennes, qui n'en finissent pas de traiter ce problème en tentant désormais de refuser l'immigration et la diversité accrue qu'elle apporte. Cette imposition d'une langue et d'une vision du monde standard

suppose d'évider toute la connaissance de sa dimension expérientielle et de lui inventer mythiquement ou par effet statistique une unité fictive au profit des dominants.

Gouverner les infrastructures de connaissance

Avec les IA génératives, la méthode de mise à l'écart de la sémantique se prolonge donc, la faille originaire des probabilités fréquentistes s'agrandit, depuis le web (dont le web sémantique n'a jamais pu exister), le moteur de recherche (et la fausse objectivité de Google) et les applis de réseaux sociaux (qui ordonne nos fils d'actualité statistiquement et commercialement).

Tous ces choix qui n'en sont pas et qui n'ont pas été explicités montrent bien une *faillite générale dans le gouvernement des infrastructures de connaissance* dès lors qu'elles se situent à l'échelle mondiale. Cette faillite était somme toute assez prévisible étant donné l'absence d'une gouvernamentalité d'internet, du web, des applis et des IA autre que les *stores* des OS dominants, ou que les instances techniques d'internet, trustées par les mêmes firmes qui profitent du bazar sémantique contemporain. Les gouvernements, lorsqu'ils acceptent de sortir du dogme libéral du laissez-faire, arrivent après la bataille et surtout pratiquent le *patch*, le correctif ponctuel, en ignorant tous les enjeux structuraux de la *dérive asémantique* de notre architecture de connaissance partagée. Comme si nous n'avions plus de contrôleurs du ciel, de contrôleurs du rail et que tout et n'importe quoi pouvait circuler quitte à rendre l'espace invivable pour les autres, quitte à saboter la performance des infrastructures communes, quitte à épuiser les ressources devenues rares. La toile des connaissances porteuse d'une potentielle intelligence collective au début 2000 a été confisquée au profit de quelques firmes pour dissoudre tout repère de connaissances partagées. Ces modèles LLM et ces firmes ne permettent pas d'avancer vers une IA interprétable ou explicable qui est pourtant devenu depuis 2015 un objectif de nombreux chercheurs. Cette opacité désarme tout contrôle technique, juridique et politique et ruine toute possibilité d'apprentissage humain et organisationnel, le machine learning ne peut plus devenir un « human-machine learning » et enregistre alors comme un fatalité ce que S. Zuboff (2019) appelle une « division of learning » qui se substitue à la « division of labor ».

Mais qu'alors y faire ? Un cahier des charges techniques qui vaut charte civilisationnelle

De nombreuses solutions différentes sont parvenues à se maintenir, à évoluer et sont opérationnelles dès maintenant. D'autres prônent désormais une hybridation étant données les limites patentes des LLM, alors qu'il serait sans doute plus productif d'abandonner carrément ces approches.

.

Nous ne traitons pas ici d'un enjeu de compétition industrielle ou financière. Nous ne traitons pas d'un retard constitutif dont l'étalon serait les performances des IA

génératives étatsuniennes ou chinoises. Nous traitons de principes d'organisation des connaissances et de design des architectures qui vont constituer notre horizon cognitif commun. Ce n'est plus la bibliothèque d'Alexandrie car ce que nous avons développé n'a plus grand rapport avec une bibliothèque mais on peut pourtant s'en inspirer pour fournir des classifications communes mais aussi pour le repère politique que cela peut constituer. Il existe un équivalent qui n'est pas européen dans ses fondations, Wikipédia, mais qui est devenu un repère pour tout utilisateur d'internet. Or, il a fallu instituer quantité de procédures pour organiser la validation par des experts, ordinaires ou savants, que tout le monde peut devenir. Ce qui compte ici, c'est avant tout la procédure collective comme vision centrale, qui permet notamment de refuser la publicité et les vidéos, tout ce qui pourrait faire dévier de la trajectoire encyclopédique de référence. Certes, il existe des conflits, des risques de captation du pouvoir, des épreuves financières, etc. Mais la vision prime et elle permet de constituer au fur et à mesure une ontologie de fait qu'on peut naviguer de façon assez ouverte (« see also »). La sémantique est présente partout, la statistique très faiblement utilisée si ce n'est pour repérer des dysfonctionnements. De fait, c'est une procédure équivalente qu'il faut encourager dans la production des modèles d'IA européens, qui devront désormais tous être à forte composante sémantique, adaptés à chaque domaine et sobres en ressources, de calcul et de données.

Ce cahier des charges doit être mis à l'ordre du jour et décliné selon toutes les couches de la chaîne de production. Car l'enjeu industriel reste important, non pas pour copier les autres mais pour se différencier au contraire, en jouant le coup d'avance, celui qui verra la fin de la pertinence des LLM, abusivement statistiques et totalement inadaptés à la résolution des problèmes réels des entreprises et des organisations, hors de tout esprit de séduction d'un grand public avec des réponses magiques et fausses.

Ce choix est porteur de valeurs et porteur de conflits. Il n'est pas adapté à des retours sur investissement immédiats et les spéculateurs ne vont pas se précipiter. Mais dès lors que Deep Seek a introduit du doute dans leur esprit avant tout cupide, il est possible de les attirer vers d'autres voies moins profitables mais garanties par des gouvernements et par des collectifs. Cependant, ce sont d'autres voies de financement qu'il faut trouver, car tout le capitalisme financier numérique a désormais perverti tout le système d'innovation, devenu avant tout financier. Les emprunts européens et les financements collaboratifs sont des pistes à valoriser dès lors que le capital-risque a montré ses effets délétères pour l'innovation et son asservissement au court terme et à la rentabilité financière. Les LLM sont une solution seulement dans un monde où l'intensité capitaliste n'est pas un problème, pas plus que ne le sont les effets climatiques, la responsabilité juridique, le droit du travail, le droit d'auteur ou le droit de la privacy, ou le débat et l'audit explicites, un monde étatsunien libertarien.

Mettre en œuvre de façon opérationnelle ce cahier des charges pour une sémantisation de l'IA

L'exigence éthique est souvent traduite par des formules creuses comme mettre « l'humain au centre », ce qui ne peut guider aucun choix technique. Il faut mettre au centre le sens, la sémantique et les valeurs de vie commune ancrées dans le droit et dans le pluralisme. Techniquement, cela se décline déjà en cahier des charges pour prévoir la nouvelle génération d'IA qui se fondera sur une sémantique qui ne pourra être fournie que par des humains, traités comme tels, et non comme esclaves, en mobilisant le moins de ressources énergétiques possibles, en s'adaptant à des environnements culturels de travail toujours précis et spécifiques.

La *barbarie computationnelle coloniale* que l'on veut nous imposer avec la diffusion d'un certain type d'IA n'est pas une fatalité. Mais pour cela, il faut accepter de faire sécession dans toutes les couches de nos systèmes de connaissance. Qui produit les puces qu'il nous faut ? Qui produit les machines qu'il nous faut ? Qui stocke les données qu'il nous faut ? Qui nettoie, annote, et valide ces données d'apprentissage ? Qui les traite et avec quelles sources de code ? Qui écrit ces normes de sécurité, de privacy, de conformité légale de tout le processus ? Qui en contrôle l'application ? Qui finance tous ces développements ? Qui en est propriétaire (ou pas) ? Qui déploie les applications qui permettent l'utilisation de tous ces systèmes ? Dans chacun de ces cas, des solutions européennes doivent être trouvées et partagées, en appuyant délibérément les organisations qui offrent des solutions compatibles avec ce cahier des charges et en bloquant les autres.

Tous les services publics doivent appliquer cette *préférence sémantique européenne* et ne plus s'obliger à des règles de concurrence déjà biaisées car ignorant délibérément les enjeux de responsabilité. Les institutions éducatives en particulier doivent être vigilantes sur toutes ces dimensions car leur mission est d'apprendre à penser, à décider, et à faire, et cela de façon autonome sans se reposer sur les performances de leurs dispositifs de dopage incontrôlés, qu'on ne peut même pas expliciter. La formation des prochaines générations à cette autonomie de jugement et à cette responsabilité est un enjeu crucial face au rouleau compresseur libertarien encapsulé dans les IA génératives et dans les LLM probabilistes opaques.

La dimension communicationnelle de cette architecture de connaissances, mise en œuvre dans les réseaux sociaux et les médias en général, doit être reprise avec les mêmes ambitions car tout notre système médiatique a dérivé en machine virale aux mains de milliardaires et de plateformes équipées d'IA génératives désormais. [D'autres cahiers des charges ont été publiés](#) sur ce plan. La complémentarité de tous ces systèmes qui forment nos façons de penser, de discuter, de décider, est très forte et indique à quel point cette mutation médiologique est un enjeu de civilisation majeur.