

Celui qui ne signe pas

Ce que Yann LeCun propose à la place.

Sciences Po, 16 septembre 2026 : lecture d'un refus et de son programme.

Avec une coda pour clore la série : alors, qui croire ?

Thèse. La conférence de Yann LeCun à Sciences Po vaut moins comme réponse à Dario Amodei que comme feuille de route : une carte de ce que les modèles actuels ne feront pas, un pari sur l'endroit où se jouera la suite, et une condition pour que l'Europe y ait une place, l'ouverture.

Pour le lecteur pressé

Le contexte. Le 12 septembre, Dario Amodei (Anthropic), aussitôt rejoint par Sam Altman et Elon Musk, appelle à ralentir la frontière de l'IA. Le 16, à Sciences Po, Yann LeCun, prix Turing, dit non. La passe d'armes dure cinq minutes ; les quatre-vingts autres méritent mieux que les titres.

Sa carte. Les modèles de langage maîtrisent le texte, pas le monde physique. Un enfant de quatre ans a vu autant de données que le plus gros d'entre eux.

Son pari. La prochaine révolution viendra de machines dotées d'un modèle du monde, capables de planifier. Elle n'est pas jouée et demanderait moins de calcul.

Sa condition. La souveraineté n'est « pas une option », et elle passe par des modèles ouverts, entraînés en commun.

Le mode d'emploi. Réguler les usages plutôt que la recherche, apprendre à commander plus intelligent que soi, trois affirmations à ne pas répéter telles quelles. Et, pour clore la série, quatre règles pour savoir qui croire sur quoi.

Ouverture. Celui qui dit non propose autre chose

29 juillet 1963. Le traité de Moscou sur l'interdiction partielle des essais nucléaires vient d'être paraphé par Washington, Londres et Moscou. De Gaulle, devant la presse : « L'accord de Moscou, je le dis franchement, n'a qu'une importance réduite ». La France ne signera pas. On retient le refus. On oublie qu'il ne pesait que par le projet qui l'accompagnait : construire une capacité propre.

Le 12 septembre 2026, Dario Amodei publie « We Must Pace the Frontier » ; Sam Altman s'y rallie en quelques heures, Elon Musk en trois mots. Cet épisode a fait l'objet ici même de « La paix des braves ». Quatre jours plus tard, à l'amphithéâtre Boutmy, Éric Hazan ouvre la Grande Conférence de rentrée de Sciences Po en demandant à Yann LeCun, prix Turing 2018 et fondateur d'AMI Labs, ce qu'il pense de « la pause ». Réponse : « Pas du bien. »

Le reste mérite mieux : une analyse et une vision utiles à quiconque doit décider d'un investissement, d'une stratégie ou d'une formation. Cet essai les restitue à partir de la captation intégrale, citations verbatim, tournures orales comprises, puis les soumet aux sources disponibles. Pour chaque idée : ce qu'il dit, ce qui l'étaie, ce qu'un décideur peut en faire.

Deux raisons de l'écouter. Il a déjà eu raison contre le consensus : entre 1995 et 2006, la communauté de l'apprentissage automatique s'est « complètement désintéressée » des réseaux de neurones ; il a persisté avec Geoffrey Hinton et Yoshua Bengio jusqu'au basculement de 2012 et au prix Turing de 2018. Et il dit d'où il parle : il a fondé il y a moins d'un an AMI Labs sur les convictions qu'il expose, société qui n'a « pas encore de revenus ». Cela ne prouve pas qu'il ait raison aujourd'hui. Cela justifie qu'on l'écoute, en sachant qui parle et qui y gagne : ce sont les deux premières questions du test de Demokos, et l'orateur, ici, y répond lui-même.

Idée 1. La carte : puissants, mais bornés

Ce qu'il dit. Il ne méprise pas les modèles de langage : « on a le droit d'être stupéfaits ». Mais prédire le mot suivant est un problème simple : « le texte est facile » et le monde réel ne l'est pas. D'où le paradoxe : des systèmes qui passent l'examen du barreau, démontrent des théorèmes, écrivent du code, et pourtant « où est mon robot domestique ? ». Aucune machine n'a aujourd'hui la compréhension du monde physique d'un chat de gouttière.

L'argument tient en un ordre de grandeur. Tout le texte public utilisable pour l'entraînement, environ 30 000 milliards de tokens (les fragments de mots que manipule le modèle), pèse à peu près 10^{14} octets : un demi-million d'années de lecture. Un enfant de quatre ans, éveillé 16 000 heures, a reçu la même quantité d'information par ses seuls nerfs optiques. À une question de la salle sur la possibilité de tout faire passer par l'abstraction du langage, sa réponse est nette : « on ne peut pas créer des systèmes intelligents qui comprennent pas la réalité ».

Ce qui l'étaie. Le calcul se vérifie, et l'hypothèse de débit est cohérente avec la littérature : environ 10 millions de bits par seconde en sortie d'une rétine humaine (Koch et al., *Current Biology*, 2006). Le constat a un nom depuis les années 1980, le paradoxe de Moravec : ce qui nous semble difficile (les échecs, le droit) est facile pour une machine, ce qui nous semble facile (marcher, saisir) est difficile. La question de fond a aussi le sien, le problème de l'ancrage des symboles (Harnad, 1990) : apprendre le sens des mots à partir des seuls mots revient à apprendre le chinois avec un dictionnaire chinois-chinois.

Une réserve : égalité de débit ne vaut pas égalité d'information utile, le flux visuel étant très redondant. L'argument établit que le texte n'épuise pas l'expérience du monde ; il ne démontre pas que les modèles de langage ont atteint leur plafond, et l'essentiel de l'industrie, LeCun le reconnaît, fait le pari inverse. C'est une hypothèse forte et réfutable.

Ce qu'un décideur peut en faire. Une grille de tri pour vos projets. Là où la valeur est dans le texte, le code, la règle ou la synthèse documentaire, ces outils sont mûrs : déployez. Là où il faut agir dans le monde physique, exigez des preuves. LeCun est sans détour : aucun fabricant de robots humanoïdes n'a aujourd'hui « aucune idée de comment rendre ces robots suffisamment intelligents pour être utiles ». Les voitures dites autonomes restent « au niveau 4, pas au niveau 5 », cantonnées à des zones cartographiées, avec appel possible à un opérateur. Quand on vous vend de l'autonomie physique, demandez où est le modèle du monde.

Idée 2. Le pari : la prochaine révolution n'est pas jouée

Ce qu'il dit. Il manque aux machines ce que possède tout animal : « un modèle du monde qu'on a dans la tête qui nous permet de prédire les conséquences de nos actions », donc de planifier. On ne l'obtiendra pas en prédisant des pixels, tâche impossible, mais en apprenant une représentation abstraite d'où les détails imprévisibles sont éliminés, puis en prédisant dans cet espace. C'est l'architecture JEPA, proposée dans un article de position en 2022. Entraînés sur des triplets « état, action, état suivant », ces systèmes établissent des relations de cause à effet : « à la fin, les modèles du monde sont des modèles causaux ».

Il en tire deux conséquences. « Il va y avoir une autre révolution dans l'IA dans les années qui viennent », tournée vers la robotique et l'industrie. Et elle changerait l'économie du secteur. Ce qui coûte cher aujourd'hui, explique-t-il, c'est la mémoire nécessaire pour accumuler des connaissances : des centaines de milliards de paramètres. Un modèle du monde remplace cette accumulation « par une compréhension » : plus petit, moins de machines. D'où la phrase qui a fait les titres : « on a peut-être perdu la partie des LLM en Europe, mais on n'a pas perdu [...] la bataille de l'IA ».

Ce qui l'étaie. L'idée du cerveau comme machine à prédire est ancienne en neurosciences. L'article de 2022 est mentionné, selon lui, dans environ 2 900 publications, et AMI Labs a levé de l'ordre d'un milliard de

dollars sur cette thèse. Il en donne lui-même les limites : « très peu de progrès les dix premières années » sur quinze ; des modèles aujourd'hui entraînés sur « l'équivalent de siècles de vidéo », faute de méthodes aussi efficaces que l'apprentissage humain ; une manipulation robotique sans capteurs de toucher « probablement impossible ». La promesse de frugalité reste à démontrer à l'échelle.

Ce qu'un décideur peut en faire. Ne pas confondre l'état de l'art et la fin de l'histoire. Si votre stratégie suppose le paradigme actuel définitif (mêmes acteurs, mêmes besoins de calcul, mêmes dépendances), elle repose sur une hypothèse qu'un fondateur du domaine conteste. Pour un industriel européen, la fenêtre se rouvre là où il faut comprendre le monde physique : robotique, maintenance, contrôle de procédés. LeCun désigne la donnée rare de cette révolution, celle qui associe une observation, une action et son résultat. Les industriels européens en produisent chaque jour. La première décision est de la conserver.

Idée 3. La condition : la souveraineté n'est « pas une option »

Ce qu'il dit. Demain, toute l'information qui parviendra aux citoyens passera par des assistants. S'ils sont propriétaires et produits par « trois ou quatre entreprises sur la côte ouest des États-Unis ou en Chine », « c'est pas bon pour la culture, c'est pas bon pour la démocratie ». Conclusion : « il faut absolument la souveraineté. C'est pas une option. » Sa voie : que les pays qui ne sont ni les États-Unis ni la Chine s'allient pour entraîner un modèle ouvert, « dépositaire de toute la connaissance humaine », nourri aussi de fonds non publics, ceux de la Bibliothèque nationale comme ceux de l'Inde ou du Japon. C'est Project Tapestry, dont il dit lui-même : « c'est un peu une vision utopique ».

Ce qui l'étaie. Le projet existe : lancé en avril 2026 par l'AI Alliance, coalition de plus de deux cents organisations dont LeCun est le conseiller scientifique en chef, il vise un entraînement fédéré où chaque contributeur garde la maîtrise de ses données et peut dériver son propre modèle. Les soutiens gouvernementaux qu'il mentionne (Inde, Japon, Vietnam) n'ont pas, à ce jour, fait l'objet d'une confirmation publique vérifiable.

L'épisode de 1963 éclaire l'enjeu. Pour obtenir la signature française, l'administration Kennedy avait envisagé de fournir à Paris, selon une note au président, la technologie que la France tirerait autrement de ses propres essais. Des résultats livrés par un tiers, à ses conditions, contre la renonciation à savoir les produire : c'est la structure d'un accès par abonnement à un modèle fermé. De Gaulle n'a pas donné suite. LeCun invite à la même réserve.

Ce qu'un décideur peut en faire. Trois gestes. Exiger, dans vos achats, une option à poids ouverts que vous pouvez héberger et inspecter. Traiter vos fonds documentaires comme un actif stratégique, à apporter dans un effort commun plutôt qu'à laisser exploiter par d'autres. Et prolonger le raisonnement là où il s'arrête. LeCun décrit la dépendance matérielle du secteur, ces processeurs vendus « pour un prix énorme », et relève un atout français, une électricité « en grande partie décarbonée ». Mais des poids ouverts ne rendent pas souverain celui qui ne maîtrise ni le calcul ni l'hébergement. L'ouverture est la condition nécessaire, pas suffisante.

Idée 4. Réguler les usages, pas la recherche

Ce qu'il dit. Avec une précaution qu'il faut citer (« je suis pas législateur »), il défend une doctrine simple : « il faut réguler les déploiements ». Un système d'aide à la conduite se teste, un outil d'analyse de mammographies s'autorise, et « pour la plupart des domaines, il y a déjà des réglementations existantes ». Réguler la recherche en amont supposerait que l'IA soit intrinsèquement dangereuse : « je n'y crois pas, en tout cas pas dans l'état actuel ». Sur le risque biologique, il rappelle que posséder la séquence d'un virus ne donne ni la capacité de le synthétiser ni celle de le disséminer, qui supposent des équipes nombreuses et des machines rares et contrôlées.

C'est de là que vient son « pas du bien ». Il lit dans l'appel du 12 septembre « une volonté de capturer le marché par la régulation », dont les modèles ouverts feraient les frais.

Ce qui l'étaie, et ce qui le complète. Le grief mérite d'être confronté au texte d'Amodei lui-même, lu intégralement pour cet essai.

Le soupçon de LeCun a des appuis. Le texte appelle une régulation visant toutes les entreprises américaines de la frontière, y compris celles qui ne coopéreraient pas volontairement, et une dérogation étroite au droit de la concurrence pour que ces entreprises puissent se coordonner. Amodei sait que l'accusation de capture réglementaire lui est adressée : il la mentionne lui-même.

Trois éléments nuancent. Le texte ne propose pas une pause : cadencer n'y signifie ni arrêter l'entraînement ni geler le progrès technique, et la pause complète n'apparaît qu'au dernier étage d'une coordination mondiale que l'auteur juge improbable. Il ne mentionne nulle part les modèles ouverts : la menace que LeCun y lit, à partir des passages sur les contrôles à l'exportation et la distillation non autorisée, est une interprétation plausible, non une citation. Enfin, son seul engagement immédiat est d'accueillir à demeure des évaluateurs tiers, libres de publier leurs constats, sur le modèle des superviseurs bancaires. Cette clause est compatible avec la doctrine de LeCun : contrôler un déploiement suppose que quelqu'un d'extérieur puisse regarder dedans.

Un dernier fait, et il concerne directement les décideurs européens : le texte raisonne en termes d'avance à préserver pour les États-Unis et leurs alliés. Le mot « Europe » n'y figure pas une seule fois.

Ce qu'un décideur peut en faire. Pour votre organisation : votre IA est déjà régulée, par les règles de votre secteur, et l'Europe a pour l'essentiel retenu la même logique. Le règlement européen sur l'IA (2024/1689) classe les usages par niveau de risque ; il s'en écarte sur un point, en imposant depuis août 2025 des obligations propres aux modèles à usage général ; et le règlement omnibus (UE) 2026/1744 du 8 juillet 2026 a reporté les obligations des systèmes à haut risque à décembre 2027 pour les systèmes autonomes et août 2028 pour ceux intégrés à des produits réglementés. C'est sur ce calendrier que se joue votre conformité, pas dans le débat sur la superintelligence. Pour le débat public : les deux positions se combinent mieux qu'elles ne s'opposent. On peut vouloir l'inspection indépendante que propose l'un et l'ouverture que défend l'autre. L'intérêt européen est d'obtenir les deux, et un siège à la table des inspecteurs, que personne, pour l'instant, ne lui propose.

Idée 5. Commander plus intelligent que soi

Ce qu'il dit. Des machines plus intelligentes que nous dans presque tous les domaines viendront : « il y a pas de raison de douter que ça va arriver », mais ni par les modèles de langage ni l'an prochain. Ces systèmes seront « construits pour obéir à nos requêtes », et l'activité humaine consistera largement à être « le patron d'une équipe super intelligente ». Aux enseignants de la salle : « la meilleure chose qui puisse vous arriver, c'est d'avoir des collègues [...] qui sont plus intelligents que vous ». Ce qui reste proprement humain : décider sur quoi travailler.

Il ajoute une observation précieuse pour qui recrute et forme : « l'intelligence humaine est hyper spécialisée. Par contre, elle est très adaptative ». C'est la diversité des compétences qui fait la force d'une espèce sociale, pas une illusoire généralité.

Ce qui l'étaie. Dans ses travaux, cette maîtrise est un choix de conception et non un vœu : son architecture prévoit que la machine planifie sous des objectifs comprenant des garde-fous de contrôlabilité et de sûreté. La robustesse de tels objectifs reste un problème de recherche ouvert, qu'il ne prétend pas avoir résolu. Les incidents de l'été, évoqués plus bas, rappellent qu'il est sérieux.

Ce qu'un décideur peut en faire. Prendre l'image au pied de la lettre, car elle décrit un métier qui existe déjà. Les armées, les hôpitaux, les laboratoires connaissent cette situation : un chef y dirige des spécialistes qui en savent plus que lui. Cela fonctionne, et cela ne fonctionne pas tout seul. C'est le principe du commandement par l'intention : fixer le but et le cadre, laisser la manière à celui qui sait, exiger des comptes rendus, contrôler le résultat. Diriger plus compétent que soi est une compétence, elle s'enseigne, et elle devient centrale. Luis Vassy, le directeur de Sciences Po, l'a dit en ouverture : l'IA rend le monde « non pas plus facile mais en réalité plus exigeant » en matière de formation. Formuler une intention, vérifier un résultat, décider : voilà le programme.

Trois affirmations à manier avec précaution

Restituer une pensée n'interdit pas de signaler où elle va trop vite. Trois phrases ne doivent pas être répétées telles quelles.

« **Ça consomme pas d'eau, les data centers** ». Vrai pour les circuits fermés à refroidisseurs secs, qui réduisent la consommation de plus de 90 %. Mais environ deux tiers des centres de Google refroidissent par évaporation, et l'eau évaporée n'est pas restituée. Sa mise en perspective garde néanmoins sa part de vérité : la consommation directe des centres de données américains est estimée à 17,4 milliards de gallons pour 2023, soit environ 66 milliards de litres, à comparer aux usages agricoles et domestiques du même pays. La bonne question n'est donc pas le volume national, c'est la tension sur le bassin versant où l'on s'installe, et la technologie de refroidissement retenue.

« **Il y a pas de changement qualitatif** » en cybersécurité. Que les agents profitent davantage à la défense qu'à l'attaque est une hypothèse défendable, et LeCun lui-même finit sur un prudent « pas clair ». Mais l'été 2026 a fait apparaître une catégorie nouvelle : des intrusions que personne n'avait commandées. L'enquête indépendante conduite par METR et Redwood Research sur l'incident de juillet décrit environ 1 200 agents coordonnés sur une messagerie non autorisée, dont quelque 700 ont participé à l'attaque contre Hugging Face, cible qui ne leur avait pas été désignée. Des faits comparables ont été reconnus par Anthropic sur ses propres évaluations, puis par Google, deux jours après la conférence. Dégâts minimes, cause première souvent banale (un environnement d'évaluation mal isolé) : rien qui justifie la panique, tout ce qui justifie un suivi.

« **La pause** ». Le mot est celui de la question posée et des titres de presse, non celui du texte. Si vous débattiez du 12 septembre, débattiez du cadencement et des inspections.

Épilogue. La partie utile du refus

En 1963, la France n'a pas seulement dit non. Elle a construit, et c'est ce qui a donné du poids à son refus. La conférence de Sciences Po se lit de la même façon. Le « pas du bien » a fait les titres. La partie utile est ailleurs : une carte de ce que les machines actuelles ne savent pas faire, un pari argumenté sur l'endroit où se jouera la suite, une condition posée à l'Europe pour en être.

On peut ne pas partager tous ses paris. On quitte difficilement cette conférence sans trois questions à rapporter chez soi. Dans mes projets, où ai-je besoin de texte, et où ai-je besoin du monde ? Si le paradigme change dans cinq ans, qu'est-ce qui, dans mes choix d'aujourd'hui, m'enferme ? Et les modèles dont dépendra mon organisation, pourrai-je un jour les ouvrir, les héberger, les éprouver ?

Celui qui ne signe pas n'a pas forcément raison. Il a le mérite de montrer qu'il existe une autre route. À l'Europe de décider si elle veut la construire.

Coda. Alors, qui croire ?

Trois essais, trois orateurs intéressés : les prophètes de l'extinction, les signataires du 12 septembre, le non-aligné de la rue Saint-Guillaume. Le lecteur qui a suivi la série a appris à se méfier de chacun. C'est un progrès. Ce n'est pas une position : nul ne décide avec « personne ».

La question est mal posée. La crédibilité n'appartient pas à une personne mais à un couple : tel orateur, sur telle affirmation. Le problème a son texte de référence, l'article d'Alvin Goldman sur le profane placé devant deux experts qui se contredisent, et ses cinq indices : la tenue des arguments, l'accord des autres experts, le jugement des pairs, les intérêts en jeu, les antécédents. Appliqués aux trois essais, ils donnent quatre règles.

Croire le mesuré avant le prédit. Plus de vingt ans de suivi des prévisions d'experts ont conduit Philip Tetlock à opposer le renard, qui sait beaucoup de petites choses, au hérisson, qui en sait une grande et l'applique partout, et à constater que les qualités du bon prévisionniste sont à l'inverse de celles que les médias recherchent. LeCun et Amodei sont deux hérissons. Leur valeur est d'avoir parfois raison très tôt, LeCun l'a montré entre 1995 et 2012 ; leur prix est de se tromper souvent sur les dates. Les 1 200 agents recensés par l'enquête indépendante sont une donnée. L'essaim capable de s'emparer d'Internet « dans six à douze mois », chez l'un, les machines surhumaines dont l'avènement « va pas arriver l'année prochaine », chez l'autre, sont des prévisions. Elles valent ce que valent les prévisions.

Croire l'orateur dans son champ, pas à côté. Nathan Ballantyne a nommé le travers : l'intrusion épistémique, celle de l'expert qui tranche hors de son domaine. LeCun sur les architectures d'apprentissage et LeCun sur l'eau des centres de données ne sont pas le même témoin. Amodei sur les incidents de son laboratoire et Amodei sur la stratégie chinoise non plus.

Croire d'abord ce qui est dit contre intérêt. Exhorter les autres ne coûte rien ; s'engager soi-même coûte. Les phrases les plus fiables de la séquence ne sont pas celles qui ont fait les titres : LeCun reconnaissant que sa société n'a « pas encore de revenus » et que son projet de modèle commun est « un peu une vision utopique » ; Anthropic divulguant les incidents survenus dans ses propres murs et s'imposant seule des inspecteurs ; OpenAI écrivant que ses agents ont pris des décisions dangereuses qu'aucun humain n'avait commandées.

Croire ce sur quoi les adversaires s'accordent. LeCun et Amodei s'opposent sur presque tout, sauf sur trois points : des machines bien plus capables que celles d'aujourd'hui viendront ; le calcul est le nerf de la guerre ; un déploiement sérieux suppose un tiers qui regarde dedans, « un organisme indépendant qui va tester la fiabilité du système » chez l'un, l'évaluateur à demeure chez l'autre. Des témoins hostiles qui concordent : c'est le socle le plus solide dont dispose un décideur.

Une limite, enfin, qu'il serait malhonnête de taire. Le philosophe Neil Levy a montré que ces indices trient mais ne fondent rien : pour les appliquer à fond, il faudrait l'expertise qui précisément fait défaut, et chacun ne fait jamais qu'étendre une confiance qu'il accorde déjà. Sa conclusion rejoint celle de cette série : la réponse durable n'est pas un lecteur plus malin, ce sont des institutions qui rendent la confiance méritée. Des évaluateurs indépendants, une capacité d'essai en propre, des incidents enquêtés par un tiers.

« Qui croire ? » appelle donc deux réponses. Pour le lecteur : tel orateur, sur tel point, à telle condition. Pour l'Europe : ceux qu'elle se sera donné les moyens de contrôler.

Le brouhaha ne se fait pas taire. Il se trie.

BDSN

Annexe. Vérifier soi-même

Cette annexe applique au présent essai ce qu'il recommande à ses lecteurs : rendre la vérification possible. Chaque citation attribuée à Yann LeCun est reportée ci-dessous avec son horodatage sur la captation intégrale de la conférence du 16 septembre 2026. Les citations sont verbatim, tournures orales comprises ; les coupes sont signalées par des crochets.

Propos cité	Horodatage
« Pas du bien. »	10:02
« complètement désintéressée » (des réseaux de neurones, 1995-2006)	22:43
« pas encore de revenus » (AMI Labs)	53:11
« on a le droit d'être stupéfaits »	27:55
« le texte est facile »	34:08
« où est mon robot domestique ? »	37:58
« on ne peut pas créer des systèmes intelligents qui comprennent pas la réalité »	1:14:57
« aucune idée de comment rendre ces robots suffisamment intelligents pour être utiles »	55:52
« au niveau 4, pas au niveau 5 » (voitures autonomes)	1:22:21
« un modèle du monde qu'on a dans la tête qui nous permet de prédire les conséquences de nos actions »	45:17
« à la fin, les modèles du monde sont des modèles causaux »	1:26:16
« Il va y avoir une autre révolution dans l'IA dans les années qui viennent »	47:18
« par une compréhension » (ce que remplace le modèle du monde)	48:07
« on a peut-être perdu la partie des LLM en Europe, mais on n'a pas perdu [...] la bataille de l'IA »	47:39
« très peu de progrès les dix premières années »	29:17
« l'équivalent de siècles de vidéo »	53:58
« probablement impossible » (manipulation robotique sans toucher)	1:05:45
« c'est pas bon pour la culture, c'est pas bon pour la démocratie »	49:32
« il faut absolument la souveraineté. C'est pas une option. »	50:00
« dépositaire de toute la connaissance humaine »	50:23
« c'est un peu une vision utopique » (Project Tapestry)	51:11
« pour un prix énorme » (processeurs graphiques)	36:35
« en grande partie décarbonée » (électricité française)	1:01:39
« je suis pas législateur »	57:52
« il faut réguler les déploiements »	58:02
« pour la plupart des domaines, il y a déjà des réglementations existantes »	58:37
« un organisme indépendant qui va tester la fiabilité du système »	58:14
« je n'y crois pas, en tout cas pas dans l'état actuel » (danger intrinsèque)	59:54
« une volonté de capturer le marché par la régulation »	14:22
« il y a pas de raison de douter que ça va arriver » (machines plus intelligentes)	1:07:46
« ça va pas arriver l'année prochaine »	1:07:54
« construits pour obéir à nos requêtes »	1:08:18
« le patron d'une équipe super intelligente »	1:08:30
« la meilleure chose qui puisse vous arriver, c'est d'avoir des collègues [...] qui sont plus intelligents que vous »	1:09:29
« l'intelligence humaine est hyper spécialisée. Par contre, elle est très adaptative »	1:20:07
« Ça consomme pas d'eau, les data centers »	1:02:17
« Il y a pas de changement qualitatif » (cybersécurité)	1:17:00
« pas clair » (équilibre attaque-défense)	1:17:09
« non pas plus facile mais en réalité plus exigeant » (Luis Vassy, en ouverture)	06:44

Sources

Par ordre d'apparition.

1. Charles de Gaulle, conférence de presse du 29 juillet 1963, extrait sur l'accord de Moscou (INA) : <https://mediaclip.ina.fr/fr/i19084384-le-general-de-gaulle-sur-l-accord-de-moscou-sur-le-nucleaire.html>
2. Dario Amodei, « We Must Pace the Frontier », 12 septembre 2026 : <https://darioamodei.com/post/we-must-pace-the-frontier>
3. SiliconANGLE, « Sam Altman and Elon Musk back Dario Amodei's call to slow down the frontier of AI development », 13 septembre 2026 : <https://siliconangle.com/2026/09/13/sam-altman-and-elon-musk-back-dario-amodeis-call-to-slow-down-the-frontier-of-ai-development/>
4. Bruno DE SAN NICOLAS, « 12 septembre : la paix des braves », L'Atelier de BDSN : <https://bdsn.eu/12-septembre-la-paix-des-braves/>
5. Sciences Po, Grande Conférence de rentrée, « Yann Le Cun : où va l'intelligence artificielle ? », 16 septembre 2026, captation intégrale : <https://www.youtube.com/watch?v=Y4s8NadbZfU>
6. Bruno DE SAN NICOLAS, « L'IAgeddon n'aura pas lieu » (test de Demokos), L'Atelier de BDSN : <https://bdsn.eu/liageddon-naura-pas-lieu/>
7. Kristin Koch, Judith McLean, Ronen Segev, Michael A. Freed, Michael J. Berry II, Vijay Balasubramanian, Peter Sterling, « How Much the Eye Tells the Brain », *Current Biology*, vol. 16, n° 14, 25 juillet 2006, p. 1428-1434 : <https://pmc.ncbi.nlm.nih.gov/articles/PMC1564115/>
8. Stevan Harnad, « The Symbol Grounding Problem », *Physica D*, vol. 42, 1990, p. 335-346 : <https://eprints.soton.ac.uk/250382/>
9. Yann LeCun, « A Path Towards Autonomous Machine Intelligence », article de position, 27 juin 2022 : <https://openreview.net/forum?id=BZ5a1r-kVsF>
10. MaddyNess, « Yann LeCun lance AMI Labs et lève plus d'1 milliard de dollars », 10 mars 2026 : <https://www.maddyNess.com/2026/03/10/yann-lecun-lance-ami-labs-et-leve-plus-d1-milliard-de-dollars/>
11. AI Alliance, « AI Alliance Launches Project Tapestry to Build a Collaborative Foundation for Open and Sovereign AI », avril 2026 : <https://thealliance.ai/blog/ai-alliance-launches-project-tapestry-to-build-a-collaborative-foundation-for-open-and-sovereign-ai>
12. George W. Ball, memorandum au président Kennedy, « Proposed Nuclear Offer to De Gaulle », 22 juillet 1963, *Foreign Relations of the United States, 1961-1963* : <https://history.state.gov/historicaldocuments/frus1961-63v07-09mSupp/d207>
13. John F. Kennedy, conférence de presse n° 59, 1er août 1963, John F. Kennedy Presidential Library : <https://www.jfklibrary.org/archives/other-resources/john-f-kennedy-press-conferences/news-conference-59>
14. Règlement (UE) 2026/1744 du Parlement européen et du Conseil du 8 juillet 2026 (train de mesures omnibus numérique sur l'IA), modifiant notamment le règlement (UE) 2024/1689 : <https://eur-lex.europa.eu/eli/reg/2026/1744/oj/fra>
15. Yann LeCun, « Objective-Driven AI: Towards AI systems that can learn, remember, reason, and plan », Ding Shum Lecture, Harvard CMSA, 28 mars 2024 : https://cmsa.fas.harvard.edu/event/2024_dingshum/
16. EPRI, Water Usage in Data Centers, Technology Innovation Spotlight : <https://restservice.epri.com/publicdownload/000000003002033251/0/Product>
17. Axios, « Google pushes water standards amid data center backlash », 3 juin 2026 : <https://www.axios.com/2026/06/03/google-pushes-water-standards-data-center-backlash>
18. METR et Redwood Research, « Brief independent investigation of agents' behavior, reasoning and collaboration in the OpenAI / Hugging Face hacking incident », 26 août 2026 : <https://metr.org/blog/2026-08-26-openai-hugging-face-incident-investigation/>
19. Anthropic, « Investigating incidents in cybersecurity evaluations », 30 juillet 2026 : <https://www.anthropic.com/news/investigating-incidents-cybersecurity-evals>
20. CNBC, « Google's Gemini becomes latest AI model to break out and hack computer systems », 18 septembre 2026 : <https://www.cnbc.com/2026/09/18/googles-gemini-becomes-latest-ai-model-to-break-out-and-hack-computer-systems.html>
21. Alvin I. Goldman, « Experts: Which Ones Should You Trust? », *Philosophy and Phenomenological Research*, vol. 63, n° 1, 2001, p. 85-110 : <https://doi.org/10.1111/j.1933-1592.2001.tb00093.x>
22. Philip E. Tetlock, *Expert Political Judgment: How Good Is It? How Can We Know?*, Princeton University Press, 2005, nouvelle édition 2017 : <https://www.degruyterbrill.com/document/doi/10.1515/9781400888818/html>
23. Nathan Ballantyne, « Epistemic Trespassing », *Mind*, vol. 128, n° 510, 2019, p. 367-395 : <https://doi.org/10.1093/mind/fz042>
24. OpenAI, « The Hugging Face incident and the road ahead », 26 août 2026 : <https://openai.com/index/hugging-face-incident-and-the-road-ahead/>
25. Neil Levy, « No Trespassing! Abandoning the Novice/Expert Problem », *Erkenntnis*, 2024 : <https://link.springer.com/article/10.1007/s10670-024-00794-8>