



Antoine Pallot

9 SEPTEMBRE 2026

ECLAIRAGE

Du refus à la collaboration : comment les grandes entreprises d'IA ont changé de discours quant aux usages militaires

grip

En février 2026, alors qu'*Anthropic* fournit déjà certains services d'intelligence artificielle (IA) au Pentagone via un partenariat avec l'entreprise *Palantir*¹, le Département de la Guerre exige que l'entreprise retire les restrictions contractuelles qu'elle avait formulées sur deux usages spécifiques de son IA : la surveillance de masse et les armes entièrement autonomes. *Anthropic* refuse². Le secrétaire à la Guerre américain Pete Hegseth désigne alors publiquement l'entreprise comme une « *menace pour la chaîne d'approvisionnement* »³, label jusqu'alors réservé aux entreprises basées dans des États que les États-Unis considèrent comme des adversaires, par exemple les entreprises chinoises *Huawei*⁴ ou *ZTE*⁵. Quelques heures plus tard, *OpenAI* annonce avoir signé le contrat que son concurrent venait de refuser⁶. La séquence est saisissante d'autant plus que ces deux entreprises n'avaient, au départ, aucune intention de travailler pour le Pentagone.

Pour comprendre cela, il faut remonter à 2018 et au projet *Maven*, un programme du Pentagone visant à intégrer l'IA dans l'analyse d'images de drones et les opérations de renseignement militaires⁷. Ce projet révèle alors une fracture au sein de la tech américaine. Certaines entreprises, comme *Palantir* ou *Anduril*, y participent sans résistance interne : leur vocation sécuritaire, voire militaire, est constitutive de leur identité⁸. Pour *Google* en revanche, la participation à *Maven* déclenche une protestation massive en interne, qui contraint l'entreprise à y renoncer et à formaliser des principes lui interdisant explicitement le développement de systèmes d'armement ou de surveillance⁹. *OpenAI* inscrit une interdiction similaire dans ses conditions

d'usage¹⁰. *Anthropic*, fondée en 2021 par d'anciens employés d'*OpenAI*, oriente ses travaux vers des modèles d'IA soumis à une grille de valeurs contraignante, une « *constitution* » que l'IA ne peut enfreindre, quelle que soit la demande de son utilisateur¹¹. Ces trois entreprises affichent donc, au moins dans un premier temps, une réticence de principe vis-à-vis du monde militaire.

Ce positionnement s'érode pourtant à partir de 2024. En janvier, *OpenAI* retire, sans l'annoncer publiquement, la restriction de ses conditions d'utilisation quant aux applications « *militaires et de guerre* ». Elle reconnaît par la suite vouloir permettre des usages de ses IA relevant de la sécurité nationale¹². *Google* supprime, en février 2025, la section entière de ses principes d'IA listant les « *applications qu'elle ne poursuivra pas* » (armes, surveillance et technologies contrevenant au droit international)¹³. *Anthropic* va même plus loin : en septembre 2024, elle noue un partenariat avec *Palantir* pour déployer ses modèles sur des réseaux classifiés de la Défense, afin d'y accélérer le traitement des données¹⁴. Elle demeure néanmoins la seule des trois à maintenir des contraintes contractuelles explicites sur les usages de son IA¹⁵. Ce partenariat s'inscrit au sein du projet *Maven*, qui a depuis donné naissance au *Maven Smart System* (MSS), sa plateforme phare. Développée par *Palantir*, cette dernière agrège de vastes flux de données de renseignement et de reconnaissance au sein d'une interface unique d'aide au ciblage et au commandement interarmées (CJADC2) sur laquelle *Claude* (l'IA principal d'*Anthropic*) est déployé depuis novembre 2024¹⁶. Selon plusieurs sources, le modèle aurait participé

à l'élaboration de l'enlèvement de Nicolas Maduro¹⁷ et des frappes en Iran¹⁸.

Cet *éclairage* propose une analyse critique des discours produits par ces trois entreprises pour justifier un rapprochement avec les institutions de défense américaines, alors même qu'elles avaient auparavant manifesté leur réticence à l'égard des applications militaires de leurs IA. Ces

discours s'organisent autour de trois registres argumentatifs : des arguments empruntés aux justifications classiques de l'industrie de défense (1); des arguments liés à l'évolution du contexte technologique et géopolitique (2); et enfin, des arguments portant sur le rôle que ces entreprises revendiquent dans la gouvernance de leurs propres systèmes (3). C'est sur ce troisième registre que les discours divergent le plus nettement.

1. La rhétorique classique des entreprises de défense

Les trois entreprises partagent une partie de leur argumentaire. Elles mobilisent plusieurs éléments des discours traditionnels de l'industrie de défense aux États-Unis : soutien aux démocraties, nécessité du leadership militaire et technologique américain, protection du bien-être de la population. Cette rhétorique se retrouve presque mot pour mot dans la communication des acteurs majeurs de l'industrie. RTX invoque ainsi le soutien aux démocraties : « *au cours de l'année, RTX a soutenu les États-Unis et leurs alliés dans la défense de la démocratie à travers le monde*¹⁹. » General Dynamics reprend le même registre sur un ton plus moral : « *Nous avons le devoir moral de soutenir notre pays et d'appuyer les alliés de notre pays dans leurs efforts pour préserver la sécurité des démocraties*²⁰. » Boeing y fait également appel : « *partout dans le monde, les forces armées américaines et leurs alliés comptent sur nos produits de défense pour garantir et protéger la liberté, ainsi que pour apporter une aide humanitaire en cas de crise*²¹. »

Reprenant ce même discours, Google écrit en février 2025 :

« Nous pensons que les démocraties doivent mener le développement de l'IA, guidées par des valeurs fondamentales, telles que la liberté, l'égalité et le respect des droits humains. Et nous pensons que les entreprises, les gouvernements et les organisations partageant ces valeurs devraient travailler ensemble afin de créer une IA qui protège les personnes, promeut la croissance mondiale et soutient la sécurité nationale »²².

Si Google est l'entreprise qui mobilise le plus ce registre²³, on le retrouve aussi dans la communication d'Open AI et d'Anthropic. OpenAI formule ainsi la même idée dans des termes quasi identiques : « *Soutenir les efforts des États-Unis et de leurs alliés pour faire progresser l'IA d'une manière qui défende les valeurs démocratiques est essentiel à notre mission*²⁴. » Anthropic,

dans l'annonce de son partenariat avec le Département de la Défense de juillet 2025, reprend le même cadre :

« Nous pensons que les démocraties doivent travailler ensemble pour s'assurer que le développement de l'IA renforce les valeurs démocratiques à l'échelle mondiale en conservant leur leadership technologique afin de se prémunir contre toute utilisation abusive à des fins autoritaires²⁵. »

Ainsi, ces trois entreprises initialement réticentes à toute application militaire de leurs IA reprennent aujourd'hui un discours porté jusque-là par les producteurs d'armement. Ce glissement rhétorique appelle toutefois une lecture critique.

Il est important de souligner que le discours légitimateur des entreprises d'armement est un discours partial qui occulte certains aspects de leur activité, notamment les intérêts économiques de ces entreprises et

leurs exportations vers de nombreux pays, y compris vers des régimes autoritaires²⁶. De la même façon, le discours des entreprises d'IA est lui aussi partiel et orienté. *OpenAI* présente ainsi son partenariat de mai 2025 avec les Émirats arabes unis, régulièrement classés par l'ONG américaine *Freedom House* parmi les régimes autoritaires²⁷, comme « enraciné dans des valeurs démocratiques²⁸ ». Cette communication passe aussi sous silence les intérêts économiques sous-jacents à l'accord : celui-ci permet notamment à *OpenAI* d'accéder à d'importantes capacités de calcul, grâce au fonds d'investissement émirati G42²⁹. *Google* a de son côté noué, au même moment, un partenariat avec *Humain*³⁰, le nouvel opérateur d'IA du fonds souverain saoudien³¹. Ces choix contredisent directement la rhétorique démocratique que ces entreprises mobilisent, et suggèrent qu'elle répond moins à une conviction qu'à un impératif commercial : sécuriser l'accès aux capitaux et à l'énergie d'un secteur en pleine expansion, y compris auprès de régimes autoritaires.

2. Les arguments liés à l'évolution du contexte technologique

Ce premier registre discursif n'est pas le seul que ces entreprises mobilisent. Un deuxième registre vient rendre compte de l'évolution de leur position en la présentant non comme un reniement éthique, mais comme une réponse logique à un contexte géopolitique et technologique transformé.

OpenAI développe à cet égard un argument singulier : les modèles d'IA développés par

le passé n'auraient pas été suffisamment robustes pour des usages sensibles. C'est leur maturité croissante qui rendrait aujourd'hui acceptables des applications auparavant exclues. L'entreprise écrit ainsi en février 2026 : « Au départ, nous n'avons pas conclu immédiatement un contrat pour un déploiement en milieu classifié, car nous estimions que nos mécanismes de protection et nos systèmes n'étaient pas encore prêts³² ».

Ainsi, si l'on en croit cette rationalisation *a posteriori*, le refus initial n'était pas la marque d'une objection de principe, mais une précaution technique provisoire. L'argument opère une conversion de registre : ce qui relevait du débat moral devient une question de maturité technologique.

Cette logique est en partie reprise dans le discours d'*Anthropic*. L'entreprise maintient aujourd'hui deux restrictions explicites : son IA ne doit pas être utilisée pour la surveillance de masse ou équiper des armes entièrement autonomes. Si la restriction concernant la surveillance de masse semble reposer sur un refus de principe, l'interdiction des armes entièrement autonomes est en revanche justifiée par l'insuffisante fiabilité actuelle des modèles d'IA pour des usages létaux³³. Sur ce point, *Anthropic* mobilise un raisonnement proche de celui d'*OpenAI* : la limite tient au niveau de maturité des systèmes plus qu'à une objection de principe. Cette ligne rouge pourrait donc être susceptible d'évoluer à mesure que les capacités techniques des systèmes progresseront.

Sans reprendre exactement le même argument, *Google* adopte la même logique, mobilisant un argument à la fois géopolitique et technologique. Ainsi, *Google* avance que « Depuis la publication de nos principes relatifs à l'IA en 2018, cette technologie a évolué rapidement. [...] Une course mondiale au leadership en matière d'IA se déroule actuellement dans un contexte géopolitique de plus en plus complexe³⁴ ». Sans préciser davantage la nature de cette complexité censée justifier le changement de politique, celui-ci apparaît comme une adaptation à un nouvel environnement.

Cette logique se double, chez *Google* comme chez *Anthropic*, d'un second argument : puisque l'IA est de toute façon développée et déployée à des fins militaires par de nombreux acteurs, mieux vaudrait y participer en s'assurant de le faire de manière responsable, plutôt que de s'en exclure. *Google* l'exprime ainsi :

« Guidés par nos principes en matière d'IA, nous continuerons à nous concentrer sur la recherche et les applications de l'IA qui s'alignent avec notre mission [...] Ces évaluations sont particulièrement importantes à mesure que l'IA est développée par de nombreuses organisations et gouvernements pour des applications dans des domaines tels que la santé, la science, la robotique, la cybersécurité, les transports, la sécurité nationale, l'énergie, le climat, etc.³⁵ »

Anthropic développe une idée similaire : considérant que la militarisation de l'IA est appelée à se développer durablement, la question n'est plus de savoir si elle doit être utilisée dans le domaine de la défense, mais dans quelles conditions elle le sera³⁶.

Ce raisonnement n'a rien d'un argument éthique. Il ne justifie pas une pratique, il la présente comme inévitable. Or, *Google* et *Anthropic* ne sont pas des observateurs extérieurs à cette dynamique : ils comptent parmi les tout premiers acteurs qui, par leurs choix technologiques et commerciaux, façonnent l'ampleur et les formes que prend aujourd'hui la militarisation de l'IA.

Cette reformulation en termes techniques mérite également d'être confrontée aux discours initiaux. Lorsque Google refuse en 2018 de renouveler sa participation au projet *Maven*, ce n'est pas au nom d'une immaturité technologique qu'elle formule ses principes, mais au nom d'un refus explicitement moral : l'entreprise s'engage à ne pas développer de technologies « *dont l'objectif principal est de causer ou de faciliter directement des dommages aux personnes*³⁷ », ni de systèmes de surveillance violant les droits humains. De même, jusqu'en janvier 2024, la politique d'usage d'*OpenAI*

interdisait explicitement le « *développement d'armes* » et les usages « *militaires et de guerre* » classés parmi les activités à haut risque de dommage physique³⁸. Ces formulations ne relevaient pas d'une réserve technique provisoire, mais d'un jugement sur la nature intrinsèquement problématique de ces usages. Présenter aujourd'hui l'évolution de ces positions comme une simple adaptation à la maturité croissante des systèmes ou à un « *contexte géopolitique* » revient donc à effacer rétrospectivement ce qui était, à l'origine, un choix éthique.

3. Le rôle des entreprises dans la gouvernance de leurs propres systèmes

C'est sur un troisième registre que les discours des trois entreprises divergent le plus : elles ne revendiquent pas le même rapport à l'État ni la même capacité ou volonté de peser sur les conditions concrètes d'utilisation de leurs technologies.

Google affirme développer l'IA « *en accord avec les principes de droit international et de droits humains largement acceptés*³⁹ » et s'engage à évaluer chaque application selon un principe général : « *si les bénéfices l'emportent substantiellement sur les risques potentiels*⁴⁰ », selon les termes de son vice-président Tom Lue en mars 2026. Il s'agit donc d'une évaluation au cas par cas et en amont des usages potentiels. Contrairement à ses concurrents, l'entreprise ne prétend pas exercer un droit de regard sur la façon

dont un utilisateur étatique emploierait ses technologies.

Anthropic non plus ne prétend pas intervenir dans l'usage ponctuel de sa technologie par les autorités militaires :

« *Anthropic est conscient que c'est le Département de la Guerre, et non les entreprises privées, qui prend les décisions militaires. Nous n'avons jamais émis d'objections à l'égard d'opérations militaires spécifiques ni tenté de limiter l'utilisation de notre technologie de manière ponctuelle*⁴¹. »

Son levier se situe ailleurs, en amont, dans la définition de catégories d'usage structurellement exclues. Deux d'entre

elles sont explicitement écartées de ses contrats avec le Département de la Guerre : la surveillance de masse et les armes entièrement autonomes. Même face aux pressions du Pentagone, l'entreprise a maintenu cette ligne, affirmant que « ces menaces ne changent pas notre position : nous ne pouvons pas en bonne conscience accéder à leur demande⁴². » Ces deux lignes rouges ne reposent cependant pas sur le même fondement, comme on l'a vu plus haut : l'une relève d'un engagement éthique, tandis que l'autre d'un seuil technique appelé à évoluer.

Le contrôle qu'*Anthropic* revendique peut produire des effets concrets. À mesure que l'IA devient centrale dans la prise de décision militaire, pouvoir exclure certains usages revient à peser sur les conditions dans lesquelles ces décisions pourront ou non s'appuyer sur l'IA, voire lui être déléguées.

Le discours d'*OpenAI* revendique quant à lui un autre type de contrôle :

« Nous conservons un contrôle total sur notre dispositif de sécurité : le déploiement se fait exclusivement dans le cloud, du personnel d'OpenAI habilité reste impliqué et nous bénéficions de protections contractuelles solides⁴³. »

Cette logique fait de l'entreprise non plus un simple fournisseur de technologie, mais un acteur dont l'implication continue dans le déploiement de ses systèmes garantirait une protection contre les risques de mésusages. Elle mobilise d'ailleurs cet argument pour se distinguer des autres entreprises :

« D'autres laboratoires d'IA ont réduit ou supprimé leurs garde-fous et s'appuient principalement sur des politiques d'utilisation comme mécanisme de protection dans leurs déploiements liés à la sécurité nationale. Nous pensons que notre approche protège mieux contre les utilisations inacceptables⁴⁴. »

Cette forme de contrôle appelle cependant également une remarque : rester impliqué dans la boucle opérationnelle ne revient pas à se réserver le droit de refuser un usage. Le contrôle qu'*OpenAI* revendique porte sur les modalités du déploiement, non sur la nature des usages eux-mêmes. Cette limite n'est pas que théorique : c'est précisément dans le cadre de ce dispositif de contrôle qu'*OpenAI* a accepté, en février 2026, le contrat avec le Département de la Guerre qu'*Anthropic* venait de refuser. Prétendre maintenir une présence continue dans le déploiement n'a donc pas empêché l'entreprise de céder sur des conditions que sa concurrente jugeait rédhibitoires. Ceci interroge la portée réelle de ce mécanisme comme garde-fou contre les « utilisations inacceptables ».

Ces trois formes de contrôle ne se valent donc pas. *Google* renonce à tout droit de regard une fois sa technologie déployée. *OpenAI* mise sur une présence continue dans la chaîne de déploiement, sans que cela lui donne pour autant la capacité de refuser un usage. *Anthropic* seule conserve un levier antérieur au déploiement lui-même, par l'exclusion contractuelle de catégories d'usage entières – un levier plus étroit qu'un droit de regard sur chaque usage, mais réel, et c'est précisément ce levier que le Pentagone a cherché à faire sauter en février 2026.

Conclusion

L'analyse du discours de ces trois entreprises d'IA montre que le rapprochement avec les institutions de défense ne s'est pas accompagné d'un abandon des références éthiques initiales, mais d'une transformation des arguments mobilisés pour les articuler à la coopération militaire. Sur le registre classique (préservation de la démocratie, leadership technologique, compétition géopolitique), les trois entreprises convergent et reprennent à leur compte une rhétorique employée habituellement par l'industrie de défense. Sa portée est cependant limitée par les faits : les partenariats noués par *OpenAI* et *Google* avec les Émirats arabes unis et l'Arabie saoudite montrent que l'invocation des « valeurs démocratiques » cède vite le pas à des impératifs commerciaux, y compris auprès de régimes autoritaires. Les différences apparaissent sur les deux autres registres. Sur celui de l'adaptation au changement, *OpenAI* efface ses objections passées en les réinterprétant comme des précautions techniques devenues caduques, tandis que *Google* et *Anthropic* les retournent en argument de présence nécessaire : puisque la militarisation de l'IA est vouée à se développer durablement, mieux vaut y participer sous conditions plutôt que de s'en exclure. Cet argument, on l'a vu, n'a rien d'éthique. Il présente comme inévitable

une dynamique que ces mêmes entreprises contribuent, par leurs propres choix, à façonner. Sur celui de la gouvernance, *Google* normalise sa participation sans revendiquer de contrôle propre. *OpenAI* fait de son implication continue dans le déploiement un mécanisme de gouvernance. *Anthropic* seule conserve un levier antérieur au déploiement, par l'exclusion contractuelle de catégories d'usage.

Ce dernier point est sans doute le plus nouveau. Aucune des trois entreprises ne fait entièrement confiance au gouvernement américain pour déployer seule leurs technologies de manière responsable, mais aucune non plus ne dispose d'un mécanisme pleinement robuste pour l'imposer. Si *Anthropic* n'intervient jamais dans l'usage ponctuel de sa technologie, elle maintient sa volonté d'en exclure certaines catégories en amont. Cette posture a d'ailleurs un coût réel : c'est précisément parce qu'*Anthropic* a maintenu ses lignes rouges contractuelles que le Pentagone a rompu avec elle et s'est tourné vers *OpenAI*. Le discours sur les garde-fous n'est donc pas que rhétorique, il a des conséquences. Reste à savoir si cette capacité de résistance perdurera à mesure que les enjeux financiers et politiques autour des contrats d'État continueront de croître.

Références

1. Entreprise spécialisée dans l'analyse de données pour le renseignement, notamment en partenariat avec la police de l'immigration américaine ICE.
2. COPP Tara et DUNCAN Ian, « [Anthropic rejects Pentagon terms for lethal use of its chatbot Claude](#) », *The Washington Post*, 26 février 2026.
3. « [...] Supply-Chain Risk [...] » [traduction libre] : Secretary of War Pete Hegseth, « [This week, Anthropic delivered a master class in arrogance and betrayal as well as a textbook case of how not to do business with the United States Government or the Pentagon.](#) », X, 27 février 2026.
4. « [FCC Designates Huawei as a National Security Threat](#) », *Federal Communications Commission*, s. d. (consulté le 31 juillet 2026).
5. « [FCC Designates ZTE as a National Security Threat](#) », *Federal Communications Commission*, s. d. (consulté le 31 juillet 2026).
6. GABBATT Adam et KERR Dara, « [OpenAI to work with Pentagon after Anthropic dropped by Trump over company's ethics concerns](#) », *The Guardian*, 28 février 2026.
7. PELLERIN Cheryl, « [Project Maven to deploy computer algorithms to war zone by year's end](#) », *U.S. Department of War*, 21 juillet 2017.
8. « [Defense](#) », *Palantir*, s. d. (consulté le 15 mai 2026) ; « [Mission](#) », *Anduril*, s. d. (consulté le 15 mai 2026).
9. WAKABAYASHI Daisuke et SHANE Scott, « [Google will not renew Pentagon contract that upset employees](#) », *The New York Times*, 1^{er} juin 2018.
10. BIDDLE Sam, « [OpenAI quietly deletes ban on using ChatGPT for "military and warfare"](#) », *The Intercept*, 12 janvier 2024.
11. BAI Yuntao et al., « [Constitutional AI : Harmlessness from AI Feedback](#) », *Anthropic*, 15 décembre 2022.
12. ZEFF Maxwell, « [OpenAI Had Banned Military Use. The Pentagon Tested Its Models Through Microsoft Anyway](#) », *Wired*, 5 mars 2026 ; BIDDLE Sam, « [OpenAI quietly deletes ban on using ChatGPT](#) », loc. cit.
13. TIKU Nitasha et DE VYNCK Gerrit, « [Google drops pledge not to use AI for weapons or surveillance](#) », *The Washington Post*, 4 février 2025.
14. « [Anthropic and Palantir Partner to Bring Claude AI Models to AWS for U.S. Government Intelligence and Defense Operations](#) », *Palantir*, 11 juillet 2024.
15. « [Statement from Dario Amodei on our discussions with the Department of War](#) », *Anthropic*, 26 février 2026.
16. MANDE Matt et ALLEN C. Gregory, « [What Is Maven Smart System, and What Does It Do?](#) », *Center for Strategic and International Studies*, 2 juin 2026.
17. RAMKUMAR Amrith, HAGEY Keach et BERGENGRUEN Vera, « [Pentagon Used Anthropic's Claude in Maduro Venezuela Raid](#) », *The Wall Street Journal*, 15 février 2026.

18. COPP Tara, DWOSKIN Elizabeth et DUNCAN Ian, « [Anthropic's AI tool Claude central to U.S. campaign in Iran, amid a bitter feud](#) », *The Washington Post*, 4 mars 2026.
19. « Throughout the year, RTX supported the United States and its allies in safeguarding democracy worldwide. » [traduction libre] : « [2023 Annual Report](#) », RTX, 2024.
20. « It is our moral imperative to support our nation and to support our nation's allies in the work they do to keep democracies safe » [traduction libre] : BOBROW Emily, « [General Dynamics CEO Phebe Novakovic Believes in Patriotism and Resilience](#) », *The Wall Street Journal*, 25 juin 2021.
21. « Around the world, the U.S. armed forces and allies depend on our defense products to secure and protect freedom and provide humanitarian aid in time of crisis. » [traduction libre] : ORTBERG Kelly, « [2025 Address to Shareholders](#) », Boeing, 24 avril 2025.
22. « We believe democracies should lead in AI development, guided by core values like freedom, equality, and respect for human rights. And we believe that companies, governments, and organizations sharing these values should work together to create AI that protects people, promotes global growth, and supports national security. » [traduction libre] : MANYIKA James et HASSABIS Demis, « [Responsible AI: Our 2024 report and ongoing work](#) », Google blog, 4 février 2025.
23. Ibid. ; « [Our AI Principles](#) », Google, s. d. (consulté le 31 juillet 2026) ; JADOTTE Marcus, « [How we're helping democracies stay ahead of digital threats](#) », Google, 11 février 2026.
24. « Supporting US and allied efforts to advance AI in a way that upholds democratic values is essential to our mission » [traduction libre] : « [OpenAI's approach to AI and national security](#) », OpenAI, 24 octobre 2024.
25. « We believe democracies must work together to ensure AI development strengthens democratic values globally by maintaining technological leadership to protect against authoritarian misuse » [traduction libre] : « [Anthropic and the Department of Defense to advance responsible AI in defense operations](#) », Anthropic, 14 juillet 2025.
26. Pour une analyse similaire portant sur les entreprises de production d'armement en Allemagne, voir : KÜNTZLER Marie, « [Regard critique sur la nouvelle image de l'industrie militaire en Allemagne](#) », Note d'analyse du GRIP, 7 août 2025.
27. Le score de 18/100 en matière de liberté a été attribué aux Émirats arabes unis par le *Freedom Index* : « [United Arab Emirates](#) », *Freedom House*, s. d. (consulté le 31 juillet 2026).
28. « [Introducing Stargate UAE](#) », OpenAI, 22 mai 2025.
29. FRIED Ina, « [OpenAI, UAE to build massive AI center in Abu Dhabi](#) », AXIOS, 22 mai 2025.
30. « [Google Cloud and PIF Advance AI Hub in Saudi Arabia](#) », Google, 13 mai 2025.
31. Le score de 9/100 en matière de liberté a été attribué à l'Arabie saoudite par le *Freedom Index* : « [Saudi Arabia](#) », *Freedom House*, s. d. (consulté le 31 juillet 2026).
32. « [Notre accord avec le Department of War américain](#) », OpenAI, 28 février 2026.
33. « [Statement from Dario Amodei on our](#)

- [discussions with the Department of War](#) », loc. cit.
34. « Since we first published our AI Principles in 2018, the technology has evolved rapidly. [...] There's a global competition taking place for AI leadership within an increasingly complex geopolitical landscape. » [traduction libre] : MANYIKA James et HASSABIS Demis, « [Responsible AI : Our 2024 report and ongoing work](#) », loc. cit.
 35. « Guided by our AI Principles, we will continue to focus on AI research and applications that align with our mission [...] These assessments are particularly important as AI is increasingly being developed by numerous organizations and governments for uses in fields like healthcare, science, robotics, cybersecurity, transportation, national security, energy, climate, and more. » [traduction libre] : MANYIKA James et HASSABIS Demis, « [Responsible AI : Our 2024 report and ongoing work](#) », loc. cit.
 36. « [Statement from Dario Amodei on our discussions with the Department of War](#) », loc. cit.
 37. « [...] whose principal purpose or implementation is to cause or directly facilitate injury to people » [traduction libre] : « [Google won't deploy AI to build military weapons](#) », The Week, 10 décembre 2018.
 38. « weapons development », « military and warfare » [traduction libre] : BIDDLE Sam,
 39. « [OpenAI quietly deletes ban on using ChatGPT](#) », loc. cit.
 40. « [...] whether the benefits substantially outweigh potential risks » [traduction libre] : LANGLEY Hugh, « [Google told staff worried about Pentagon AI deals that the company is "leaning more" into national security contracts](#) », Business Insider, 19 mars 2026.
 41. « Anthropic understands that the Department of War, not private companies, makes military decisions. We have never raised objections to particular military operations nor attempted to limit use of our technology in an ad hoc manner » [traduction libre] : « [Statement from Dario Amodei on our discussions with the Department of War](#) », loc. cit.
 42. « [...] these threats do not change our position: we cannot in good conscience accede to their request » [traduction libre] : « [Statement from Dario Amodei on our discussions with the Department of War](#) », loc. cit.
 43. « [Notre accord avec le Department of War américain](#) », loc. cit.
 44. Ibid.

L'auteur

Antoine Pallot est assistant de recherche au GRIP. Il est titulaire d'un master en relations internationales et d'un master spécialisé en droit international de l'Université libre de Bruxelles. Il a consacré son mémoire de fin d'études à l'application du principe de précaution aux utilisations militaires de l'IA.

Pour citer cette publication

PALLOT Antoine, « Du refus à la collaboration : comment les grandes entreprises d'IA ont changé de discours quant aux usages militaires », *Éclairage du GRIP*, 9 septembre 2026.



Le GRIP bénéficie du soutien
du Service de l'Éducation
permanente de la Fédération
Wallonie-Bruxelles.

Photo de couverture :

De gauche à droite : Demis Hassabis, cofondateur et directeur général de *Google Deepmind* ; Dario Amodei, cofondateur et PDG d'*Anthropic* ; Sam Altman, cofondateur et PDG d'*OpenAI*. - Crédit : [Flickr - World Economic Forum](#), [Flickr - TechCrunch](#), [Flickr - World Economic Forum](#).

Les opinions exprimées dans le présent document ne reflètent pas nécessairement une position du GRIP dans son ensemble.

Tous droits réservés. © Groupe de recherche et d'information sur la paix et la sécurité

Groupe de recherche et d'information sur la paix et la sécurité

Mundo-Madou -Avenue des Arts, 7-8
1210 Saint-Josse-ten-Noode, Belgique

Tél. : +32 (0) 473 982 820 - admi@grip.org -
www.grip.org





Le GRIP a pour objectif la recherche et l'information sur les questions de paix et de sécurité internationales. Ses travaux portent sur :

- les conflits armés, leur gestion, leur prévention et leur résolution, notamment grâce aux outils du maintien de la paix, ainsi que leurs impacts humains et écologiques ;
- les politiques de sécurité et de défense, notamment au niveau européen ou en Afrique ;
- la production, le commerce, l'usage et le contrôle des armements.

Le GRIP entend contribuer à un monde moins armé et plus sûr en faisant la promotion des valeurs de paix, du règlement pacifique des différends et de la sécurité humaine. Il prend en compte dans ses travaux les conséquences humaines, sociales, économiques et écologiques des activités militaires et de production d'armement.

Le GRIP considère que les enjeux de sécurité et de défense doivent faire l'objet d'un débat démocratique et pluraliste. Il œuvre pour préserver l'accès des citoyen-ne-s, des médias et des décideur-euse-s politiques à une recherche indépendante sur ces questions.

5 BONNES RAISONS DE SOUTENIR LE GRIP

Le GRIP a pour mission d'étudier les conflits et les conditions de la paix. Il le fait dans l'optique de donner aux citoyens, à la société civile et aux élus accès à des analyses indépendantes permettant aux décideurs comme au grand public de renforcer leurs capacités critiques face à des enjeux complexes où s'entremêlent des intérêts politiques et économiques et des conceptions normatives et éthiques parfois contradictoires. En faisant un don au GRIP, vous participez au renforcement de ses moyens et œuvrez à :

1. Développer une recherche indépendante sur la paix ;
2. Consolider les capacités en tant que force de proposition auprès des décideurs politiques ;
3. Garantir l'accès en langue française à une recherche rigoureuse et accessible au public ;
4. Former une relève à qui il incombera de relever les défis de demain ;
5. Préserver l'activité Édition du GRIP qui permet de mettre de l'avant les combats des acteurs au service de la paix qu'ils soient journalistes, médiateurs ou militants des droits de la personne.

Le GRIP ne saurait accomplir efficacement sa mission d'information et de sensibilisation du public sans le soutien de donateurs motivés par la défense de la paix comme bien commun.

En soutenant le GRIP, vous contribuez au renforcement d'une recherche indépendante et de qualité au service de la société civile sur de nombreux sujets sensibles relatifs aux droits humains, aux libertés fondamentales ou encore à la sécurité des personnes. Vous permettez aussi aux chercheurs du GRIP de s'investir dans la formation d'une relève étudiante, en fournissant un encadrement propice à la transmission des savoirs et des compétences nécessaires à l'analyse critique des enjeux de société.

→ Rejoignez-nous sur www.grip.org

→ Devenez donateur :

IBAN : BE87 0001 5912 8294
BIC/SWIFT : BPO TBE B1