

DEUXIÈME ÉDITION

L'IA au service de la souveraineté *pour la diplomatie multilatérale*



GAINFOREST · YOUTH NEGOTIATORS ACADEMY



MMXXVI

— AVANT DE COMMENCER

Ce document est vivant.

Il fait suite au guide de 2024 rédigé pour le Climate Youth Negotiator Programme, en conserve les entretiens et les prompts, et les actualise pour des outils qui agissent désormais autant qu'ils répondent. Saisir les possibilités et reconnaître les limites, pour renforcer votre plaidoyer.



Table des matières

1	Introduction	4
2	Que sont les grands modèles de langage ?	6
3	Éclairages des parties prenantes	9
4	Les LLM au service des négociations	11
5	Utiliser les modèles	13
6	Agents	17
7	Agents pour les courriels	20
8	Agents pour la logistique	22
9	Agents pour d'autres activités	23
10	Souveraineté, et propriété des données et des décisions	24
11	Déployer votre propre IA souveraine	26
12	L'empreinte environnementale	28
13	Conclusion et recommandations	30
14	Annexe : synthèse des entretiens	32
15	Glossaire de l'IA	34
16	Références	38

Introduction

L'espace de négociation de la CCNUCC concentre des données, des documents, des historiques et des terminologies en évolution constante. C'est un lieu où la précision du langage et une solide maîtrise d'enjeux multiformes peuvent influencer sensiblement sur les résultats. C'est là que des technologies telles que les grands modèles de langage peuvent procurer un avantage.

1.1 Pourquoi devons-nous nous en soucier ?

Pour les jeunes négociateurs et négociatrices climatiques qui représentent leur pays auprès de la Convention-cadre des Nations Unies sur les changements climatiques, les enjeux sont considérables. Ils portent les préoccupations et les aspirations d'une génération plus jeune, qui héritera des conséquences des décisions prises aujourd'hui. Pourtant, ils rencontrent souvent des difficultés pour accéder à de vastes volumes de données climatiques, comprendre les nuances du langage des politiques publiques ou communiquer efficacement leurs positions.

Les modèles de langage peuvent aider sur ces trois plans. Ils traitent, comprennent et génèrent du texte à partir de très vastes ensembles d'écrits, ce qui les rend particulièrement utiles pour quatre fonctions.

Analyse de l'information. Synthétiser de grands volumes de données et de recherches climatiques en enseignements concis, afin que les négociateurs puissent suivre les résultats les plus récents.

Appui à la rédaction des politiques. Aider les négociateurs à formuler leurs points dans un langage qui trouve un écho, respecte les usages de la convention et conserve néanmoins sa force distinctive.

Communication interculturelle. Traduire et contextualiser l'information, ce qui rend l'environnement de négociation plus inclusif.

Simulation et formation. Proposer des scénarios de négociation simulés afin que les jeunes négociateurs puissent s'exercer avant d'entrer sur la scène mondiale.

La technologie elle-même ne représente qu'une partie de l'enjeu. Bien utiliser ces outils témoigne d'une évolution plus large, dans laquelle les jeunes conjuguent leur engagement en faveur de l'action climatique avec les outils qui façonnent la manière dont les décisions sont prises.

1.2 Comment utiliser ce guide

La première moitié explique ce que sont ces systèmes, ce dont les négociateurs ont besoin et comment demander à un modèle de produire quelque chose d'utile. La seconde moitié porte sur les agents, qui lisent des ensembles complets de docu-

ments et exécutent une tâche en plusieurs étapes. Les derniers chapitres traitent de la propriété de vos données, de la manière d'exécuter un modèle que vous contrôlez et de ce qu'il coûte à la planète. Chaque chapitre contient des prompts que vous pouvez copier, et aucun bagage technique n'est nécessaire.¹

¹ La première édition a été rédigée en 2024 pour le Climate Youth Negotiator Programme. Les entretiens du chapitre 3 et de l'appendice proviennent de ce travail.

Que sont les grands modèles de langage ?

2.1 Comment nous en sommes arrivés là

En 1957, le psychologue Frank Rosenblatt a lancé le programme du perceptron au Cornell Aeronautical Laboratory, et l'US Navy a présenté une machine à perceptron l'année suivante. Un perceptron est un modèle mathématique simplifié d'une cellule cérébrale : il combine des entrées, leur applique des poids et produit une sortie. Il s'agissait d'un modèle abstrait plutôt que d'une simulation du cerveau, mais les réseaux neuronaux modernes en descendent en partie.

Les progrès ont ensuite marqué le pas pendant des décennies, limités par les données disponibles, le matériel disponible et les méthodes d'entraînement disponibles. Le domaine a repris son élan lorsque ces trois contraintes se sont desserrées simultanément. En 2012, Alex Krizhevsky, Ilya Sutskever et Geoffrey Hinton ont entraîné **AlexNet** sur 1,2 million d'images étiquetées au moyen d'unités de traitement graphique, et ont fait état d'un taux d'erreur de 17,0 pour cent dans le top cinq sur ImageNet.¹

En 2017, des chercheurs de Google ont introduit le **Transformer**, qui repose sur l'attention plutôt que sur la récurrence. L'attention calcule des poids entre les mots d'un passage en fonction du contexte, ce qui permet au modèle de relier des parties éloignées d'un texte. Les Transformers s'entraînent bien en parallèle, et ils soutiennent presque tous les systèmes présentés dans ce guide.

OpenAI a ensuite appliqué le Transformer au préentraînement génératif. GPT et GPT-2 ont montré qu'un modèle entraîné uniquement à prédire le mot suivant pouvait acquérir de nombreuses capacités sans modèle distinct entraîné pour chaque tâche. Les chercheurs ont observé des **lois de mise à l'échelle**, c'est-à-dire que les performances s'amélioraient de façon prévisible à mesure que le calcul, les données et la taille du modèle augmentaient, à condition que la taille du modèle et les données d'entraînement restent équilibrées.

Le 30 novembre 2022, OpenAI a publié ChatGPT sous la forme d'une version de recherche gratuite, et l'interface conversationnelle a mis ces capacités à la portée de tous.

Deux autres évolutions ont suivi. En septembre 2024, OpenAI a annoncé o1, et en janvier 2025, DeepSeek a publié R1. Tous deux consacrent un calcul supplémentaire au moment de répondre, en générant des notes de travail plus longues avant leur réponse, ce qui aide pour les problèmes qui comportent plusieurs étapes. Et les modèles ont commencé à être reliés à des outils et à des boucles, ce qui fait l'objet du chapitre 6.

¹ Ce résultat tenait autant à l'architecture, à l'augmentation des données et à l'ensemble de données qu'au matériel.

2.2 La recette

Fondamentalement, un modèle de langage est un type d'intelligence artificielle conçu pour comprendre, générer et traiter le langage humain. On peut le concevoir comme un vaste cerveau numérique, entraîné sur des textes issus de livres, d'articles, de sites web et d'autres sources, qu'il utilise pour générer un texte d'apparence humaine à partir des motifs qu'il reconnaît.

ChatGPT, comme d'autres modèles de la série GPT d'OpenAI ou Llama de Meta, est construit en deux étapes : le préentraînement, puis le post-entraînement. Le post-entraînement désignait auparavant une seule phase d'affinage. Il comprend désormais deux parties, et la seconde explique pourquoi les modèles récents peuvent traiter des problèmes longs.

Préentraînement. Le modèle est entraîné sur une très grande quantité de texte provenant d'internet. Il ne sait pas précisément quels documents figuraient dans son ensemble d'entraînement. L'objectif est d'apprendre la grammaire, des faits sur le monde, une certaine capacité de raisonnement et, surtout, à prédire le mot suivant dans une phrase. Ce faisant, il absorbe une quantité étonnante d'informations, depuis des faits quotidiens triviaux jusqu'à des notions complexes.

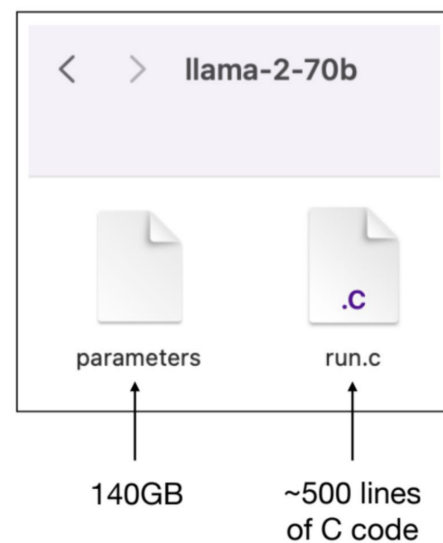
Affinage supervisé. La première moitié du post-entraînement. Le modèle est entraîné davantage sur un ensemble de données plus restreint, rédigé et évalué par des examinateurs humains qui suivent des lignes directrices publiées. Les examinateurs notent les sorties possibles pour une série d'entrées d'exemple, et le modèle généralise à partir de ce retour d'information pour répondre à un large éventail de demandes des utilisateurs. C'est ce qui lui apprend à suivre une instruction plutôt qu'à simplement poursuivre votre phrase.

Apprentissage par renforcement. La seconde moitié, et la plus récente. Au lieu de copier une réponse écrite, le modèle est récompensé lorsqu'il parvient à un bon résultat; il apprend ainsi quelles façons de traiter un problème tendent à être efficaces. C'est ce qui apprend à un modèle à raisonner sur un problème long et à maintenir une chaîne d'étapes cohérente avant de répondre.

2.3 Limites

Les chercheurs continuent de démêler le fonctionnement détaillé de ces modèles. L'opinion générale est que le préentraînement incorpore un large éventail de connaissances sur le monde, et que le post-entraînement façonne la manière dont ces connaissances sont utilisées.

Ce qui demeure vrai, c'est qu'un modèle de langage prédit le mot suivant dans une séquence à partir de motifs statistiques. Cette méthode saisit une grande partie des connaissances humaines, mais elle entraîne aussi des problèmes. Les modèles reproduisent les biais des textes sur lesquels ils ont été entraînés, et produisent avec assurance des affirmations fausses, ce que le domaine appelle hallucination. Une phrase probable n'est pas une phrase vérifiée.



Sur disque, un modèle est constitué de deux fichiers : les paramètres qu'il a appris et un court programme qui les exécute. Llama 2 70B représente 140 Go de paramètres et environ 500 lignes de C.

2.4 D'où viennent les biais

Ces modèles sont entraînés sur ce que contient internet, et internet ne représente pas le monde de manière équilibrée. Cela importe davantage pour un négociateur climat que pour presque tout autre lecteur.

Prenons le relevé de la nature elle-même. La Global Biodiversity Information Facility est le plus grand dépôt ouvert d'observations d'espèces sur Terre, et environ 83 pour cent de ses enregistrements proviennent d'Amérique du Nord et d'Europe. Les 17 pour cent restants couvrent tout le reste du monde, y compris les lieux les plus riches en biodiversité de la planète.

Le même déséquilibre traverse les textes. Il a été écrit en ligne beaucoup plus de contenus en anglais, en français et en espagnol qu'en quechua, en twi ou en tok pisin, si bien qu'un modèle a vu beaucoup plus de textes du premier groupe. Il traduira, résumera et rédigera avec fluidité pour une délégation travaillant dans une langue bien représentée, et il échouera plus souvent, et avec davantage d'assurance, pour une langue qui ne l'est pas.

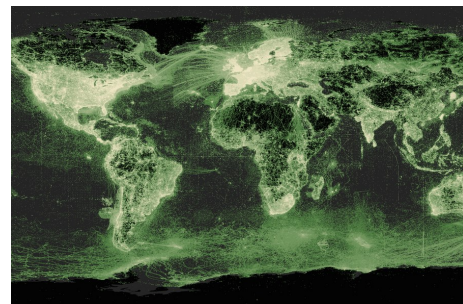
C'est aussi de là que provient une grande partie des hallucinations. Interrogé sur un lieu bien documenté, un modèle répond à partir d'un corpus dense. Interrogé sur un lieu peu documenté, il produit les mêmes phrases assurées à partir d'un matériau bien plus limité et comble les lacunes avec ce que le motif suggère. La sortie ne paraît pas moins certaine. Elle l'est seulement.

Rien de cela n'est imputable à quiconque de manière simple. La collecte et la publication de données exigent des infrastructures et des financements dont de nombreux pays ne disposent pas, et certaines communautés refusent de transmettre des connaissances relatives à leurs terres et à leurs espèces, ce qui constitue une position légitime plutôt qu'un échec. La conséquence s'impose néanmoins : ces outils fonctionnent le mieux pour les lieux déjà les mieux servis par les données, et le moins bien pour les lieux où les enjeux sont les plus élevés.

Le même écart se retrouve dans l'identité de celles et ceux qui utilisent ces outils. L'indice d'Anthropic sur l'utilisation de Claude par pays suit étroitement le revenu national, et vingt-cinq pays se situent dans un niveau inférieur où l'usage mesuré est presque nul.² Parmi eux figurent les Tonga, Samoa, Nauru et Palau, qui sont parties à la même négociation et comptent parmi les plus exposés à son issue.

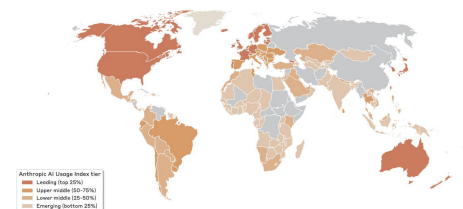
■ REMARQUE Ce que cela signifie dans la salle

Considérez la fluidité d'un modèle au sujet de votre région comme une mesure de ce qui a été écrit à son propos, et non de ce qu'il sait. Moins votre pays, votre langue ou votre écosystème est représenté, plus vous devez vérifier rigoureusement ce qu'il vous indique.



Toutes les observations d'espèces dans GBIF, cartographiées. Les zones lumineuses indiquent l'emplacement des enregistrements, non celui de la vie. En chiffres bruts, la Belgique et la Suède apparaissent plus riches en biodiversité que le Brésil et Madagascar.

² Anthropic, *Anthropic Economic Index*, 2025. L'utilisation croît avec le PIB par personne en âge de travailler selon une puissance d'environ 0,69.



Utilisation de Claude par personne en âge de travailler, par pays. Comparez-la à la carte des enregistrements d'espèces ci-dessus.

Éclairages des parties prenantes

Les entretiens menés avec de jeunes négociateurs climatiques de plusieurs pays font apparaître un ensemble de difficultés communes et une vision partagée de ce qui serait utile.¹

¹ Six négociateurs du Libéria, du Paraguay, du Pérou, du Nigéria, du Liban et de l'Indonésie. Le résumé complet figure en annexe.

3.1 Principaux défis et points de friction

Barrières linguistiques. L'anglais n'est pas la langue première de la plupart des personnes présentes dans la salle.

Langage technique. Le jargon de la UNFCCC exige des années d'assimilation (Libéria, Paraguay).

Transfert des connaissances. Les connaissances et informations techniques sont partagées de manière inégale, et une grande partie se perd d'un cycle à l'autre.

Mémoire historique. La connaissance des négociations passées est difficile à acquérir (Liban).

Communication de sujets complexes. Les échanges avec les négociateurs chevronnés sont difficiles à mener sous pression (Pérou, Nigéria), et il faut acquérir une légitimité avant que le fond ne soit entendu.

Compréhension des positions des Parties. Les dynamiques de négociation évoluent rapidement et sont difficiles à interpréter.

Contraintes de temps. Le temps manque toujours pour traiter les informations qui comptent.

Expression d'idées complexes. S'exprimer en anglais avec rapidité et précision demeure un obstacle persistant.

3.2 Outils et ressources souhaités

Outils linguistiques. Un appui linguistique et grammatical plus sophistiqué, conçu pour les textes de la UNFCCC plutôt que pour la rédaction professionnelle générale.

Recherche rapide d'informations. Des plateformes capables d'analyser de longs documents, d'en extraire les éléments pertinents et de confirmer quel texte est la version en vigueur.

Formation personnalisée. Des ressources d'apprentissage adaptées à la région, au groupe et au niveau d'expérience d'un négociateur.

Les outils effectivement utilisés au moment des entretiens étaient Grammarly pour la rédaction et Google Drive pour la collaboration. Un négociateur a indiqué n'utiliser aucun outil de traduction, au motif que le langage de la négociation ne résiste pas à la traduction.

Les LLM au service des négociations

Les négociations climatiques sont multiformes et portent sur des sujets allant des émissions de carbone et de la biodiversité aux incidences économiques et aux dynamiques sociopolitiques. Les aborder exige trois choses à la fois.

Un accès rapide à l'information. La science du climat évolue en permanence, et les négociateurs doivent disposer des données les plus récentes à portée de main.

Une rédaction efficace des accords. La précision du langage peut déterminer la réussite ou l'échec d'un accord.

Une communication multilingue. Une plateforme mondiale exige un dialogue interculturel et multilingue.

4.1 Où les modèles sont utiles

Communication. Rédiger et traduire des documents, afin que l'information et les accords soient communiqués avec exactitude entre des parties ayant des langues maternelles différentes.

Collecte d'information. Réunir et synthétiser de vastes volumes de documents sur les questions climatiques, notamment des travaux scientifiques, des documents d'orientation et des accords antérieurs, afin que les négociateurs restent informés.

Travail sur les scénarios. Examiner les conséquences possibles de différentes décisions de politique publique, ce qui aide un négociateur à explorer des positions et des résultats avant de s'engager dans l'un d'eux.¹

Analyse des arguments. Synthétiser les arguments avancés par différentes parties pendant les négociations, ce qui aide un négociateur à discerner les points essentiels et les contre-arguments.

¹ On parle parfois de modélisation prédictive, ce qui est excessif. Les modèles de langage n'exécutent pas de modèles climatiques ou économiques. Ils reformulent et comparent des scénarios, et tout chiffre qu'ils produisent doit être vérifié auprès d'une source réelle.

4.2 Pièges potentiels

Les avantages sont réels, mais certaines réserves comptent en diplomatie climatique.

Simplification de questions complexes. Ces systèmes condensent, et la condensation supprime parfois le qualificatif qu'une partie a défendu pendant trois sessions pour le faire insérer.

Absence d'intelligence émotionnelle. Un modèle ne lit pas les courants émo-

tionnels et politiques d'une négociation, et il ne sait pas ce qui a été convenu dans un couloir.

Absence d'expertise humaine. La maîtrise apparente du registre de la politique climatique n'est pas l'expertise en politique climatique. S'appuyer sur un modèle pour l'une ou l'autre peut produire des informations incomplètes ou inexactes.

Sécurité et confidentialité. Les discussions portent souvent sur des éléments sensibles. Tout contenu collé dans un système hébergé quitte votre délégation, et les conditions selon lesquelles il est stocké sont rarement lues.

Considérations éthiques. Soyez transparent quant à l'assistance automatisée, maintenez la responsabilité des décisions entre les mains des personnes, et surveillez les risques d'utilisation abusive.

Mauvaise communication et mauvaise interprétation. Un modèle peut mal comprendre l'intention qui sous-tend une demande et produire un texte qui ressemble à votre position sans l'être.

Utiliser les modèles

5.1 Quel modèle, et selon quelles conditions

Avant de choisir un assistant, il faut savoir qu'il en existe deux catégories, et que la différence importe davantage que la marque.

Poids fermés. ChatGPT, Claude et Gemini fonctionnent sur les serveurs d'entreprises et sont accessibles au moyen d'un compte. Ce sont les modèles les plus puissants et les plus rapides à prendre en main. Les niveaux gratuits permettent de rédiger, traduire et résumer. Les niveaux payants coûtent environ vingt dollars par mois et ajoutent la prise en charge de documents plus longs, le téléversement de fichiers et les fonctions d'agent du chapitre 6. Tout ce que vous saisissez est transmis à l'entreprise et, selon vos paramètres, peut être conservé, examiné ou utilisé pour entraîner le modèle suivant.

Poids ouverts. Llama, Mistral, Qwen, Gemma et DeepSeek publient leurs paramètres, de sorte que le même modèle peut être exécuté par vous, par votre institution ou par un fournisseur que vous choisissez plutôt que par un fournisseur qui vous est imposé. Ollama en installe un sur un ordinateur portable en quelques minutes. Polly, exploité avec GainForest et la Youth Negotiators Academy sur polly.ai4cop.org, est un modèle hébergé à poids ouverts qui ne conserve aucune trace de ce que vous envoyez. Le chapitre 11 traite de l'exécution de votre propre modèle.

Une règle pour commencer. Utilisez un modèle fermé hébergé pour les éléments qui sont déjà publics, et un modèle à poids ouverts pour tout ce qui ne l'est pas. Ne collez pas dans un service que vous ne contrôlez pas une position que votre délégation n'a pas encore déposée. Le chapitre 10 explique ce qui est en jeu dans ce choix, et il vaut la peine de le lire avant d'ouvrir un compte plutôt qu'après.

Quel que soit votre choix, installez l'application mobile. La saisie vocale est utile lorsque vous vous déplacez entre les séances, et dicter une idée approximative dans votre propre langue produit souvent un meilleur brouillon que taper une phrase soigneusement formulée en anglais.

5.2 Prompt engineering 101

L'instruction que vous saisissez s'appelle un prompt. En rédiger un consiste à créer des instructions précises et bien définies pour vous aider à accomplir une tâche, et la qualité ainsi que la précision de ce que vous écrivez façonnent la réponse que vous recevez.

Un prompt peut contenir l'un quelconque des éléments suivants, sans devoir nécessairement contenir les quatre.

Instruction. La tâche que vous voulez que le modèle accomplisse.

Contexte. Les informations de fond qui orientent le modèle vers une meilleure réponse.

Données d'entrée. Le texte, la question ou le document sur lequel vous voulez qu'il travaille.

Indication de sortie. Le type ou le format de sortie que vous souhaitez.

5.3 Ce qui fait un bon prompt

Six gestes accomplissent l'essentiel du travail, quel que soit le modèle que vous utilisez.

Soyez clair, direct et détaillé. Dites qui vous êtes, ce que vous voulez et à quoi la sortie doit ressembler. Un modèle ne peut pas déduire la salle dans laquelle vous vous trouvez.

Montrez un exemple. Un seul échantillon de la sortie souhaitée enseigne plus vite qu'un paragraphe qui la décrit.

Demandez des citations. Dites au modèle de citer le passage sur lequel il s'appuie. Ancrer une réponse dans un document retrouvé est la protection la plus solide contre une invention assurée.

Balisez des parties. Encadrez chaque élément par des balises, <context>, <task>, <text>, afin que le modèle puisse distinguer votre instruction du document que vous avez collé.

Laissez-le réfléchir d'abord. Demandez le raisonnement avant la réponse. Les modèles commettent moins d'erreurs sur toute tâche comportant plus d'une étape lorsqu'ils la traitent ouvertement.

Donnez-lui un rôle. « Vous êtes chercheur sur les pertes et préjudices » produit un texte différent, et généralement meilleur, que la même demande sans locuteur identifié.

Puis continuez. Affiner un prompt est la manière dont la qualité de l'échange s'améliore, et la deuxième tentative est presque toujours meilleure que la première.

Une habitude mérite d'être prise dès le départ. Demandez des sources, et demandez au modèle de signaler les parties dont il n'est pas sûr. Une réponse que vous ne pouvez pas vérifier est une réponse que vous ne pouvez pas utiliser en plénière.

L'ANATOMIE D'UNE INTERVENTION

Pour la tâche que les négociateurs accomplissent le plus souvent, cinq éléments font la différence : qui parle et où, la question précise en jeu, les éléments probants qui la sous-tendent, le cadre auquel elle se rattache, et ce qui devrait être fait.

¹En tant que jeune délégué de la Malaisie à la COP29, rédigez une intervention de deux minutes ²sur les pertes et préjudices qui touchent les enfants en Malaisie. Mettez en évidence ³des préoccupations telles que l'élévation du niveau de la mer menaçant les communautés côtières, les perturbations de l'éducation dues aux incidences des changements climatiques, et les problèmes de santé mentale chez les enfants. ⁴Formulez des recommandations précises pour donner la priorité au financement destiné à protéger les droits des enfants et à renforcer leur participation aux discussions sur le climat ⁵dans le cadre de l'Accord de Paris.

¹ Rôle et contexte ² Thème central ³ Éléments probants ⁴ Appel à l'action
⁵ Contenu technique

Supprimez l'un de ces cinq éléments et l'on sent ce qui manque. Sans le rôle, le modèle rédige pour personne; sans les éléments probants, il produit du sentiment; et sans la dernière ligne, il ne rattache le texte à aucun accord sur lequel quelqu'un puisse agir.

5.4 Prompts pour différents cas d'utilisation

1. Analyse de données

TÉLÉVERSEMENT DE FICHIER

À partir de cette feuille de calcul, tracez l'augmentation moyenne des émissions au cours des dernières années, indiquez votre méthode et énumérez les hypothèses que vous avez formulées concernant les données manquantes.

2. Traduction

N'IMPORTE QUEL MODÈLE

Traduisez ce paragraphe en espagnol, en conservant le registre formel utilisé dans les textes de décision.

3. Résumé

RECHERCHE WEB ACTIVÉE

Résumez l'intervention du LDC Group lors du dialogue technique du Global Stocktake le 10 juin 2023, et citez les passages sur lesquels vous appuyez.

4. Collecte d'informations pour élaborer une déclaration de position ou préparer un discours

N'IMPORTE QUEL MODÈLE

Vous êtes le rédacteur de discours d'un négociateur de l'ONU sur les changements climatiques. Vous écrivez au sujet des efforts du Royaume-Uni en matière de renforcement des capacités. Je suis le négociateur. Quelles cinq questions me poseriez-vous pour extraire les informations dont vous avez besoin afin de rédiger une déclaration?

5. Anticiper les contre-arguments et préparer des réfutations

N'IMPORTE QUEL MODÈLE

Les pertes et préjudices sont une question essentielle dans les négociations climatiques. Vous êtes chercheur sur les pertes et préjudices. Présentez les principales raisons pour lesquelles les pays n'accordent pas d'importance aux pertes et préjudices. Produisez un contre-argument solide à ces points, démontrant la nécessité d'un fonds pour les pertes et préjudices.

6. Rédiger une déclaration liminaire ou un discours d'ouverture

N'IMPORTE QUEL MODÈLE

QUEL MODÈLE

Vous êtes rédacteur de discours et préparez une déclaration à partir des questions ci-dessus. Produisez une déclaration de 300 mots présentant les efforts du Royaume-Uni en matière de renforcement des capacités climatiques. La déclaration doit refléter le style rédactionnel utilisé par le Gouvernement britannique.

7. Décoder le jargon

N'IMPORTE QUEL MODÈLE

Expliquez ce que signifie « moyens de mise en œuvre » dans ce paragraphe, en langage clair, et donnez un exemple de la manière dont cette expression a été utilisée dans une décision antérieure.

8. Trouver le document

RÉCUPÉRATION

Quel texte de décision contient le mandat relatif à ce programme de travail ? Donnez-moi la cote du document et citez le paragraphe dans lequel elle apparaît.

9. Traduire pour votre propre public

N'IMPORTE QUEL MODÈLE

Traduisez cette intervention en Bahasa Malaysia pour un public communautaire, en conservant le sens et en supprimant le jargon de la UNFCCC.

Agents

Tout ce qui précède fonctionne de la même manière. Vous saisissez une demande, vous lisez une réponse, vous décidez quoi en faire. Quatre éléments ont changé et, ensemble, ils ajoutent une seconde manière de travailler.

Le contexte s'est allongé. Un modèle peut désormais contenir simultanément l'intégralité d'un texte de négociation, ses révisions antérieures et les notes d'information de votre délégation.

La récupération est devenue normale. Un système peut rechercher dans un ensemble de documents ou sur le web avant de répondre, et citer ce qu'il a trouvé. C'est ce qui rend une réponse vérifiable.

Les modèles ont reçu des outils. Un modèle peut être relié à un index de recherche, à une feuille de calcul, à un calendrier ou à un dossier de fichiers, de sorte qu'il puisse agir sur ceux-ci plutôt que seulement les décrire.

Les systèmes ont reçu des boucles. Au lieu de répondre une seule fois, un système peut planifier, franchir une étape, examiner le résultat, puis recommencer jusqu'à ce que la tâche soit accomplie.

Un modèle relié à des outils, à une mémoire et à une boucle « répéter jusqu'à achèvement », fonctionnant dans le cadre d'autorisations que quelqu'un lui a accordées, correspond à ce que l'on entend par agent. Il n'a ni jugement ni responsabilité. Il a une tâche et un ensemble d'autorisations.

■ LIGNE ROUGE La ligne

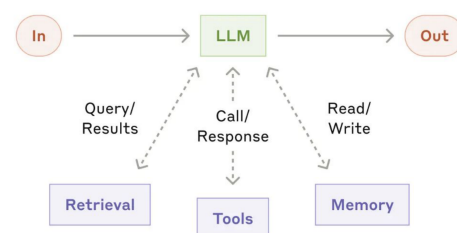
Un agent peut préparer une position. Il ne peut envoyer, soumettre ni publier quoi que ce soit pour une délégation sans qu'une personne nommément désignée n'autorise cet acte.

6.1 Trois manières de travailler

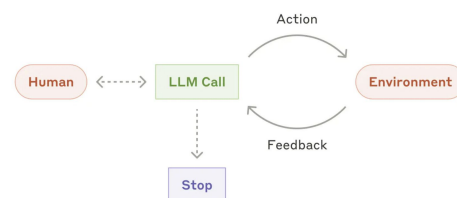
Chat. Vous posez une question, il répond, vous vérifiez. Le plus adapté pour la rédaction, la traduction, l'explication et la répétition.

Récupération. Le système recherche d'abord dans un ensemble défini de documents et répond avec des citations. À utiliser pour tout élément factuel relatif à la négociation : historique du texte, positions des Parties, décisions antérieures.¹

Agentique. Le système exécute une tâche en plusieurs étapes à l'aide d'outils. À utiliser pour les travaux répétitifs produisant un résultat vérifiable, par exemple le suivi des modifications entre révisions ou la constitution d'une note quotidienne.



Un modèle augmenté. Le même modèle, auquel on donne quelque chose à rechercher, des outils à appeler et un endroit où conserver ce qu'il a appris.



La boucle. Le modèle agit, lit ce qui est revenu, puis recommence jusqu'à ce qu'il s'arrête. Une personne fixe la tâche et les autorisations, et elle est le seul élément du schéma qui puisse être tenu responsable.

¹ Si un système vous dit ce que dit un paragraphe d'un texte de décision et ne peut pas vous montrer le paragraphe, traitez-le comme une rumeur.

6.2 Cas d'utilisation

1. Suivi du texte

FICHIERS JOINTS

Compare la révision 3 de ce projet de décision avec la révision 2. Dresse la liste de chaque modification substantielle, cite les deux versions et signale les modifications qui touchent les lignes rouges de ma délégation.

2. Cartographie des positions

RÉCUPÉRATION

À partir de ces soumissions, construis un tableau des positions des Parties sur l'Article 6.4 avec une citation et une référence pour chaque ligne. Indique toute position que tu as inférée.

3. Recherche historique

RÉCUPÉRATION

Retrace la manière dont l'expression «responsabilités communes mais différenciées» a été qualifiée dans les textes de décision depuis 2015, en citant chaque décision.

4. Lecture nocturne

FICHIERS JOINTS

Voici les quatre textes publiés aujourd'hui. Dis-moi ce qui a avancé, ce qui est resté bloqué et quelles sont les trois pages que je dois lire avant demain.

5. Synthèse quotidienne

AGENT PROGRAMMÉ

Chaque soir, collecte les textes publiés dans la journée pour les axes que je suis, résume ce qui a changé et envoie-moi la liste.

6. Analyse des crochets

FICHIERS JOINTS

Pour chaque option entre crochets dans ce paragraphe, explique qui en bénéficie, qui s'y oppose et à quoi pourrait ressembler une formulation de repli.

7. Répétition

CHAT

Joue le rôle d'un négociateur qui s'oppose à ce texte. Conteste mon intervention et ne sois pas complaisant.

8. Retrotraduction

CHAT

Traduis ma formulation en français, puis retraduis-la indépendamment en anglais, et dis-moi ce qui a dérivé.

9. Passation

FICHIERS JOINTS

À partir de mes notes dans cette session, rédige le document de passation pour la personne qui reprendra ce dossier au prochain cycle.

10. Travail sur les données

ACCÈS AUX FICHIERS

À partir de cette feuille de calcul sur les émissions, trace la tendance pour ces pays depuis 2010, indique ta méthode et liste chaque hypothèse concernant les données manquantes.

6.3 Travailler avec un agent en toute sécurité

Un agent qui agit à partir d'une mauvaise lecture produit des conséquences, et non un mauvais paragraphe.

Vérifiez ce que vous répétez. Ouvrez un document cité et confirmez que la

citation s’y trouve avant de l’utiliser en salle.

Accordez la plus petite autorisation suffisante. Un système qui lit vos fichiers n’a pas également besoin d’envoyer des courriels.

Traitez tout document externe comme non fiable. Un fichier soumis peut contenir des instructions visant l’agent.

Décidez qui approuve un texte assisté par l’IA. Avec plusieurs outils dans la chaîne, personne d’autre ne le fera.

Agents pour les courriels

Le courriel demeure un mode essentiel de communication professionnelle, et trois facteurs le rendent difficile.

Gestion du volume. Les boîtes de réception débordent, et les hiérarchiser prend du temps dont vous ne disposez pas.

Rédaction efficace. Un courriel clair, concis et orienté vers l'action exige un savoir-faire.

Barrières multilingues. La communication mondiale suppose de travailler dans des langues qui ne sont pas les vôtres.

Un assistant connecté à votre boîte aux lettres lit désormais lui-même le fil de discussion; accordez-lui donc un accès en lecture et laissez-le rédiger. L'envoi reste sous votre contrôle.

1. Suivi de réunion

Résume les points clés et les engagements pris lors de la récente réunion de négociation, attribue chacun d'eux à la personne qui l'a formulé, et rédige un courriel de suivi courtois aux participants afin d'encourager la collaboration sur les actions convenues.

2. Planification

Génère un courriel diplomatique et flexible proposant plusieurs dates et heures pour le prochain cycle de réunions, exprimées dans les fuseaux horaires de tous les participants.

3. Notes d'information quotidiennes et comptes rendus

Crée un courriel quotidien concis de note d'information mettant en évidence les progrès réalisés lors des dernières sessions, y compris les nouvelles propositions, les points convenus et les domaines de divergence nécessitant un examen plus approfondi.

4. Hiérarchisation intelligente

Élabore un courriel à l'équipe de négociation présentant une stratégie de hiérarchisation des points de l'ordre du jour selon leur urgence, leur incidence potentielle sur les objectifs climatiques et la faisabilité de parvenir à un accord.

5. Triage

Classe mes messages non lus en quatre catégories : nécessite une décision de ma part, nécessite une réponse, doit être transféré, ne nécessite aucune action. Rédige des réponses uniquement pour le troisième groupe et laisse-les en brouillon.

6. Vérification du registre

Réécris ce courriel dans un registre diplomatique formel sans modifier aucun

engagement de fond, puis dresse la liste des changements effectués afin que je puisse vérifier.

Ces systèmes adoptent par défaut une voix chaleureuse et conciliante qui peut concéder des éléments absents de votre texte initial. Demandez ce qui a été modifié, puis lisez la liste.

Agents pour la logistique

Organiser les calendriers, créer des listes de contrôle et réserver les déplacements peut devenir accablant pendant une session. Trois problèmes reviennent régulièrement.

Gestion du temps. Hiérarchiser les tâches et leur attribuer du temps n'a rien d'évident.

Oubli de tâches. Des éléments essentiels disparaissent de la liste de contrôle.

Réservation de vols et d'hôtels. Trouver des options qui correspondent à votre calendrier et à votre budget prend des heures.

Il s'agit du travail le moins risqué du guide et du meilleur point de départ avec un agent. Les tâches se répètent, les résultats sont faciles à vérifier et presque rien ici n'est confidentiel.

1. Réservations

Recense les vols et hôtels les plus économiques et les plus pratiques pour un trajet de [ville de départ] à [ville de destination] aux [dates], avec une préférence pour les vols du matin et un hébergement proche du lieu de la réunion. Présente-moi les arbitrages plutôt qu'une seule recommandation.

2. Création de listes de contrôle

Génère une liste de contrôle complète pour une session internationale de deux semaines, comprenant les documents nécessaires, l'accréditation, les vêtements adaptés à une météo variable, le matériel de travail essentiel et des copies hors ligne des textes clés.

3. Planifier la journée

Examine mon calendrier et le programme de la session, puis propose un emploi du temps qui équilibre réunions, plages de travail individuel, exercice physique et temps personnel, compte tenu de ces échéances et engagements.

4. Suivi des progrès

Crée un modèle pour suivre les progrès des axes de travail que je suis, comprenant les jalons, les échéances et un système permettant de consigner les mises à jour et les prochaines étapes.

La logistique est aussi le premier domaine dans lequel un agent reçoit de véritables autorisations, sur votre calendrier, votre messagerie et vos fichiers. C'est raisonnable, et c'est là que commence l'habitude d'accorder un peu plus d'accès à chaque fois. Le système qui réserve vos vols n'a pas besoin de rédiger vos positions.

Agents pour d'autres activités

Les négociations sur le climat ne se résument pas à des chiffres et à des faits. Elles sont traversées par les émotions humaines, les différences culturelles et le bien-être mental.

Tensions sur la santé mentale. Les réalités des changements climatiques et la pression de la négociation pèsent sur les délégués.

Malentendus culturels. Les interprétations erronées issues des différences culturelles créent des fractures et ralentissent le dialogue.

C'est la seule partie du guide où une simple fenêtre de conversation l'emporte sur un agent. Ces échanges appellent un modèle avec lequel dialoguer, non un système doté d'autorisations.

1. Note culturelle

Je participe à une négociation des Nations Unies sur les changements climatiques avec une personne originaire du Kenya. À quels éléments culturels devrais-je être attentif afin de faire progresser ma négociation ?

2. Reconnaissance des émotions

Je me sens abattu en ce moment ; pouvez-vous me dire quelque chose d'encourageant ?

3. Débriefing

Voici ce qui s'est passé lors de cette séance et la manière dont je l'ai géré. Qu'est-ce qu'un négociateur plus expérimenté aurait fait différemment ?

Deux mises en garde. Une note culturelle produit une généralisation sur un pays, alors que vous rencontrerez une personne ; considérez donc ce que vous obtenez comme une orientation plutôt que comme une prédiction. Par ailleurs, un modèle peut être un lieu raisonnable où déposer une pensée à une heure du matin, mais ce n'est ni un clinicien, ni un collègue, et ce n'est pas confidentiel. Si la pression dépasse la simple fatigue, adressez-vous à votre chef de délégation, aux contacts d'appui de votre programme ou à un professionnel.

Souveraineté, et propriété des données et des décisions

Ces outils peuvent contribuer à rééquilibrer l'espace multilatéral. Ils lèvent la barrière linguistique qui détermine qui prend la parole avec assurance, condensent des documents que personne n'a le temps de lire, et donnent à une délégation de deux personnes une partie de la portée d'une délégation de cinquante personnes. C'est un acquis important. Cela soulève aussi des questions d'intégrité qu'il convient de se poser avant de saisir quoi que ce soit.

Qui détient les données. Le texte que vous collez dans un assistant hébergé quitte votre délégation. Selon le compte et ses paramètres, il peut être conservé, examiné par du personnel, ou utilisé pour entraîner le prochain modèle. Les positions nationales, les instructions reçues de la capitale, et tout élément qui vous a été communiqué à titre confidentiel n'y ont pas leur place.

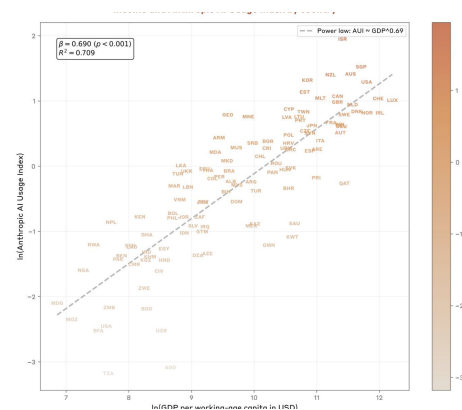
Qui détient la décision. Un modèle produit un texte qui ressemble à une position. S'il rédige et que personne ne vérifie, le raisonnement qui sous-tend la position de votre partie commence à exister en dehors de votre délégation. Toute communication sortante doit être assumée par une personne investie d'un mandat.

Ce que vous ne pouvez pas voir. OpenAI et Anthropic ne restituent pas les traces de raisonnement. Vous obtenez une réponse, parfois un résumé de la manière dont elle a été produite, tandis que les étapes effectives restent chez le fournisseur. Vous ne pouvez pas auditer le cheminement, le reproduire, ni montrer à un collègue pourquoi le système a répondu ce qu'il a répondu.

Poids ouverts et poids fermés. Les poids d'un modèle sont les nombres qu'il a appris lors de l'entraînement. ChatGPT, Claude et Gemini conservent les leurs sur les serveurs de leurs entreprises, et ce sont les plus puissants et les plus faciles à utiliser. Llama, Mistral, Qwen, Gemma et DeepSeek publient les leurs, de sorte que vous pouvez exécuter le modèle sur votre propre machine ou sur un serveur régional. Les modèles ouverts accomplissent bien la traduction, la synthèse et la rédaction, et restent en retrait sur les problèmes multi-étapes les plus difficiles. Les poids que vous détenez ne peuvent pas être renchérissés ni retirés au milieu d'une session.

Une IA alignée sur la souveraineté signifie que la délégation, la région ou la communauté fixe les conditions : quel modèle est utilisé, quelles données y sont introduites, qui peut inspecter ce qui s'est passé, et le droit de refuser. Dans le processus de la COP, les parties qui disposent des capacités les plus limitées sont généralement celles qui ont le plus en jeu, ce qui rend la question pratique plutôt que philosophique.

Règle de travail. Utilisez le modèle hébergé le plus performant pour les documents



Utilisation de Claude en fonction du revenu, par pays. L'adoption suit de près le PIB par personne en âge de travailler, ce qui revient à dire que ces outils arrivent en dernier là où ils sont le plus nécessaires.

publics, la lecture de contexte, la rédaction et les exercices de préparation. Utilisez des poids ouverts, ou un déploiement régi par un contrat que votre institution a examiné, pour tout élément touchant à une position qui n'est pas encore publique.

Ce guide existe pour maintenir cet arbitrage visible. Les chapitres précédents montrent ce que ces systèmes font bien. Celui-ci porte sur ce à quoi vous renoncez en échange.

Déployer votre propre IA souveraine

Le chapitre précédent a exposé l'arbitrage. Celui-ci porte sur la réponse la plus robuste à cet arbitrage, à savoir exploiter le modèle soi-même.

L'auto-hébergement signifie que rien ne quitte les locaux. Il n'y a pas de politique de conservation à laquelle faire confiance, pas de conditions d'utilisation susceptibles de changer, et pas de compte pouvant être suspendu au milieu d'une session. Il permet aussi de développer les compétences au sein de votre institution plutôt que de les louer.

Jusqu'à récemment, cela signifiait accepter un modèle beaucoup moins performant. Ce n'est plus le cas.

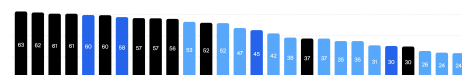
Kimi K3 et GLM-5.3 obtiennent tous deux un score de 60 dans l'Artificial Analysis Intelligence Index, trois points derrière Claude Opus 5 à 63, et Qwen3.8 suit à 58. Dix-sept des 27 modèles de l'indice ont des poids ouverts, même si le haut du classement reste encore majoritairement propriétaire.¹ Il y a un an, le principal modèle ouvert accusait treize points de retard sur le principal modèle propriétaire ; cet écart s'est donc rapidement réduit et pourrait se réduire encore.

Le prix constitue la seconde moitié de l'argument. DeepSeek V4 Pro 0813 obtient un score de 53 pour environ un quart de dollar par tâche, dans ce quadrant. Claude Opus 5 obtient dix points de plus, à un prix environ dix fois supérieur. Ce rapport est la raison pratique pour laquelle une délégation peut se permettre d'exploiter son propre modèle.

11.1 Quel modèle exploiter

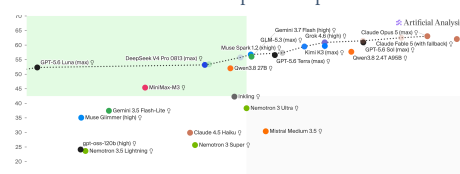
MODÈLE	LABORATOIRE	POIDS	INDICE
Kimi K3 (max)	Moonshot	Ouverts, usage restreint	60
GLM-5.3 (max)	Z.ai	Ouverts, MIT	60
Qwen3.8 2.4T A95B	Alibaba	Ouverts, usage restreint	58
DeepSeek V4 Pro 0813	DeepSeek	Ouverts	53
Qwen3.8 27B	Alibaba	Ouverts	52
MiniMax-M3	MiniMax	Ouverts, usage restreint	45
Nemotron 3 Super	NVIDIA	Ouverts	26
gpt-oss-120b (high)	OpenAI	Ouverts	24

État des lieux au 20 août 2026. Lisez la licence avant de vous engager. Plusieurs des modèles ouverts les plus performants restreignent l'usage commercial, ce qui



Les 27 modèles de l'Artificial Analysis Intelligence Index, au 20 août 2026, classés par rang. Le noir indique les modèles propriétaires, le bleu les poids ouverts. Le principal modèle ouvert, Kimi K3, occupe la cinquième place avec 60, à trois points du premier.

¹ Le graphique compte GLM-5.3 comme propriétaire. Z.ai publie ses poids sous licence MIT, ce guide le compte donc comme ouvert, ce qui explique pourquoi la figure montre seize barres bleues et le texte indique dix-sept.



Le même indice rapporté au coût par tâche. La zone ombrée correspond au quadrant le plus intéressant, score élevé et faible coût, et les modèles ouverts en occupent la majeure partie.

ne pose généralement pas de problème pour une délégation, mais en pose un pour un fournisseur qui construit un service par-dessus l'un d'eux.

11.2 Combien cela coûte

Un NVIDIA DGX Spark est une machine de bureau dotée de 128GB de mémoire unifiée, actuellement vendue 4 699 dollars. Deux unités se connectent directement l'une à l'autre sans commutateur, offrant 256GB pour environ 9 400 dollars. Cette paire exécute un modèle épars de la catégorie DeepSeek à 55 à 60 jetons par seconde avec un contexte d'un million de jetons, ce qui suffit à servir une délégation au travail.

Une seule unité exécute un modèle de 27B à 35B assez rapidement pour soixante utilisateurs simultanés, couvrant la traduction, la synthèse et la rédaction pour tout un bureau. Trois ou quatre unités, à 384 à 512GB, hébergent les plus grands modèles ouverts sans les élaguer.

Comparé à un abonnement pour chaque délégué, le matériel s'amortit en un cycle ou deux, et il continue de fonctionner une fois le financement terminé.

11.3 Comment le mettre en place

Servez le modèle avec vLLM ou SGLang. Les deux présentent la même API qu'OpenAI, de sorte que tout outil qui dialogue avec ChatGPT dialoguera avec votre propre serveur après un changement d'adresse et de clé. Pour un premier essai sur un ordinateur portable, Ollama ou llama.cpp sont plus simples. NVIDIA publie des guides de configuration pour le Spark, et les recettes communautaires pour connecter deux, trois et quatre unités sont publiques et reproductibles.

Si exploiter votre propre modèle n'est pas réaliste avant la prochaine session, la solution de repli est un service hébergé assorti d'un accord écrit de conservation nulle des données. Traitez cela comme une mesure transitoire plutôt que comme une destination.

11.4 Une réserve à connaître

L'indice ci-dessus est une moyenne de neuf évaluations, et les modèles ouverts ne sont pas en retrait de manière uniforme dans chacune d'elles. L'écart est le plus marqué pour les tâches de raisonnement les plus difficiles et pour les hallucinations, domaine dans lequel les modèles propriétaires restent nettement en tête. Pour un négociateur, ce second point décide de tout. Orientez le modèle vers vos propres documents, demandez des citations, et ouvrez-les.

■ REMARQUE Nous sommes heureux d'aider

GainForest et la Youth Negotiators Academy conseillent gratuitement les délégations à ce sujet : dimensionnement du matériel, choix d'un modèle, mise en service et formation des personnes qui l'utiliseront. Écrivez à team@gainforest.net.

L'empreinte environnementale

Un négociateur climat qui utilise un outil consommant de l'énergie et de l'eau devrait pouvoir dire en quelle quantité. La réponse honnête comporte deux volets, qui vont dans des directions différentes.

12.1 Votre propre utilisation est limitée

Google a mesuré une requête textuelle médiane adressée à Gemini à 0,24 Wh d'énergie, 0,03 gramme de CO₂e et 0,26 millilitre d'eau.¹ Cela correspond à quelques secondes d'utilisation d'un ordinateur portable. Même une journée de travail avec un usage intensif reste inférieure à un court trajet en voiture.

Polly indique cette donnée pour chaque compte, et le calcul est publié plutôt qu'affirmé.² Il part des paramètres actifs du modèle :



ÉNERGIE PAR JETON

2 FLOP par paramètre actif → diviser par le débit de la puce → multiplier par la puissance de la puce

EAU ET CARBONE

1,1 litre par kWh pour le refroidissement → 480 grammes de CO₂e par kWh sur un réseau moyen

La lecture de votre requête coûte environ un quart de ce que coûte la rédaction d'une réponse, et la relecture d'une requête mise en cache coûte un dixième; les trois opérations sont donc comptabilisées séparément. Sur le modèle épars qu'exécute Polly, qui active 13 milliards de paramètres par jeton, un échange typique représente environ 0,13 Wh et 0,14 millilitre d'eau. Ce chiffre se situe juste en dessous de la médiane de Google.

Considérez chacun de ces chiffres comme exact à un facteur deux près. Ils n'incluent pas l'entraînement du modèle, l'eau utilisée pour produire l'électricité, votre propre appareil, ni le carbone incorporé dans le matériel.

12.2 L'utilisation de l'industrie ne l'est pas

Les chiffres intéressants se situent à l'échelle du secteur plutôt qu'à celle de la personne. La demande d'électricité des centres de données augmente plus vite que les réseaux ne se décarbonent, de nouvelles capacités sont construites dans des régions soumises au stress hydrique, et les entreprises qui communiquent les chiffres par

¹ Google, *Measuring the environmental impact of delivering AI at Google scale*, arXiv :2508.15734, August 2025. Ce chiffre inclut la capacité inactive et les frais généraux des centres de données, que la plupart des estimations publiées omettent.

² La méthode et chaque constante figurent dans `footprint.ts` dans `pi-village-core`. Un chiffre qu'un négociateur ne peut pas interroger est pire que l'absence de chiffre.

requête sont les mêmes dont les émissions totales ont augmenté depuis qu'elles ont commencé à construire ces systèmes.

Les deux volets sont vrais simultanément. Votre propre utilisation est négligeable dans l'arrondi, et l'agrégat constitue une revendication sérieuse et croissante sur l'énergie, l'eau et les terres. Énoncer le premier sans le second relève du type de comptabilité qu'un négociateur n'accepterait pas d'une Partie.

12.3 Que faire

Choisissez un modèle plus petit lorsqu'un modèle plus petit suffit. La traduction et la synthèse n'exigent pas un modèle de pointe, et un modèle épars qui active quelques milliards de paramètres coûte une fraction du coût d'un modèle dense.

Demandez aux fournisseurs leurs chiffres, y compris pour l'eau. Demandez à partir de quel réseau la requête est servie. Ce sont des questions d'achat raisonnables, et elles sont rarement posées.

Indiquez votre propre chiffre lorsque vous utilisez ces outils dans votre travail. Une délégation qui publie l'empreinte de sa propre utilisation de l'IA est légitime à demander aux autres d'en faire autant.

Conclusion et recommandations

13.1 Forces et faiblesses

Forces. Grande efficacité dans le traitement et la synthèse des vastes volumes de textes sur lesquels repose la politique climatique. Capacité à produire des rapports, à rédiger des textes de politique générale et à simuler des négociations en langage naturel, ce qui permet d'économiser du temps et des ressources. Capacité à assister simultanément de nombreuses personnes, des décideurs aux chercheurs et aux militants.

Faiblesses. Biais potentiels dans les modèles, qui faussent la manière dont l'information est traitée s'ils ne sont pas surveillés et corrigés. Dépendance à l'égard de la qualité et de l'étendue des données d'entraînement, qui limite la capacité d'un modèle à traiter des questions climatiques nuancées. Difficulté à interpréter avec exactitude un langage juridique et technique complexe sans supervision humaine.

Possibilités. Stratégies de communication personnalisées permettant de mobiliser différentes parties prenantes en faveur de l'action climatique. Intégration avec d'autres outils d'analyse des données environnementales, à l'appui de politiques mieux éclairées. Éducation et diffusion plus larges de l'information relative aux politiques climatiques, la rendant accessible aux non-spécialistes.

Menaces. Désinformation, lorsqu'un modèle génère un contenu incorrect ou trompeur à partir de sources peu fiables. Dépendance excessive, qui freine le développement de l'expertise locale en matière de politique climatique. Préoccupations éthiques concernant la transparence et la responsabilité dans la prise de décision automatisée.

Les agents ajoutent deux autres menaces qu'il convient de nommer. Les délégations dotées de ressources feront fonctionner des systèmes bien gouvernés, tandis que les petites délégations utiliseront des outils gratuits aux conditions peu claires, ce qui élargit l'écart sous l'apparence d'un accès égal. Et, lorsque plusieurs personnes et plusieurs outils interviennent dans une chaîne, la responsabilité d'un texte devient difficile à retracer, à moins que quelqu'un ne décide à l'avance où elle se situe.

13.2 Recommandations

Trois éléments méritent d'être mis en place : des outils de traduction du langage technique de la UNFCCC, des plateformes d'analyse des données historiques de négociation, et des environnements dans lesquels les jeunes négociateurs disposent de la technologie dont ils ont besoin pour contribuer. À ceux-ci s'en

ajoutent six autres.

Culture numérique. Formation visant spécifiquement à naviguer dans les bases de données et plateformes internationales relatives aux politiques publiques.

Appui en temps réel. Outils fournissant des clarifications pendant les négociations plutôt qu'après celles-ci.

Formation virtuelle à la négociation. Environnements simulés permettant de s'exercer à des scénarios de négociation.

Plateformes d'appui financier. Plateformes de financement participatif ou de micro-subventions répondant aux obstacles financiers à la participation aux négociations, soulevés par des négociateurs du Libéria et du Paraguay.

Formation à la communication interculturelle. Programmes répondant aux malentendus recensés par des négociateurs du Pérou et du Nigéria.

Programmes de mentorat. Réseaux reliant les jeunes négociateurs à des négociateurs expérimentés, afin que les connaissances se transmettent d'un cycle à l'autre.

Pour un négociateur, quatre habitudes couvrent l'essentiel. Utiliser ces systèmes pour la préparation, la langue et la mémoire. Vérifier tout élément que l'on répètera à voix haute. Ne jamais laisser un système s'exprimer au nom de sa Partie. Continuer à effectuer soi-même une part suffisante des recherches pour pouvoir encore s'en acquitter sans aide.

Pour une délégation, consigner par écrit les outils autorisés, ce qui peut y être collé, et qui approuve un texte assisté par l'IA. Une page suffit, et elle coûte moins cher qu'un incident.



*En mobilisant le potentiel de l'IA, la jeunesse d'aujourd'hui
ne se contente pas d'hériter de l'avenir, elle le façonne
activement; que le prochain chapitre de l'histoire de notre
planète soit écrit avec les traits audacieux de l'innovation et
du leadership des jeunes.*



Les questions et commentaires sont les bienvenus à l'adresse team@gainforest.net.

Annexe : synthèse des entretiens

Synthèse consolidée des entretiens avec de jeunes négociateurs et négociatrices sur le climat, menés afin de comprendre les difficultés auxquelles ils sont confrontés et la manière dont les modèles de langage peuvent appuyer leurs travaux.

Personnes interrogées. Six négociateurs et négociatrices, originaires du Liberia, du Paraguay, du Pérou, du Nigéria, du Liban et de l'Indonésie. Leur expérience variait, certains étant nouveaux dans le processus.

14.1 Difficultés et points de blocage

Dans le rôle exercé. Obstacles linguistiques lorsque l'anglais n'est pas la première langue. Le jargon technique de la UNFCCC (Liberia, Paraguay). Lacunes dans les connaissances techniques et dans les modalités de partage de l'information. Méconnaissance de l'historique des négociations (Liban). Difficulté à dialoguer avec des négociateurs chevronnés sur des sujets complexes (Pérou, Nigéria). Difficulté à s'adapter rapidement à la dynamique des négociations.

Dans des situations particulières. Les premières expériences à la COP ont été intimidantes en raison d'un manque de préparation (Paraguay, Pérou). La prise de parole en public dans un environnement à forts enjeux rend les sujets complexes difficiles à formuler.

Dans la collecte d'informations. Le site web de la UNFCCC est la principale source et il est complexe à parcourir. Le temps disponible pour traiter les informations essentielles est insuffisant (Paraguay). La diffusion repose sur les communications internes et les réseaux personnels.

Dans la communication. Il est difficile d'exprimer rapidement et avec exactitude des idées complexes en anglais. Les niveaux de maîtrise de l'anglais varient et créent des obstacles (Indonésie).

14.2 Outils

Outils utilisés. Grammarly pour l'aide à la rédaction (Paraguay). Google Drive pour la collaboration documentaire (Paraguay). Aucun recours à des outils de traduction, au motif que la langue est difficile à traduire (Paraguay).

Outils souhaités. Des outils linguistiques et grammaticaux plus sophistiqués pour les textes de la UNFCCC (Paraguay, Liberia). Des plateformes permettant un examen efficace des documents et l'extraction des informations clés (Liban,

Nigéria). Des ressources d'apprentissage adaptées (Liberia).

14.3 Traitement de l'information et prise de décisions

Suivi de l'actualité. Par l'intermédiaire de collègues et de ressources partagées, ainsi que par des bulletins d'information et des plateformes en ligne (Liberia, Liban).

Approche des informations complexes. Travail d'équipe et mobilisation de l'expertise d'autrui (Pérou), ainsi que recherches approfondies sur les sujets peu familiers.

Collaboration. Des communautés de soutien entre négociateurs sont activement encouragées (Pérou, Nigéria). Des initiatives visant à améliorer les compétences linguistiques sont en cours (Indonésie).

14.4 Idées et solutions

Améliorations souhaitées. Simplification des documents officiels, dans un souci d'inclusivité (Liban, Nigéria).

Tâches répétitives méritant d'être automatisées. Récupération de détails historiques relatifs aux négociations (Liban).

Vision d'avenir. L'IA et les outils numériques sont perçus comme des appuis potentiels aux négociations.

14.5 Possibilités offertes par les modèles de langage

Avantages. Traduction et synthèse, qui répondent aux obstacles linguistiques et condensent l'information.

Automatisation. Synthèse à partir de documents volumineux.

Réserves. La dépendance à l'égard de la technologie pourrait affaiblir les compétences essentielles en matière de recherche (Liberia, Paraguay).

Glossaire de l'IA

Termes que vous entendrez dans la salle, définis simplement.

15.1 Termes fondamentaux

Intelligence artificielle (IA)	Systèmes conçus pour accomplir des tâches associées à l'intelligence humaine, telles que la prédiction, la classification, la génération ou la planification.
Apprentissage automatique	Méthode par laquelle un système apprend des motifs récurrents à partir d'exemples au lieu de recevoir explicitement chaque règle.
Modèle	Système mathématique appris qui associe des entrées à des sorties.
Grand modèle de langage (LLM)	Modèle entraîné sur de vastes collections de textes afin de prédire et de générer du langage.
IA générative	IA qui produit de nouveaux textes, images, sons, codes ou autres résultats.
Système agentique / agent d'IA	Modèle relié à des outils, à une mémoire, à des instructions et à des boucles d'action afin de pouvoir accomplir plusieurs étapes.
Automatisation	Utilisation d'un système pour accomplir une tâche avec une intervention humaine directe limitée.
Augmentation	Utilisation de l'IA pour appuyer le travail humain plutôt que le remplacer.
Humain dans la boucle	Une personne examine, approuve ou intervient dans un processus assisté par l'IA.
Humain sur la boucle	Une personne supervise un système sans examiner chaque action.
Humain hors de la boucle	Un système agit sans qu'une personne examine chaque décision ou action.

15.2 Données et connaissances

Données d'entraînement	Exemples utilisés pour ajuster les paramètres d'un modèle pendant l'entraînement.
Jeu de données	Collection organisée de données utilisée pour l'entraînement, les tests ou l'analyse.
Provenance des données	Informations indiquant d'où proviennent les données et comment elles ont été collectées ou modifiées.
Traçabilité des données	Historique de la manière dont les données circulent dans les systèmes et les transformations.
Données structurées	Données organisées en champs, lignes, catégories ou schémas définis.
Données non structurées	Matériaux tels que texte, images, audio ou vidéo, sans structure tabulaire fixe.
Données synthétiques	Données générées artificiellement et conçues pour ressembler à des données réelles.
Vérité de terrain	Étiquette ou résultat de référence considéré comme correct aux fins d'évaluation.
Base de connaissances	Collection organisée d'informations qu'un système peut retrouver.
Connaissances paramétriques	Informations encodées dans les paramètres appris d'un modèle.
Connaissances factuelles	Informations sur des faits ou des relations qu'un modèle peut avoir apprises à partir de ses données.

	d'entraînement.
Génération augmentée par récupération (RAG)	Système qui récupère des documents ou des enregistrements avant de générer une réponse.
Connaissances locales	Connaissances ancrées dans l'expérience d'un lieu ou d'une communauté particulière.
Savoirs autochtones	Systèmes de connaissances élaborés et maintenus par les peuples et communautés autochtones.
Souveraineté des données	Principe selon lequel les données sont régies conformément aux lois, aux droits et à l'autorité des personnes ou du lieu qu'elles concernent.
Minimisation des données	Collecte et conservation uniquement des données nécessaires à une finalité déclarée.
Données sensibles	Données dont la divulgation ou l'utilisation abusive pourrait causer un préjudice.

15.3 Entrées et sorties des modèles

Entrée / prompt	Instruction ou matériau fourni à un modèle.
Sortie	Texte, prédiction, classification, image ou action produit par un système.
Inférence	Exécution d'un modèle entraîné pour produire une sortie.
Fenêtre de contexte	Quantité de texte ou d'autre matériau qu'un modèle peut prendre en compte en une seule fois.
Jeton	Unité de texte ou d'autres données traitée par un modèle de langage.
Plongement	Représentation numérique d'un texte, d'images ou d'autres données, utilisée pour comparer des relations.
Classification	Attribution d'un élément à une ou plusieurs catégories.
Prédiction	Estimation d'un résultat probable à partir des données disponibles.
Prévision	Prédiction relative à un état ou à un événement futur.
Recommandation	Sortie qui suggère une ligne d'action.
Optimisation	Recherche du meilleur résultat selon un objectif et des contraintes définis.
Fonction objectif	Quantité formelle qu'un système cherche à maximiser ou à minimiser.
Variable de substitution	Mesure indirecte utilisée à la place d'un objectif plus difficile à mesurer.
Référence d'évaluation	Test utilisé pour comparer les performances des modèles.
Évaluation	Mesure d'un modèle ou d'un système au regard de critères définis.
Incertitude	Degré selon lequel une sortie peut être erronée ou incomplète.
Score de confiance	Estimation numérique associée à une prédiction ou à une classification.
Étalonnage	Relation entre la confiance prédite et l'exactitude réelle.
Explicabilité	Capacité de fournir un compte rendu compréhensible de la manière dont une sortie a été produite.
Interprétabilité	Mesure dans laquelle le comportement interne d'un système peut être compris.
Boîte noire	Système dont le fonctionnement interne est difficile à examiner.
Hallucination	Sortie de modèle fluide mais non étayée, fausse ou fabriquée.
Biais	Erreur systématique ou performance inégale associée aux données, à la conception ou à l'utilisation.
Équité	Approches permettant d'évaluer et de traiter les traitements ou résultats inégaux.
Robustesse	Capacité de fonctionner de manière fiable dans des conditions modifiées ou imparfaites.
Attaque antagoniste	Entrée conçue délibérément pour provoquer la défaillance d'un système ou un comportement inattendu.

15.4 Modèles et entraînement

Réseau neuronal	Système informatique composé de couches connectées qui transforment des entrées en sorties.
Perceptron	Modèle mathématique précoce d'une cellule cérébrale unique, et fondement de recherches ultérieures sur les réseaux neuronaux.
Transformer	Architecture de réseau neuronal qui utilise l'attention pour traiter les relations entre jetons.
Poids	Nombres qu'un modèle a appris pendant l'entraînement. En pratique, les poids sont le modèle.
Poids ouverts	Poids publiés de sorte que chacun puisse exécuter, examiner ou héberger le modèle lui-même.
Poids fermés	Poids conservés sur les serveurs du fournisseur et accessibles uniquement par l'intermédiaire de son service.
Préentraînement	Entraînement initial sur un vaste corpus de données, généralement en prédisant un contenu manquant ou subséquent.
Ajustement fin	Entraînement supplémentaire sur un jeu de données ou une tâche plus restreint.
Post-entraînement	Processus appliqués après le préentraînement pour améliorer l'utilité, le comportement ou l'alignement.
Apprentissage par renforcement (RL)	Entraînement au moyen d'un retour d'information qui récompense certaines actions et en pénalise d'autres.
Apprentissage par renforcement à partir de retours humains (RLHF)	Apprentissage par renforcement utilisant des préférences ou des évaluations humaines.
Modèle de raisonnement	Modèle ou système optimisé pour consacrer des calculs supplémentaires à des problèmes en plusieurs étapes.
Chaîne de pensée	Texte de raisonnement intermédiaire généré pendant la résolution d'un problème.
Lois d'échelle	Relations observées entre les performances d'un modèle, sa taille, les données d'entraînement et la puissance de calcul.
Mélange d'experts (MoE)	Architecture contenant plusieurs sous-réseaux spécialisés, dont seuls certains sont utilisés pour chaque entrée.

15.5 Gouvernance et responsabilité

Alignement	Conception d'un système de sorte que son comportement reflète des objectifs ou des valeurs spécifiés.
Sécurité de l'IA	Recherches et pratiques visant à réduire les comportements préjudiciables ou incontrôlés des systèmes.
Gouvernance de l'IA	Règles, institutions, processus et pratiques qui façonnent le développement et l'utilisation de l'IA.
IA responsable	Développement et utilisation de l'IA en tenant compte des incidences sociales, éthiques, juridiques et environnementales.
Conception sensible aux valeurs	Conception de technologies accordant une attention explicite aux valeurs humaines et aux intérêts concernés.
IA alignée sur la souveraineté	IA conçue et gouvernée pour préserver la capacité d'action d'une communauté ou d'un lieu défini.
Autorisations granulaires	Restrictions concernant les outils, fichiers, données ou actions auxquels un système peut accéder.
Bac à sable	Environnement isolé dans lequel un système peut fonctionner avec un accès limité.
Point d'approbation humaine	Autorisation humaine requise avant qu'une action ne se poursuive.
Journal d'audit	Enregistrement des entrées, sorties, actions et modifications d'un système.
Fiche de modèle	Documentation décrivant l'utilisation prévue, les limites et l'évaluation d'un modèle.

Fiche de système	Documentation couvrant un système déployé plus large, y compris les risques et les mesures d'atténuation.
Dérive du modèle	Variation des performances à mesure que les données, les comportements ou les conditions évoluent au fil du temps.
Gestion des versions	Suivi des modifications apportées aux modèles, aux données, aux prompts et aux logiciels.
Retour à une version antérieure	Rétablissement d'un système ou d'un déploiement à une version antérieure.
Auto-hébergement	Exécution d'un modèle sur du matériel que vous ou votre institution contrôlez.
Conservation nulle des données	Engagement contractuel selon lequel un fournisseur ne stocke rien de ce que vous lui envoyez.

Références

1. Le perceptron de Rosenblatt et la démonstration de 1958, Cornell Chronicle
2. Krizhevsky, Sutskever, Hinton, *ImageNet Classification with Deep Convolutional Neural Networks*, NeurIPS 2012
3. Vaswani et al., *Attention Is All You Need*, Google Research 2017
4. Radford et al., *Language Models are Unsupervised Multitask Learners*, OpenAI 2019
5. Kaplan et al., *Scaling Laws for Neural Language Models*, OpenAI 2020
6. Hoffmann et al., *Training Compute-Optimal Large Language Models*, NeurIPS 2022
7. OpenAI, *Introducing ChatGPT*, 30 novembre 2022
8. OpenAI, *Learning to Reason with LLMs*, 12 septembre 2024
9. DeepSeek, *DeepSeek-R1*, 20 janvier 2025
10. METR, *Measuring AI Ability to Complete Long Tasks*, 2025
11. Anthropic, *Anthropic Economic Index*, 2025, utilisation par pays et par tranche de revenu
12. Global Biodiversity Information Facility, données d'occurrence et biais géographique de l'échantillonnage, gbif.org et le blog des données du GBIF
13. Artificial Analysis, *Intelligence Index by Open Weights and Proprietary*, 20 août 2026
14. Artificial Analysis, *Intelligence Index vs. Cost per Intelligence Index Task*, 20 août 2026
15. NVIDIA, guide utilisateur et aperçu matériel de DGX Spark
16. Google, *Measuring the environmental impact of delivering AI at Google scale*, arXiv :2508.15734, août 2025
17. Epoch AI, *How much energy does ChatGPT use?*, 2025
18. GainForest, footprint.ts dans pi-village-core, le calcul qui sous-tend les chiffres par compte de Polly
19. GainForest, *Generative AI + Climate Guide 2024*, établi pour le Climate Youth Negotiator Programme en amont de COP29