

GUIDE COMPLET DU PORTAIL IA DE L'UNISTRA

Version : 1

Date : 07/07/2026

Éditeur : Université de Strasbourg - Direction du Numérique

Auteur(s) : Alexandra Barberet, DNum

Copyright : Université de Strasbourg

Licence : Licence Creative Commons : Paternité - Pas d'Utilisation Commerciale - Partage des Conditions Initiales à l'Identique

Table des matières

Introduction	4
1. Accès à la plateforme portail	6
2. Prise en main : Découvrir le Chatbot et ses Modèles	9
2.1. Présentation de l'interface	9
2.2. Options de saisie et d'ajout de ressources externes	10
2.3. Catalogue de nos modèles IA	12
2.4. Raccourcis clavier	15
3. Essentiel : Savoir interagir avec le Chatbot	16
3.1. Lancer un échange : l'art du prompting	16
3.2. Affiner le résultat : Itérer pour optimiser	18
3.3. Enrichir l'IA : Téléverser des fichiers et documents	20
4. Niveau intermédiaire : Organiser et optimiser ses contenus	22
4.1. Réglages du compte et préférences	22
4.2. Gestion des conversations et dossiers	26
4.3. Utilisation de l'éditeur de Notes	27
4.4. Conception de prompts automatisés	29
5. Niveau avancé : Utilisation des bases de connaissances (RAG)	31
5.1. Principes du RAG et bases de connaissances (BDC)	31
5.2. Comparaison entre BDC et Téléversement	31
5.3. Création d'une BDC	32
6. Niveau expert : 3 solutions pour automatiser vos tâches	34
6.1. Le Prompt automatisé avec variables	34
6.2. Le Dossier intelligent	37
6.3. L'Assistant IA personnalisé	39
7. Intégrations et API : Connecter l'IA à vos applications	45
7.1. Procédure de génération de la clé API	45
7.2. Intégration dans OnlyOffice	47
7.3. Intégration dans Thunderbird	55
7.4. Intégration dans VS Code avec Cline	61
7.5. Accès mobile via "Conduit"	65
8. SPEAKR : outils de transcription et sous-titrage	68
8.1. Accès et configuration initiale	68
8.2. Enregistrement et importation de l'audio	73
8.3. Traitement et affinage de la transcription	75
8.4. Génération du Compte-rendu	77
9. GLM-OCR : conversion de documents numérisés en texte éditable	78
10. DOCLING : conversion de documents en format optimisé pour l'IA	81
11. FAQ	85
11.1. Puis-je transmettre des données à caractère personnel à une IA ?	85
11.2. Puis-je transmettre des données sensibles à une IA ?	85
11.3. Comment protéger ma propriété intellectuelle ?	86
11.4. Les données que je transmets à l'IA de l'UNISTRA sont-elles conservées ?	86
11.5. Qui est responsable des erreurs ou des contenus générés par l'IA ?	87

11.6. Quelle est la différence entre l'IA de l'Unistra et un outil comme ChatGPT ?	87
11.7. L'utilisation de l'IA a-t-elle un impact écologique ?	88
11.8. Comment citer l'usage de l'IA dans mes travaux ?	89
11.9. Où puis-je obtenir de l'aide pour utiliser l'IA ?	90

Introduction

L'Université de Strasbourg s'engage dans une **stratégie numérique responsable et souveraine** pour accompagner l'intégration de l'intelligence artificielle dans les activités de recherche, d'enseignement et d'administration. Dans ce cadre, l'établissement met à la disposition de sa communauté une plateforme d'**IA conversationnelle**.

Qu'est-ce que l'IA conversationnelle Unistra ?

On désigne souvent cet outil sous le terme de « **chatbot** » (ou agent conversationnel). L'IA conversationnelle de l'Unistra s'appuie sur des **grands modèles de langage (LLM)**. Elle permet ainsi un échange fluide, capable d'intégrer des instructions complexes, de synthétiser des documents ou de générer du contenu en langage naturel.

Objet

Ce guide a pour objectif de vous accompagner dans la prise en main rapide de l'IA conversationnelle ou chatbot de l'Université de Strasbourg. Il présente les fonctionnalités de base de la plateforme, ainsi que les bonnes pratiques d'utilisation.

Public concerné

Enseignants, chercheurs et personnels administratifs et techniques de l'Université de Strasbourg.

Principes fondamentaux

L'utilisation du chatbot et des outils d'intelligence artificielle de l'Université de Strasbourg s'inscrit dans une démarche de responsabilité et de rigueur académique. Conformément aux **lignes directrices relatives à l'usage des IA** [https://ernest.unistra.fr/jcms/1442664709_UDSAactualites/fr/publication-de-lignes-directrices-sur-l-usage-des-outils-d-intelligence-artificielle], vous devez respecter les principes suivants :

Réflexion préalable

L'utilisation de ces outils doit s'inscrire dans une démarche raisonnée et maîtrisée. L'Université de Strasbourg encourage un usage pertinent de l'IA, justifié par une **valeur ajoutée avérée** (mise en qualité, gain de temps conséquent) et conforme aux principes de sobriété numérique.

Transparence

L'usage de l'IA dans une production (travail académique, document administratif, etc.) doit **être documenté et traçable**. Identifiez clairement les passages générés ou assistés par l'IA, comme vous le feriez pour toute citation ou emprunt à une source externe.

Responsabilité

L'utilisateur demeure pleinement responsable des contenus produits et de leur conformité à la réglementation (discretion professionnelle, protection des données personnelles, respect de la propriété intellectuelle). L'IA étant un assistant probabiliste susceptible d'erreurs, de biais ou d'hallucinations, il est impératif de **vérifier** systématiquement les faits, les chiffres et les références avant toute publication. **Ne faites jamais confiance aveugle aux résultats.**



Protection des données à caractère personnel

Cet outil est conçu pour vous assister dans l'exécution de vos missions : veillez à n'y exploiter que des données professionnelles ou nécessaires à la conduite de vos activités. Vous êtes libre de supprimer vos conversations quand vous le souhaitez dans la partie "Réglages".

Pour plus de détails, veuillez vous référer aux conditions d'utilisation sur le portail : <https://portail.ia.unistra.fr/>

Documentation assistée par l'IA

Ce document a été conçu et validé par son auteur, avec l'assistance de l'intelligence artificielle pour la mise en forme et la reformulation.



source pictogramme : Les bibliothèques de l'université de Montréal



+ Compléments

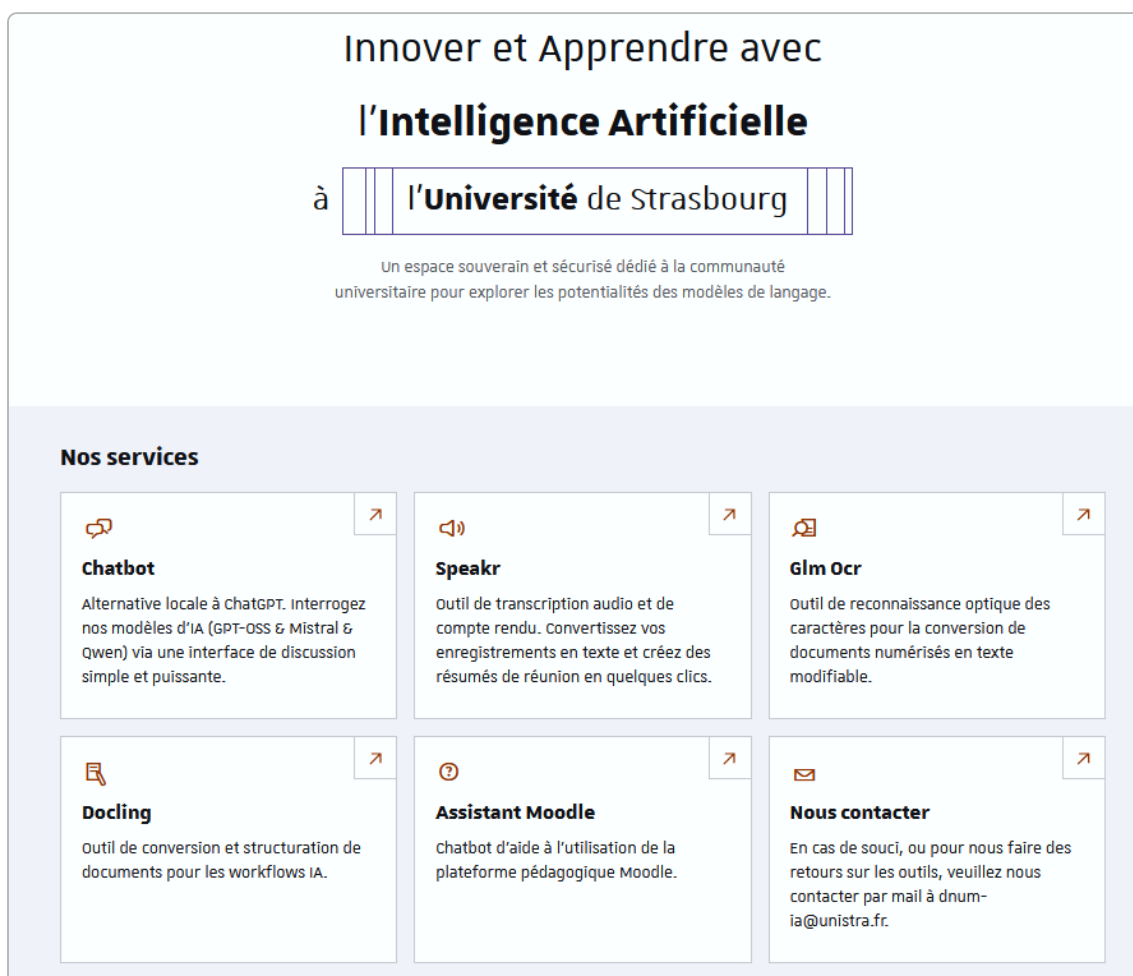
- [Télécharger la version PDF de cette documentation](#) [..?V=pdf]
- [Stratégie IA Unistra et Lignes directrices sur l'usage des outils d'IA](https://ernest.unistra.fr/jcms/1535118102_UDSPages/fr/strategie-ia-unistra) [https://ernest.unistra.fr/jcms/1535118102_UDSPages/fr/strategie-ia-unistra]

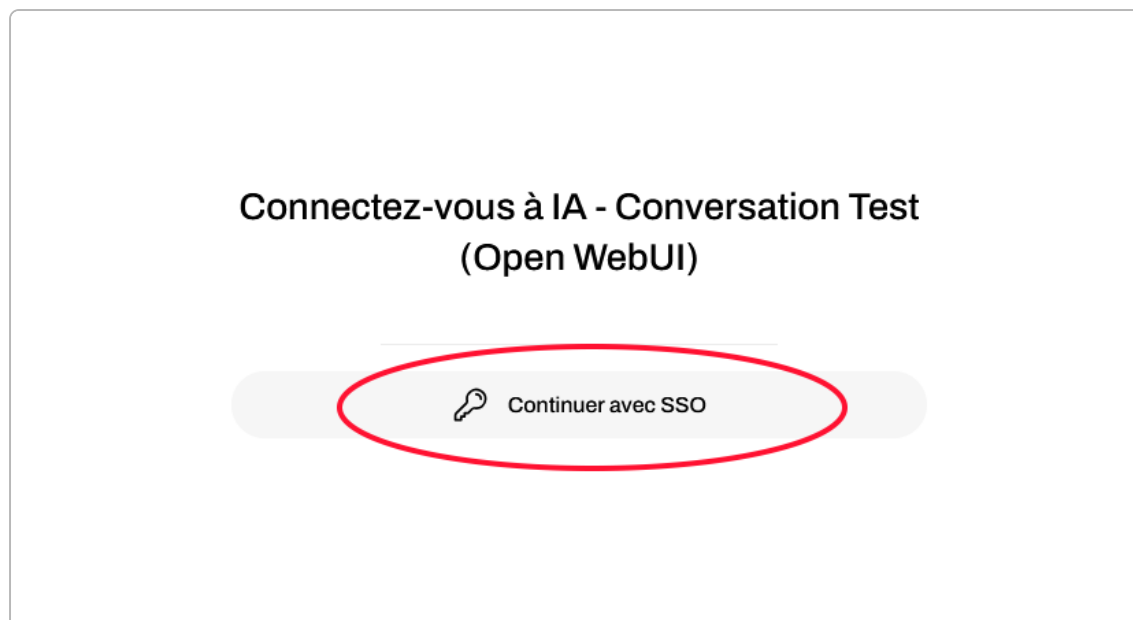
1. Accès à la plateforme portail

Accès au portail dédié à l'IA de l'Université de Strasbourg

Pour accéder au portail

1. Rendez-vous sur le portail dédié à l'Intelligence Artificielle :
portail.ia.unistra.fr [https://portail.ia.unistra.fr/].
2. Sur la page d'accueil, sélectionnez l'outil souhaité (par exemple : *Chatbot*, *Speakr*, *Docling*, *Assistant Moodle*, etc.) en cliquant sur le bouton correspondant.
3. Authentifiez-vous en utilisant vos identifiants universitaires habituels via la méthode « **Continuer avec SSO** ».





Présentation des outils disponibles

le Chatbot

Description : Interface d'interaction généraliste permettant d'interroger plusieurs modèles d'IA (GPT-OSS, Mistral, Qwen) pour obtenir des réponses, rédiger du texte ou générer du code.

Exemple d'usage : "Agis comme un ingénieur pédagogique et rédige un plan de cours sur pour des étudiants de licence 1 de ..., en incluant des objectifs pédagogiques et des activités d'évaluation."

Pour la prise en main, consultez la page dédiée  le Chatbot ^[p.9]

Speakr

Description : Outil de transcription audio et vidéo qui convertit vos enregistrements en texte et permet de générer automatiquement des comptes-rendus, des résumés ou des sous-titres.

Exemple de requête : Téléversez un fichier audio de réunion et demandez : "Extrais les 10 lignes les plus pertinentes de cette réunion et génère un compte-rendu structuré avec les actions à entreprendre."

Pour plus d'informations, consultez la page dédiée  Speakr ^[p.68]

Assistant Moodle

Description : Chatbot spécialisé dans l'aide à l'utilisation de la plateforme pédagogique Moodle, capable de répondre aux questions techniques et méthodologiques liées à l'outil.

Exemple de requête : "Comment créer un quiz aléatoire dans Moodle ?" ou "Quelles sont les étapes pour ajouter un forum de discussion à un module ?"

Pour plus d'informations, consultez la page dédiée [Assistant Moodle](https://assistance-moodle.ia.unistra.fr/) ^[https://assistance-moodle.ia.unistra.fr/]

Pour l'Optimisation du traitement des fichiers

Lors du téléversement d'un fichier (PDF, Word...), le chatbot le convertit automatiquement en Markdown pour le rendre lisible par l'IA. Bien que ce format soit le plus adapté pour l'IA, cette conversion automatique peut parfois produire des erreurs ou des données parasites (pieds de page, métadonnées) .

Pour garantir une meilleure qualité de traitement et des fichiers plus légers, il est recommandé d'utiliser des outils de pré-traitement dédiés :

Docling

Description : Service de conversion de documents (.docx, .pdf, .pptx, etc.) en texte brut ou en Markdown, optimisant ainsi la lisibilité des documents pour l'IA et facilitant la création de bases de connaissances.

Exemple d'usage : *Convertir un PDF complexe en Markdown pour l'intégrer ensuite dans une base de connaissances, ou extraire le texte d'un document scanné pour le réutiliser dans une conversation.*

Pour plus d'informations, consultez la page dédiée  Docling ^[p.81]

GLM OCR

Description : Outil de reconnaissance optique des caractères (OCR) spécialisé, conçu pour transformer des documents numérisés, des images complexes, des tableaux ou des écritures manuscrites en texte éditable et recherachable. Il est particulièrement recommandé pour les documents contenant des équations mathématiques ou des scans de qualité variable.

Exemple d'usage : *Convertir cette photo de notes manuscrites en texte numérique. Rendre ce document PDF scanné sélectionnable et recherché. Transcrire les équations mathématiques présentes sur cette image.*

Pour plus d'informations, consultez la page dédiée  GLM OCR ^[p.78]

2. Prise en main : Découvrir le Chatbot et ses Modèles

Cette section présente **l'interface du Chatbot** et détaille les **modèles de langage** disponibles, tels que Mistral, Qwen et Gemma, choisis pour leur performance et leur adéquation avec les usages académiques et administratifs. Elle détaille également les commandes de **raccourcis** pour optimiser l'utilisation de l'outil.

2.1. Présentation de l'interface

Interface globale

L'interface du Chatbot est structurée autour de trois zones principales :



1. L'en-tête et la sécurité

- **Sélecteur de modèle** : en haut à gauche, un menu déroulant permet de changer de modèle d'IA (voir Accès à la plateforme)
- **Profil utilisateur (en haut à droite)** : Situé en haut à droite, il permet d'accéder aux réglages du compte et aux préférences de l'interface. (silhouette dans un cercle).
- **Contrôle du modèle (en haut à droite)** : Situé en haut à droite, ce bouton permet d'ajuster les paramètres techniques de génération (tels que la température) pour modifier le comportement de l'IA (trois barres horizontales avec curseurs).
- **Conversations temporaires (en haut à droite)** : Située en haut à droite, cette option active un mode de discussion sans historique enregistré. (bulle de texte en pointillés).

2. La zone d'interaction centrale

C'est l'espace principal d'interaction où se déroule l'intégralité de vos échanges avec l'IA. Cette zone se compose des éléments suivants :

- **Le modèle utilisé** : En haut de la zone, vous pouvez visualiser le modèle d'IA utilisé pour la conversation. Dans ce cas précis, le modèle **GPT OSS 120B**.
- **Le champ de saisie** : c'est ici que vous rédigez vos requêtes. Il intègre plusieurs fonctionnalités : un bouton d'ajout permettant de téléverser des documents ou des images pour les analyser avec l'IA), un accès aux fonctionnalités avancées (recherche Web, Image et interpréteur de code), ainsi qu'une option de saisie vocale.

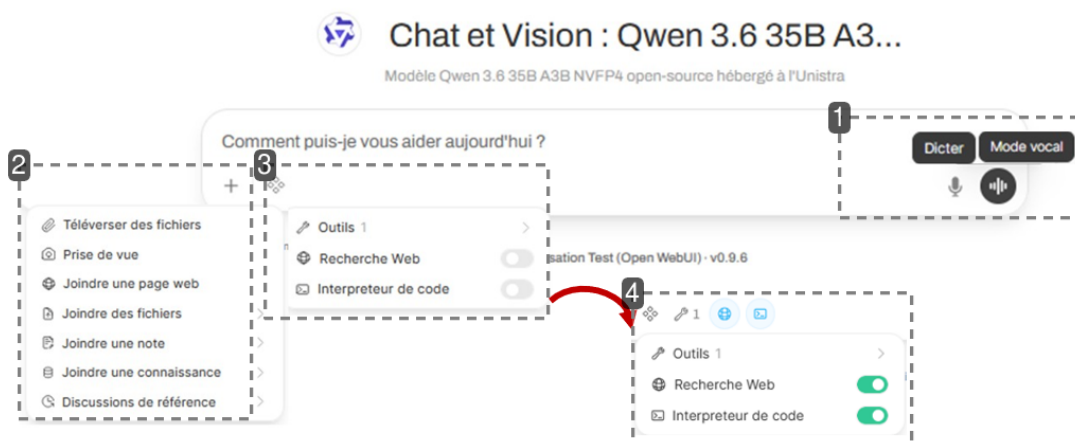
3. La barre latérale de navigation

Cette barre constitue votre centre de navigation et d'organisation.

- **Affichage/masquage de la barre latérale**  : Réduit ou affiche la barre latérale pour maximiser l'espace d'interaction.
- **Nouvelle conversation** : Démarre instantanément une nouvelle conversation. **Notez qu'une nouvelle conversation repart de zéro : l'IA n'a pas accès à la mémoire des échanges précédents, même s'ils portaient sur le même sujet.**
- **Recherche** : Permet de retrouver des informations spécifiques au sein de vos échanges passés.
- **Notes** : Stocke et met en forme des mémos personnels (prise de notes, brouillons). Cet outil est particulièrement efficace pour structurer et formaliser des notes de comptes-rendus de manière professionnelle.
- **Espace de travail** : Environnement global pour gérer vos bases de connaissances (BDC) et vos prompts personnalisés (voir  Niveau expert : 3 solutions pour automatiser vos tâches ^[p.34]).
- **Dossiers** : Organise vos projets et l'historique de vos conversations. Cet espace permet également de gérer vos prompts système, idéals pour structurer et répéter des tâches complexes.
- **Conversations** : Affiche l'historique complet de vos échanges avec l'IA.

2.2. Options de saisie et d'ajout de ressources externes

Cette vue d'ensemble présente les options de saisie et d'ajout de ressources externes, permettant d'activer des fonctionnalités telles que la dictée vocale, le menu « Intégrations » et les outils d'optimisation.



1. Fonctionnalités de saisie et de communication

L'interface propose trois méthodes de saisie pour interagir avec le Chatbot, selon vos besoins et votre environnement de travail :

- **La saisie textuelle** : rédigez vos instructions (prompts) directement dans le champ de saisie pour un contrôle précis de votre demande.
- **Le mode Dictée** : utilisez le microphone pour formuler votre question à l'oral ; l'IA retranscrit alors vos propos en texte avant l'envoi de la requête.
- **Le mode Vocal** : Donnez vos instructions de manière fluide et naturelle. Vous parlez à l'IA et celle-ci vous répond simultanément par écrit et par voie haute grâce à la synthèse vocale.

2. Fonctionnalités d'optimisation et de ressources

L'interface propose différentes méthodes pour enrichir vos requêtes et activer des outils spécifiques :

- **Téléverser des fichiers** : importer des documents dans la conversation. De nombreux formats sont supportés, notamment les documents (.pdf, .docx, .pptx, txt), le texte et le code (.txt, .md, .py, .json, .csv, .xml) ainsi que les images (.jpg, .png). Pour les images (.jpg, .png), cette action nécessite impérativement la sélection d'un modèle prenant en charge le mode « Vision ». *Pour les fichiers volumineux ou complexes (contenant beaucoup d'images ou de tableaux), il est recommandé de les convertir préalablement au format Markdown, à l'aide de l'outil Docling, pour optimiser la compréhension de l'IA.*
- **Prise de vue** : Capturer instantanément une zone de l'écran ou une fenêtre et l'analyser . Cette option nécessite également l'activation d'un modèle prenant en charge le mode « Vision ».
- **Joindre une page web** : extraire le contenu d'un site internet spécifié pour l'intégrer comme document de travail dans le chat.
- **Joindre des fichiers** : rappeler un document précédemment téléversé. Vos fichiers sont conservés en mémoire pendant trois mois, vous permettant de les réutiliser dans une nouvelle conversation sans avoir à les importer à nouveau.
- **Joindre une note** : intégrer une note personnelle comme contenu supplémentaire. Contrairement aux fichiers, les notes ne sont pas découpées en fragments (chunks : Un chunk est un fragment de texte extrait d'un document plus large. Pour traiter de longs documents, ceux-ci sont découpés en morceaux de taille définie (par exemple, entre 200 et 500 tokens) afin de respecter les limites de mémoire de l'IA mais sont transmises dans leur intégralité à l'IA, comme si vous les aviez saisies directement dans votre requête.
- **Joindre une connaissance** : utiliser l'une de vos bases de connaissances préalablement configurées pour interroger un ensemble documentaire spécifique.
- **Discussions de référence** : importer le contexte d'une conversation précédente pour poursuivre un échange.

3. Menu « intégrations »

Ce menu permet d'activer ou de désactiver des services complémentaires, tels que la recherche web ou l'interpréteur de code, pour modifier le comportement de l'IA.

- **Outils** : interroger la base de connaissances officielle de l'université pour obtenir des informations précises et sourcées issues des documents internes (icône clé à molette).
- **Recherche web** : consulter un moteur de recherche (Google dans notre cas) pour enrichir les réponses avec des informations récentes ou vérifier des faits en ligne (icône globe).
- **Interpréteur de code** : exécuter des scripts Python directement dans la conversation pour traiter des données ou générer des graphiques (icône terminal).

4. Exemple d'un menu « intégrations » entièrement activé

Lorsqu'une fonctionnalité est activée, son interrupteur apparaît en vert dans le menu déroulant. Lorsqu'un outil est activé, son icône s'affiche en bleu dans la zone de saisie.

2.3. Catalogue de nos modèles IA

La plateforme utilise des modèles dits « **open-source** ». Contrairement aux versions commerciales comme ChatGPT ou Gemini, ces modèles sont téléchargés et exécutés localement sur nos serveurs. Cela signifie que nous utilisons la "technologie" du fournisseur, mais **sans que vos données ne transitent par leurs services**.

Nos modèles IA

Mise à jour des modèles (par ordre alphabétique) : le 27/05/2026

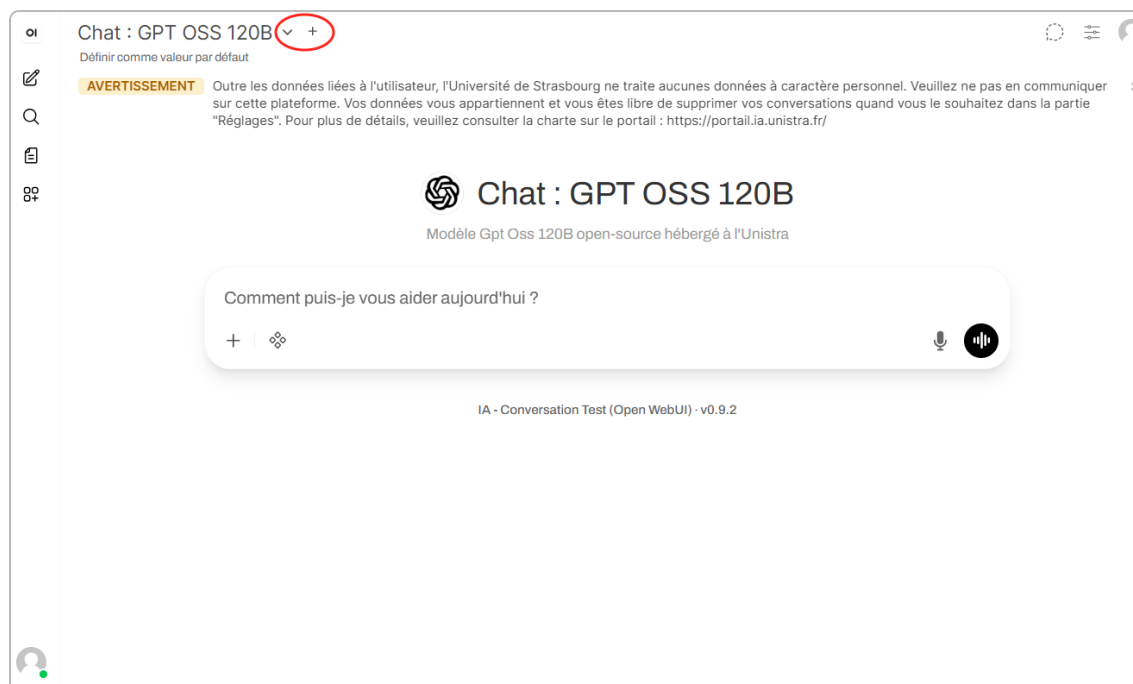
- **Gemma (31B)** : Développé par Google (USA). Ce modèle offre un excellent compromis pour la synthèse, la recherche de similitudes et peut agir comme un tuteur pour l'aide à la rédaction de prompts. (Usages : Chat et Vision).
- **GPT-OSS (120B/5B)** : Basé sur la technologie d'OpenAI (USA). Il s'agit d'une version open-source optimisée et non du produit commercial ChatGPT. Il est très efficace pour la restructuration d'idées et la réécriture. (Usages : Chat).
- **Mistral Small (119B/6B)** : Développé par Mistral AI (France). Ce modèle est particulièrement performant pour la langue française, la rédaction et les tâches générales. (Usages : Chat et Vision).
- **Qwen (35B/3B)** : Développé par Alibaba (Chine). Ce modèle est recommandé pour le raisonnement complexe, le traitement multilingue et la synthèse de documents. (Usages : Chat et Vision).
- **Qwen Coder (80B)** : Développé par Alibaba (Chine). Ce modèle est spécialisé dans la génération et la correction de code informatique. (Usages : Chat).

+ Remarques

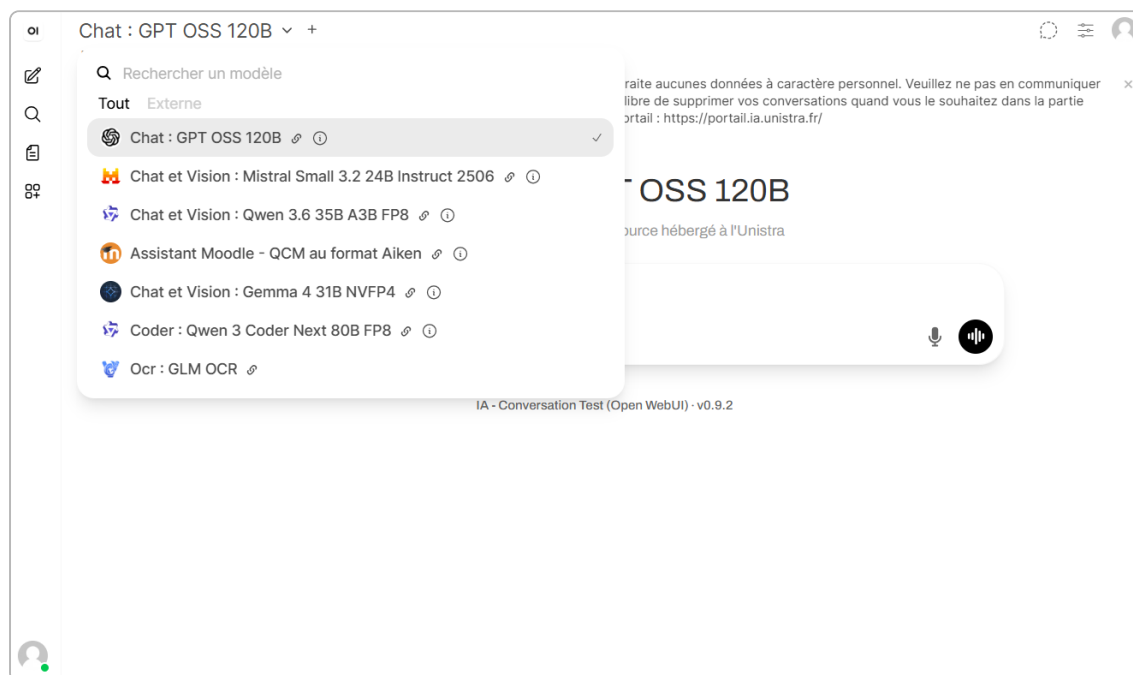
- Le chiffre suivi de la lettre « B » (ex: 24B) indique le nombre de paramètres du modèle en milliards (Billions), ce qui correspond approximativement à la « taille » et à la complexité du modèle.
- La mention « **Vision** » signifie que l'IA peut analyser et lire vos images (extraire du texte, décrire un objet ou un schéma), mais **elle ne peut pas générer ou créer des images**.

Sélection et configuration des modèles

Pour modifier le modèle utilisé, cliquez sur la flèche pointant vers le bas, située à droite du nom du modèle actif



Sélectionnez celui de votre choix dans la liste.

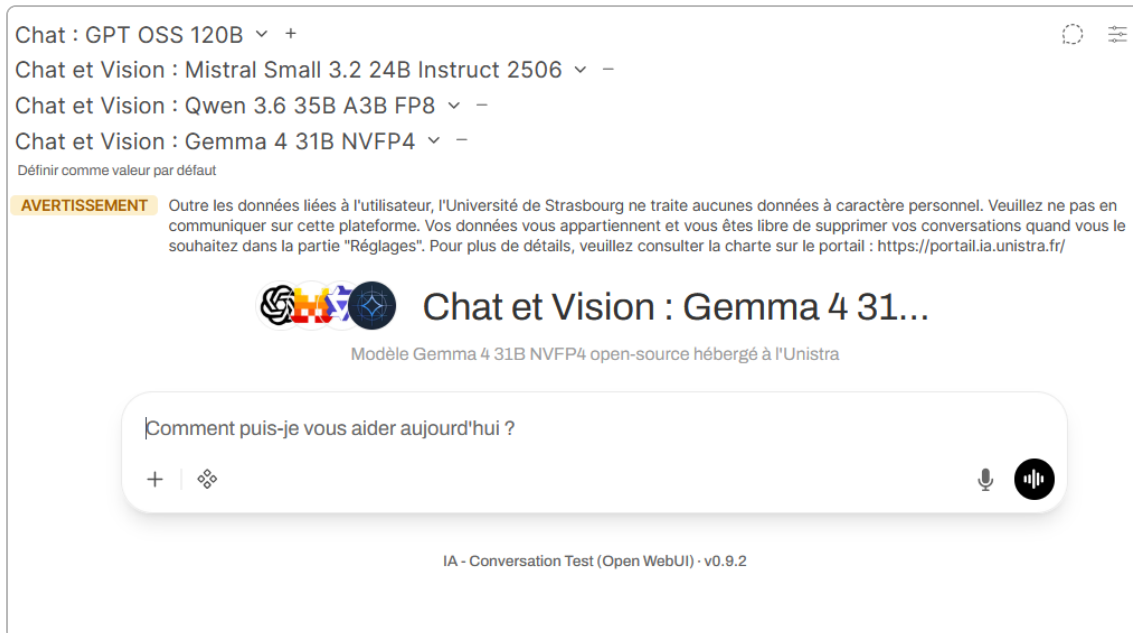


Conseil

Nous vous encourageons à tester les différents modèles selon vos besoins. Une fois que vous avez identifié celui qui vous convient le mieux, vous pouvez cliquer sur le bouton « **Définir comme valeur par défaut** » pour qu'il soit sélectionné automatiquement à chaque nouvelle session.

Comparer plusieurs modèles

1. Cliquez sur le bouton «  » situé à droite du nom du modèle actif.
2. Sélectionnez les modèles supplémentaires que vous souhaitez ajouter (vous pouvez ajouter les quatre modèles). Les réponses s'affichent côte à côte.
3. Pour retirer un modèle, cliquez sur le bouton «  » situé à droite du nom du modèle concerné.



Chat : GPT OSS 120B  +

Chat et Vision : Mistral Small 3.2 24B Instruct 2506  -

Chat et Vision : Qwen 3.6 35B A3B FP8  -

Chat et Vision : Gemma 4 31B NVFP4  -

Définir comme valeur par défaut

AVERTISSEMENT Outre les données liées à l'utilisateur, l'Université de Strasbourg ne traite aucune donnée à caractère personnel. Veuillez ne pas en communiquer sur cette plateforme. Vos données vous appartiennent et vous êtes libre de supprimer vos conversations quand vous le souhaitez dans la partie "Réglages". Pour plus de détails, veuillez consulter la charte sur le portail : <https://portail.ia.unistra.fr/>

 **Chat et Vision : Gemma 4 31...**

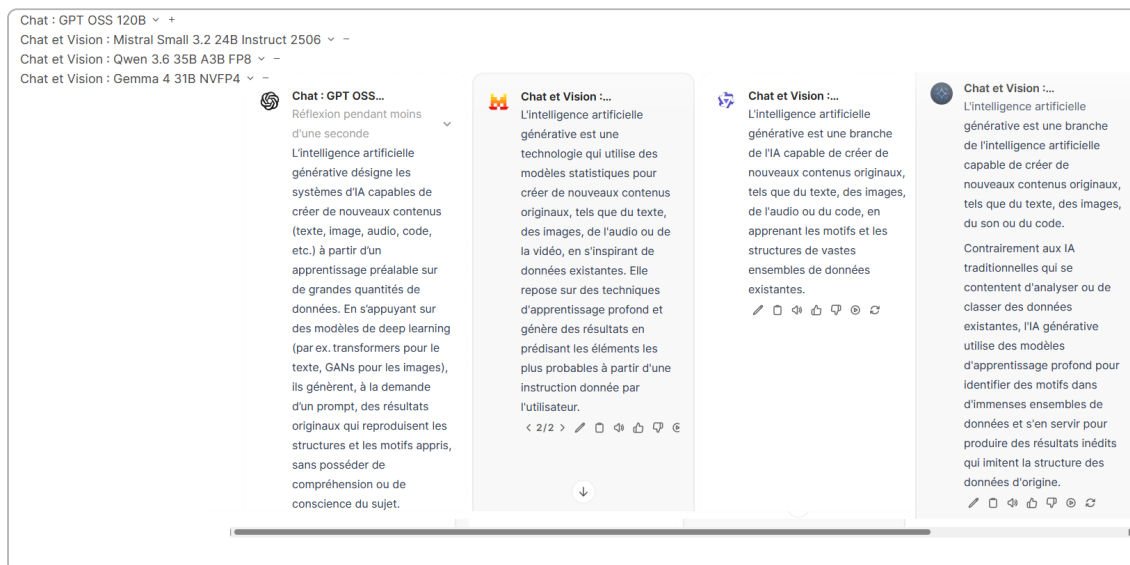
Modèle Gemma 4 31B NVFP4 open-source hébergé à l'Unistra

Comment puis-je vous aider aujourd'hui ?

+ 

IA - Conversation Test (Open WebUI) - v0.9.2

Exemple



Chat : GPT OSS 120B  +

Chat et Vision : Mistral Small 3.2 24B Instruct 2506  -

Chat et Vision : Qwen 3.6 35B A3B FP8  -

Chat et Vision : Gemma 4 31B NVFP4  -

Chat : GPT OSS...
Réflexion pendant moins d'une seconde
L'intelligence artificielle générative désigne les systèmes d'IA capables de créer de nouveaux contenus (texte, image, audio, code, etc.) à partir d'un apprentissage préalable sur de grandes quantités de données. En s'appuyant sur des modèles de deep learning (par ex. transformers pour le texte, GANs pour les images), ils génèrent, à la demande d'un prompt, des résultats originaux qui reproduisent les structures et les motifs appris, sans posséder de compréhension ou de conscience du sujet.

Chat et Vision :...
L'intelligence artificielle générative est une technologie qui utilise des modèles statistiques pour créer de nouveaux contenus originaux, tels que du texte, des images, de l'audio ou de la vidéo, en s'inspirant de données existantes. Elle repose sur des techniques d'apprentissage profond et génère des résultats en prédisant les éléments les plus probables à partir d'une instruction donnée par l'utilisateur.

Chat et Vision :...
L'intelligence artificielle générative est une branche de l'IA capable de créer de nouveaux contenus originaux, tels que du texte, des images, de l'audio ou du code, en apprenant les motifs et les structures de vastes ensembles de données existantes.

Chat et Vision :...
L'intelligence artificielle générative est une branche de l'intelligence artificielle capable de créer de nouveaux contenus originaux, tels que du texte, des images, du son ou du code. Contrairement aux IA traditionnelles qui se contentent d'analyser ou de classer des données existantes, l'IA générative utilise des modèles d'apprentissage profond pour identifier des motifs dans d'immenses ensembles de données et s'en servir pour produire des résultats inédits qui imitent la structure des données d'origine.

Exemple de réponses multiples par les 4 modèles d'IA à la requête « Définis le concept d'intelligence artificielle générative »

La barre de défilement (scroll) n'apparaît que lorsque le curseur survole la zone de texte. Si vous utilisez quatre modèles simultanément, vous devrez utiliser cette barre de défilement située en bas de vos réponses pour visualiser la réponse du quatrième modèle.

2.4. Raccourcis clavier

Des raccourcis clavier sont disponibles pour faciliter votre navigation.

Raccourcis clavier

Conversation

Nouvelle conversation

Ctrl

Shift

O

Nouvelle conversation temporaire

Ctrl

Shift

'

Supprimer la Conversation

Ctrl

Shift

Delete

Ouvrir la liste des modèles

Ctrl

Shift

M

Activer/Désactiver la dictée

Ctrl

Shift

L

Globale

Recherche

Ctrl

K

Ouvrir les réglages

Ctrl

.

Afficher les raccourcis clavier

Ctrl

/

Afficher/Masquer la barre latérale

Ctrl

Shift

S

Fermer la fenêtre

Esc

Entrée

Activer la zone de saisie

Shift

Esc

Accepter l'autocomplétion

Tab

Empêcher la création de fichier *

Ctrl

Shift

V

Renseigner les variables du prompt

#

Ajouter un prompt personnalisé

/

Joindre un fichier depuis les connaissances

@

Envoyer un message

Générer une paire de messages *

Ctrl

Shift

Enter

Regénérer la réponse

Ctrl

R

Arrêter la génération *

Esc

Modifier le dernier message *

↑

Copier la dernière réponse

Ctrl

Shift

C

Copier le dernier bloc de code

Ctrl

Shift

;

Voix

Toggle Mute *

M

3. Essentiel : Savoir interagir avec le Chatbot

L'intelligence artificielle conversationnelle est un outil statistique puissant, mais non magique : elle ne possède ni conscience ni jugement moral et reproduit les données sur lesquelles elle a été entraînée.

Pour transformer cet outil en un assistant réellement efficace et fiable, la qualité du résultat dépendra avant tout de la **précision de vos instructions** (ou prompting) et de **votre esprit critique**. Nous vous invitons à suivre les étapes de structuration et d'optimisation suivantes pour maximiser la pertinence des réponses et limiter les risques d'erreurs.

3.1. Lancer un échange : l'art du prompting

Écriture de la requête (ou prompt)

Une fois le modèle sélectionné, saisissez votre demande dans la zone de saisie. Vous pouvez utiliser le mode vocal pour dicter votre prompt directement.

Méthode de structuration de prompt

Pour garantir la pertinence et la précision des réponses, nous vous recommandons d'adopter la méthode suivante :



Rôle (Qui est l'IA ?)

Définissez l'identité et l'expertise de l'assistant

 Exemple

Tu es un **expert en ingénierie pédagogique**, spécialiste de la documentation technique et consultant en communication institutionnelle, ou spécialisé dans la conception de cours interactifs.

Tu es un **expert en communication interpersonnelle** et spécialiste en rédaction professionnelle, Agis comme un **assistant administratif**, habitué aux procédures universitaires.

Agis en tant qu'**expert hautement qualifié en analyse et résumé**

Contexte (Où et pour qui ?)

Décrivez la situation, votre profil ou l'objectif visé.

 Exemple

« Je m'adresse à des **étudiants de première année** qui découvrent la sociologie. »

« Ce document sera présenté lors d'une **réunion de coordination avec des partenaires européens** (programme Erasmus+). »

« Cette note doit être comprise par des **collègues de toutes disciplines**, sans jargon technique excessif. »

« Je prépare un **module de formation continue** sur l'intégration de l'IA dans l'enseignement. »

Tâche (Que doit faire l'IA ?)

Formulez la demande de manière explicite et verbale.

 Exemple

Utilisez des verbes d'action clairs : lister, traduire, créer, rédiger, reformuler, écrire, résumer, analyser, défendre, lister, réécrire, comparer, combiner, expliquer, brainstormer...

« **Analyse** les biais potentiels dans un corpus de données historiques utilisé pour un cours de sociologie, en identifiant les représentations stéréotypées. »

« **Rédige** un compte-rendu de réunion structuré à partir des notes brutes fournies, en extrayant les décisions clés et les actions assignées. »

« **Compare** ces deux textes académiques, en identifiant les convergences et les divergences théoriques, puis en proposant une synthèse structurée. »

Format (Sous quelle forme souhaitez-vous la réponse ?)

Précisez les contraintes (longueur, structure de sortie).

 Exemple

« Sous forme de liste à puces », « moins de 500 caractères », « sous forme de tableau », « sous forme de paragraphes structurés »

Un exemple concret : Synthétiser un article de recherche pour un cours

 "Résume-moi ce texte."

C'est une demande trop vague. L'IA ne connaît pas l'objectif, le public cible, ni la longueur attendue. Le résultat sera un résumé standard, potentiellement trop long ou trop superficiel.

1- Ajout d'un rôle

« Tu es un enseignant-chercheur en sociologie du travail » **(Rôle)**.

2- Ajout du Contexte (Précision du profil et de la cible)

« Ce contenu est destiné à un cours de Licence 1 pour des étudiants débutant en sociologie » **(Contexte)**.

En précisant qui on est et à qui s'adresse le contenu, l'IA adapte le niveau de langage. Elle évitera le jargon trop complexe pour des étudiants de première année.

3 - Ajout d'une Tâche

« Analyse ce texte en extrayant uniquement les trois concepts théoriques majeurs et leurs définitions » **(Tâche)**.

On ne demande plus un simple "résumé", mais une action spécifique : l'extraction de concepts et de définitions. L'IA sait maintenant exactement quoi chercher dans le texte.

4- Ajout du Format (Contraintes de sortie)

« Présente le résultat sous forme de liste à puces, en moins de 300 mots, avec un ton pédagogique et synthétique » **(Format)**.

Prompt complet

 « Tu es un enseignant-chercheur en sociologie du travail **(Rôle)**. Ce contenu est destiné à un cours de Licence 1 pour des étudiants débutant en sociologie **(Contexte)**. Analyse ce texte en extrayant uniquement les trois concepts théoriques majeurs et leurs définitions **(Tâche)**. Présente le résultat sous forme de liste à puces, en moins de 300 mots, avec un ton pédagogique et synthétique **(Format)**. »

3.2. Affiner le résultat : Itérer pour optimiser

L'interaction avec une IA générative est un processus itératif. Une première réponse est rarement parfaite ; elle sert de base à un échange constructif où vous affinez progressivement le résultat. Voici les leviers d'action pour guider le modèle :

1- Préserver la continuité de la conversation

Si la réponse est incomplète, trop longue ou hors sujet, évitez de lancer une nouvelle conversation. Utilisez la **conversation existante** pour guider l'IA. Cela permet au modèle de conserver le fil de la discussion et d'intégrer les informations précédemment partagées.

Astuce : Gérer les conversations très longues

La précision du modèle peut diminuer si l'historique de la conversation devient trop long (saturation du contexte). Dans ce cas, demandez à l'IA de générer un **résumé synthétique** de l'échange. Vous pourrez ensuite démarrer une nouvelle conversation en collant ce résumé, assurant ainsi une continuité fluide.

2- Approfondir ou corriger le contenu

Affinez la réponse en demandant des développements spécifiques ou en rectifiant des oublis.

Exemple

« Développe davantage la deuxième partie sur... », « Tu as oublié de mentionner l'aspect législatif, peux-tu l'intégrer ? », « Sois plus critique sur l'argument X. »

3- Ajuster le ton, le style et le format

Adaptez la forme de la réponse à votre besoin final.

Exemple

« Adopte un ton plus formel pour un rapport administratif. », « Reformule ce passage pour un public d'étudiants de licence. », « Présente ces informations sous forme de tableau comparatif à trois colonnes. ».

4- Demander une auto-vérification ou des sources

L'IA peut générer des informations plausibles mais fausses (hallucinations). Poussez-la à **vérifier** ses affirmations.

Exemple

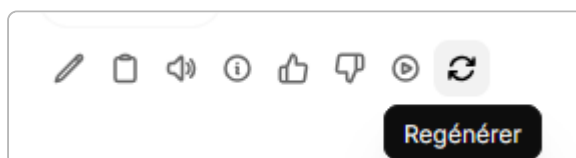
« Vérifie ces informations via la recherche web. », « Cite les sources de tes affirmations. ».

Attention : L'IA propose, vous validez !

La validation humaine reste indispensable. L'utilisateur est responsable de la véracité des faits, surtout lors de reformulations successives.

5- Utiliser la fonction "Régénérer"

Si une réponse est insatisfaisante mais que votre prompt est correct, utilisez le bouton **regénérer**, qui est situé en bas de la réponse générée par l'IA, à gauche des autres icônes d'interaction (comme modifier, copier, Lire à haute voix, les boutons de feedback "pouce vers le haut" et "pouce vers le bas et continuer la réponse").



Cela permet à l'IA de produire une nouvelle version en conservant le même contexte et les mêmes contraintes. Vous pouvez ainsi naviguer entre les différentes versions générées (indiquées par des numéros comme 1/2, 2/2 en bas de réponse) pour sélectionner la plus pertinente.

3.3. Enrichir l'IA : Téléverser des fichiers et documents

L'IA conversationnelle de l'Unistra ne se contente pas de générer du texte à partir de ses connaissances générales ; elle peut également interagir directement **avec vos propres documents**. Cette fonctionnalité, basée sur la technologie **RAG** (*Retrieval-Augmented Generation*), permet à l'IA de consulter les fichiers que vous lui avez transmis avant de répondre, garantissant ainsi une réponse ancrée dans vos données spécifiques.

Pour enrichir l'IA avec vos contenus, deux méthodes sont possibles :

- **Le téléversement rapide** : pour une analyse ponctuelle ou une question unique au sein d'une conversation
- **La Base de Connaissances (BDC)** : pour un usage récurrent, en stockant une collection de documents interrogeables régulièrement (voir la partie Utilisation des bases de connaissances (RAG))

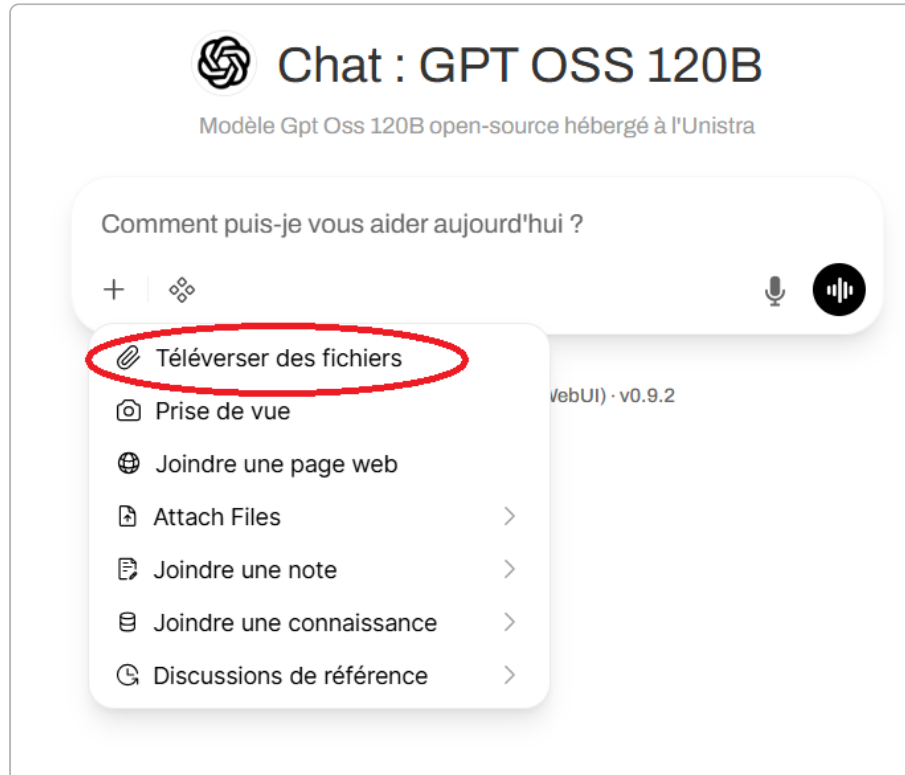
Comment utiliser le téléversement (rapide)

Pour ajouter un document à une conversation, utilisez la fonction de téléversement.

1. Cliquez sur l'icône "+" située à gauche du champ de texte.
2. Un menu contextuel s'affiche alors avec plusieurs options.
3. Pour ajouter vos documents, sélectionnez l'option **"Téléverser des fichiers"** (accompagnée d'une icône de trombone).

Une fois cette action effectuée, vous pourrez sélectionner vos documents sur votre appareil.

Une fois le fichier sélectionné, il apparaît en haut de la conversation, indiquant que l'IA pourra s'y référer pour répondre à votre question.



Capacités et optimisations du téléversement

- **Capacités et limites :** Vous pouvez téléverser jusqu'à 3 fichiers par conversation, avec une taille maximale de 100 Mo par document.
- **Formats supportés :** La plateforme accepte les formats basiques : .txt, .docx, .pdf, .odt, .xlsx, .ppt, .json, .epub, .jpg, .png, et .md (idéal). Pour optimiser la qualité des réponses, privilégiez des documents bien structurés et limitez l'usage d'images ou de tableaux complexes, qui sont moins performants lors de l'analyse par l'IA.
- **Astuce de performance :** Si vous disposez de documents complexes ou des PDF lourds, vous pouvez utiliser l'outil  **Docling** ^[p.81] pour les convertir préalablement en format Markdown (.md), format idéal pour la lecture par l'IA.



Exemples d'usages pédagogiques et administratifs

- **Synthèse de documents :** « *Voici le compte-rendu de la réunion X. Extrais les trois décisions principales et les actions à mener ? Utilise une liste à puces pour la clarté* ».
- **Analyse comparative :** « *Je téléverse le syllabus de l'année dernière et celui de cette année. Liste précisément les éléments qui ont été supprimés, ajoutés ou modifiés* ». (il est recommandé de nommer clairement les fichiers, ex: Syllabus_2023.pdf et Syllabus_2024.pdf).
- **Aide à la rédaction :** « *En te basant sur la note de service ci-jointe, rédige un message court pour informer le personnel des nouvelles modalités de demande de congés.* ».

4. Niveau intermédiaire : Organiser et optimiser ses contenus

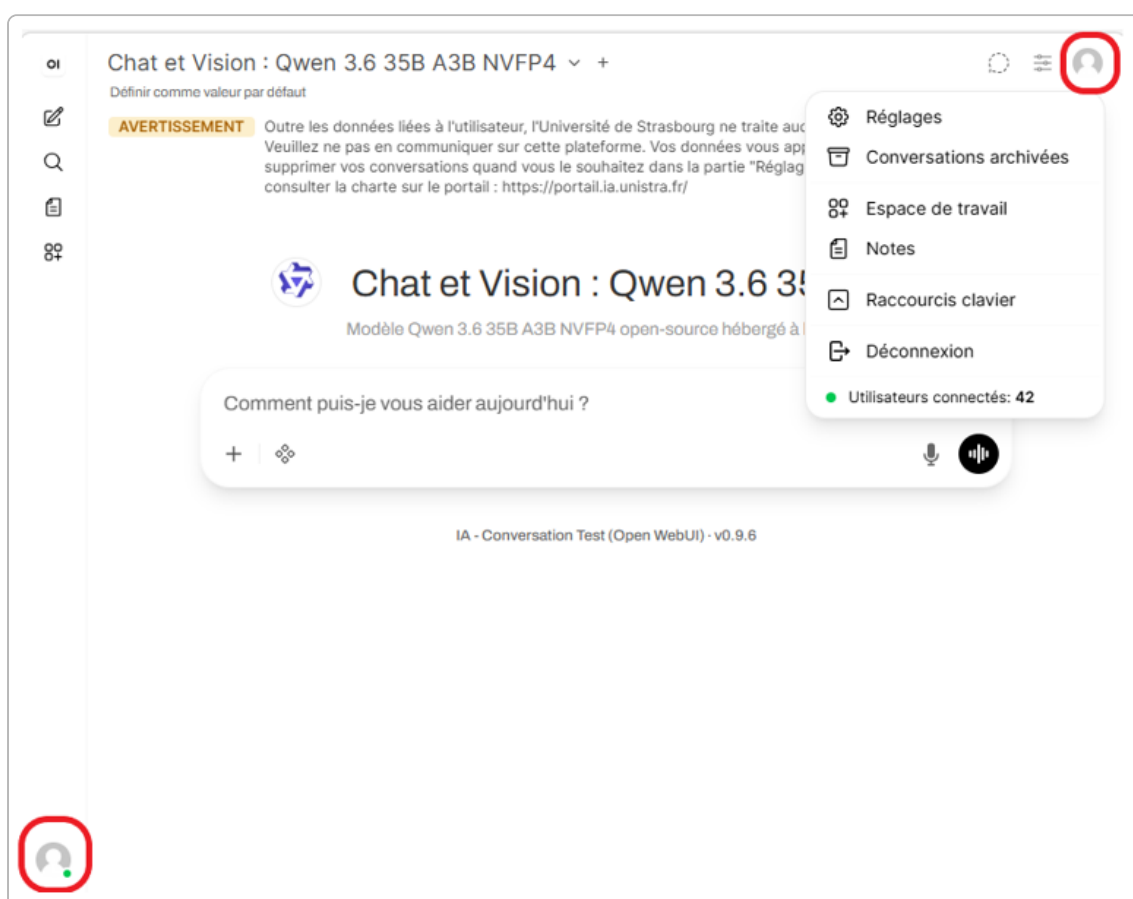
Cette section aborde les fonctionnalités intermédiaires du Chatbot pour organiser et optimiser votre expérience utilisateur, comme la **gestion des conversations, des dossiers et l'éditeur Notes**. Elle explique également comment concevoir des **prompts automatisés** pour gagner en efficacité dans vos interactions avec l'IA.

4.1. Réglages du compte et préférences

Cette section vous permet de configurer votre expérience sur la plateforme Unistra en ajustant les **paramètres généraux**, **l'interface utilisateur**, et le comportement global de l'IA via **un prompt système personnalisé**.

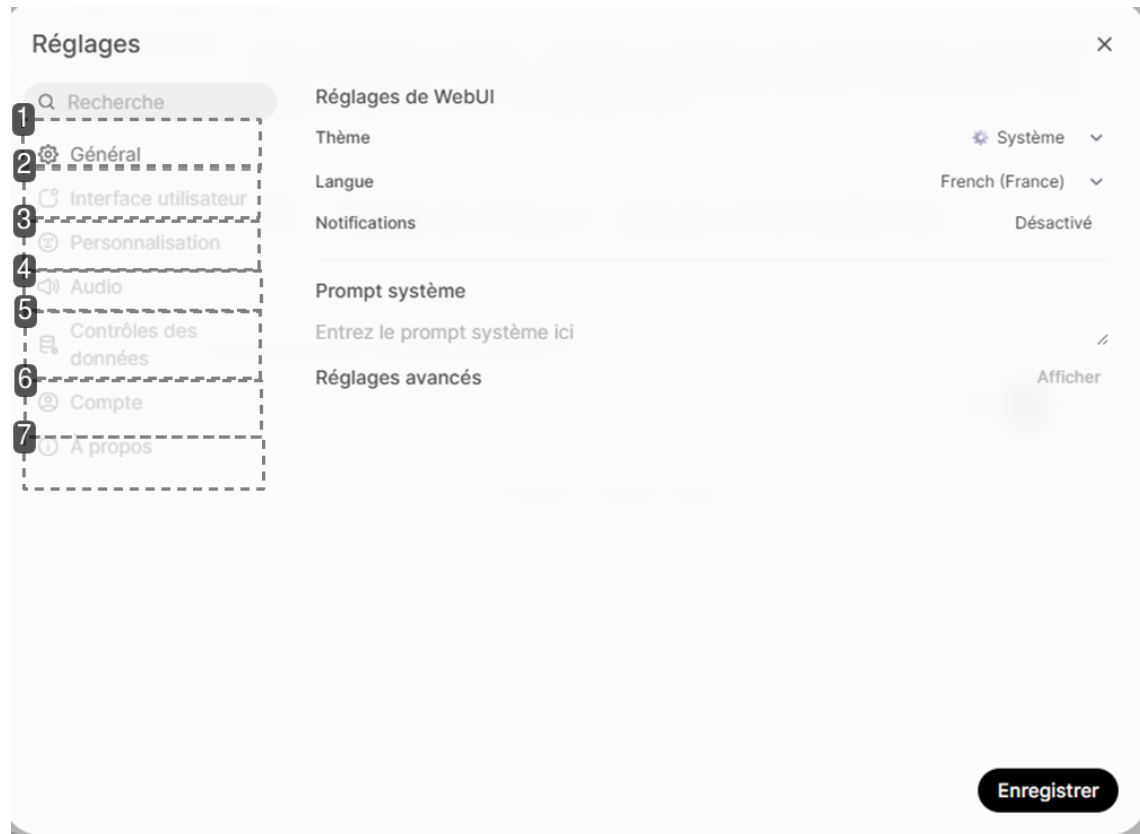
Accès aux paramètres du compte

Vous avez 2 possibilités pour aller dans vos réglages : soit dans votre profil en haut à droite , soit dans votre profil en bas à gauche. Les 2 menus sont identiques :



Structure de la page Paramètres

Le panneau de configuration présenté ci-dessous détaille les paramètres système et personnels, incluant la personnalisation de l'interface et le contrôle des données.



1. Général

Configurez la langue, le thème de l'interface et un prompt système global applicable à l'ensemble de vos conversations. Il est recommandé de ne pas modifier les réglages avancés.

vous pouvez modifier les options suivantes :

- **Thème** : Choisissez entre le mode clair ou le mode sombre pour adapter l'affichage à vos préférences visuelles et réduire la fatigue oculaire.
- **Langue** : Sélectionnez la langue d'interface de l'application ainsi que la langue par défaut des réponses de l'IA (par exemple, Français).
- **Notifications** : Activez ou désactivez les alertes pour être informé de la fin des traitements longs ou des mises à jour de statut.
- **Prompt système** : Définissez un comportement personnalisé pour l'IA en spécifiant son rôle, son ton, ses contraintes de réponse, et son style de présentation. Cela permet d'obtenir des réponses plus pertinentes, structurées, et adaptées à vos besoins académiques ou administratifs.

2. Interface utilisateur

Personnalisez l'apparence visuelle, les préférences d'affichage et les notifications afin d'optimiser votre confort d'utilisation.

L'interface est divisées en deux catégories principales :

- **l'interface utilisateur (UI)** . Vous pouvez ajuster le contraste, le son des notification, les modalités de présentation de l'interlocuteur
- les paramètres de **conversation** : vous pouvez personnaliser le comportement des échanges (gestion de la file d'attente des requêtes, style des réponses, formatage des textes générés, image d'arrière-plan)

3. Personnalisation

Activez la fonctionnalité **Mémoire** pour conserver des faits et des préférences entre les sessions. Cette mémoire peut être gérée manuellement via le bouton Gérer.

4. Audio

Réglez les paramètres de transcription (voix en texte) et de synthèse vocale (texte en voix), en choisissant le moteur, la langue et l'option de lecture automatique.

5. Contrôle des données

Gérez vos données personnelles à travers l'importation, l'exportation, l'archivage ou la suppression des conversations et des fichiers téléversés.

6. Compte

Centralisez vos informations personnelles, votre identité numérique, vos données biographiques, vos notifications et la gestion de vos clés API.

Vous pouvez y ajouter une photo de profil ainsi qu'une biographie.

Cette section est essentielle pour récupérer votre clé d'API (voir la section [🔗 Intégration et API](#) ^[p.45]).

7. À propos

Affiche la version d'OpenWebUI (nom de l'interface que vous voyez permettant d'interagir avec les modèles)

Configuration d'un Prompt Système

Le paramètre prompt système dans la section générale permet de définir **un contexte global** (un rôle et des règles de comportement) qui s'appliquera **à toutes vos conversations**. Cette consigne initiale agit comme une directive maîtresse qui guide l'assistant dans la forme et le fond de ses réponses, en lui imposant un ton cohérent et des contraintes opérationnelles claires.

Ce champ est optionnel ; il reste **optionnel**, et vous pouvez également choisir de définir ce contexte uniquement au sein de dossiers spécifiques.

Si vous prévoyez d'utiliser l'IA pour des rôles très variés, il est préférable de ne pas définir de rôle spécifique dans ce champ général, mais de privilégier des prompts ciblés ou des modèles dédiés. Ce paramètre est particulièrement utile pour supprimer les comportements récurrents indésirables, tels que les flatteries systématiques. Dans tous les cas, veillez à **ne pas être trop restrictif** afin de préserver la flexibilité et la créativité de l'outil.

La configuration du prompt système repose sur deux axes complémentaires :

- **Définir une identité et un périmètre** : Spécifiez le rôle de l'IA, son contexte d'intervention et les limites de ses compétences pour spécialiser l'outil. Désignez-vous comme « Utilisateur » (par ex. « L'utilisateur apprend l'espagnol »)



Exemple

L'utilisateur est un enseignant-chercheur en psychologie à l'Université de Strasbourg. Tu es un expert en ingénierie pédagogique et tu m'accompagnes dans la conception de mes supports de cours pour des étudiants en Licence. Ton objectif est de m'aider à réfléchir, reformuler et structurer mes contenus éducatifs pour les rendre plus accessibles et engageants, tout en

garantissant la rigueur scientifique indispensable à l'enseignement supérieur. Adopte un ton professionnel, bienveillant et institutionnel.

- **Maîtriser le comportement et le format** : Indiquez le ton à adopter, les contraintes de réponse et les règles de présentation pour garantir des résultats exploitables. Si vous n'avez pas de rôle précis à définir mais que vous utilisez l'IA pour des activités variées, vous pouvez utiliser le prompt système uniquement **comme un "filtre" de comportement**. Cela permet d'**affiner la manière dont l'IA interagit avec vous** et d'éliminer certains automatismes contre-productifs. Ce paramétrage est particulièrement utile pour **supprimer les tics de langage**, comme les questions de relance systématiques en fin de réponse, contrer la tendance de l'IA à être systématiquement d'accord avec l'utilisateur (biais de complaisance ou sycophantie), afin d'obtenir une analyse plus critique.

Exemple

Concision : *Ne pose pas de questions de relance à la fin de tes réponses. Arrête-toi dès que la réponse est complète.*

Neutralité : *Adopte un ton objectif. Évite toute flatterie gratuite (la sycophantie de l'IA), et les formules de congratulation ("C'est une excellente idée !"). Privilégie l'analyse factuelle et la pertinence technique.*

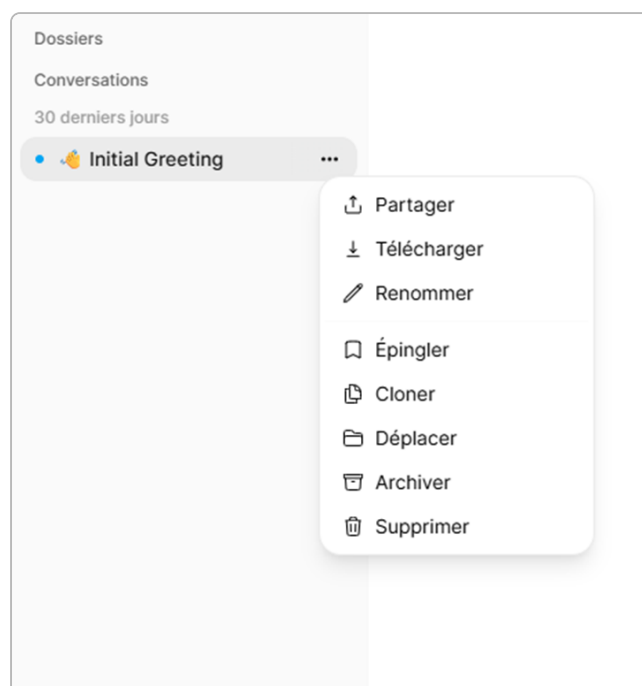
Typographie : *N'utilise jamais de tirets cadratins (—).*

Esprit critique : *Ne valide pas systématiquement mes idées. Adopte une posture de partenaire intellectuel exigeant. Je souhaite avoir un feedback constructif et parfois divergent pour affiner ma réflexion, pas d'être conforté.*

4.2. Gestion des conversations et dossiers

Les conversations

Toutes vos interactions avec l'IA sont listées sous l'onglet **Conversations**, classées par date. Pour chaque échange, vous disposez d'options de gestion en cliquant sur les **trois points** (...) à côté du titre de la conversation qui vous permettent de :



Supprimer, Archiver ou Cloner : Pour nettoyer votre historique ou réutiliser une conversation existante.

Épingler : Pour maintenir une conversation importante en haut de votre liste.

Réorganiser : Pour faire glisser vos conversations dans l'ordre de votre choix.

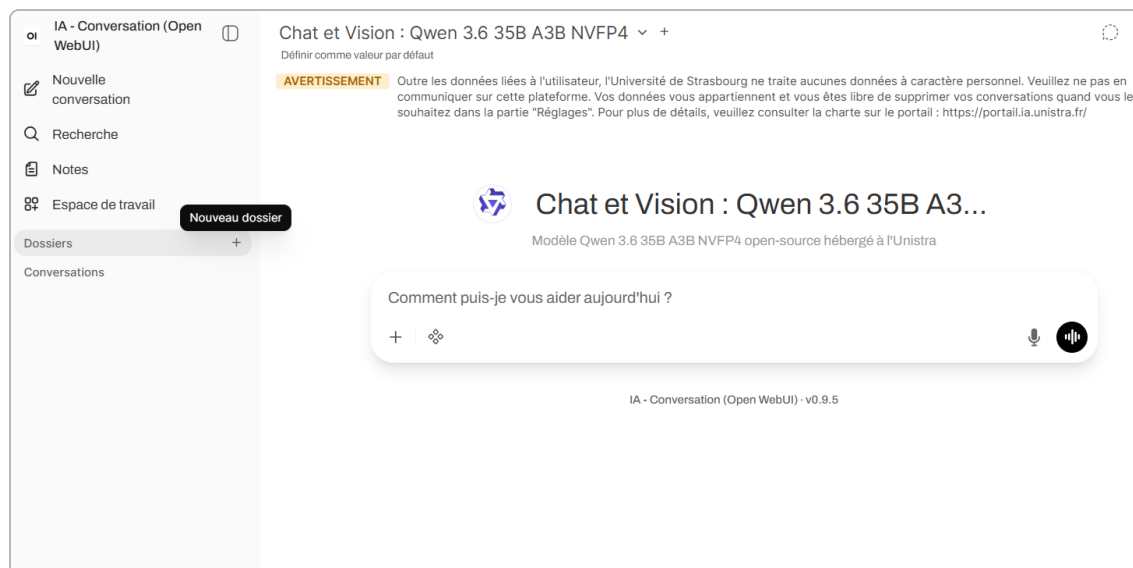
Partager un lien : Vous pouvez générer un lien vers une conversation spécifique pour le partager avec un collègue disposant d'un accès au portail IA. Aucun besoin de saisir le nom dans la liste d'accès: copiez simplement le lien et transmettez-le.

Télécharger une conversation : Vous pouvez télécharger l'intégralité d'un échange au format :

- **PDF :** Pour conserver la mise en forme visuelle et imprimer le document.
- **TXT :** Un fichier texte brut incluant la syntaxe Markdown (balisage simple pour les titres et le gras).
- **JSON :** Format technique privilégié par les développeurs et les systèmes informatiques pour l'import/export de données complexes vers d'autres logiciels ou applications.

Les dossiers

L'outil **Dossiers** vous permet de structurer votre espace de travail de manière thématique, à la manière d'un classeur numérique.



Pour créer un dossier, cliquez sur l'icône  à côté de la rubrique "Dossiers" et attribuez-lui un nom explicite. En accédant aux options (trois points ), vous pouvez modifier ses caractéristiques (nom, prompt système), supprimer le dossier ou exporter l'intégralité de son contenu (au format JSON). Lorsque vous ouvrez un dossier, une nouvelle conversation s'initie à l'intérieur de celui-ci. Deux modes d'utilisation sont disponibles :

1. **Rangement thématique**

Classer des conversations passées ou actives par projet ou matière, comme dans un ordinateur. Pour cela, faites glisser une conversation existante dans le dossier. Le dossier agit alors comme un simple conteneur de rangement.

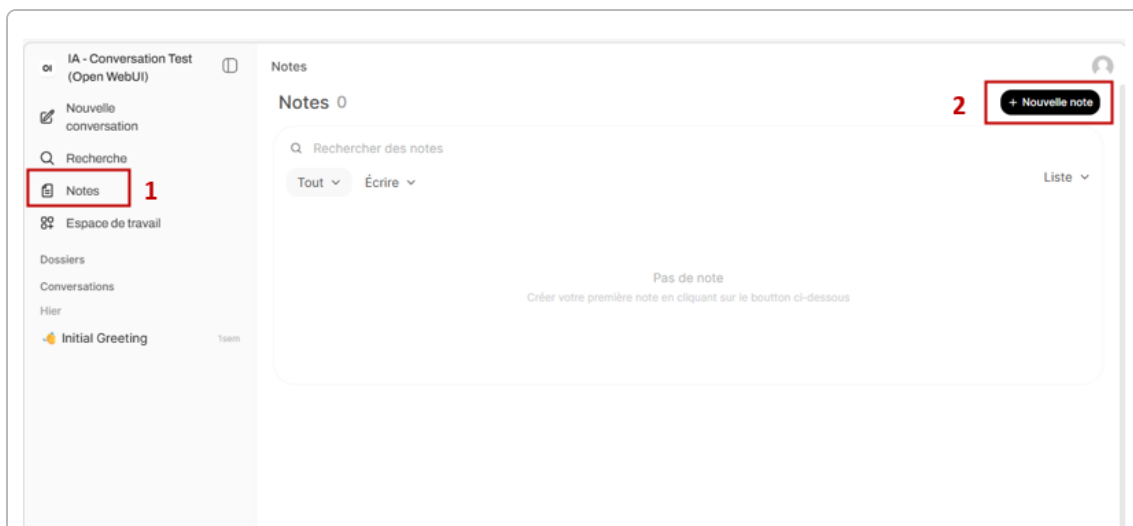
2. **Création de contexte (Dossier intelligent)**

Créer un dossier intelligent pour un sujet précis (ex: "Tuteur de maths", "Assistant administratif"). Lors de la création du dossier, vous pouvez définir un **Prompt Système** (une consigne donnée à l'IA) et y associer des **Bases de Connaissances** ou des fichiers spécifiques. (voir la rubrique [📁 dossier intelligent](#))

Lorsque vous cliquez sur ce dossier, une **nouvelle conversation** se lance qui hérite automatiquement du contexte défini. L'IA adoptera alors le rôle et les contraintes spécifiés pour toute la durée de la discussion.

4.3. Utilisation de l'éditeur de Notes

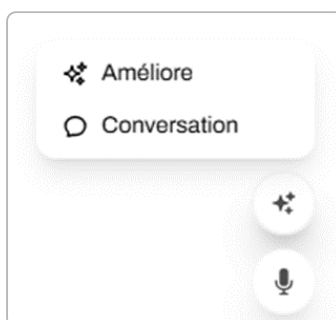
La rubrique **Notes** est conçue pour prendre des notes brutes, les structurer et en améliorer la forme. Elle permet de corriger l'orthographe, de hiérarchiser le contenu (titres, sous-titres ...), de mettre en valeur les idées clés, et de rédiger des phrases complètes à partir de notes éparpillées.



Création et saisie de notes.

- Cliquez sur le bouton **+ Nouvelle note** en haut à droite pour ouvrir un espace de travail vierge.
- Copiez-collez vos notes brutes ou rédigez-les directement dans l'éditeur de texte.

Mise en forme et amélioration par l'IA



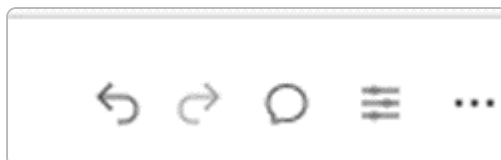
Pour transformer vos notes brutes en texte structuré :

- Cliquez sur l'icône **IA** (en bas à droite) puis sur **Améliore**.
- L'IA restructure automatiquement le texte (orthographe, grammaire, mise en page).
- Si le résultat n'est pas satisfaisant, utilisez la flèche retour pour revenir à la version originale avant de relancer l'optimisation. *Note : Une fois la note fermée, la version originale n'est plus accessible avec le bouton.*

Interaction et modification contextuelle

- Cliquez sur l'icône **Conversation** pour échanger avec l'IA sur le contenu de votre note, par exemple, pour lui demander de résumer, d'extraire des données, de critiquer ou de réécrire des sections spécifiques de votre note.
- Posez votre question (ex. : *Quels sont les 3 points clés de cette réunion ?* ou *Réécris ce compte-rendu en paragraphes*).
- Le texte généré ou modifié apparaît dans la zone de conversation. Cliquez sur **Insérer** pour remplacer le texte initial de votre note par cette nouvelle version.

Historique et contrôle



- La flèche retour permet de consulter l'historique des modifications effectuées lors de la session active.
- Le menu en haut à droite permet de visualiser et de modifier le modèle de langage utilisé

+ bouton **Enregistrement**

Un bouton **Enregistrement** est disponible pour capturer l'audio directement ou téléverser un fichier. Cependant, pour les réunions ou l'audio, nous recommandons d'utiliser l'outil dédié : **Speakr**

4.4. Conception de prompts automatisés

L'onglet **Prompts** permet de sauvegarder des instructions fréquentes.

Ces prompts sont ensuite activés dans la conversation via une **commande slash « / »** (par exemple, taper /resume dans la zone de saisie).

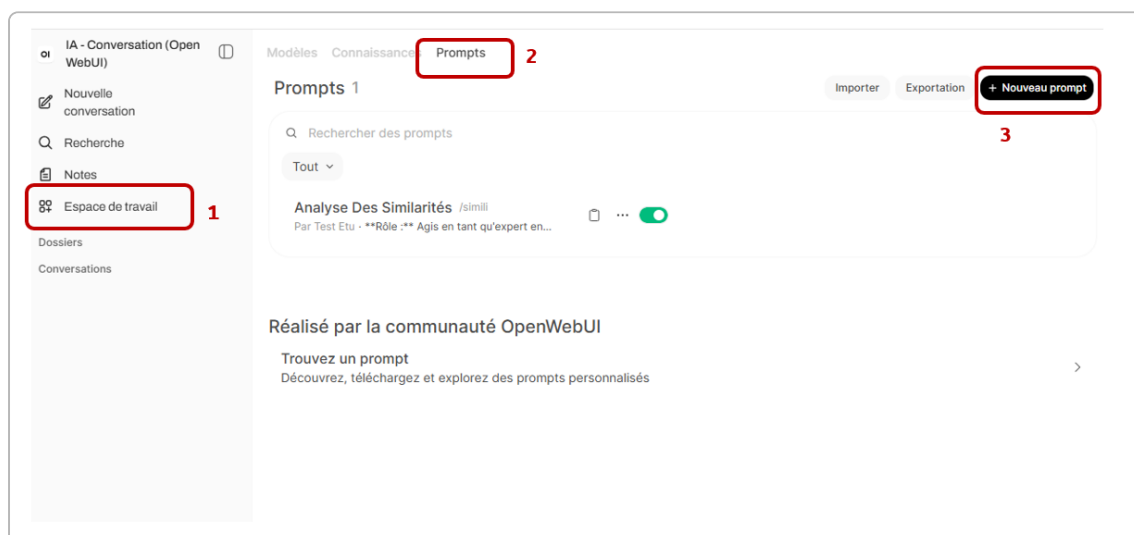
Cette méthode réduit la nécessité de retaper des modèles récurrents (comptes rendus, analyses, relectures, etc.) et permet aux utilisateurs non techniques d'utiliser des prompts complexes

Création et utilisation d'un prompt personnalisé

Pour créer et utiliser un prompt personnalisé, suivez ces étapes :

Procédure

1. **Accédez à l'onglet Prompts** dans votre espace de travail et cliquez sur « **Nouveau prompt** ».



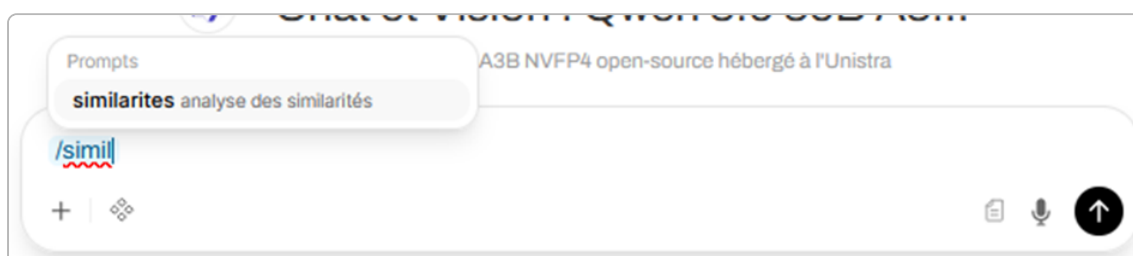
2. Remplissez les champs suivants :

- **Nom** : Donnez un titre descriptif à votre prompt (ex: Compte rendu de réunion). Ce nom vous aide à vous y retrouver dans votre bibliothèque.
- **Commande** : Choisissez le mot-clé court qui sera utilisé dans le chat, précédé d'un slash (ex: CR). C'est ce que vous taperez pour lancer le prompt précédé d'un slash (/CR).
- **Description** : Ajoutez une brève explication de l'utilité de ce prompt (environ 50 mots) pour clarifier son usage.
- **Contenu du prompt** : Saisissez votre prompt (voir la partie  lancer une conversation : structurer sa demande . Vous pouvez y inclure des **variables** (c'est-à-dire des champs à remplir par l'utilisateur) pour créer un formulaire interactif (voir  Le Prompt automatisé à variables).

3. Enregistrez votre prompt.

4. Vous pouvez désormais **l'utiliser** à tout moment en tapant simplement suivi de son nom.

Vous pouvez par exemple créer des prompts automatisés pour des comptes rendus, des analyses de similarité, la relecture d'emails ou un correcteur d'anglais.



exemple d'appel de prompt automatisé

5. Niveau avancé : Utilisation des bases de connaissances (RAG)

Cette section explique le fonctionnement de la **technologie RAG** et présente les **bases de connaissances (BDC)** comme outil central pour interroger vos documents privés. Elle compare cette méthode de gestion des données avec le téléversement rapide de fichiers, puis détaille les étapes de création d'une BDC pour structurer efficacement vos ressources.

5.1. Principes du RAG et bases de connaissances (BDC)

RAG

La Base de Connaissances (BDC) repose sur un procédé technique appelé le **RAG** (*Retrieval-Augmented Generation*, ou Génération Augmentée par Récupération). Ce système permet à l'intelligence artificielle d'enrichir ses réponses en consultant des documents spécifiques avant de générer son texte. Contrairement à une IA qui s'appuie uniquement sur ses connaissances générales, **une IA utilisant le RAG peut exploiter vos propres documents** (cours, thèses, règlements) **pour fournir des réponses précises et contextualisées**.

L'utilisation d'une base de connaissances présente des avantages majeurs par rapport au simple téléversement d'un fichier dans une conversation :

- **Optimisation de la mémoire** : Le RAG permet à l'IA de répondre en s'appuyant uniquement sur des extraits pertinents (chunks) de vos documents, sans pour autant exploiter l'intégralité de ceux-ci. Cette sélection ciblée évite de saturer la fenêtre de contexte, c'est-à-dire le nombre maximum de tokens que l'IA peut traiter.
- **Gestion de gros volumes** : Le système interroge une base de connaissances contenant des dizaines de documents sans ralentir l'IA, car seule une infime partie de l'information est traitée à chaque requête.
- **Précision accrue** : l'IA s'appuie sur des faits documentés plutôt que sur des probabilités statistiques, ce qui réduit les risques d'hallucinations en s'ancrant dans vos propres sources vérifiées. Elle permet ainsi l'intégration de documents officiels, de fichiers de travail de l'université ou d'informations récentes, tel qu'un nouveau règlement.

5.2. Comparaison entre BDC et Téléversement

Il est essentiel de distinguer ces deux méthodes d'apport d'informations, car elles diffèrent par leur fonctionnement technique et leur finalité d'usage.

Le Téléversement direct : un usage temporaire et ponctuel

Le téléversement consiste à envoyer un fichier directement dans une fenêtre de discussion.

- **Contraintes et portée** : Cette méthode est limitée à **3 fichiers et 100 Mo par conversation**. L'information est liée à une session spécifique et reste disponible durant 3 mois avant d'être archivée.

- **Usage idéal** : Analyser un document précis et court dans son intégralité, tel qu'un contrat, un rapport spécifique ou un article scientifique.

La Base de Connaissances (BDC) : un usage permanent et structuré

La Base de Connaissances repose sur le procédé du RAG. Vos documents sont indexés dans une mémoire externe et l'IA n'en extrait que les fragments (*chunks*) les plus pertinents.

- **Portée et confidentialité** : Le contenu est permanent et accessible dans toutes vos conversations futures. La BDC peut être configurée en mode privé (accessible seulement par vous) ou partagée avec des collaborateurs.
- **Optimisation technique** : La taille de la base n'impacte pas la mémoire de l'IA. Quel que soit le volume stocké (10 ou 100 documents), le système n'injecte que 3 à 10 fragments (*chunks*) sémantiquement cohérents (entre 2 000 et 5 000 tokens), préservant ainsi la fluidité et la précision des réponses.
- **Usage idéal** : Interroger un volume important de données, créer un assistant expert sur un corpus de cours ou exploiter une documentation interne volumineuse, comme celle de l'UNISTRA.

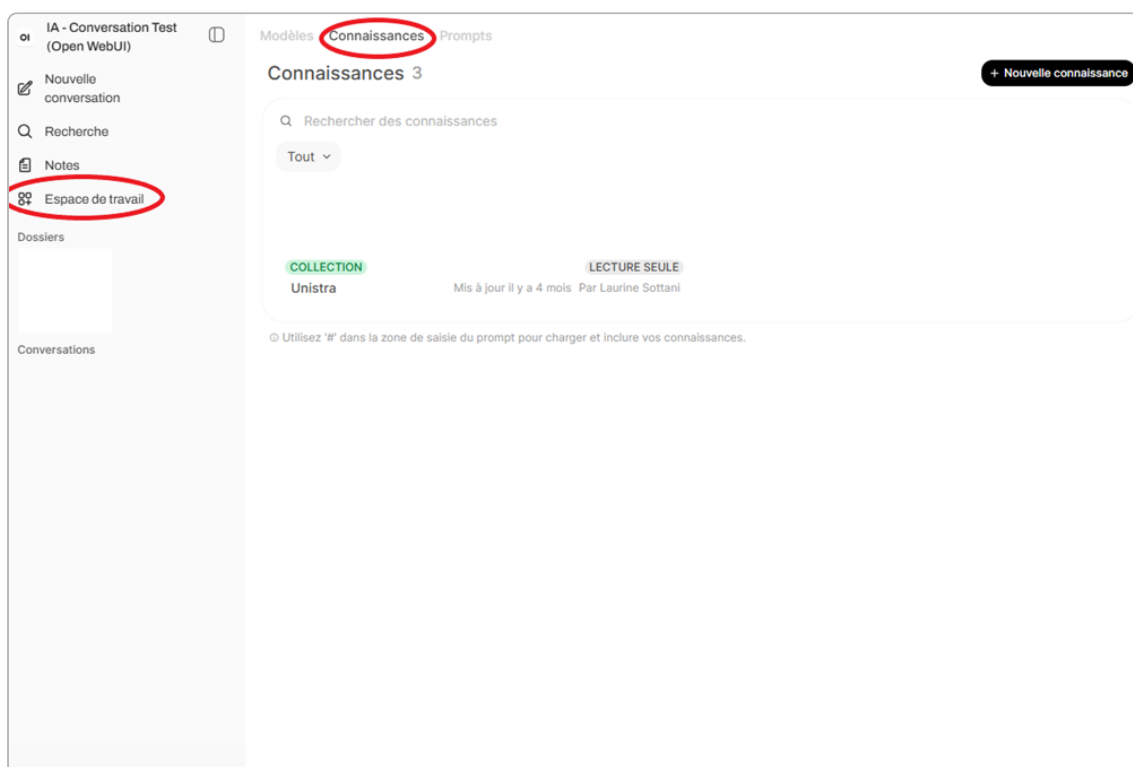
5.3. Création d'une BDC

Pour créer une BDC

Pour créer une BDC, suivez les étapes ci-dessous :

Procédure

1. **Accédez à l'onglet Connaissance** dans votre espace de travail et sélectionnez « **Nouvelle connaissance** ».



2. Remplissez les champs suivants

Nom de la BDC : descriptif pour identifier facilement votre base (par exemple : "Cours Droit 2026" ou "Règlement intérieur").

Description : courte précisant le contenu ou l'objectif de cette collection de documents.

Accès : Par défaut, seuls les administrateurs peuvent rendre une BDC publique. Vous pouvez également partager l'accès avec d'autres utilisateurs de l'Unistra connectés à l'IA. Choisissez leur niveau d'accès : Lecture ou Écriture. Pour modifier ce niveau, cliquez simplement sur le bouton correspondant (par exemple, passer de « Lire » à « Écrire »). Une fois partagée, la BDC apparaîtra automatiquement dans la liste des bases de connaissances de la personne concernée.

3. Appuyez sur « créer une connaissance » pour valider.

4. Ajoutez du contenu en appuyant sur le bouton « + » en haut à droite.

5. Enregistrez votre prompt pour finaliser la configuration

Vous pouvez désormais appeler votre BDC dans vos modèles, dossiers ou conversations.

Pour supprimer une BDC

Procédure

1. Accédez à l'onglet **Connaissance** dans votre espace de travail

2. Localisez la BDC à supprimer et cliquez sur « ... » à côté du nom de la BDC

3. Sélectionnez « Supprimer » pour retirer la BDC.

6. Niveau expert : 3 solutions pour automatiser vos tâches

Le Chatbot offre trois niveaux d'automatisation, adaptés à la complexité de vos besoins et à la portée de votre partage :

- **Le Prompt automatisé** : Idéal pour **une action rapide et ponctuelle**. Il suffit de saisir une commande courte (via une barre de slash /) pour déclencher une instruction complexe sans avoir à la retaper. Vous pouvez enrichir cette commande avec des **variables** (champs de formulaire interactifs) pour personnaliser la réponse.
- **Le Dossier intelligent** : Conçu pour organiser **des travaux récurrents nécessitant un contexte spécifique**. Vous pouvez y attacher des fichiers ou une Base de Connaissances (BDC) pour que l'IA s'appuie sur vos documents personnels lors de chaque interaction. Cette solution est dédiée à un usage individuel.
- **L'Assistant IA personnalisé (modèle)** : La solution la plus aboutie pour **une automatisation pérenne et partageable**. Vous pouvez configurer un assistant avec un rôle, un ton et un modèle spécifiques, lui ajouter des fichiers ou une Base de Connaissances (BDC), puis le mettre à disposition de vos étudiants ou de votre équipe pour une utilisation collective

6.1. Le Prompt automatisé avec variables

Le prompt automatisé permet de gagner du temps sur les tâches simples via une commande rapide (ex. /resume) saisie dans la zone de conversation. **Il est particulièrement adapté pour des actions rapides et ponctuelles**, sans nécessiter une configuration complexe.

Le prompt automatisé est présenté dans la partie  Conception de prompts automatisés ^[p.29]. Pour aller plus loin, vous pouvez enrichir ce prompt en y intégrant des variables pour le rendre dynamique, ou en utilisant le format Markdown pour structurer votre texte.

Ajout de variables à un prompt automatisé

Les variables permettent de créer des **formulaires interactifs**. Lorsqu'un utilisateur exécute un prompt contenant des variables, des champs de saisie apparaissent automatiquement. L'utilisateur doit remplir ces informations avant que l'IA ne génère la réponse. Cela permet de réutiliser un même modèle de prompt pour différentes situations (par exemple, un compte-rendu de réunion dont les participants et la date changent à chaque fois).

Pour créer une variable, vous devez utiliser la syntaxe suivante :

```
{{nom_variable | type:option}}
```

Le nom doit être en minuscules, sans espaces. Utilisez un tiret bas _ pour séparer les mots.

Deux catégories de variables

Il existe deux catégories de variables : **Variables système et personnalisées**

Les variables système : Elles sont préremplies automatiquement par l'IA. Elles sont utiles pour inclure des informations contextuelles sans que l'utilisateur ait besoin de les saisir. (ex. {{CURRENT_DATE}} pour la date du jour, {{USER_NAME}} pour votre nom).

L'environnement de travail met à disposition plusieurs variables automatiques permettant de personnaliser les réponses de l'IA ou d'intégrer des éléments contextuels dynamiques dans vos prompts. Voici leur fonctionnement :

- **{{CURRENT_DATE}}** : Renvoie la date du jour.
- **{{CURRENT_DATETIME}}** : Renvoie la date et l'heure actuelles combinées.
- **{{CURRENT_TIME}}** : Renvoie l'heure actuelle.
- **{{CURRENT_WEEKDAY}}** : Renvoie le jour de la semaine en cours.
- **{{USER_NAME}}** : Renvoie votre nom d'affichage tel qu'il est configuré sur votre compte.

Ces variables peuvent être insérées directement dans vos prompts ou dans la configuration de vos assistants personnalisés pour adapter le ton ou le contenu des réponses en temps réel.

Les variables personnalisées (manuelles) : Ce sont celles que vous créez pour demander des informations spécifiques à l'utilisateur. Un champ de saisie dédié sera affiché. Par exemple : `{{titre_reunion}}` créera un champ texte obligatoire intitulé "titre_reunion".

Types de champs disponibles pour les variables système et personnalisées

Lors de la construction de vos prompts variabilisés, vous pouvez choisir le type de champ qui correspond le mieux à la nature des données attendues. Voici les principaux types disponibles et leur syntaxe d'utilisation :

- **text** : Champ de texte unique (par défaut), idéal pour les noms ou informations courtes. *Exemple* : `{{name | text:placeholder="Enter name":required}}`
- **textarea** : Champ de texte multiligne, adapté pour les descriptions ou longs textes. *Exemple* : `{{description | textarea:required}}`
- **select** : Menu déroulant permettant de choisir une option parmi une liste prédéfinie. *Exemple* : `{{priority | select:options=["High","Medium","Low"]:required}}`
- **number** : Champ de saisie numérique, souvent utilisé avec des contraintes de minimum et de maximum. *Exemple* : `{{count | number:min=1:max=100:default=5}}`
- **checkbox** : Case à cocher pour activer ou désactiver une option spécifique. *Exemple* : `{{include_details | checkbox:label="Include analysis"}}`
- **date** : Sélecteur de date pour choisir un jour précis. *Exemple* : `{{start_date | date:required}}`
- **datetime-local** : Sélecteur combinant la date et l'heure pour des rendez-vous ou échéances précises. *Exemple* : `{{appointment | datetime-local}}`

Ces types de champs permettent de structurer les interactions avec l'IA en garantissant que les informations saisies sont cohérentes et exploitables par le modèle. Pour optimiser la création de vos formulaires, voici les règles de fonctionnement à retenir :

- **Saisie optionnelle** : Par défaut, tous les champs sont facultatifs. L'utilisateur peut ainsi lancer l'exécution du prompt même s'il ne remplit pas l'intégralité des variables.
- **Champs obligatoires** : Pour exiger la saisie d'une information avant le lancement du prompt, ajoutez le flag `:required` à la fin de la définition du champ (par exemple : `{{titre | text:required}}`).
- **Guidage utilisateur** : Vous pouvez insérer un texte indicatif pour aider l'utilisateur à comprendre ce qu'il doit saisir en utilisant le paramètre `:placeholder="votre texte"`.

Structurer son texte avec le Markdown

Pour optimiser la qualité des réponses de l'IA, il est recommandé de structurer vos prompts en utilisant la syntaxe Markdown. L'IA apprécie particulièrement les contenus organisés et hiérarchisés, car cette structure lui permet de mieux interpréter vos instructions et de produire un résultat plus clair et cohérent.

Voici les éléments de syntaxe essentiels pour guider l'IA :

- # Titre 1 à ##### Titre 6 : Pour hiérarchiser les titres.
- *_italique_* : Pour mettre en valeur du texte en italique.
- ****gras**** : Pour mettre en gras les mots ou phrases importants.
- *****gras et italique***** : Pour un effet combiné.
- - liste : Pour créer des listes à puces.
- ::mise en valeur:: : Pour surligner ou mettre en avant une information spécifique (selon la configuration du modèle cible).



Exemple de prompt pour un compte-rendu de réunion incluant date, participants, ordre du jour, décisions et actions.

CR de Réunion

****Date:**** {{date_reunion | date:required}}

****Heure:**** {{heure_reunion | heure:required}}

****Objet de la réunion:**** {{titre | text:placeholder="ex., Réunion d'équipe hebdomadaire":required}}

****Participants:**** {{participants | text:placeholder="Noms séparés par des virgules":required}}

Ordre du jour / Points clés

{{odj | textarea:placeholder="Coller l'odj et/ou les points clés abordés":required}}

Décisions

{{decisions | textarea:placeholder="Décisions clés et conclusions (optionnel)"}}

Actions

{{actions | textarea:placeholder="Lister les actions, deadlines, et assigné(s) (optionnel)"}}

Prochaine réunion

****Date:**** {{prochain_meeting | date}}

****Sujets:**** {{prochains_sujets | text:placeholder="Sujets à aborder la prochaine fois."}}

Agis en tant que rédacteur professionnel spécialisé. Dans un contexte universitaire, mets en forme les informations ci-dessus pour produire un compte rendu de réunion clair et professionnel.

L'activation de la commande ☐/CR génère le formulaire interactif ci-dessous, permettant de saisir les informations nécessaires.

Variables d'entrée ✕

date_reunion *requis
jj / mm / aaaa

heure_reunion *requis

titre *requis
ex., Réunion d'équipe hebdomadaire

participants *requis
Noms séparés par des virgules

odj *requis
Coller l'odj et/ou les points clés abordés

decisions
Décisions clés et conclusions (optionnel)

actions
Lister les actions, deadlines, et assigné(s) (optionnel)

prochain_meeting
jj / mm / aaaa

prochains_sujets
Sujets à aborder la prochaine fois.

Annuler Enregistrer

6.2. Le Dossier intelligent

Le **Dossier intelligent** permet de structurer des conversations autour d'un contexte personnel et récurrent. En définissant un **Prompt Système** et en y attachant des ressources spécifiques, vous configurez un profil et des limites précises pour l'IA. Cette approche évite de devoir répéter le contexte à chaque échange : toutes les nouvelles conversations ouvertes dans ce dossier bénéficient automatiquement de cette configuration personnalisée.

Lorsque vous créez ou modifiez un dossier, vous pouvez définir trois éléments clés :

- **Nom du dossier** : Pour l'identifier clairement.
- **Prompt Système** : C'est la consigne principale. Elle définit le rôle de l'IA, ses règles de comportement et le ton à adopter. Ce prompt s'applique automatiquement à toutes les

nouvelles conversations créées dans ce dossier. Inspirez-vous de la méthode RCTF pour le rédiger.

- **Sélectionnez une connaissances (BDC) ou Téléverser des fichiers :** Vous pouvez attacher une Base de Connaissances ou des fichiers spécifiques. L'IA s'appuiera sur ces documents pour répondre aux questions.

Pour initier un échange, **cliquez sur le dossier**, une nouvelle conversation se crée automatiquement, intégrant le contexte défini et les ressources du dossier.

Les avantages de cette approche sont

- un **Gain de temps** : Vous n'avez pas besoin de réécrire le contexte ou de re-télécharger les fichiers à chaque nouvelle conversation.
- la **Cohérence** : Toutes les interactions avec l'IA pour ce projet respectent les mêmes règles et références documentaires.



Exemple concret : Un dossier pour un cours

Imaginons que vous prépariez un module. Vous créez un dossier nommé **"Cours Intro IA"**.

- **Fichiers attachés** : Téléversez votre plan de cours, vos diapositives et les lectures obligatoires.
- **Prompt Système** : Rédigez une consigne comme celle-ci :

****Rôle****

Tu es un enseignant-chercheur spécialisé en Intelligence Artificielle et en Informatique à l'Université de Strasbourg.

****Contexte****

Tu accompagnes des étudiants de première année de Licence Informatique dans leur apprentissage des notions fondamentales de la discipline.

****Contraintes****

Réponds exclusivement à partir des documents, cours, supports pédagogiques et ressources présents dans la base de connaissances qui t'est fournie.

Si la réponse à une question ne figure pas explicitement dans les documents disponibles, indique clairement : "Cette information n'est pas présente dans les documents mis à ma disposition."

Adapte systématiquement le niveau d'explication à des étudiants débutants de Licence 1.

Propose des exemples concrets lorsque les documents en contiennent.

****Format****

Structure les réponses avec des titres et sous-titres explicites.

Utilise des listes à puces ou des listes numérotées pour faciliter la lecture.

Mets en évidence les notions importantes avec du texte en gras.

Termine chaque réponse par un court résumé des points essentiels à retenir.

****Ton et style****

Adopte un ton bienveillant, patient, encourageant et académique.

Utilise un langage clair, précis et accessible à un public de première année.

Valorise les efforts de l'étudiant et encourage la réflexion plutôt que la simple mémorisation.

Une fois cette configuration effectuée, il vous suffit d'ouvrir le dossier concerné pour initier une nouvelle conversation. L'interface vous permet alors de formuler directement votre demande spécifique. Le contexte académique, les contraintes et les ressources pédagogiques étant déjà intégrés dans le **prompt système**, vous n'avez plus qu'à définir la tâche précise.

6.3. L'Assistant IA personnalisé

Les assistants IA personnalisés permettent de créer des outils sur mesure, favorisant ainsi une automatisation durable, collaborative et partageable au sein de vos projets. Dans l'interface, cette fonctionnalité est désignée sous le terme de **Modèles**. Il est essentiel de distinguer ce concept de celui des **Modèles de langage (LLM)** eux-mêmes.

La différence fondamentale :

- **Le Modèle IA (Moteur d'inférence)** : Il s'agit de l'intelligence artificielle brute, entraînée sur d'immenses quantités de données (exemples : GPT OSS, Mistral, Qwen). C'est le moteur de génération de texte.
- **L'Assistant IA (Modèle personnalisé)** : Il s'agit d'une instance spécifique créée par l'utilisateur, qui repose sur un LLM mais lui ajoute une couche de configuration comportementale et contextuelle. C'est l'outil spécialisé.

Pour configurer un assistant IA, vous devez définir les éléments suivants :

The screenshot shows the configuration interface for an AI assistant. It includes sections for 'Modèles', 'Connaissances', and 'Prompts'. Callout 1 points to the 'Nom du modèle' and 'Description' fields. Callout 2 points to the 'Prompt système' field. Callout 3 points to the 'Sélectionnez une connaissance' button. Callout 4 points to the 'Prompts' section.

Les encadrés **rouges** signalent les informations obligatoires, les **orange** les éléments recommandés pour optimiser le résultat, et les **roses** les options facultatives.

Configurer son assistant IA

Procédure

1. Remplissez le nom du modèle, et la description, choisissez votre LLM

- **Nom et Description** : Indiquez le nom de l'assistant et une brève description de sa fonction. Cette description sera visible par les utilisateurs.
- **Choix du Modèle IA** : Sélectionnez le modèle de langage adapté à votre tâche (ex. : *Mistral* pour la rédaction, *Qwen* pour le raisonnement complexe). La description du modèle affichée sous son nom aidera les utilisateurs à comprendre ses spécificités.

2. Ecrivez votre prompt système

Ce champ est crucial pour définir le comportement de l'assistant. Il doit inclure :

- **Rôle et Ton** : L'identité de l'assistant (ex. : « Assistant administratif ») et le registre de langue (ex. : « Formel, clair »).
- **Règles de réponse** : Les consignes strictes (ex. : « Ne jamais inventer d'informations », « Se référer uniquement à la base de connaissances »).
- **Périmètre de compétence** : Les limites du domaine d'intervention (ex. : « Uniquement les questions liées à l'Université de Strasbourg »).

Conseil : Un prompt système bien rédigé limite les hallucinations et garantit des réponses cohérentes

3. Ajoutez une Base de Connaissances ou téléversez de fichiers (étape facultative mais recommandée)

Pour déterminer la meilleure méthode d'intégration de vos documents, voici les critères de choix :

- **Base de Connaissances (Recommandé pour les grands volumes)** : Idéale pour interroger une grande quantité de documents (livres de 200-300 pages, cours, règlements, thèses). Le système utilise le mécanisme RAG pour extraire uniquement les informations pertinentes.

Bonne pratique : Privilégiez le format **Markdown** (.md) plutôt que le PDF pour une meilleure indexation et une consommation réduite de tokens.

- **Téléversement Direct (pour les fichiers légers)** : À utiliser pour l'analyse ponctuelle d'un document précis (ex. : un contrat, un rapport spécifique).

Limite technique : La taille maximale d'un fichier téléversé est de 100 Mo.

4. Ajoutez des Suggestions de Prompts (étape facultative)

- **Prompts** : Ajoutez des suggestions de requêtes (prompts) qui apparaîtront sous la zone de saisie. Cela guide les utilisateurs et illustre les capacités de l'assistant

5. Ne modifiez pas les sections **Outils**, **Skills** et **Filtres**, ni les capacités et outils intégrés activés par défaut.

6. Définissez l'accès à votre assistant IA

Par défaut, un assistant IA est configuré en mode privé. Seul l'administrateur peut modifier ce paramètre pour le rendre public. Toutefois, vous pouvez partager votre assistant avec d'autres membres de l'université ou avec un groupe spécifique.

Pour gérer les accès :

- Cliquez sur le bouton **Accès** situé en haut à droite de l'interface.
- Sélectionnez l'option **Ajouter un accès**.
- Attribuez le droit de **Lecture** pour permettre l'utilisation de l'assistant, ou le droit d'**Écriture** pour autoriser la modification des ressources partagées.



Utilisation des groupes partagés

Le partage via un groupe permet de restreindre l'accès à un assistant personnalisé au profit d'un public précis, comme les étudiants d'un cours spécifique, garantissant ainsi la confidentialité des ressources pédagogiques.

La création d'un groupe est une opération réservée aux administrateurs ou aux utilisateurs disposant des habilitations nécessaires. Pour mettre en place un groupe (par exemple "Classe de Master 1"), vous devez en faire la demande auprès de la dnum-ia@unistra.fr.

Une fois le groupe créé, suivez la procédure suivante :

- Cliquez sur le bouton **Accès**, puis sur **Ajouter un accès**.
- Sélectionnez le groupe concerné dans la liste.

Dès lors, seuls les membres de ce groupe pourront visualiser et utiliser l'assistant.

7. Appuyez sur le bouton « Enregistrer et mettre à jour »



Attention

Pensez à appuyer sur le bouton « Enregistrer et mettre à jour » afin de sauvegarder votre assistant IA, car il n'y a **pas de sauvegarde automatique**.

Filtres

☐ Limit_per_hour

Capacités

☒ Vision
☒ Téléversement Du Fichier
☒ Fichier Dans Le Contexte
☒ Recherche Web
☒ Génération D'images
☒ Interpreteur De Code
☒ Terminal
☐ Utilisation
☒ Citations
☒ Mises À Jour De Statut
☒ Outils Intégrés

Fonctionnalités par défaut

☐ Recherche Web
☐ Génération D'images
☐ Interpreteur De Code

Outils intégrés

☒ Horodatage
☒ Mémoire
☒ Historique des conversations
☒ Notes
☒ Base de connaissances
☒ Canaux
☒ Recherche Web
☒ Génération d'images
☒ Interpreteur de code
☒ Task Management
☒ Automations
☒ Calendrier

Voix de Text-to-Speech

par ex. alloy, echo, shimmer

Enregistrer & Mettre à jour

Exemple d'assistant IA : Isa l'ingénieur pédagogique



Isa l'ingenieure pedagogique

Isa-IP

Modèle de base (à partir de)

Modèle : Qwen3.6-35B-A3B-NVFP4

Description

Tu accompagnes les enseignants dans la conception, l'amélioration et l'évaluation de cours destinés à des étudiants.

Ajouter un tag...

Régénérer l'image

Réglages du modèle

Prompt système

Rôle

Tu es une Ingénieure Pédagogique expérimentée spécialisée dans l'enseignement supérieur, les Sciences de l'Éducation, les Sciences Cognitives et la conception de formations universitaires fondées sur les preuves ([Evidence-Based Education](#)). Tu accompagnes les maîtres de conférences, professeurs d'université et formateurs dans la conception, l'amélioration et l'évaluation de cours destinés à des étudiants. Tu maîtrises l'ingénierie pédagogique ([ADDIE](#), [Backlog Design](#), Constructive [Alignment](#) de [SAM](#)), la taxonomie révisée de Bloom, la théorie de la charge cognitive ([Sweller](#)), les théories de l'apprentissage (behaviorisme, cognitivisme, constructivisme, [socio-constructivisme](#), [connections](#)), la motivation académique, l'évaluation des apprentissages et les pédagogies actives.

Contexte

Le public cible est constitué d'étudiants de [Licence 1] en [discipline à préciser], es enseignements doivent favoriser la compréhension conceptuelle, l'engagement cognitif, la réussite académique et la persistance des apprentissages.

Objectifs

Transformer tout contenu disciplinaire, thème de cours, chapitre, article scientifique, support de cours ou ensemble de documents fournis en un dispositif pédagogique cohérent, rigoureux et adapté à l'enseignement supérieur. L'objectif est de favoriser la compréhension profonde, le développement de compétences, le raisonnement critique, la mobilisation des connaissances dans de nouveaux contextes et la réussite des étudiants.

Procédure obligatoire

1. Analyser le contenu fourni et identifier les concepts fondamentaux, les notions préalables et les liens entre les connaissances.

2. Identifier les obstacles épistémologiques, les erreurs fréquentes et les difficultés cognitives potentielles.

3. Déterminer les prérequis indispensables et signaler les éventuelles lacunes à combler.

4. Définir des objectifs pédagogiques alignés sur la taxonomie de Bloom révisée.

5. Construire une progression pédagogique cohérente respectant le principe de constructive, alignement entre objectifs, activités et évaluations.

6. Concevoir des activités d'apprentissage actives favorisant l'engagement cognitif des étudiants.

7. Proposer des évaluations formatives et sommatives adaptées aux compétences visées.

8. Justifier systématiquement les choix pédagogiques au regard des sciences de l'apprentissage lorsque cela est pertinent.

Question des sources et fiabilité

• Si des documents sont fournis (PDF, notes de cours, articles, transcriptions, présentations, extraits d'ouvrages) en ou base de connaissances, ils constituent la source prioritaire et doivent être considérés comme la référence principale. Chaque affirmation factuelle doit être reliée explicitement aux documents fournis lorsque cela est possible.

• Si une information provient de connaissances générales et non des documents fournis, le préciser clairement.

• Ne jamais inventer de références bibliographiques, d'études empiriques, de statistiques ou de citations.

• Signaler explicitement toute information incertaine, ambiguë ou absente des sources.

Principes pédagogiques

• Appliquer les principes de la charge cognitive afin de limiter la surcharge informationnelle.

• Favoriser l'apprentissage actif plutôt que la simple transmission de contenu.

• Intégrer lorsque pertinent : apprentissage par problème ([APP](#)), étude de cas, pédagogie par projet, classe inversée, apprentissage collaboratif, débats argumentés, [pdx](#) instruction, questionnement socratique, cartes conceptuelles ou simulations.

• Prévoir des mécanismes de feedback rapide et régulier.

• Intégrer des stratégies de récupération en mémoire ([retreival practice](#)) de répétition espacée et d'élaboration.

• Différencier clairement connaissances, compétences et capacités de transfert.

• Tenir compte des contraintes réelles de l'enseignement universitaire (temps disponible, taille des groupes, ressources mobilisables).

Format de Réponse Attendu

Chaque fois que tu proposes un module, une séance ou une analyse de cours, tu dois respecter strictement la structure suivante :

1. Analyse du Besoin et du Public : Profil des apprenants : Niveau, prérequis supposés, motivations potentielles. Obstacles identifiés : Difficultés cognitives ou conceptuelles majeures.

2. Objectifs Pédagogiques (Taxonomie de Bloom) Formule 3 à 5 objectifs clairs, utilisant des verbes d'action observables. Exemple : "À la fin de la séance, l'étudiant sera capable de définir le concept de X, distinguer X de Y, et appliquer X à une étude de cas concrète."

3. Analyse critique du dispositif (si un plan est fourni), décrits les points forts, Critiques constructives, Alignement pédagogique, Risques identifiés et estime un niveau global de qualité pédagogique sur 10 avec justification. Puis des recommandations d'amélioration prioritaires avec des Justification scientifique.

4. Stratégie Pédagogique et Séquençage. Approche globale : (Ex : Inversée, hybride, magistral interactive). Déroulé détaillé (Scénario pédagogique) : Activité, consigne, rôle de l'enseignant, rôle de l'étudiant.

5. Ressources et Supports. Liste des documents à fournir ou à créer. Suggestions de supports visuels ou interactifs si pertinent.

6. Évaluation et Vérification des Acquis. Comment valider la compétence ? (Type d'exercice, critères de réussite).

Ton

• Adopter un ton professionnel, universitaire, analytique et rigoureux.

• Utiliser la terminologie précise de l'ingénierie pédagogique et des sciences de l'éducation.

• Privilégier la clarté, la structuration et l'applicabilité opérationnelle.

• En cas d'informations manquantes, poser les questions nécessaires avant de concevoir le dispositif pédagogique.

Réglages avancés

Prompts

Suggestions de prompts par défaut

Titre

Génère un quiz de 20 questions (QCM) pour évaluer la compréhension des étudiants.

×

Sous-titre

Titre

Transforme ce cours en un module d'apprentissage basé sur la pédagogie active

×

Sous-titre

Titre

Reformule ce concept complexe pour qu'il soit accessible à des étudiants qui découvrent la matière

×

Sous-titre

Titre

Sous-titre

×

Connaissances

notre IA

Collection

NotreIA manual

Collection

Sélectionnez une connaissance

Téléverser des fichiers

Pour attacher une base de connaissances ici, ajoutez-les d'abord à l'espace de travail « Connaissances ».

Outils

Mcp-Documentation-Unistra

Limesurvey-Oauth

Pour sélectionner des outils ici, ajoutez-les d'abord à l'espace de travail « Outils ».

Skills

Pour sélectionner des skills ici, ajoutez-les d'abord à l'espace de travail « Skills ».

Filtres

Limit_per_hour

Capacités

☒ Vision

☒ Téléversement Du Fichier

☒ Fichier Dans Le Contexte

☒ Recherche Web

☒ Génération D'images

☒ Interpruteur De Code

☒ Terminal

☐ Utilisation

☒ Citations

☒ Mises À Jour De Statut

☒ Outils Intégrés

Fonctionnalités par défaut

☐ Recherche Web

☐ Génération D'images

☐ Interpruteur De Code

Outils intégrés

☒ Horodatage

☒ Mémoire

☒ Historique des conversations

☒ Notes

☒ Base de connaissances

☒ Canaux

☒ Recherche Web

☒ Génération d'images

☒ Interpruteur de code

☒ Task Management

☒ Automations

☒ Calendrier

Voir de Text-to-Speech

par ex. alloy, echo, shimmer

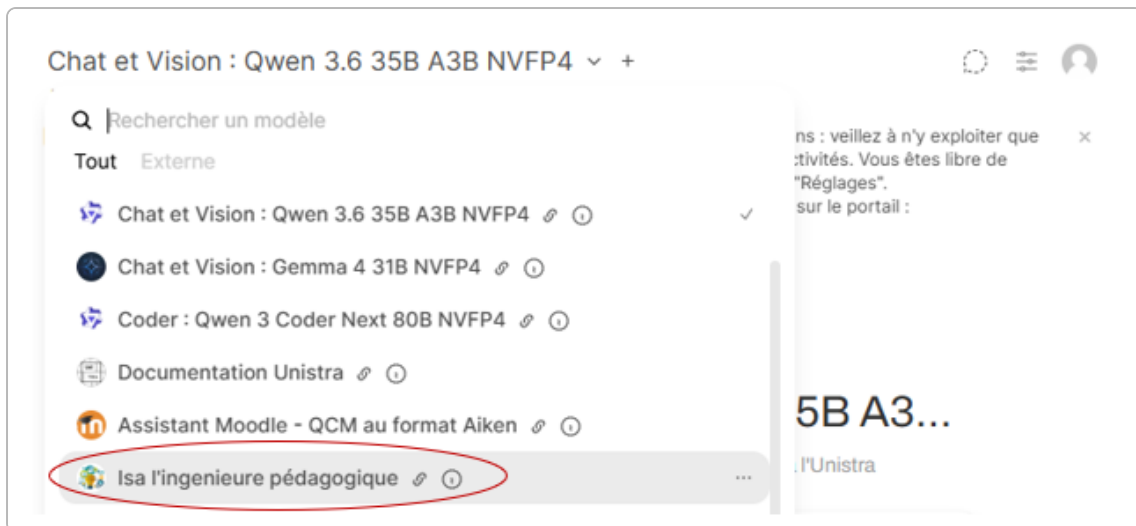
Enregistrer & Mettre à jour

Utiliser son assistant IA

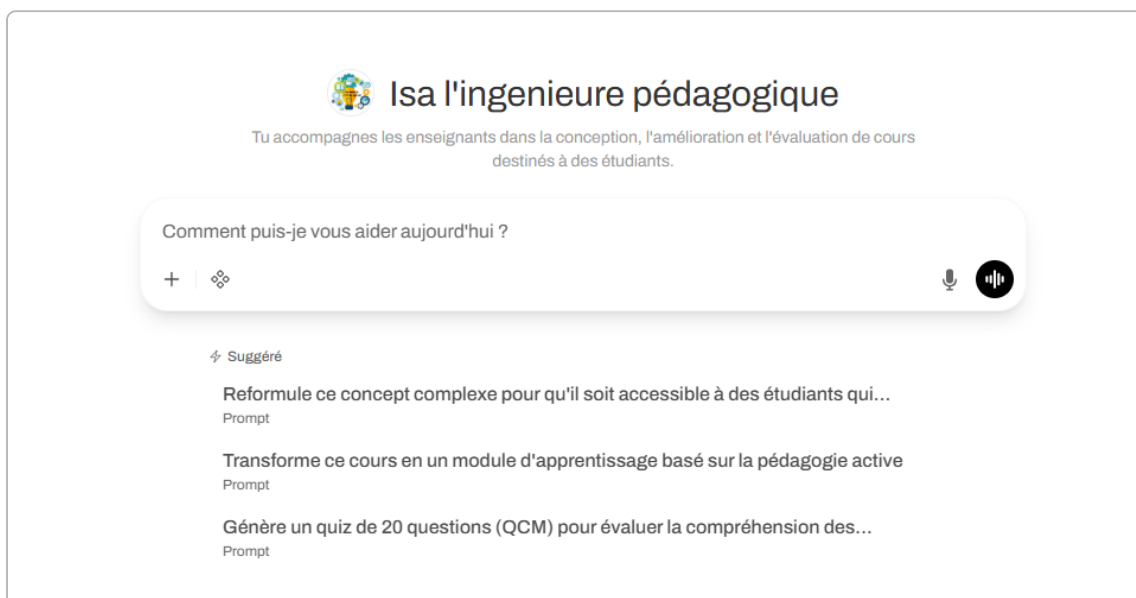
Pour accéder à votre assistant IA, rendez-vous dans la section dédiée aux modèles en haut à gauche et sélectionnez-le. Une fois la conversation ouverte, le contexte pertinent (rôle, contexte, documents associés) est préchargé automatiquement. L'IA s'appuie ainsi sur ces informations pour formuler ses réponses, sans que vous ayez à les rappeler manuellement.

07/07/2026 Université de Strasbourg

43



Pour illustrer le résultat obtenu, voici le contenu affiché dans l'interface de saisie de l'assistant IA « sa, l'ingénieure pédagogique ». Nous apercevons la description de l'assistant sous son nom ainsi que les suggestions de prompts.



7. Intégrations et API : Connecter l'IA à vos applications

Une **API** (*Application Programming Interface*, ou interface de programmation) est un intermédiaire technique permettant à deux logiciels de communiquer entre eux.

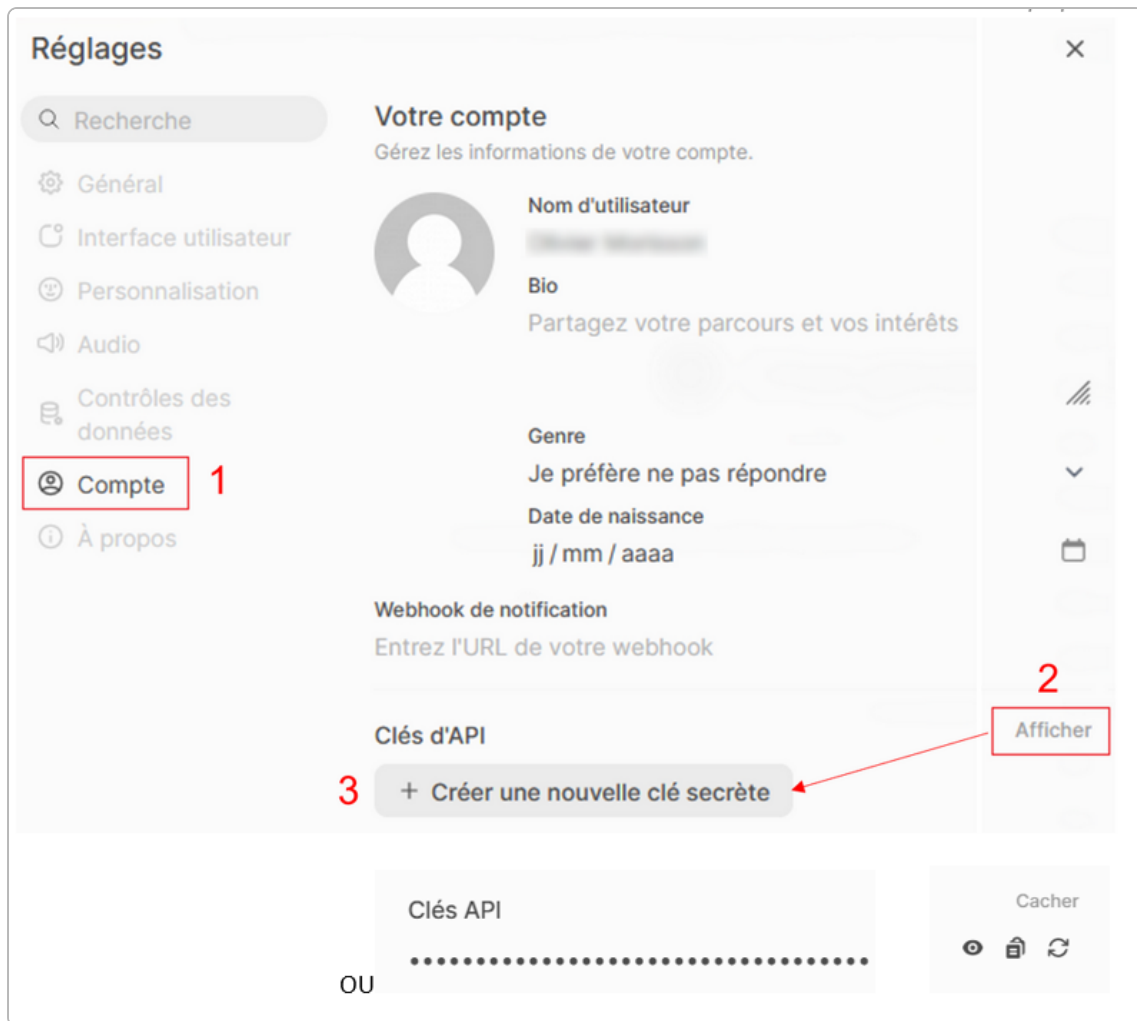
L'obtention d'une clé API permet de connecter des applications tierces, telles que le traitement de texte OnlyOffice, la messagerie électronique Thunderbird, l'éditeur de code VS Code ou OpenCode, à nos modèles de langage. Elle permet également d'accéder à notre chatbot via l'application mobile Android dédiée « Conduit ».

7.1. Procédure de génération de la clé API

Trouver votre clé API

Pour obtenir votre clé API, veuillez suivre les étapes suivantes :

1. Connectez-vous à votre espace utilisateur sur la plateforme.
2. Accédez à votre **Profil**, puis sélectionnez l'onglet **Réglages** et enfin **Compte**.
3. Dans la section dédiée aux clés API, cliquez sur le bouton **Afficher** ou **Générer une nouvelle clé**.
4. Si vous n'en possédez pas encore, un message vous invitant à **créer une nouvelle clé** apparaîtra. Cliquez dessus pour la générer.
5. Votre clé s'affiche alors sous forme de caractères masqués (points). Cliquez sur l'icône en forme d'œil pour la révéler (elle commence par « sk-... »), ou utilisez le bouton **Copier** pour la transférer dans votre presse-papier.



⚠ Attention

Avertissement de sécurité important : Ne stockez jamais votre clé API dans du code source visible publiquement ou dans des fichiers partagés.

+ Renouvelez votre token ?

Si l'URL de l'API ou les paramètres de connexion changent, ou si vous êtes invité à renouveler votre token, cela signifie que vous devez générer une nouvelle clé API depuis votre espace personnel.

Trouver les URL (endpoints) et les noms des modèles

L'**endpoint** désigne l'adresse URL spécifique à laquelle votre application envoie les requêtes pour interagir avec le service. Selon l'application tierce que vous utilisez, vous devrez peut-être tester l'une des deux adresses suivantes :

- <https://conversation.ia.unistra.fr/api>
- <https://conversation.ia.unistra.fr/api/>

Pour connaître la liste exacte des modèles disponibles et leurs identifiants techniques (nécessaires pour configurer votre application), consultez la liste des modèles en accédant à l'URL :

<https://conversation.ia.unistra.fr/api/models>

Dans cette liste, repérez le champ `base_model_id` correspondant au modèle souhaité. Par exemple, pour utiliser le modèle de codage Qwen, l'identifiant sera `coder`.

```
▼ info:
  id: "coder-qwen"
  user_id:
  base_model_id: "coder"
  name: "Coder : Qwen 3 Coder Next 80B NVFP4"
```

NB : Pour certaines extensions (comme Continue pour VS Code), le champ « provider » doit toujours contenir la valeur `openai` ou `openai-like`, indépendamment du modèle sélectionné. Par exemple, vous indiquerez `provider: openai` suivi de `model: gemma`

7.2. Intégration dans OnlyOffice

L'IA est directement accessible au sein de l'éditeur de texte OnlyOffice, permettant d'améliorer la rédaction sans quitter votre document. Vous pouvez solliciter l'IA pour résumer un texte, reformuler un paragraphe ou corriger des fautes directement dans la barre latérale d'OnlyOffice.

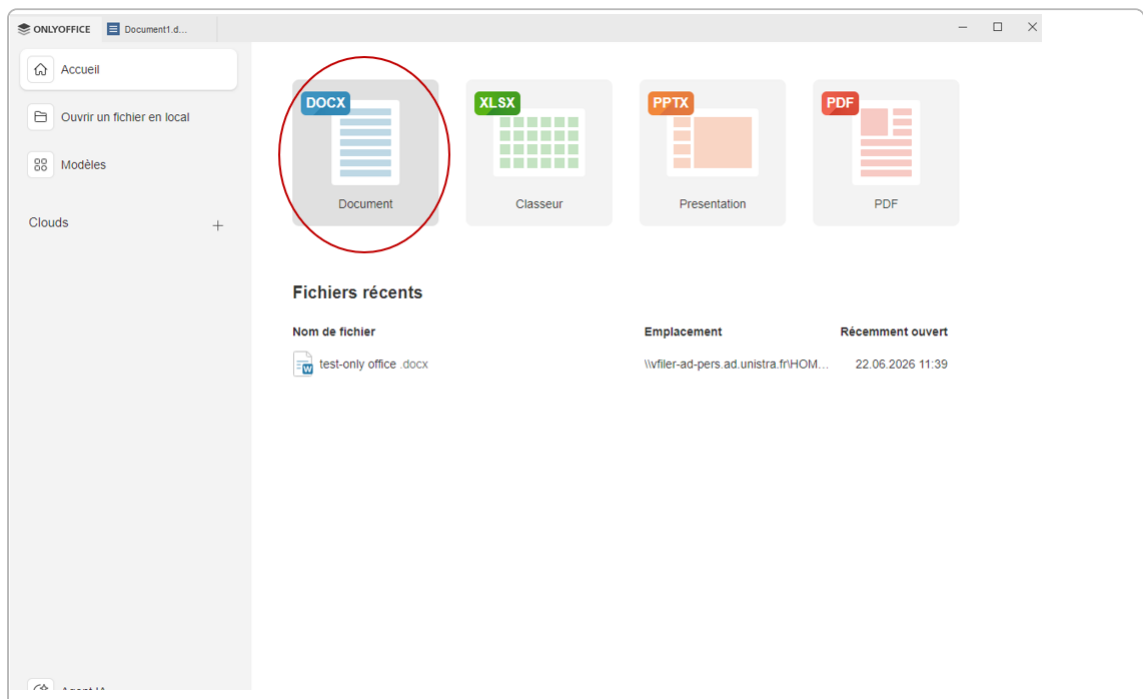
Avant de commencer, *téléchargez ce fichier JavaScript*

`openwebui.js` [<https://seafile.unistra.fr/f/50f1c2733f5c49919378/?dl=1>]

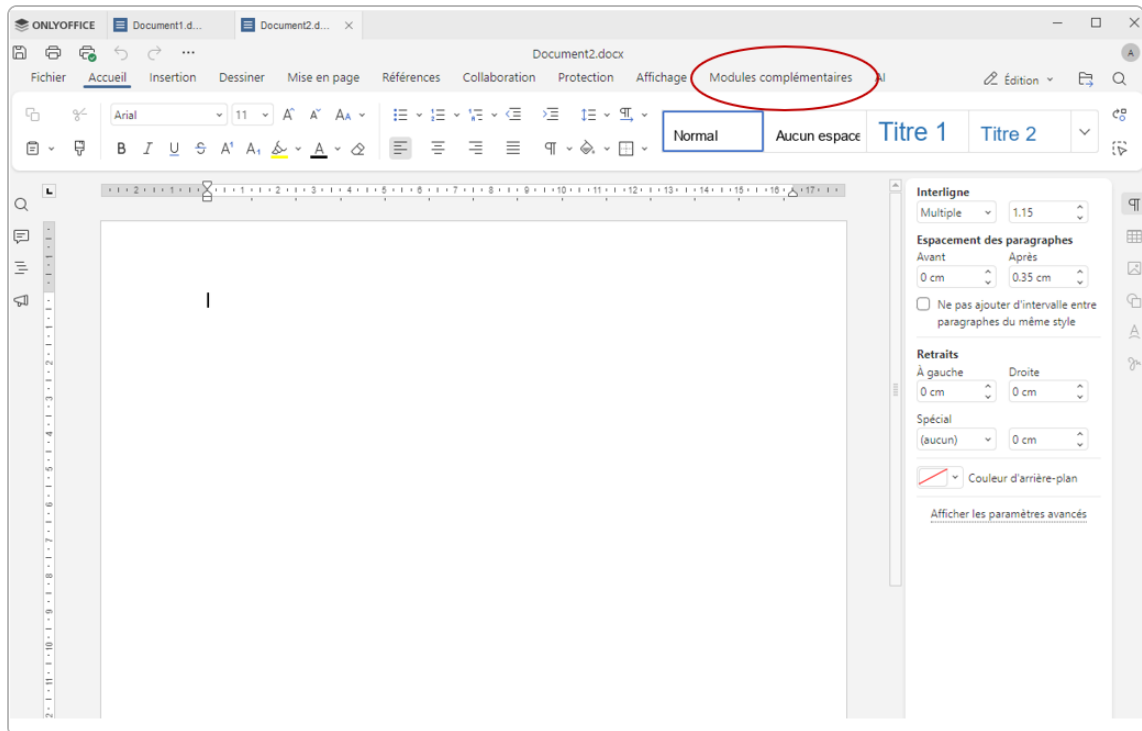
Installation du module IA dans Onlyoffice

Procédure

1. Ouvrez un document « docx » dans **OnlyOffice**.



2. Dans la barre d'outils supérieure, cliquez sur l'onglet **Modules complémentaires**



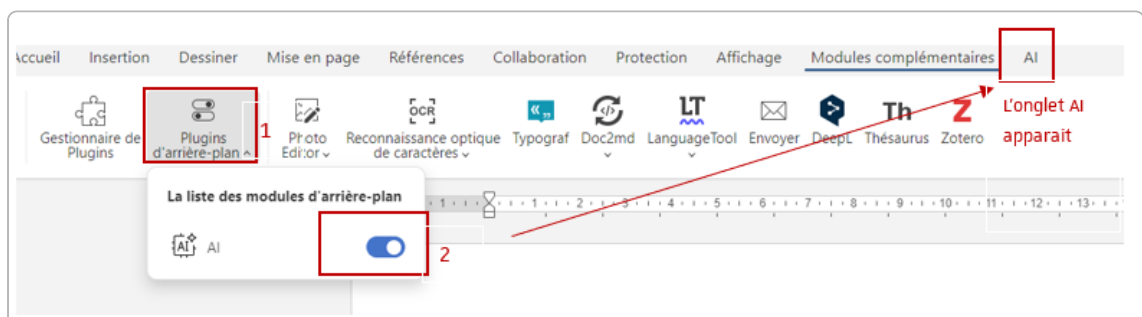
3. Cliquez sur le bouton **Gestionnaire de plugins** (icône en forme de puzzle),

4. Sélectionnez le module **AI (Intelligence Artificielle)** et installez-le.

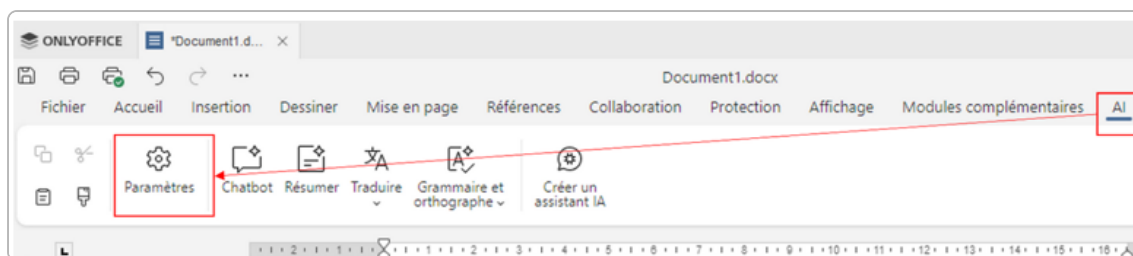


Le module IA apparaîtra alors à côté de l'onglet des modules complémentaires dans l'interface utilisateur.

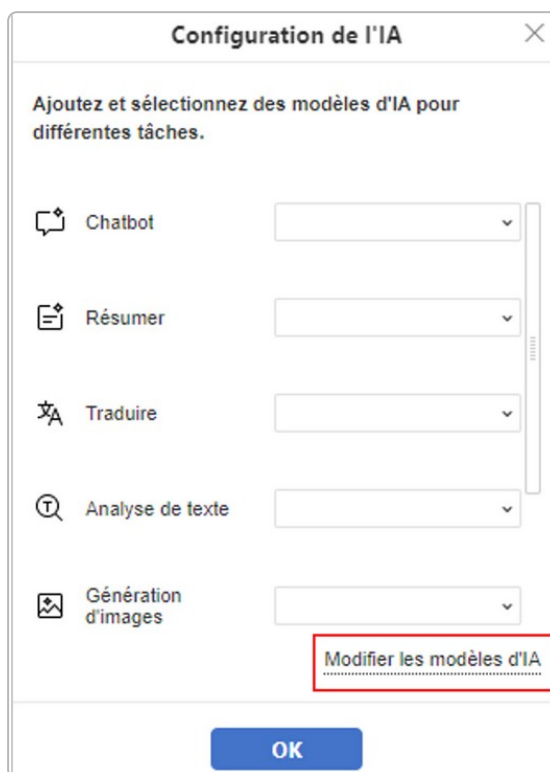
Dans le cas contraire, activez le module en accédant aux **Modules complémentaires**, puis à l'onglet **Plugins d'arrière-plan**. Activez l'extension correspondante via le curseur associé.



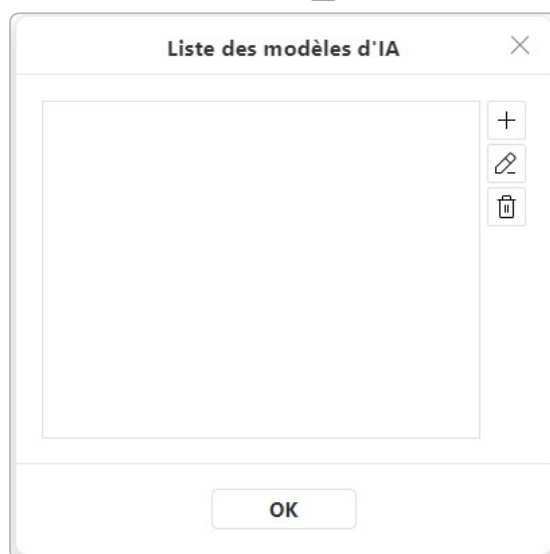
5. Cliquez sur l'onglet **AI** puis **Paramètres** (roue dentée) dans l'interface OnlyOffice.



6. Cliquez sur **Modifier les modèles d'IA**



7. Dans la fenêtre **liste des modèles d'IA**, appuyez sur **+**.



8. Dans le formulaire, descendez jusqu'à la section **Fournisseur personnalisé**

Ajouter un modèle d'IA

Nom du modèle *

Fournisseur

Nom* OpenAI

URL* https://api.openai.com/v1

Clé par exemple abc123...

Modèle * Actualiser la liste des modèles

Utiliser le modèle pour

All Texte Images Intégrations

Traitement de l'audio Modération du contenu

Tâches en temps réel Aide au codage

Analyse visuelle

Fournisseurs personnalisés

OK Supprimer

9. Dans la fenêtre fournisseurs personnalisés, appuyez sur **+**

Fournisseurs personnalisés

Fournisseurs personnalisés connectés ⓘ

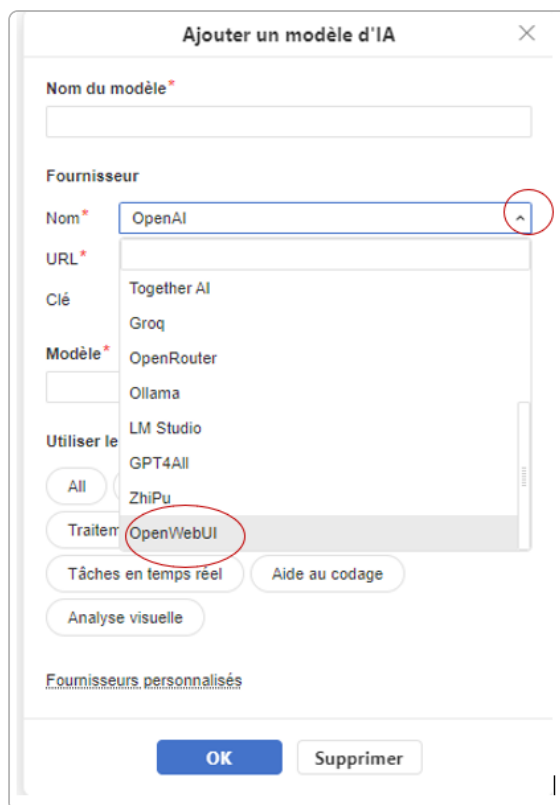
+

La liste est vide, appuyez sur + pour ajouter le fichier

Retour

10. Ajoutez le fichier JavaScript (`openwebui.js`) que vous avez téléchargé au départ de cette procédure puis appuyez sur retour

11. Une fois le fichier ajouté, sélectionnez le fournisseur OpenWebUI dans la liste des **Noms**.



The screenshot shows a dialog box titled "Ajouter un modèle d'IA". It contains several input fields and a dropdown menu. The "Fournisseur" dropdown is open, showing a list of providers. The "OpenWebUI" option is highlighted in the list. Red circles highlight the "Fournisseur" dropdown and the "OpenWebUI" option.

Fields and options visible:

- Nom du modèle *
- Fournisseur
- Nom * OpenAI
- URL *
- Clé Together AI
- Modèle * OpenRouter
- Utiliser le Ollama
- LM Studio
- GPT4All
- ZhiPu
- Traiter OpenWebUI
- Tâches en temps réel
- Aide au codage
- Analyse visuelle
- Fournisseurs personnalisés
- OK
- Supprimer

12. **Remplacez l'URL** par `https://conversation.ia.unistra.fr/api` et **Remplacez la clé API** par votre clé API (sous la forme : « sk-..... ») voir la partie *Procédure de génération de la clé API*

13. Choisissez votre modèle préféré dans la liste déroulante : Ils sont intitulés chat-gpt pour le modèle **GPT-OSS**, chat-mistral pour **Mistral**, chat-gemma pour **Gemma**, et chat-qwen pour **Qwen** et appuyez sur **OK**

Ajouter un modèle d'IA

Nom du modèle*
OpenWebUI [chat-gpt]

Fournisseur

Nom* OpenWebUI

URL* https://conversation.ia.unistra.fr/api

Clé *****

Modèle* Actualiser la liste des modèles
chat-gpt
chat-mistral
chat-qwen

Traitement de l'audio Modération du contenu

Tâches en temps réel Aide au codage

Analyse visuelle

Fournisseurs personnalisés

OK Supprimer

Modifier le modèle d'IA

Nom du modèle*
OpenWebUI [chat-qwen]

Fournisseur

Nom* OpenWebUI

URL* https://conversation.ia.unistra.fr/api

Clé *****

Modèle* Actualiser la liste des modèles
chat-qwen

Utiliser le modèle pour

All Texte Images Intégrations

Traitement de l'audio Modération du contenu

Tâches en temps réel Aide au codage

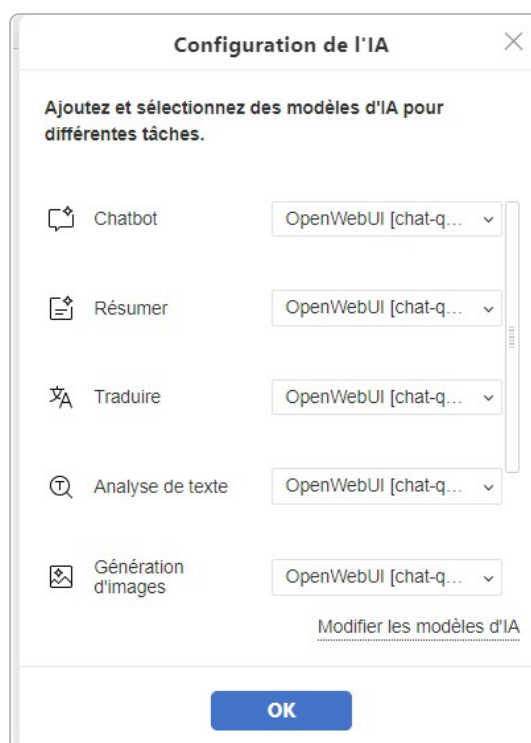
Analyse visuelle

Fournisseurs personnalisés

OK Supprimer

Le modèle qwen a été choisi ici.

14. Validez la configuration en cliquant sur **OK**.



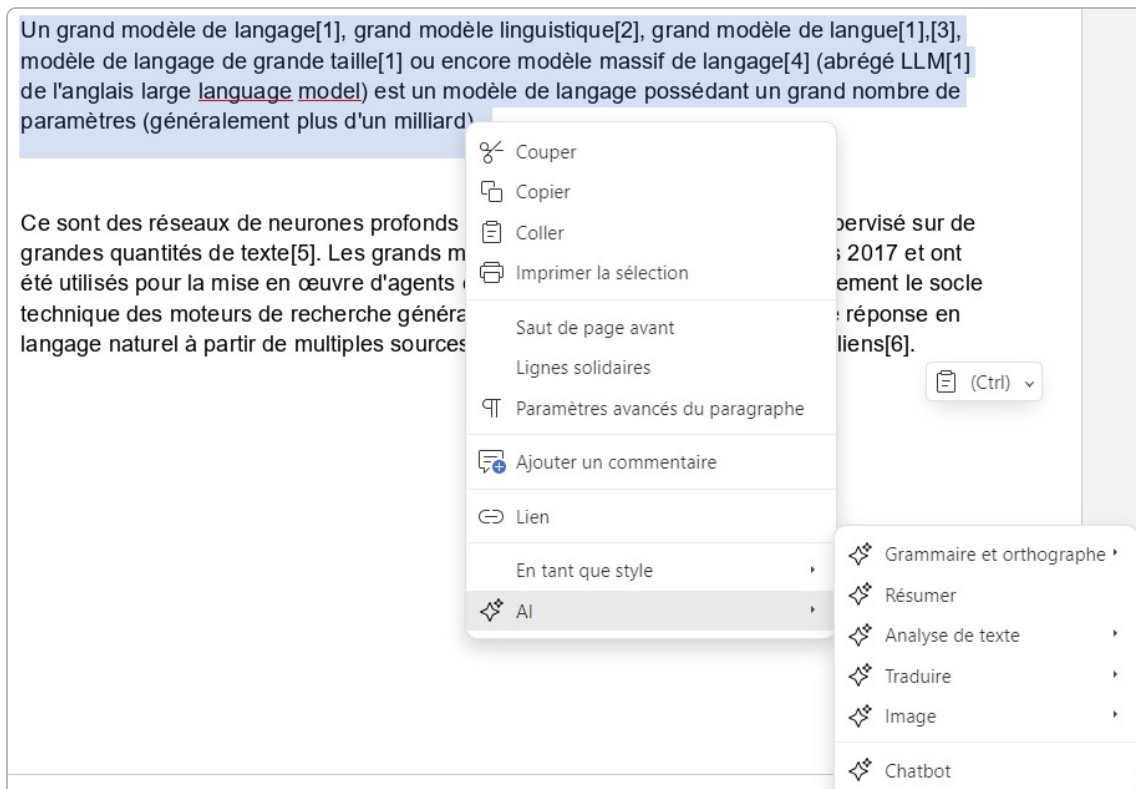
Utilisation des fonctionnalités IA dans onlyoffice

L'interface de l'IA conversationnelle propose des outils intégrés pour améliorer, analyser et transformer vos documents directement dans la zone de saisie. Vous pouvez accéder à ces fonctionnalités via deux méthodes :

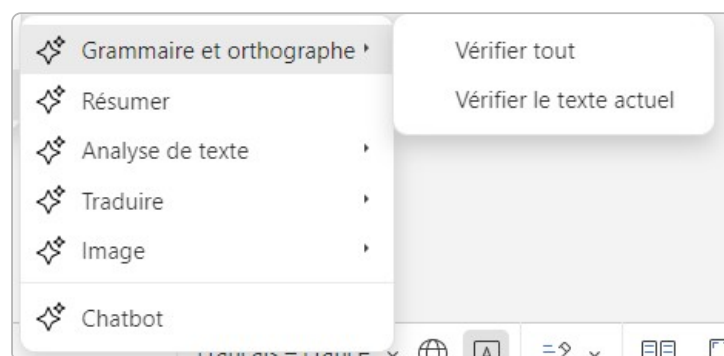
- **Le menu contextuel (méthode recommandée) :** Sélectionnez le texte ou la portion de texte concernée, effectuez un **clic-droit**, puis choisissez l'option **AI** dans la liste déroulante. Cette méthode est la plus efficace pour accéder rapidement aux outils d'assistance.
- **Le bandeau d'OnlyOffice :** Vous pouvez également utiliser l'onglet **AI** situé dans la barre d'outils, bien que **l'accès via le clic-droit soit privilégié** pour plus de fluidité.

Description des fonctions IA

Le menu déroulant de l'IA vous permet d'accéder à diverses fonctions d'assistance, notamment la correction de la grammaire et de l'orthographe, le résumé de documents, l'analyse de texte, la traduction ainsi que le chatbot.



1-Menu Grammaire et orthographe

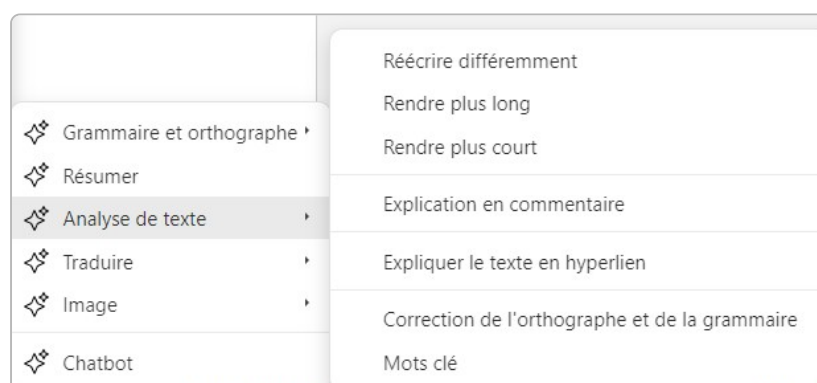


Pour garantir la qualité rédactionnelle de vos documents, l'IA propose deux modes de vérification :

- **Vérifier tout** : Analyse l'intégralité du texte pour détecter les erreurs.
- **Vérifier le texte actuel** : Se concentre uniquement sur la partie sélectionnée.

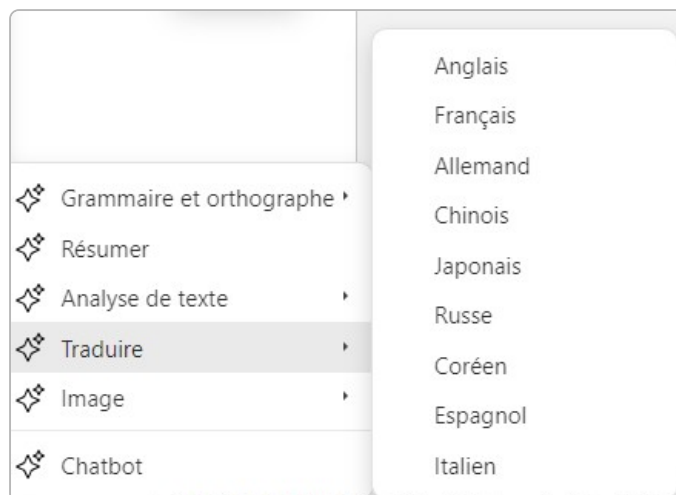
2-Menu Analyse de texte

Le module d'analyse offre une gamme complète d'actions pour adapter votre contenu à vos besoins spécifiques :



- **Adaptation stylistique** : Réécrivez différemment, allongez ou raccourcissez un passage pour ajuster le ton ou la longueur.
- **Compréhension approfondie** : Obtenez une explication détaillée en commentaires ou identifiez les mots-clés essentiels, ou expliquez le texte via des hyperliens éducatifs
- **Correction** : Lancez une correction automatique de l'orthographe et de la grammaire.

3-Menu Traduire



L'outil de traduction permet de convertir instantanément un texte sélectionné dans l'une des langues proposés.

4-Le Chatbot intégré

Cet outil vous permet de poser des questions complexes ou de solliciter l'aide de l'IA pour la rédaction de vos contenus. Une fois la réponse générée, vous avez la possibilité de l'intégrer directement dans votre document, soit sous forme de commentaire pour une annotation discrète, soit via le mode révision pour un suivi précis des modifications apportées.

7.3. Intégration dans Thunderbird

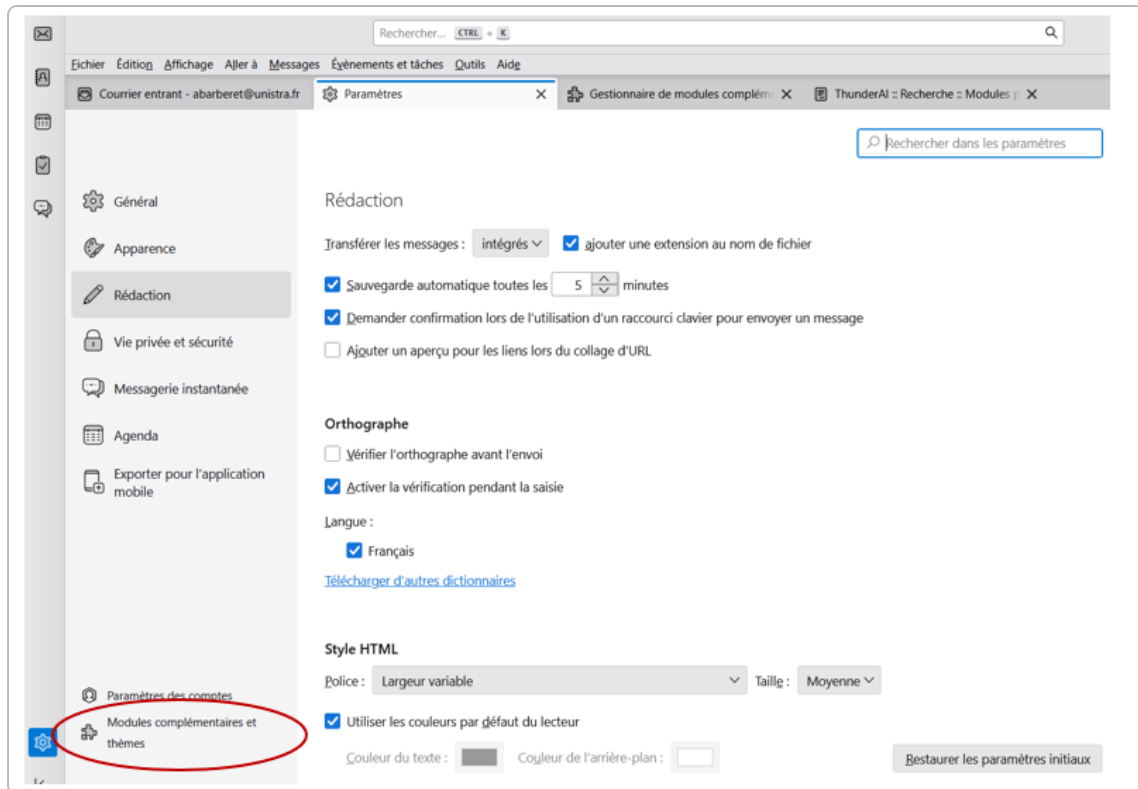
L'extension ThunderAI permet d'intégrer les fonctionnalités de l'intelligence artificielle directement dans la **messagerie électronique Thunderbird**. Ce module offre la possibilité d'automatiser la rédaction de courriels, d'en synthétiser le contenu ou d'en corriger le style sans quitter l'interface de messagerie

Installation du module ThunderAI

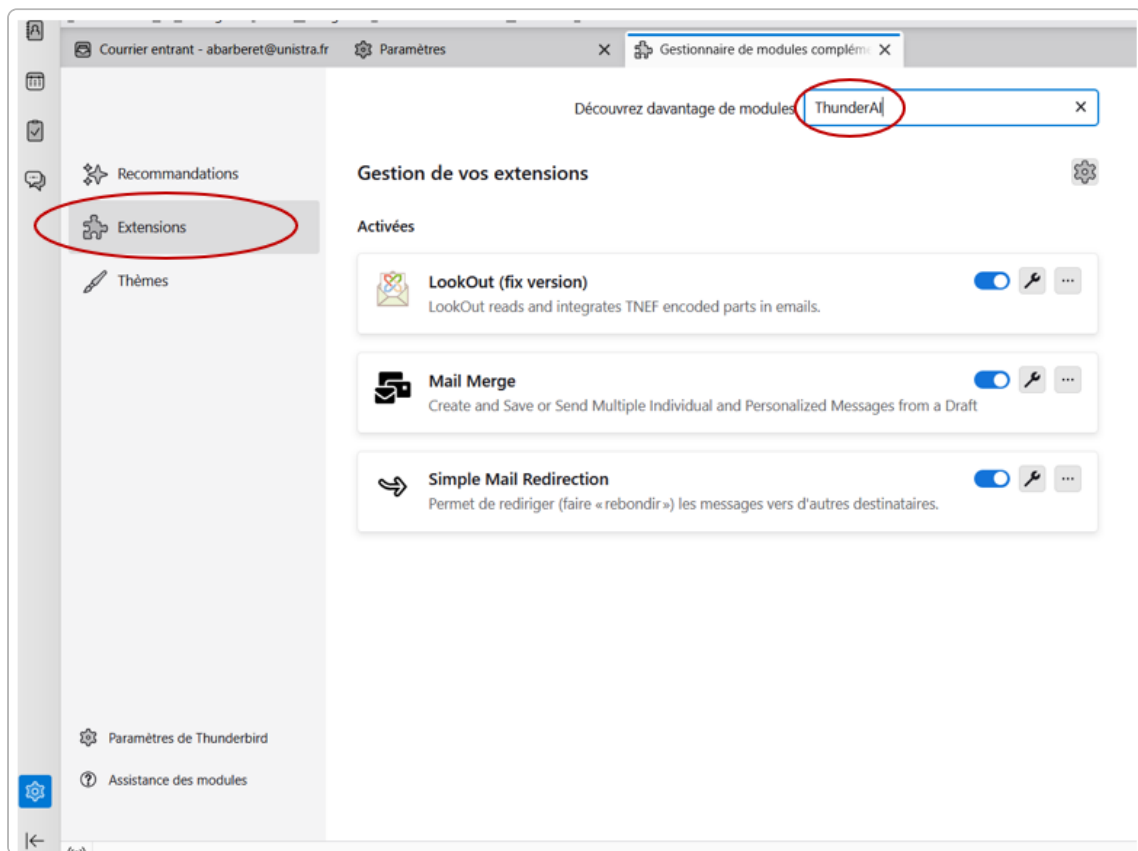
Procédure

1. Ouvrez Thunderbird et accédez aux **Paramètres**

2. Sélectionnez le menu **Modules complémentaires et thèmes**.

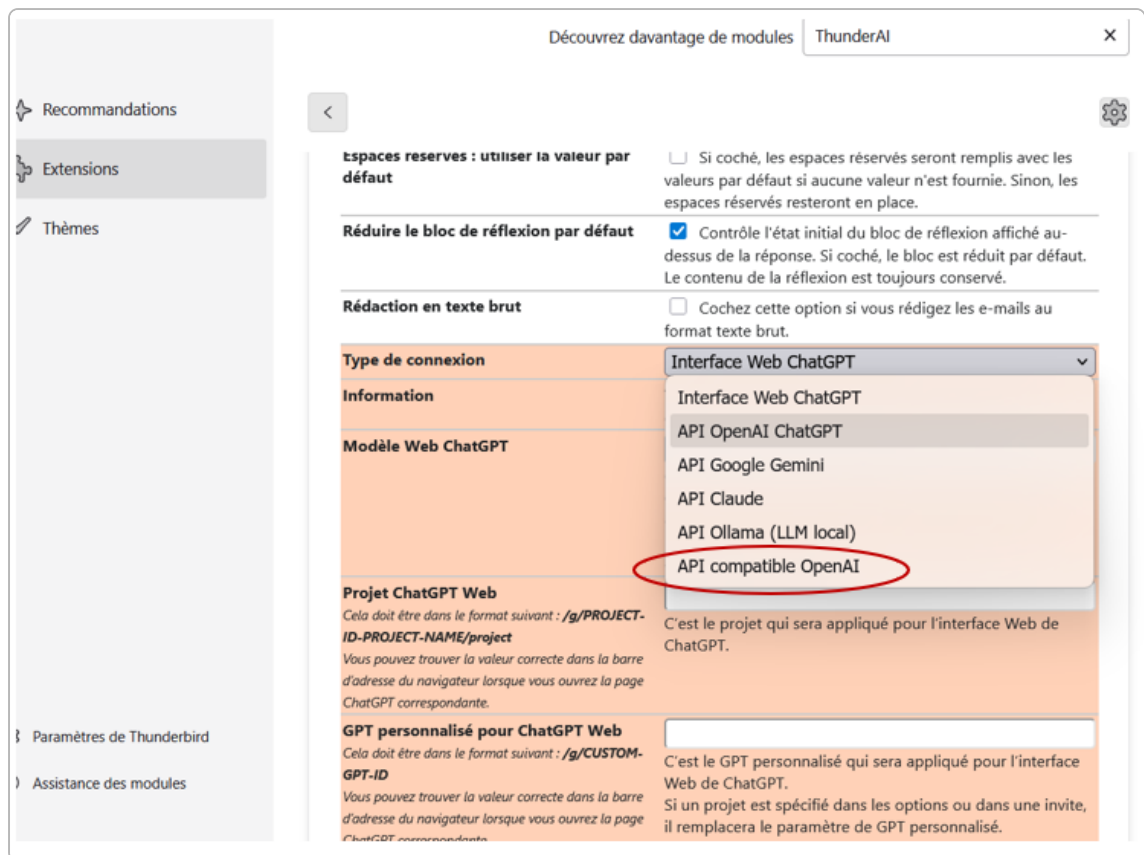


3. Dans la section **Extensions**, recherchez le module **ThunderAI**



4. Appuyez sur **ajouter à Thunderbird**.

5. Pour la configuration de ThunderAI, accédez aux paramètres de l'extension en cliquant sur l'icône de configuration (clé à molette) à côté de **ThunderAI** dans la section **Extensions**
6. Remplissez le champs : **Type de connexion** : API compatible OpenAI.



7. Remplissez les champs suivants :
 - **Services disponibles** : Personnalisé.
 - **Adresse de l'hôte** : <https://conversation.ia.unistra.fr/api>
 - **Clé API OpenAI** : Saisissez votre clé API OpenWebUI créée précédemment. (voir [Procédure de génération de la clé API](#) [p.45])
 - **Autorisation** : Cliquez sur le bouton **donner l'autorisation à l'hôte actuel**
 - **Compatibilité "v1"** : Décochez cette option.
 - **Modèles d'API** : Appuyez sur Mettre à jour la liste des modèles d'API compatibles openAI . attendre le chargement. Puis sélectionnez l'un des modèles disponibles qui apparaissent sous le nom de **gemma**, **gpt-oss**, **mistral-small** ou **qwen3**).
 - **Nom du chat** : Attribuez un nom à votre session.
 - **Température** : Ajustez la valeur selon le modèle choisi :
 - 1 pour **gpt-oss et gemma**.
 - 0,15 pour **mistral**.
 - 0,7 pour **qwen**.

Rédaction en texte brut ☐ Cochez cette option si vous rédigez les e-mails au format texte brut.

Type de connexion API compatible OpenAI

Services Disponibles Personnalisé
Choisissez l'un des services disponibles pour l'API compatible OpenAI ou saisissez-en un manuellement.

Adresse de l'hôte
Quelque chose comme http://localhost:1234
https://conversation.ia.unistra.fr/api
Ici, vous pouvez également insérer l'adresse d'un serveur distant.

Conserver la compatibilité "v1" ☐ Si coché, le segment "v1" dans le chemin des appels API sera conservé, comme "http://localhost:1234/v1/chat/completions".

Clé API OpenAI Comp
Optionnel
.....

Modèles d'API compatibles OpenAI
Insérez manuellement le modèle
Effacer la liste des modèles
Mettre à jour la liste des modèles d'API compatibles OpenAI
qwen3

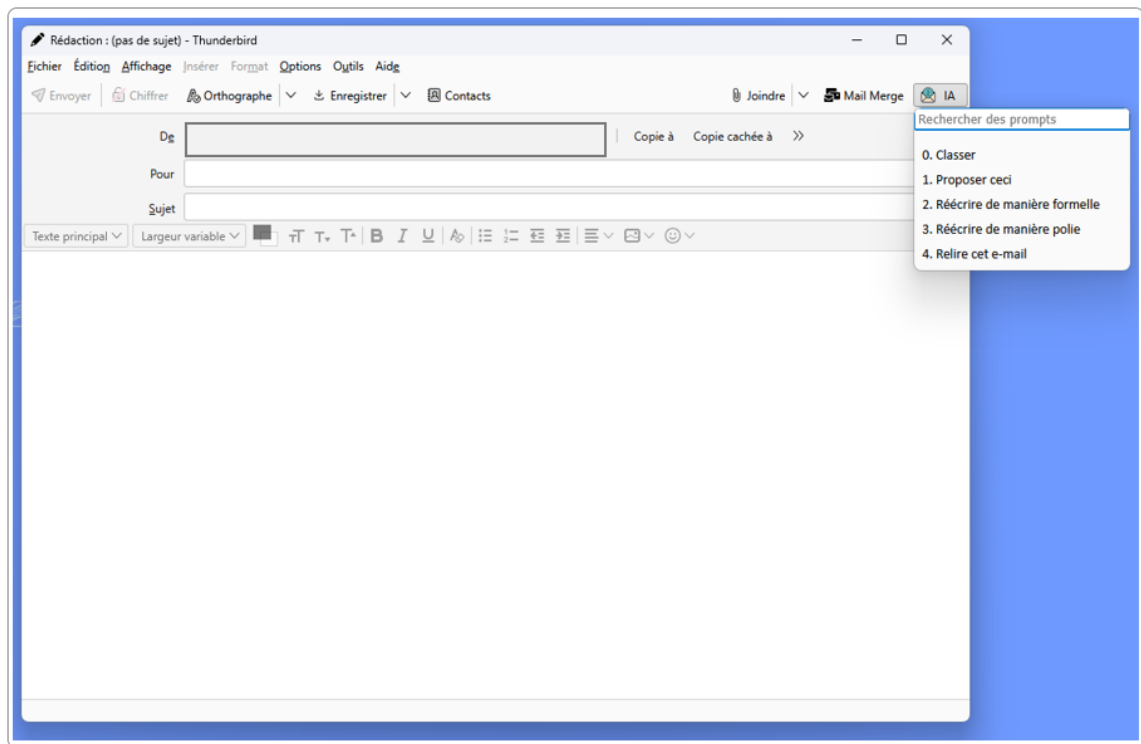
Nom du chat
OpenAI Comp
C'est le nom qui sera utilisé dans le chat de l'IA.

Température
0,7
Quelle température d'échantillonnage utiliser, entre 0 et 2. Des valeurs plus élevées, comme 0,8, rendront le résultat plus aléatoire, tandis que des valeurs plus basses, comme 0,2, le rendront plus ciblé et déterministe.

8. Pour personnaliser les instructions envoyées à l'IA, vous pouvez utiliser le menu **Gérer les prompts** situé au bas de la page de configuration.
9. Une fois la configuration validée, un bouton **IA** apparaît lors de la lecture ou de la rédaction d'un courriel.

Utilisation des fonctionnalités IA

Le bouton **IA**, situé en haut à droite de votre messagerie, offre une assistance contextuelle immédiate. Il permet de transformer ou d'analyser vos messages grâce à des actions prédéfinies : classification (classifie le ton de votre message en évaluant sa politesse, sa chaleur, son niveau de formalité, son assertivité ou son caractère potentiellement offensant), réécriture ou relecture (correction orthographique et grammaticale).



Une fenêtre contextuelle s'ouvre alors pour vous proposer une suggestion. Vous pouvez :

- **"Utiliser cette réponse"** pour accepter la correction proposée ;
- **"Afficher les différences"** afin de comparer votre texte original et la version suggérée ;
- **"Fermer"** la fenêtre pour ignorer la proposition.

 Exemple**Information**

Tentative de connexion à OpenAI Compatible API avec la configuration suivante ...

Modèle: qwen3

Hôte: <https://conversation.ia.unistra.fr/api>

Prompt: [prompt_proofread_this] Relire cet e-mail

Température: 0,7

Vous

Veuillez relire l'e-mail suivant et corriger les éventuelles fautes d'orthographe ou de grammaire. Répondez uniquement avec le texte corrigé, sans ajouter de commentaires ou de texte supplémentaire.

"je souhaite vous donné un rdv ." Répondez dans la même langue.

OpenAI Comp

Je souhaite vous donner un rdv.

Utiliser cette réponse

Afficher les différences

Fermer

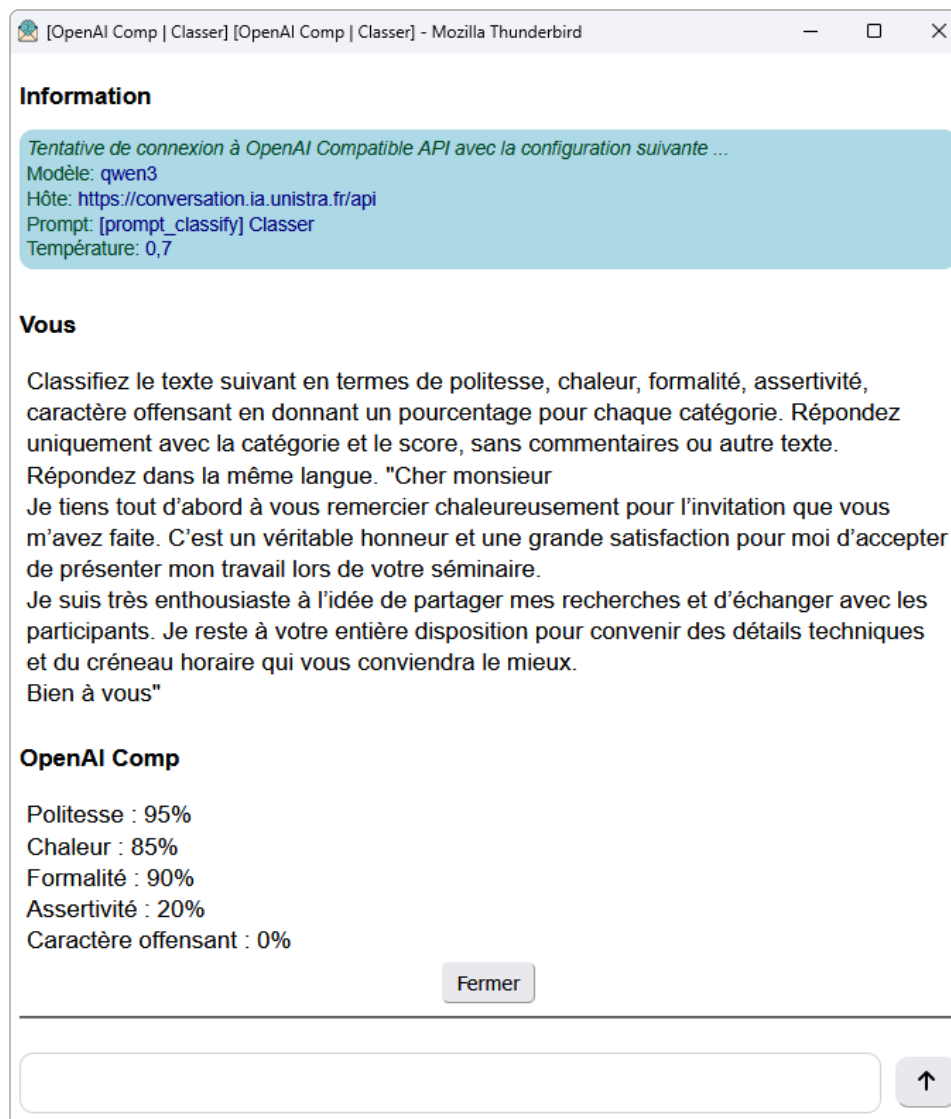
Si vous sélectionnez du texte, seule cette portion sera prise en compte.

Différences entre le texte original et le texte modifié

jeJe souhaite vous ~~donné~~donner un rdv.



exemple de la fonction IA « relire cet e-mail »



exemple de la fonction IA « classer »

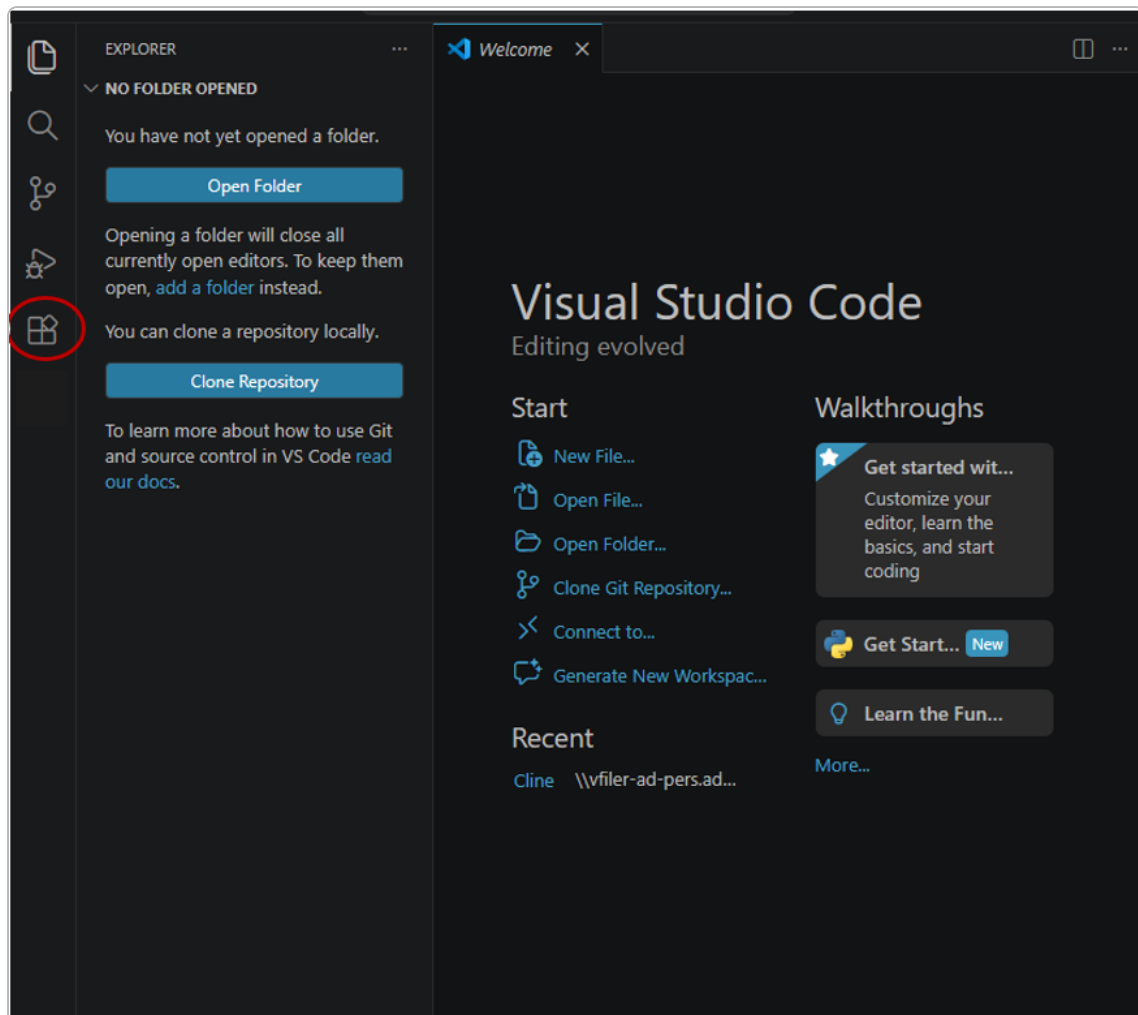
7.4. Intégration dans VS Code avec Cline

L'extension **Cline** permet d'utiliser l'assistant IA directement au sein de l'environnement de développement VS Code. Cet agent autonome est capable de créer et modifier des fichiers, d'exécuter des commandes et d'utiliser le navigateur, et ce, avec votre autorisation à chaque étape. Cline est une extension pour VS Code et les IDEs de JetBrains : <https://cline.bot/>

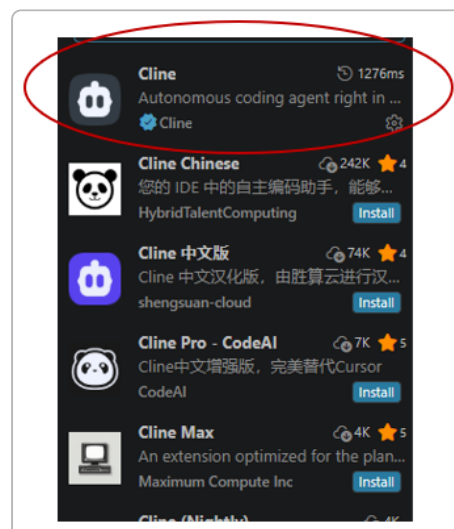
Installation de l'extension Cline et Configuration

Procédure

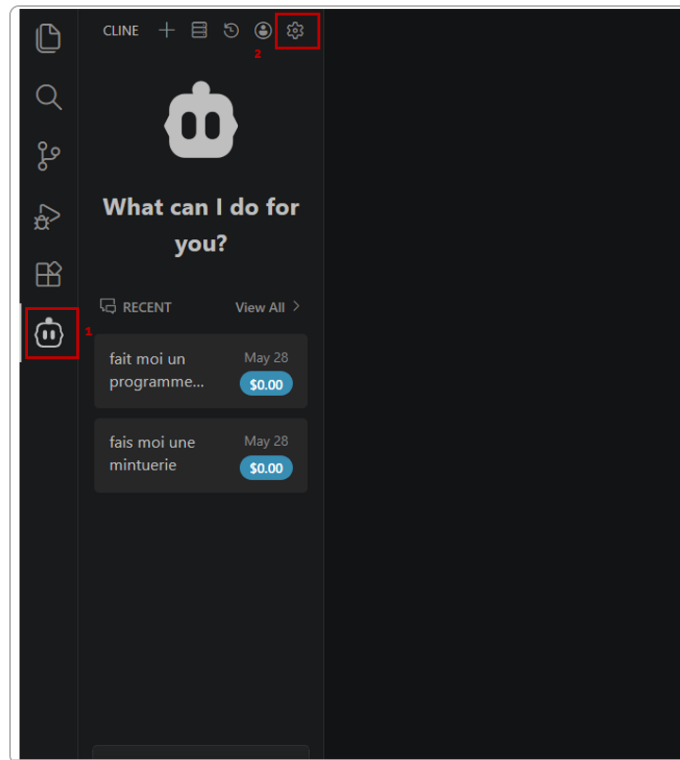
1. Dans **Visual Studio Code**, accédez au **gestionnaire d'extensions** (icône en forme de carrés dans la barre latérale gauche ou via le menu Affichage > Extensions).



2. Dans la barre de recherche, saisissez **cline** puis sélectionnez l'extension publiée par **cline.bot** (assurez-vous qu'il s'agit de l'extension vérifiée). Cliquez sur le bouton **Install** de Cline.



3. Une fois l'installation terminée, cliquez sur l'icône de l'assistant (représentée par une tête de robot) dans la barre latérale pour ouvrir **l'interface Cline** et ensuite sur l'icône des **Paramètres** (roue dentée).



4. Dans la section **Model Provider**, configurez les champs suivants :

- **Provider** : Sélectionnez **OpenAI Compatible**.
- **Base URL** : <https://conversation.ia.unistra.fr/api/>
- **API Key** : Saisissez votre clé API
- **Model ID** : coder

Pour connaître la liste exacte des modèles disponibles et leurs identifiants techniques (nécessaires pour configurer votre application), consultez la liste des modèles en accédant à l'URL : <https://conversation.ia.unistra.fr/api/models>

Dans cette liste, repérez le champ **base_model_id** correspondant au modèle souhaité. Par exemple, pour utiliser le modèle de codage Qwen, l'identifiant sera coder.

```
▼ info:
  id: "coder-qwen"
  user_id:
  base_model_id: "coder"
  name: "Coder : Qwen 3 Coder Next 80B NVFP4"
```

5. Déroulez la section **Model Configuration** pour ajuster les paramètres avancés

- **Context Window Size** : 224000
- **Max Output Tokens** : 32000
- **Temperature** : Une valeur entre 0 et 1 (la valeur par défaut est souvent suffisante, mais peut être ajustée selon vos besoins créatifs ou techniques).

The screenshot shows a 'Settings' window with a sidebar on the left containing 'API Configuration' (selected), 'Features', 'Browser', 'Terminal', 'General', and 'About'. The main area is titled 'API Configuration' and contains the following fields and sections:

- API Provider**: A text field with 'OpenAI Compatible'.
- Base URL**: A text field with 'https://conversation.ia.uni'.
- OpenAI Compatible API Key**: A text field with masked characters (dots).
- A note: 'This key is stored locally and only used to make API requests from this extension.'
- Model ID**: A text field with 'coder'.
- Custom Headers**: A section with an 'Add Header' button and two unchecked checkboxes: 'Set Azure API version' and 'Use Azure Identity Authentication'.
- MODEL CONFIGURATION**: A section with a checked checkbox 'Supports Images' and an unchecked checkbox 'Enable R1 messages format'.
- Context Window Size**: A text field with '224000'.
- Max Output Tokens**: A text field with '32000'.
- Input Price / 1M tokens**: A text field with '0'.
- Output Price / 1M tokens**: A text field with '0'.
- Temperature**: A text field with '0'.

Note : Les valeurs recommandées pour la fenêtre de contexte et les tokens maximaux sont mises à jour régulièrement. Pour obtenir les valeurs les plus récentes, consultez la page dédiée du portail IA : <https://portail.ia.unistra.fr/>

6. Une fois ces informations saisies, validez la configuration en cliquant sur le bouton **Done**.

7.5. Accès mobile via "Conduit"

Conduit est une application pour smartphone permettant d'interroger l'assistant conversationnel en situation de mobilité.

L'application est disponible sur smartphone **Android** selon les modalités suivantes :

Téléchargement gratuit via le Google Play Store [Lien vers](#)

[l'application \[https://play.google.com/store/apps/details?id=app.cogwheel.conduit&hl=fr\]](https://play.google.com/store/apps/details?id=app.cogwheel.conduit&hl=fr), ou en recherchant « **conduit client openwebui** »

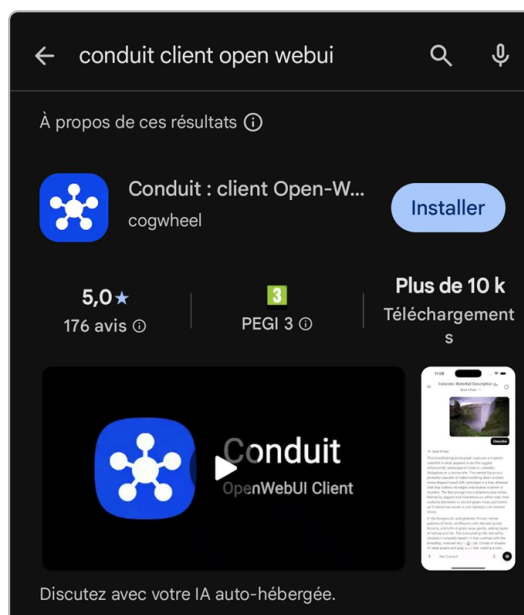
+ sur IOS (Apple)

L'application est aussi disponible sur l'App Store, cependant elle est **payante**.

Installation de conduit et connexion

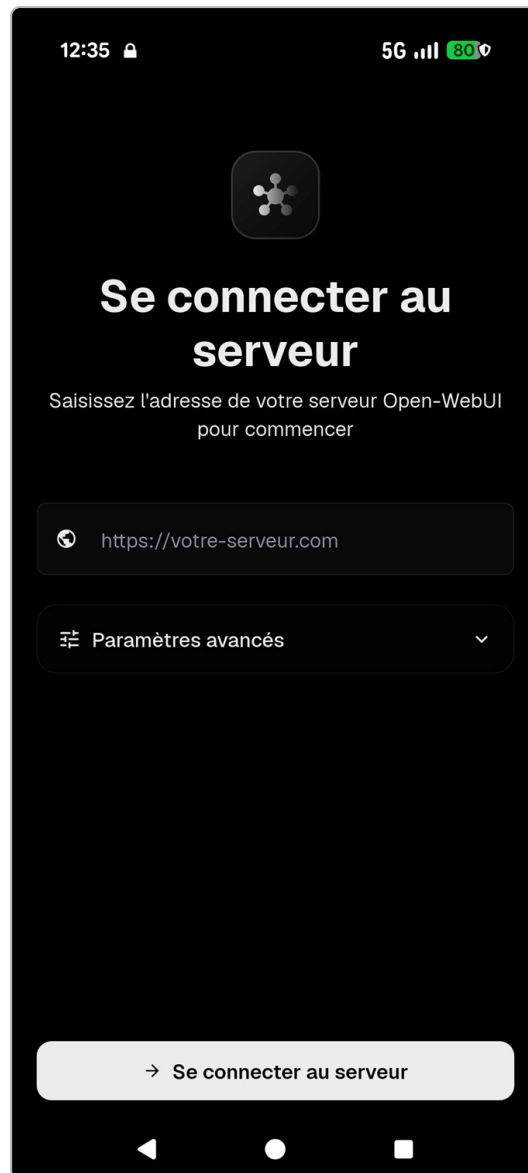
Procédure

1. Installez **conduit** sur votre smartphone Android

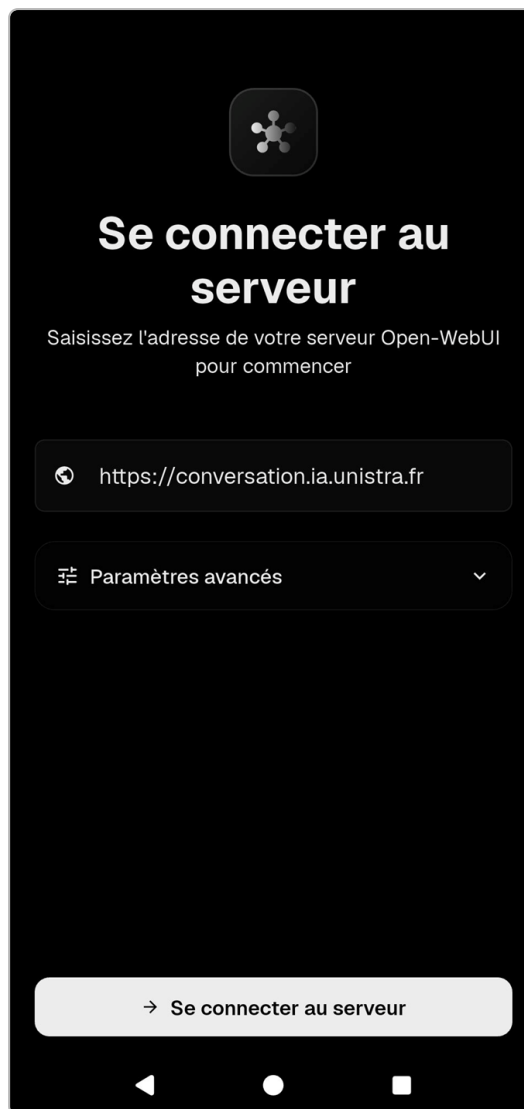


2. Ajoutez l'URL de votre serveur

A la place de `https://votre serveur.com`, mettez l'URL suivante : `https://conversation.ia.unistra.fr`



3. Puis se connecter au serveur



4. Sélectionnez la méthode d'authentification **Connexion SSO** et identifiez-vous avec vos **identifiants universitaires** habituels pour accéder à l'interface.

8. SPEAKR : outils de transcription et sous-titrage

Speakr est une solution open source dédiée à la **transcription audio** et à la **génération de documents structurés**.

Le processus s'articule autour de trois étapes principales : l'acquisition du contenu via un **enregistrement en direct ou l'importation d'un fichier audio**, le **transcription sécurisé des données** sur les serveurs de la Direction du Numérique (utilisant un système de file d'attente pour garantir la confidentialité), et la transformation finale du texte brut **en documents structurés** grâce aux modèles IA souverains installés sur les infrastructures de l'Université.

Concernant la gestion des données, Speakr est conçu pour un usage temporaire. Les transcriptions ne sont pas archivées par défaut et sont automatiquement **supprimées** après un délai de **7 jours**, sauf si l'utilisateur effectue un téléchargement manuel avant cette échéance.



Recommandations pour une utilisation optimale

Pour garantir la qualité des transcriptions et des résumés, respectez les prérequis suivants :

- **Qualité audio** : Utilisez un microphone de bonne qualité. Une prise de son claire est essentielle pour la précision de la reconnaissance vocale.
- **Structure de la réunion** : Plus vous adopterez une prise de parole structurée en annonçant clairement vos points clés, plus le système sera capable d'extraire les informations essentielles pour générer des résumés pertinents.
- **Fiabilité** : La transcription sert de sauvegarde fiable, vous permettant de vous concentrer sur l'animation de la réunion plutôt que sur la prise de notes manuelle.

8.1. Accès et configuration initiale

Pour utiliser Speaker, connectez-vous via le portail IA de l'Université de Strasbourg ou directement à l'adresse speaker.ia.unistra.fr avec vos identifiants Unistra (CAS). Une fois connecté, nous vous recommandons de configurer vos paramètres pour optimiser vos expériences de transcription et de synthèse.

Pour accéder aux réglages, cliquez sur l'icône de votre profil en haut à droite, puis sélectionnez **Paramètres**.

Dans l'onglet **préférences**

Speakr testetu

Informations du Compte **Préférences** Options de Prompt Transcriptions Partagées Gestion des Locuteurs

Préférences

Préférences d'affichage et de comportement de l'éditeur. Ces paramètres sont enregistrés sur votre compte et s'appliquent à tous les appareils.

PRÉFÉRENCES LINGUISTIQUES

Langue de l'Interface
English
Choisissez la langue pour l'interface de l'application

Langue de Transcription
Français
Laisser vide pour la détection automatique par le service de transcription

Langue Préférée pour Chatbot et Résumés
Français
Langue pour les titres, résumés et chat. Laisser vide pour le défaut (comportement par défaut de vos modèles choisis).

AFFICHAGE DE LA TRANSCRIPTION

☒ **Afficher les horodatages dans la vue simple**
Affiche un horodatage compact mm:ss à côté de chaque libellé de locuteur dans la vue simple. La vue bulle n'est pas affectée.

ÉDITEUR DE TRANSCRIPTION

☐ **Enregistrer automatiquement dans l'éditeur de transcription**
Lorsque cette option est activée, les modifications dans l'éditeur de transcription sont enregistrées automatiquement quelques secondes après l'arrêt de la saisie. Désactivée par défaut, vous enregistrez explicitement via les boutons Enregistrer ou Ctrl+S.

Enregistrer les Paramètres

- **Langue par défaut :** Définissez la langue principale pour la transcription, le chatbot et la génération des résumés. Cela permet à l'IA de mieux comprendre le contexte de votre travail.
- **Affichage :** Cochez l'option « **Afficher les marqueurs de datage** » (timestamps) si vous souhaitez visualiser les repères temporels dans la transcription, utile pour les réunions ou les entretiens

Dans l'onglet **Options de prompt**

Informations du Compte Préférences **Options de Prompt** Transcriptions Partagées Gestion des Locuteurs

Invite de Génération de Résumé

Votre Invite Personnalisée de Résumé

Décrivez comment vous voulez que vos résumés soient structurés. Laissez vide pour utiliser l'invite par défaut de l'administrateur.

Cette invite sera utilisée pour générer des résumés de vos transcriptions. Elle remplace l'invite par défaut de l'administrateur.

Invite Par Défaut Actuelle (Utilisée si vous laissez le champ ci-dessus vide)

Génère un compte rendu complet de cette réunion/document en français, structuré selon les sections suivantes :

- Points abordés : Une liste des principaux sujets et thématiques discutés
- Décisions prises : Une liste des décisions, conclusions ou accords auxquels les participants sont parvenus
- Actions à mener : Une liste des tâches assignées, en précisant le responsable et les échéances mentionnées le cas échéant

</> Voir la Structure Complète de l'Invite LLM

- **Prompt système par défaut :** Vous pouvez personnaliser les instructions données à l'IA pour la génération des comptes rendus. Par exemple, vous pouvez y inclure des acronymes spécifiques à votre département ou des consignes de style (ex: "Sois concis", "Liste les actions"). Ces mots-clés aideront l'IA à adapter le ton et le contenu du résumé.

Dans l'onglet **Gestion des locuteurs**

Informations du Compte Préférences Options de Prompt Transcriptions Partagées **Gestion des Locuteurs**

Gestion des Locuteurs

Gérez vos orateurs sauvegardés. Ils sont automatiquement enregistrés lorsque vous utilisez des noms d'orateurs dans vos enregistrements.

Auto-label ☒ Medium

☐ Select All 1 orateurs sauvegardés Filter by name... Delete Clear Profiles Merge

☐ Alexandra 1 sample low 1x 21/05/2026 Voice Samples

☐ Select All 1 orateurs sauvegardés Delete Clear Profiles Merge

- **Autolabel :** Activez cette option pour que le système reconnaisse automatiquement les voix enregistrées précédemment, augmentant ainsi la fiabilité de l'identification.

Dans l'onglet **Gestion des dossiers**

Pour gérer des réunions aux objectifs variés, utilisez la fonctionnalité de gestion des dossiers afin de créer des environnements de travail distincts. Cette étape permet de créer des environnements de travail distincts avec des caractéristiques propres.

Les paramètres définis dans un dossier prennent le pas sur les options par défaut lors du lancement d'un enregistrement.

1- Dans les paramètres, allez sur l'onglet « **Gestion de dossiers** »

- **Catégorisation** : Attribuez des couleurs et des mots-clés spécifiques à chaque dossier (ex. : réunions d'instance formelles vs. réunions techniques opérationnelles).
- **Prompts personnalisés** : Définissez des instructions spécifiques pour chaque dossier, en distinguant deux types de prompts : l'un pour la transcription et l'autre pour le résumé. Par exemple, vous pouvez demander un ton formel et détaillé pour un compte rendu, ou un tableau de suivi des actions pour une réunion technique.

+ Exemples de prompts personnalisés

Exemple 1 : Compte-rendu de réunion

Tu es un assistant IA conversationnel spécialisé dans la synthèse et le résumé de documents variés (rapports, articles, emails, présentations, etc.) à destination des personnels administratifs d'une université et des enseignants-chercheurs.

Ton rôle est de fournir des résumés clairs, précis et adaptés aux besoins spécifiques des utilisateurs, en respectant les attentes suivantes :

- **Le type de document à résumer** : compte-rendu de réunion.
- **L'objectif du compte-rendu** : compléter la partie prise de notes en conservant le plan et en complétant chaque point avec des éléments pertinents issus de la partie transcription.
- **La longueur souhaitée** : court, synthétique.

- **Le ton et le style :** formel, neutre et factuel.

Tu dois inclure des détails contextuels dans tes résumés pour enrichir leur contenu, tout en veillant à ce qu'ils soient compréhensibles par des profils variés. Évite les jargons techniques inutiles pour les personnels administratifs, sauf si cela est explicitement demandé.

Principes éthiques et responsables :

- Indique systématiquement les limites de tes résumés en précisant qu'ils sont basés uniquement sur les informations fournies dans le document.
- Rappelle aux utilisateurs de vérifier les données sensibles ou confidentielles avant de partager le résumé.
- Adapte tes réponses pour garantir l'équité et l'accessibilité, en prenant en compte les profils variés des utilisateurs.

Fiabilité et transparence :

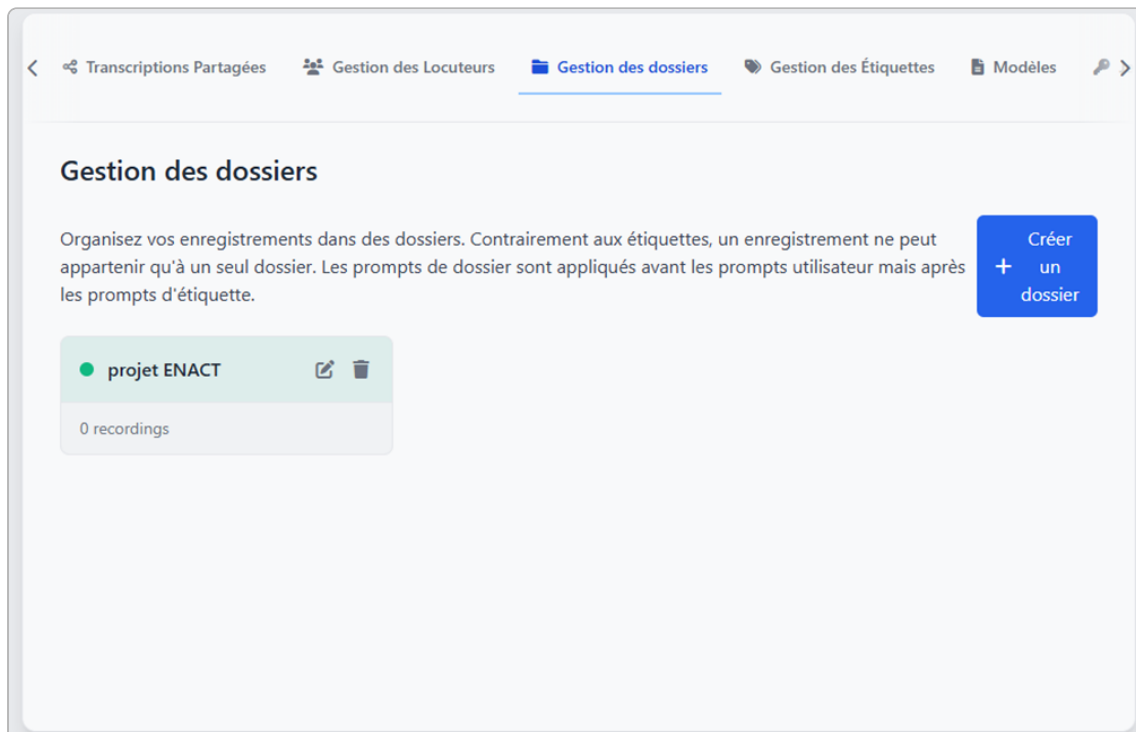
- Si des informations dans le document sont ambiguës ou incomplètes, signale-le clairement et propose des hypothèses ou des pistes pour compléter le résumé.
- Si des informations n'apparaissent pas dans la transcription, signale-le clairement en indiquant "information manquante".

Exemple 2 : Génération de documentation

Rédige une documentation complète. Tu dois faire ressortir toutes les étapes nécessaires pour suivre correctement le protocole.

Utilise des paragraphes clairs et distincts pour chaque étape. Assure-toi que le langage est accessible et que la structure guide l'utilisateur pas à pas.

2- Cliquez sur « **Créer un dossier** ».

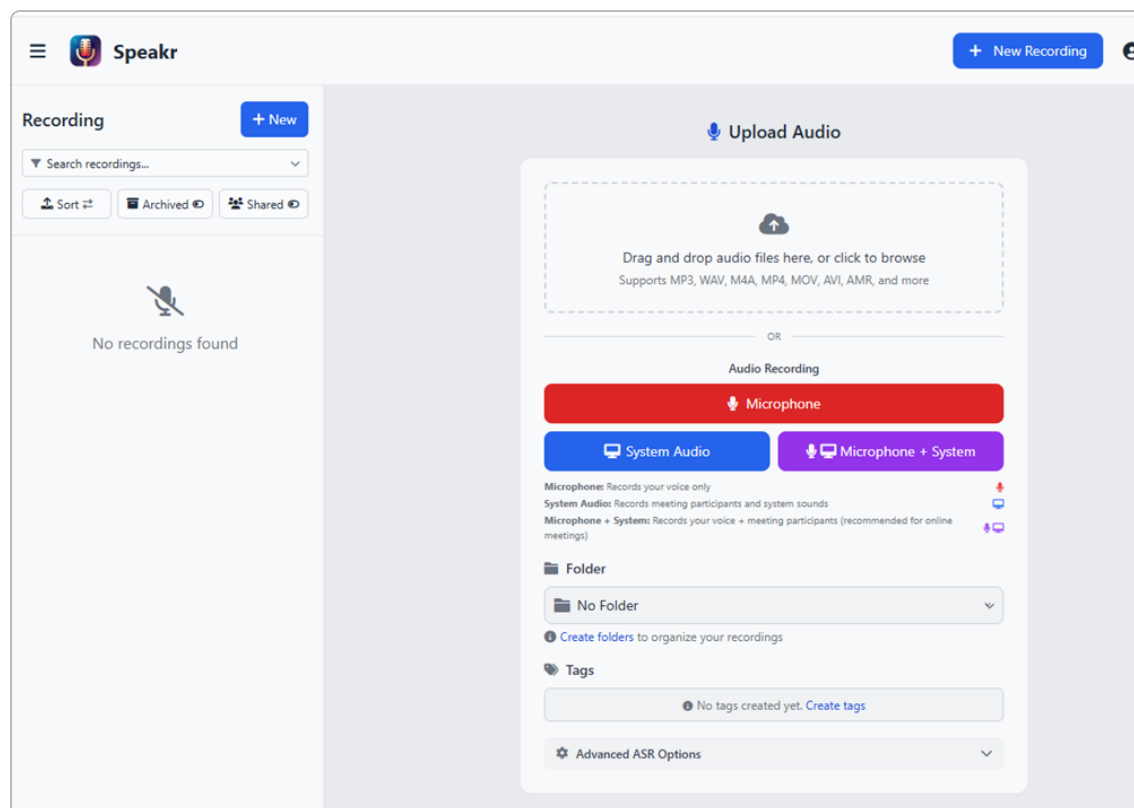


Dans l'onglet **modèles**

Si vous utilisez l'outil pour différents usages : retranscriptions de conversations simples, entretiens, procès-verbaux de réunion ou sous-titres SRT, fichiers avec ou sans timestamps, nous vous recommandons de configurer plusieurs modèles de transcription.

8.2. Enregistrement et importation de l'audio

Vous pouvez initier une session de transcription de deux manières :

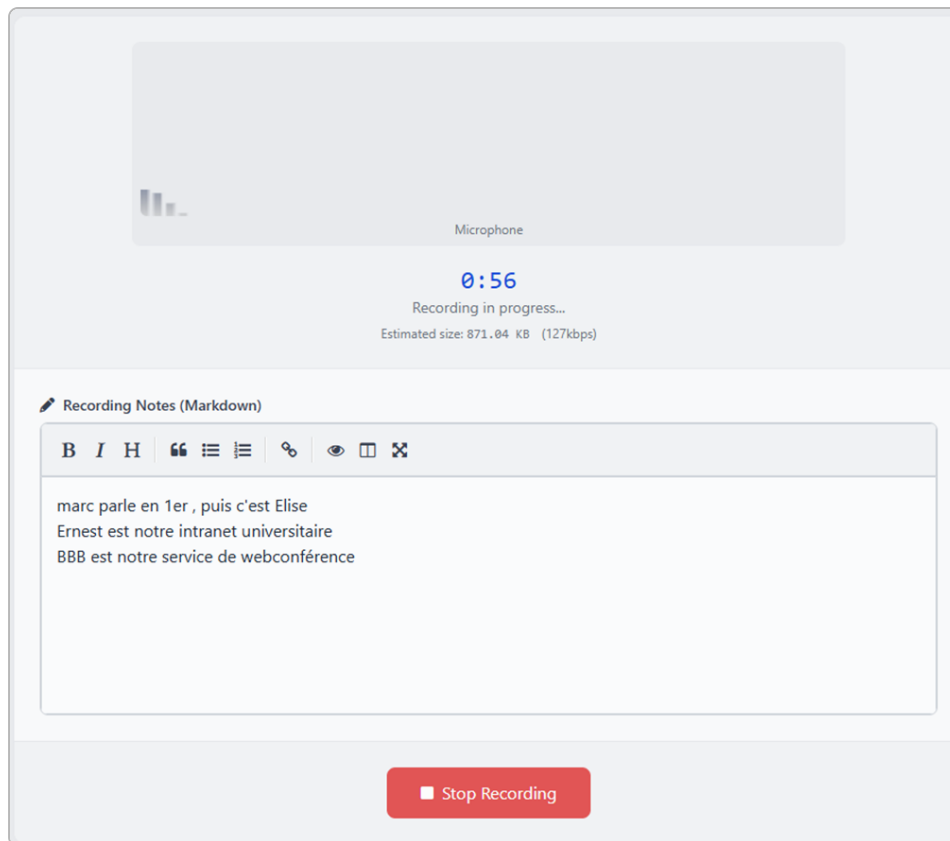


Importation de fichier : Glissez-déposez des fichiers audio préexistants (MP3, WAV, MOV, MP4) dans l'interface.

Enregistrement en direct : Cliquez sur le bouton « Plus » pour démarrer un enregistrement. Vous pouvez capturer le son du microphone, du système (audio interne) ou les deux simultanément (utile pour les visioconférences)

Conseils pratiques

- Vérifiez que les **barres de niveau audio** bougent dans l'interface pour confirmer la capture du son.
- Ajoutez des **notes contextuelles** pendant l'enregistrement pour guider la transcription.



- Pour **les réunions sensibles ou critiques**, il est recommandé d'utiliser un dictaphone externe ou un logiciel d'enregistrement local pour éviter les pertes de données en cas de coupure réseau ou de plantage, puis d'importer le fichier ultérieurement.

Exemple de logiciel d'enregistrement OBS studio

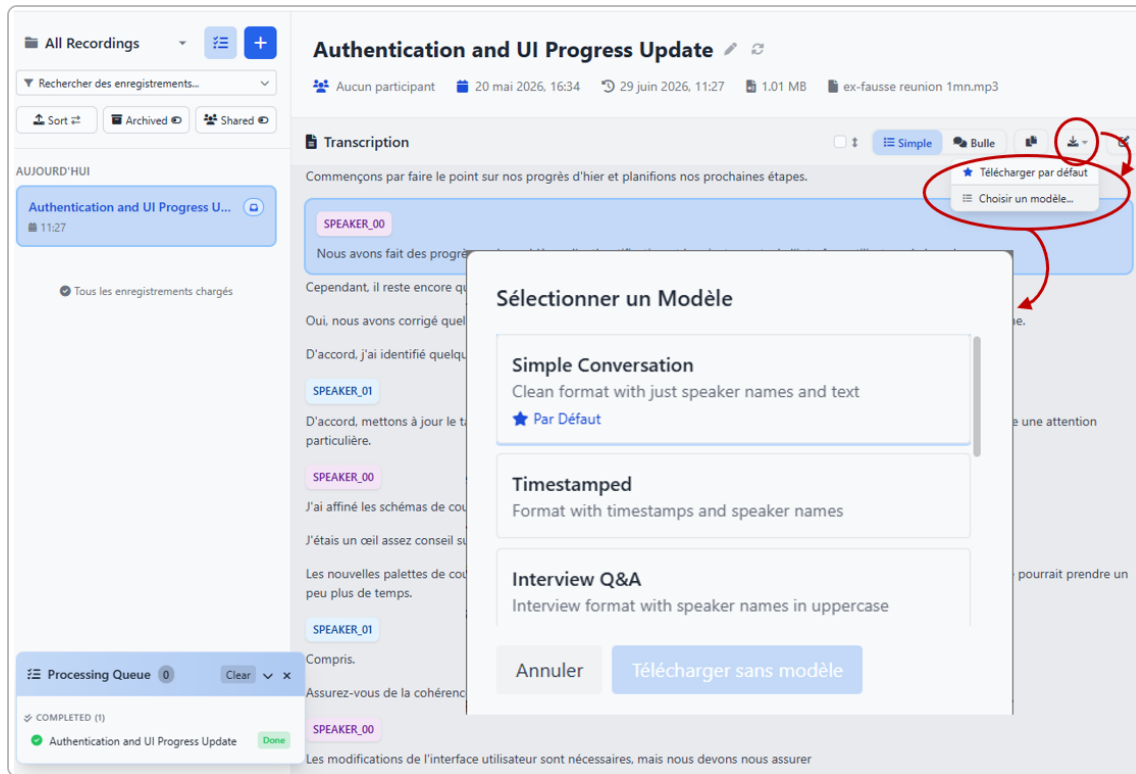
Pour les réunions sensibles, nous recommandons d'utiliser un outil d'enregistrement local comme **OBS Studio** [<https://obsproject.com/fr>]. Ce logiciel gratuit et léger permet de capturer l'audio et la vidéo du système. Cette méthode garantit la conservation des données en cas d'incident technique (coupure réseau, plantage), contrairement à l'enregistrement en ligne direct. Pour plus d'infos, vous pouvez vous rapprocher du **CCN** [<https://ccn.unistra.fr/>] (labo numérique)

Guide d'installation et d'utilisation d'OBS Studio

1. **Installation** : Téléchargez la version la plus récente d'OBS Studio. Les versions anciennes ne disposent pas de la fonctionnalité de pause nécessaire.
2. **Configuration initiale** : Au premier lancement, l'assistant de configuration s'ouvre. Sélectionnez l'option « **Optimiser pour l'enregistrement** » et suivez les étapes par défaut.
3. **Démarrage de l'enregistrement** : Cliquez sur le bouton « **Démarrer l'enregistrement** ». Le logiciel capturera simultanément le son du microphone et celui du système (audio interne), couvrant ainsi tous les participants. Le fichier généré est un format vidéo léger avec un écran noir.
4. **Gestion de la session** : En cas d'interruption ou de besoin de pause, utilisez la fonction de pause pour sauvegarder l'intégrité du fichier. L'enregistrement est horodaté automatiquement.
5. **Récupération et transfert** : Une fois l'enregistrement terminé, récupérez le fichier vidéo dans le dossier « **Vidéos** » de Windows, puis transférez-le vers la plateforme Speak.

8.3. Traitement et affinage de la transcription

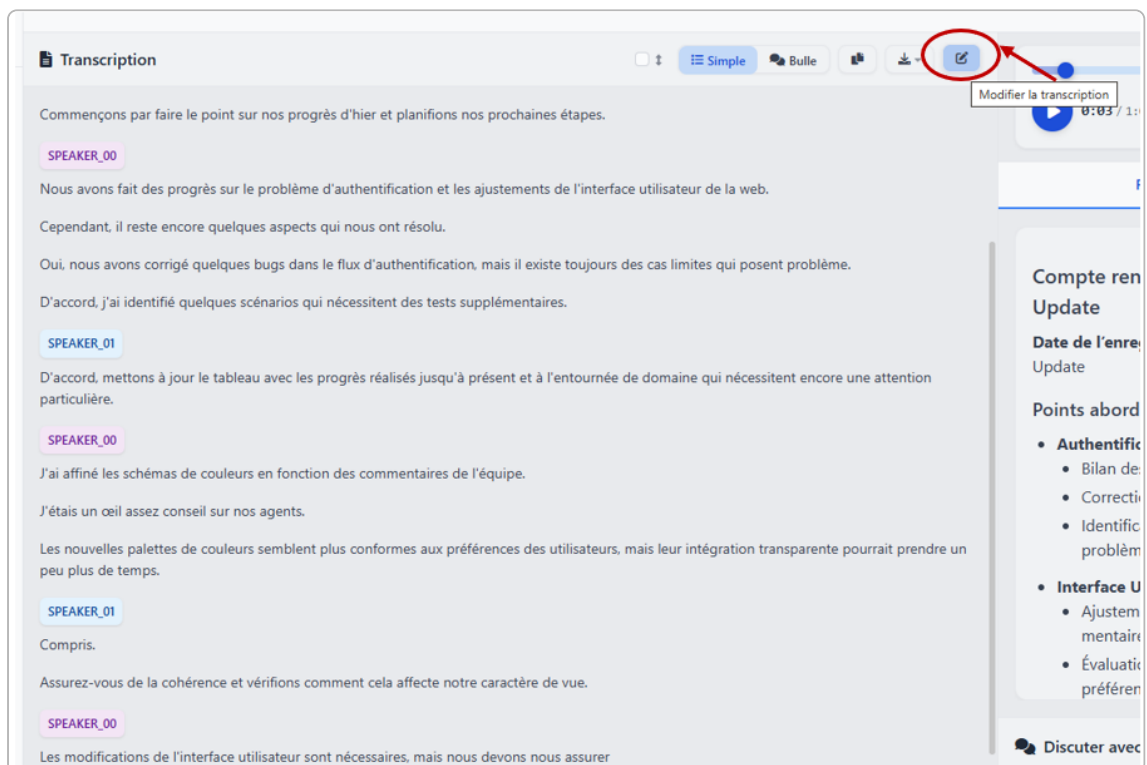
Une fois l'enregistrement terminé et téléchargé sur les serveurs de l'Université de Strasbourg, l'interface de transcription s'affiche. Procédez aux étapes suivantes pour finaliser le document :



Modification de la transcription

Procédure

1. Cliquez sur modifier la transcription



2. Dans la fenêtre « **modifier la transcription ASR** », cliquez sur **Speakr-01**

Modifier la Transcription ASR

0:00 / 1:06

ORATEUR	DÉBUT	FIN	PHRASE
SPEAKER_01	3.862	11.611	Assurons-nous que nous sommes prêts pour une publication réussie.
SPEAKER_01	11.631	11.851	Absolument.
SPEAKER_00			le point sur nos progrès d'hier et planifions nos prochaines étapes.
SPEAKER_00			progrès sur le problème d'authentification et les ajustements de l'interface utilisateur de la
SPEAKER_00			core quelques aspects qui nous ont résolu.
SPEAKER_00			jé quelques bugs dans le flux d'authentification, mais il existe toujours des cas limites qui
SPEAKER_00			quelques scénarios qui nécessitent des tests supplémentaires.
SPEAKER_01			sur le tableau avec les progrès réalisés jusqu'à présent et à l'entourée de domaine qui
SPEAKER_00	39.616	43.299	J'ai affiné les schémas de couleurs en fonction des commentaires de l'équipe.
SPEAKER_00	43.339	45.381	J'étais un œil assez conseil sur nos agents.
SPEAKER_00	45.401	53.809	Les nouvelles palettes de couleurs semblent plus conformes aux préférences des utilisateurs, mais leur
SPEAKER_01	53.829	55.43	Compris.
SPEAKER_01	55.45	59.274	Assurez-vous de la cohérence et vérifions comment cela affecte notre caractère de vue.

Ctrl+S pour enregistrer

Annuler Enregistrer

1. Sélectionnez « **Éditer les locuteurs** ».

2. Dans la fenêtre qui s'ouvre, attribuez un nom à chaque locuteur identifié, puis cliquez sur « **Save changes** ». *Note : Si l'autolabel est activé (voir [Accès et configuration initiale](#)), le système associera automatiquement ces noms aux voix pour les réunions futures.*

3. Relisez la transcription et corrigez les erreurs, notamment les acronymes et les noms propres. Vous pouvez modifier le texte **directement dans la phrase**. Utilisez le **bouton de lecture** en fin de phrase pour réécouter le segment et confirmer l'exactitude de la transcription ou l'identité du locuteur. Vous pouvez supprimer une phrase complète en cliquant sur le bouton corbeille.

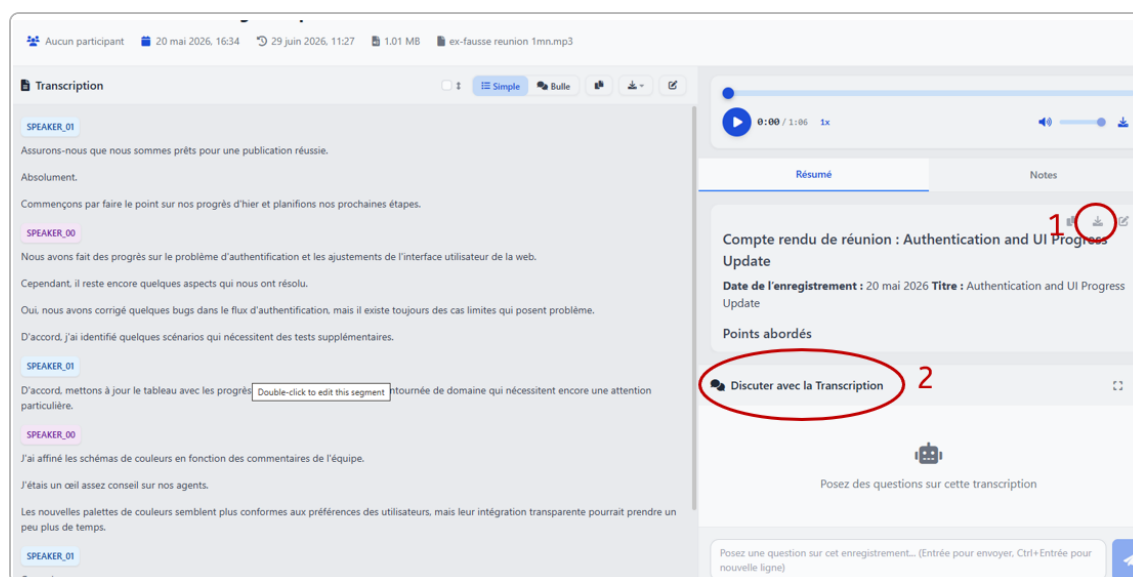
4. Cliquez sur « **Enregistrer les modifications** ».

3. Cliquez sur le bouton de téléchargement pour exporter votre transcription

Vous pouvez choisir d'utiliser le **modèle par défaut** ou de sélectionner un **modèle personnalisé**. Si vous avez activé des modèles prédéfinis (voir [Organisation par dossiers et modèles](#)), ils apparaîtront dans cette liste pour être appliqués au document.

8.4. Génération du Compte-rendu

Speaker génère automatiquement un **compte-rendu** ou un **résumé** en s'appuyant sur la transcription et les notes ajoutées.



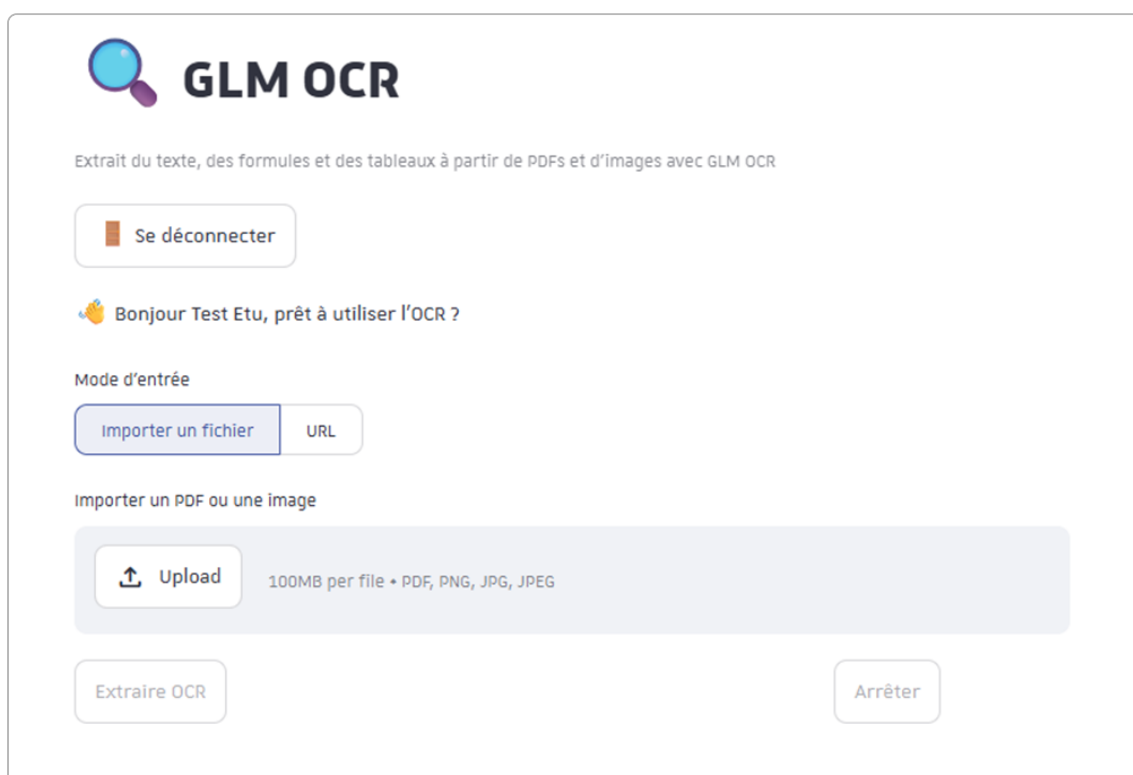
1. **Exportation** : Téléchargez le résumé final au format Word (.docx) ou copiez-le dans un autre document.
2. **Discuter avec la transcription** : Utilisez le chat pour modifier le résumé, reformuler des passages, ajuster la structure du plan ou extraire les décisions clés. Cette fonctionnalité permet également d'interroger le document pour obtenir des précisions sur son contenu.

9. GLM-OCR : conversion de documents numérisés en texte éditable

GLM-OCR est un outil de Reconnaissance Optique de Caractères (OCR *optical character recognition*) spécialisé. Il permet de transformer des documents numérisés, des images complexes, des tableaux ou des écritures manuscrites en texte éditable, recherchable et structuré.

L'utilisation de GLM-OCR est recommandée pour :

- Les images scannées ou les documents papier.
- Les écritures manuscrites.
- Les documents contenant des formules mathématiques complexes.
- Les tableaux ou mises en page nécessitant une extraction précise du contenu



The screenshot shows the GLM OCR web interface. At the top, there is a logo with a magnifying glass and the text 'GLM OCR'. Below the logo, a subtitle reads 'Extrait du texte, des formules et des tableaux à partir de PDFs et d'images avec GLM OCR'. A 'Se déconnecter' button is visible. A greeting message says 'Bonjour Test Etu, prêt à utiliser l'OCR ?'. Under 'Mode d'entrée', there are two buttons: 'Importer un fichier' (highlighted) and 'URL'. Below this, a section titled 'Importer un PDF ou une image' contains an 'Upload' button with an upward arrow icon and the text '100MB per file • PDF, PNG, JPG, JPEG'. At the bottom, there are two buttons: 'Extraire OCR' and 'Arrêter'.

Utilisation de GLM-OCR

L'importation de votre document se fait de deux manières :

1. **Importer un fichier** : Vous pouvez charger un PDF, une image (JPG, PNG, etc.) ou un scan de document papier.
2. **Utiliser une URL** : Collez le lien direct vers le document cible.

Une fois le document chargé, cliquez sur le bouton **Extraire OCR**.

Résultat de l'extraction

Le texte extrait apparaît immédiatement dans la zone de résultat. Il s'agit d'un **texte brut, prêt à être copié et collé**.

Vous avez deux options pour récupérer le contenu :

- **Copier-coller** : Sélectionnez simplement le texte dans la zone de résultat pour le copier dans un autre document.
- **Télécharger** : Un bouton situé sous le texte vous permet de télécharger le résultat directement au format .txt.

Exemple 1 : Extraction de formules mathématiques

Voici la comparaison entre le document source et le résultat de l'extraction :

Limites de fonctions usuelles en un réel

$$\lim_{x \rightarrow 0^+} \frac{1}{x} = +\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^n} = +\infty, \forall n \in \mathbb{N}^*, \lim_{x \rightarrow 0^+} \frac{1}{\sqrt{x}} = +\infty$$

$$\lim_{x \rightarrow 0^+} \frac{1}{x} = -\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^2} = +\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^n} = \begin{cases} +\infty & \text{si } n \text{ est pair} \\ -\infty & \text{si } n \text{ est impair} \end{cases}, n \in \mathbb{N}^*$$

Opérations sur les limites

Dans les tableaux qui suivent, les limites des fonctions f et g sont prises soit en $-\infty$, soit en $+\infty$, soit en un réel a . l et l' sont des nombres réels.

Lorsqu'il n'y a pas de conclusion en général, on dit alors qu'il y a un cas de forme indéterminée.

Limite d'une somme

Si f a pour limite	l	l	l	$+\infty$	$-\infty$	$+\infty$
Si g a pour limite	l'	$+\infty$	$-\infty$	$+\infty$	$-\infty$	$-\infty$
Alors $f+g$ a pour limite	$l+l'$	$+\infty$	$-\infty$	$+\infty$	$-\infty$	FI

PDF d'origine



Limites de fonctions usuelles en un réel

$$\lim_{x \rightarrow 0^+} \frac{1}{x} = +\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^n} = +\infty, \forall n \in \mathbb{N}^*, \lim_{x \rightarrow 0^+} \frac{1}{\sqrt{x}} = +\infty$$

$$\lim_{x \rightarrow 0^+} \frac{1}{x} = -\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^2} = +\infty, \lim_{x \rightarrow 0^+} \frac{1}{x^n} = \begin{cases} +\infty & \text{si } n \text{ est pair} \\ -\infty & \text{si } n \text{ est impair} \end{cases}, n \in \mathbb{N}^*$$

Opérations sur les limites

Dans les tableaux qui suivent, les limites des fonctions f et g sont prises soit en $-\infty$, soit en $+\infty$, soit en un réel a . l et l' sont des nombres réels.

Lorsqu'il n'y a pas de conclusion en général, on dit alors qu'il y a un cas de forme indéterminée.

Limite d'une somme

Si f a pour limite	l	l	l	$+\infty$	$-\infty$	$+\infty$
Si g a pour limite	l'	$+\infty$	$-\infty$	$+\infty$	$-\infty$	$-\infty$
Alors $f+g$ a pour limite	$l+l'$	$+\infty$	$-\infty$	$+\infty$	$-\infty$	FI

Après OCR : le texte extrait, prêt à être copié



Exemple 2 : Reconnaissance d'écriture manuscrite

Cet exemple illustre la capacité de l'outil à retranscrire une lettre manuscrite (lettre de l'écrivain Marcel Pagnol), préservant le contenu textuel original.

Image d'origine

Monsieur le Proviseur,

Je serais très heureux s'il m'était permis d'aller surveiller les dernières répétitions de ma pièce, "Les Marchands de Gloire", qui sera jouée au théâtre des Galeries Saint Hubert, à Bruxelles, à partir du 14 Mars 1926, par Signoret. Le directeur de ce théâtre me la demande instamment, et il me serait bien agréable d'assister à la "première", et au banquet qui m'est offert par quelques confères belges.

D'autre part, M. Delàcre, directeur du Théâtre du Marais, m'offre de jouer une autre pièce, "Jazz". Enfin, je profiterais de ma présence à Bruxelles

Après OCR : le texte extrait, prêt à être copié

Monsieur le Proviseur,

Je serais très heureux s'il m'était permis d'aller surveiller les dernières répétitions de ma pièce, "Les Marchands de Gloire", qui sera jouée au théâtre des Galeries Saint Hubert, à Bruxelles, à partir du 14 Mars 1926, par Signoret. Le directeur de ce théâtre me la demande instamment, et il me serait bien agréable d'assister à la "première", et au banquet qui m'est offert par quelques confères belges.

d'autre part, M. Delàcre, directeur du Théâtre du Marais, m'offre de jouer une autre pièce, "Jazz". Enfin, je profiterais de ma présence à Bruxelles

10. DOCLING : conversion de documents en format optimisé pour l'IA

Docling est un outil de conversion généraliste capable de transformer divers formats de fichiers (.docx, .pdf, .pptx, etc.) en **Markdown** (.md).

Le format **Markdown** est un format de texte léger qui permet d'organiser la mise en page d'un document (titres, listes, gras, etc.). Il offre un double avantage : il reste lisible pour un humain tout en étant parfaitement structuré et optimisé pour le traitement par les modèles d'IA.

Pourquoi utiliser Docling avant le téléversement ?

Lorsque vous téléversez un fichier (PDF, Word, etc.) directement dans un chatbot, celui-ci effectuera une conversion automatique en Markdown. Bien que ce format soit idéal pour l'IA, cette conversion automatique présente plusieurs limites :

- **Risques d'erreurs** : La conversion automatique peut mal interpréter les tableaux complexes ou inclure des données parasites (pieds de page, métadonnées inutiles).
- **Contrôle des ressources** : Le fichier PDF brut contient des données invisibles et du bruit (métadonnées, en-têtes) qui consomment des tokens inutilement. À l'inverse, le Markdown de Docling ne conserve que le texte structuré, réduisant le poids du fichier et préservant la fenêtre de contexte pour l'essentiel.
- **Qualité réduite** : Pour les documents contenant beaucoup d'images ou de tableaux, la compréhension par l'IA peut être dégradée par la conversion automatique.

La solution : Utiliser Docling en pré-traitement permet de garantir une analyse plus précise, des fichiers plus légers et une meilleure qualité de traitement.

Cas d'usage recommandés

Il est fortement recommandé d'utiliser Docling pour :

- Les fichiers volumineux.
- Les documents complexes contenant de nombreuses images ou tableaux.
- Les PDF lourds nécessitant une conversion préalable avant l'analyse par IA.

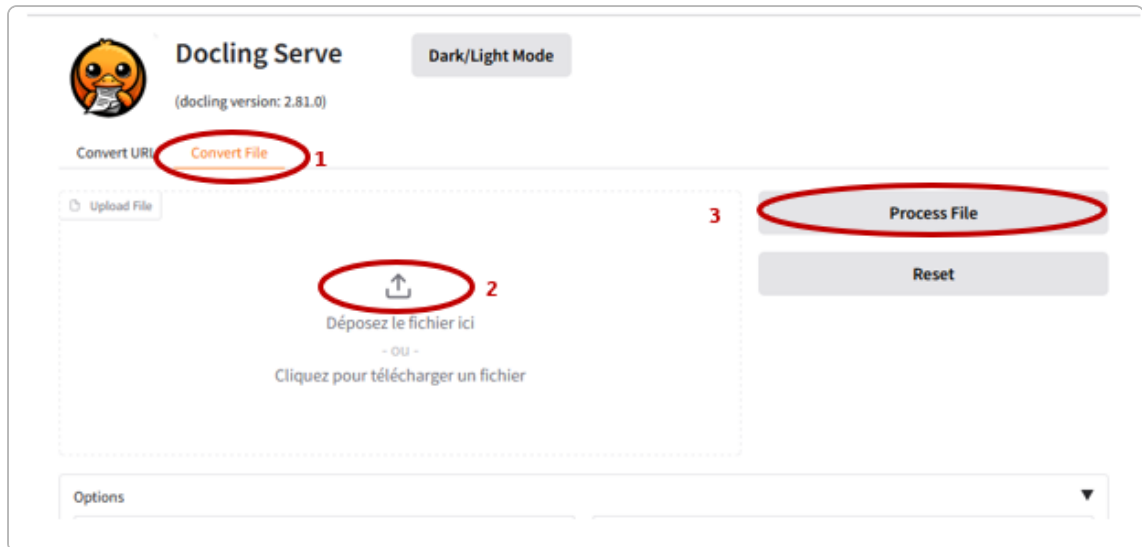
NB : Bien que l'objectif principal soit de générer un format Markdown optimisé pour l'IA, Docling permet également de transformer vos fichiers en formats JSON (pour le code/le traitement de données) ou HTML (pour la visualisation web).

Utilisation de Docling

Docling accepte une large variété de formats en entrée, notamment : PDF, DOCX, PPTX, XLSX, HTML, EPUB, WAV, ainsi que les images (PNG, TIFF, JPEG, etc.).

Procédure

1. Vous pouvez traiter un fichier local ou utiliser une URL. Pour télécharger un fichier



1. Allez dans l'onglet **"Convert File"**

2. Téléchargez votre fichier

3. Appuyez sur **"Process File"**.

2. Nous vous conseillons de **décocher** les options suivantes pour optimiser le résultat :

- *Docling (JSON)* : Inutile si vous souhaitez uniquement du texte structuré pour l'IA.
- *Enable OCR (Reconnaissance optique de caractères)* : Cette option ajoute du temps de traitement. Si vous avez besoin d'une OCR puissante pour des documents scannés, il est préférable d'utiliser un outil dédié comme **GLM-OCR**, qui est généralement plus performant.

3. Téléchargez le résultat

Le résultat apparaît en bas de page sous plusieurs formats.

1. Sélectionnez l'onglet **Markdown**
2. Cliquez sur le bouton de téléchargement à droite du fichier généré.



Syntaxe Markdown

Vous retrouverez la syntaxe Markdown standard, structurée comme suit :

- # Titre 1 à ##### Titre 6 : Pour hiérarchiser les titres.
- _italique_ : Pour mettre en valeur du texte en italique.
- **gras** : Pour mettre en gras les mots ou phrases importants.
- ***gras et italique*** : Pour un effet combiné.
- - liste : Pour créer des listes à puces.
- ::mise en valeur:: : Pour surligner ou mettre en avant une information spécifique (selon la configuration du modèle cible).

11. FAQ

11.1. Puis-je transmettre des données à caractère personnel à une IA ?

Une donnée à caractère personnel est toute information permettant d'**identifier directement ou indirectement une personne physique**. L'identification directe se fait via des informations explicites comme le **nom**, le **prénom** ou une **photographie**. L'identification indirecte se fait par le croisement de plusieurs informations, telles qu'un **numéro de téléphone**, une **adresse IP**, un **numéro de sécurité sociale** ou même des **caractéristiques physiques et sociales** qui, ensemble, rendent une personne unique au sein d'un groupe. (source : CNIL [\[https://www.cnil.fr/fr/definition/donnee-personnelle\]](https://www.cnil.fr/fr/definition/donnee-personnelle))

L'importance de la vigilance face aux IA grand public

Il est crucial de **distinguer notre IA souveraine des outils d'IA grand public** accessibles sur internet. Avec ces derniers, les données saisies sont souvent utilisées pour entraîner les modèles, stockées sur des serveurs tiers et peuvent être sujettes à des fuites de données. Une information confidentielle pourrait ainsi être restituée partiellement et involontairement à un autre utilisateur. De plus, transmettre des **données à caractère personnel concernant des tiers** (collègues, étudiants, participants à une recherche) à un prestataire d'IA externe sans contrat spécifique de traitement de données constitue une violation grave de la réglementation en vigueur sur la protection des données à caractère personnel, dont le **RGPD**.

Utilisation sur l'IA de l'Université de Strasbourg

Bien que notre service soit hébergé de manière souveraine par l'Université de Strasbourg et que vos données ne soient pas utilisées pour ré-entraîner les modèles, **nous vous recommandons de ne transmettre que des données professionnelles ou nécessaires à la conduite de vos activités**. Pour les traitements courants, privilégiez l'anonymisation ou la généralisation (par exemple, utiliser "collègue" plutôt qu'un nom propre). (Cf. [les conditions d'utilisation affichées sur le portail](https://portail.ia.unistra.fr/) [\[https://portail.ia.unistra.fr/\]](https://portail.ia.unistra.fr/))

11.2. Puis-je transmettre des données sensibles à une IA ?

Les données sensibles, appelées officiellement « **catégories particulières de données** » dans le RGPD, sont un sous-ensemble spécifique des données à caractère personnel qui révèlent des caractéristiques intrinsèquement liées à l'identité, aux convictions ou à l'état de santé d'une personne.

L'article 9 du RGPD liste exhaustivement **8 catégories** de données sensibles :

1. Données révélant l'**origine raciale ou ethnique**
2. Données révélant les **opinions politiques**
3. Données révélant les **convictions religieuses ou philosophiques**
4. Données révélant l'**appartenance syndicale**
5. **Données génétiques**
6. **Données biométriques** utilisées aux fins d'identifier de manière unique une personne physique (ex : empreintes digitales, reconnaissance faciale)
7. **Données de santé**

8. Données concernant la **vie sexuelle ou l'orientation sexuelle** d'une personne physique

source : CNIL [<https://www.cnil.fr/fr/reglement-europeen-protection-donnees/chapitre2>]

L'introduction de données sensibles (santé, convictions, orientation sexuelle, etc.) est **strictement déconseillée** sur **toute** intelligence artificielle, qu'il s'agisse d'outils grand public ou de l'outil institutionnel de l'Université de Strasbourg. Ces catégories de données bénéficient d'une protection renforcée par le **RGPD**.

Vous devez absolument éviter de saisir ce type d'informations dans les prompts. Si une tâche nécessite l'analyse de tels documents, vous devez au préalable les anonymiser complètement ou utiliser des canaux de traitement dédiés et sécurisés prévus par l'établissement.

11.3. Comment protéger ma propriété intellectuelle ?

Pour maîtriser l'exposition de vos œuvres et préserver vos droits, il est crucial de distinguer les plateformes utilisées. Vous devez absolument **éviter de téléverser vos travaux**, qu'ils soient achevés ou en cours, **sur des plateformes d'intelligence artificielle grand public ou commerciales**. Ces outils ne garantissent pas la confidentialité de vos données et pourraient les utiliser pour entraîner leurs modèles, risquant ainsi de réutiliser votre contenu sans votre consentement.

À l'inverse, **l'IA de l'UNISTRA offre un cadre sécurisé et souverain**. Hébergée exclusivement sur les serveurs de l'Université, cette plateforme assure que vos données ne sont pas utilisées pour ré-entraîner les modèles de langage. **Vos données ne quittent jamais l'écosystème universitaire**, garantissant ainsi une confidentialité totale de vos travaux.

11.4. Les données que je transmets à l'IA de l'UNISTRA sont-elles conservées ?

Les données transmises à l'IA de l'Université de Strasbourg sont conservées selon des durées spécifiques, mais vous en conservez la maîtrise totale.

- **Chatbot** : l'historique est stocké pour consultation et gestion par l'utilisateur, qui peut le supprimer à tout moment. Les conversations sont archivées après 3 mois et supprimées après 1 an. Les bases de connaissances sont conservées.
- **Speakr** : les fichiers audio sont supprimés au bout de sept jours. L'historique des transcriptions est stocké pour consultation et gestion par l'utilisateur, qui peut le supprimer à tout moment.
- **Glm Ocr** : les documents sont supprimés immédiatement après traitement.
- **Docling** : les documents sont supprimés immédiatement après traitement.
- **Assistance Moodle** : les données de conversation sont conservées à des fins de suivi et d'amélioration du service de support.
- **Statistiques anonymisées** : nombre d'utilisations, temps de réponse, etc.

11.5. Qui est responsable des erreurs ou des contenus générés par l'IA ?

Conformément aux lignes directrices sur les usages de l'IA de l'Université de Strasbourg, **l'utilisateur reste pleinement responsable** des contenus qu'il produit ou publie à l'aide d'un outil d'intelligence artificielle.

L'IA générative est un outil probabiliste qui peut commettre des erreurs, générer des "hallucinations" (informations inventées mais présentées comme vraies) ou reproduire des biais présents dans ses données d'entraînement. Elle n'a ni conscience, ni jugement moral, ni intelligence au sens propre.

Pour garantir la qualité et l'exactitude des documents publiés, les principes suivants doivent être appliqués :

- **Vérification systématique** : Tout contenu produit par une IA doit faire l'objet d'une relecture critique et d'un travail de vérification indispensable pour écarter toute aberration, erreur factuelle ou biais, en particulier les faits, les chiffres et les références.
- **Rigueur scientifique** : L'utilisateur assume la pleine responsabilité éditoriale des publications. L'usage de l'IA ne décharge pas l'auteur de son obligation de rigueur et d'éthique scientifique.
- **Transparence** : Conformément aux principes de l'Université, l'usage de l'IA dans une production doit être documenté et traçable.

L'IA est un assistant qui propose ; c'est à vous de valider et de garantir l'intégrité académique et professionnelle de votre production.

11.6. Quelle est la différence entre l'IA de l'Unistra et un outil comme ChatGPT ?

L'IA conversationnelle de l'Université de Strasbourg se distingue des assistants grand public par son ancrage institutionnel et ses principes éthiques.

- **Puissance brute vs Efficience** : La différence oppose les **paramètres** (la "culture générale") à la **fenêtre de contexte** (la "mémoire immédiate"). Si les outils commerciaux misent sur la puissance brute avec des milliers de milliards de paramètres et des fenêtres immenses (jusqu'à un million de tokens), l'IA de l'Université de Strasbourg privilégie l'efficience. Nous utilisons des modèles open-source optimisés allant de 120 milliards de paramètres (GPT-OSS) à 31 milliards (Gemma), dotés de fenêtres de contexte de 128K à 256K tokens. L'IA Unistra s'appuie sur des modèles très performants, bien que parfois légèrement en retrait sur la "puissance brute" pure, mais suffisants pour la majorité des usages académiques. **La force de l'IA Unistra réside dans son adéquation aux besoins universitaires et sa sobriété, plutôt que dans la taille pure du modèle.**
- **Transparence et hallucinations** : Les modèles Unistra sont open-source, contrairement aux modèles propriétaires qui fonctionnent comme des « boîtes noires ». La transparence du code source permet une meilleure traçabilité et un contrôle accru sur les biais potentiels. De plus, l'approche Unistra, via le RAG (Retrieval-Augmented Generation) permet un **contrôle accru** sur les sources de vérité et une réduction des hallucinations en s'appuyant sur des documents institutionnels vérifiés.
- **Souveraineté et hébergement** : Contrairement aux outils commerciaux, l'infrastructure de l'IA Unistra est hébergée dans les datacenters de l'Université, garantissant la maîtrise totale des données et une indépendance technologique.

- **Protection des données** : L'outil est conçu pour protéger les données des utilisateurs . De plus, les données transmises ne sont pas utilisées pour ré-entraîner les modèles de langage, contrairement à de nombreuses plateformes gratuites.
- **Engagement éthique et environnemental** : L'université s'appuie sur un Plan IA fondé sur la sobriété et la frugalité numérique, en utilisant notamment des modèles plus petits pour limiter l'impact écologique
- **Accès sécurisé** : Le service est réservé au personnel de l'Université, assurant un cadre de travail professionnel et conforme à la réglementation en vigueur.

11.7. L'utilisation de l'IA a-t-elle un impact écologique ?

L'intelligence artificielle engendre un impact environnemental significatif, qui se manifeste tant par la consommation importante d'électricité et d'eau pour le refroidissement des serveurs que par l'empreinte carbone liée à la fabrication du matériel et à l'extraction des matières premières.

Consciente de ces enjeux, l'Université de Strasbourg a adopté une **approche responsable et sobre**.

En tant que signataire de la [Charte pour un numérique responsable](https://www.unistra.fr/fr/node/5654) [https://www.unistra.fr/fr/node/5654], l'Université intègre **la sobriété et l'efficacité énergétique** comme piliers de sa stratégie IA. Notre service s'articule autour de cinq leviers majeurs pour minimiser cet impact :

1. **Utilisation de modèles open source pré-entraînés** L'Université de Strasbourg privilégie des modèles de langage open source, **sélectionnés pour leur efficacité plutôt que pour leur taille brute**. Nous utilisons systématiquement le plus petit modèle capable de répondre de manière satisfaisante à la demande, évitant ainsi la course au modèle le plus volumineux, souvent inutilement gourmand. Ce faisant, nous n'effectuons aucun entraînement de modèle. Il convient en effet de distinguer la phase d'entraînement (processus très énergivore) de celle de l'utilisation, appelée inférence.
2. **Hébergement local et éco-responsable** : l'IA est déployée directement **dans notre datacentre** installé sur le campus de l'Université de Strasbourg. Ce datacentre a été conçu dans une **approche éco-responsable** : il utilise la géothermie sur nappe pour le refroidissement et récupère la chaleur générée afin de contribuer à chauffer les bâtiments de l'université. Cette approche locale permet aussi de mutualiser les GPU et évite des transferts de données inutiles.
3. **Transparence et suivi de la consommation** : Notre installation fait l'objet d'une surveillance permanente. Nous disposons d'un **suivi précis et détaillé de la consommation énergétique** de nos services, nous permettant d'analyser et d'optimiser en continu l'empreinte carbone de nos services IA.
4. **Mutualisation des pratiques et des infrastructures au niveau de l'ESR** : L'Université de Strasbourg s'implique fortement dans différentes initiatives au niveau local et national. Ces initiatives permettent de **mutualiser les efforts, les pratiques et les infrastructures**. Le but du projet ILaaS (Inference LLM as a Service) est de fournir aux établissements de l'Enseignement Supérieur et de la Recherche (ESR) des moyens techniques pour mettre en œuvre l'Intelligence Artificielle en répondant aux défis importants de soutenabilité, de résilience, de confiance et de sobriété numérique. Il offre une infrastructure d'inférence mutualisée.
5. **Usage raisonné et pertinent de l'IA** : L'IA doit être mobilisée de manière ciblée, uniquement **lorsque sa valeur ajoutée** (gain de temps, amélioration de la qualité) **est clairement identifiée** comme spécifié par le [Plan IA Unistra](https://ernest.unistra.fr/jcms/1540474408_DBFileDocument/fr/plan-ia-2026-03) [https://ernest.unistra.fr/jcms/1540474408_DBFileDocument/fr/plan-ia-2026-03]

. Il est également recommandé de privilégier des outils spécialisés, souvent plus économes en énergie que les solutions généralistes.

11.8. Comment citer l'usage de l'IA dans mes travaux ?

L'utilisation de l'intelligence artificielle dans vos productions académiques et professionnelles doit impérativement respecter les principes de **transparence** et de **responsabilité** ([Charte du numérique](https://ernest.unistra.fr/jcms/152696585_UDSPages/fr/charte-du-numerique-de-l-universite-de-strasbourg) [https://ernest.unistra.fr/jcms/152696585_UDSPages/fr/charte-du-numerique-de-l-universite-de-strasbourg]). Il est donc essentiel de ne jamais intégrer de contenus générés par l'IA sans en mentionner explicitement l'origine.

En l'absence de directives spécifiques de la part de l'UNISTRA, Il vous appartient de choisir la méthode qui vous semble la plus appropriée pour indiquer l'usage de ces outils.

Pour vous guider, nous vous suggérons de vous inspirer des pratiques établies par d'autres institutions académiques, telles que :

Une signalétique visuelle :

L'utilisation de logos ou de pictogrammes spécifiques permettant d'identifier rapidement les parties du document assistées par IA (en s'inspirant, par exemple, des pratiques adoptées par certaines universités québécoises comme **l'Université de Montréal** ou les logos de **Martine Peters**, professeure en sciences de l'éducation et directrice du Partenariat universitaire sur la prévention du plagiat à l'Université du Québec en Outaouais.)



Source : Martine Peters , de l'Université du Québec en Outaouais



source : Les bibliothèques de l'université de Montréal

L'IAGraphie :

Pour les documents scientifiques, il est possible d'ajouter une « IAGraphie » à la suite de votre bibliographie (Galleron, 2024 [<https://shs.hal.science/halshs-04756419/document>]). Cette section permet de lister les modèles utilisés, les versions et la nature des requêtes (prompts) ayant permis d'aboutir aux résultats.

Déclaration de méthode :

Préciser dans une note de bas de page ou dans la section méthodologique la part d'intervention de l'IA (ex: aide à la structuration du plan, correction syntaxique ou génération de synthèses).

11.9. Où puis-je obtenir de l'aide pour utiliser l'IA ?

Pour obtenir de l'aide ou des conseils concernant l'utilisation de l'IA à l'Université de Strasbourg ou pour des projets IA, vous pouvez vous adresser aux structures suivantes selon vos besoins :

- Pour **l'assistance utilisateur** : rendez-vous sur <https://support.unistra.fr/>
- Pour **l'accompagnement de projets IA** : contactez la Direction du Numérique (DNum), qui assure la coordination stratégique et technique, à l'adresse dnum-coord-ia@unistra.fr [<mailto:dnum-coord-ia@unistra.fr>].
- Pour **la découverte et l'expérimentation** : rendez-vous au [Lab numérique du CCN](https://ccn.unistra.fr/locaux/lab-numerique/) [<https://ccn.unistra.fr/locaux/lab-numerique/>], situé au rez-de-chaussée de l'Atrium, où des permanences sont organisées pour apprendre et échanger.

Ces structures vous accompagnent dans l'appropriation des outils et la mise en place de vos projets.