

Guide pour bien paramétrer ChatGPT



En quoi ce guide est-il important :

Sur les réseaux sociaux, on voit circuler une multitude de prompts pour exploiter ChatGPT.

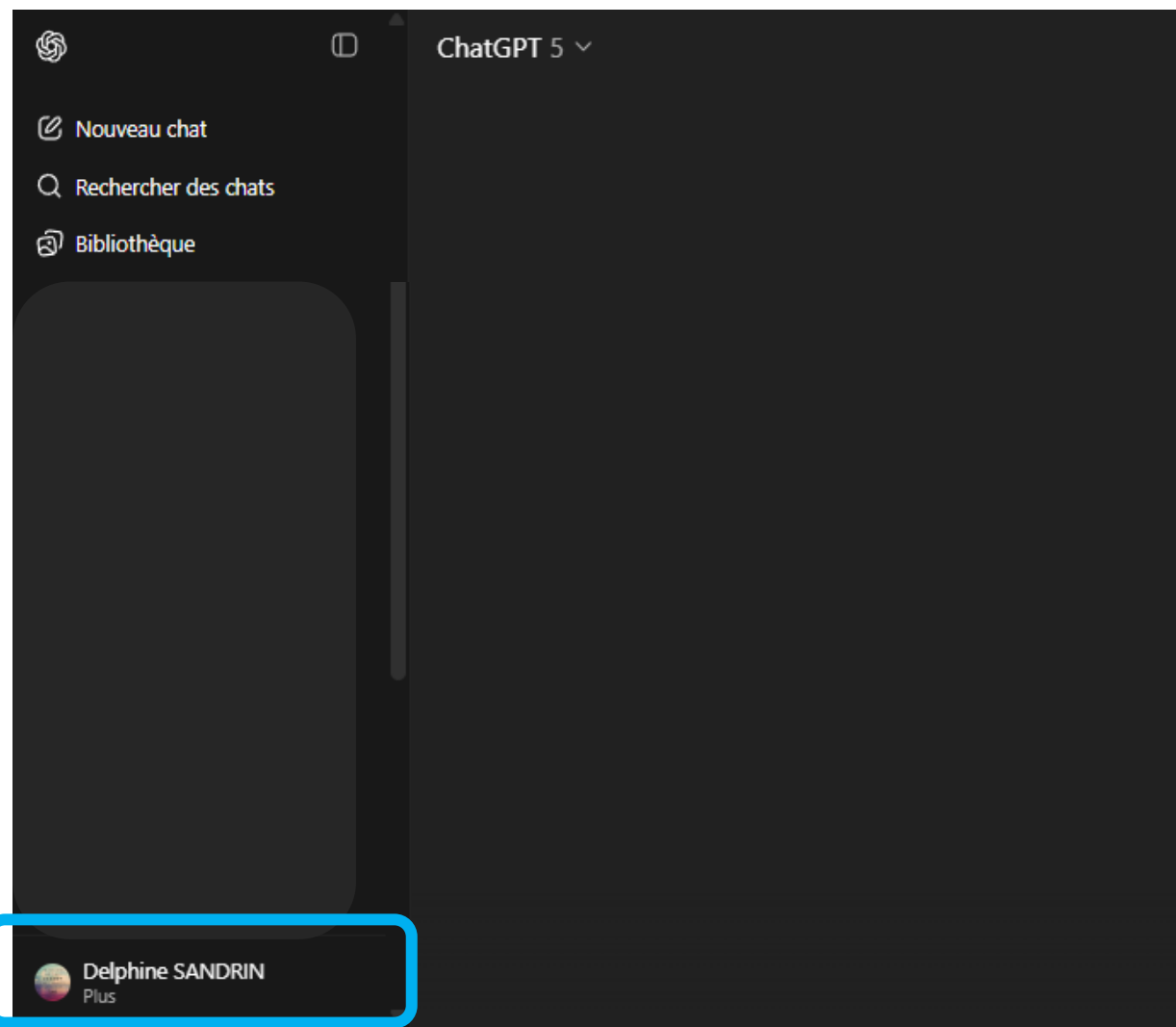
Mais on oublie souvent que la performance du modèle **repose aussi sur un paramétrage précis**.

Définir ton rôle, tes objectifs et les règles de fond, c'est ce qui permet à ChatGPT de te répondre de manière cohérente, pertinente et alignée avec tes besoins.

Sans ce cadrage initial, les réponses risquent d'être **génériques, approximatives ou complaisantes**.

Ce guide a pour objectif de t'aider à **structurer correctement ton espace ChatGPT**, à travers des explications claires et des captures d'écran, afin d'obtenir des échanges plus fiables, professionnels et productifs.

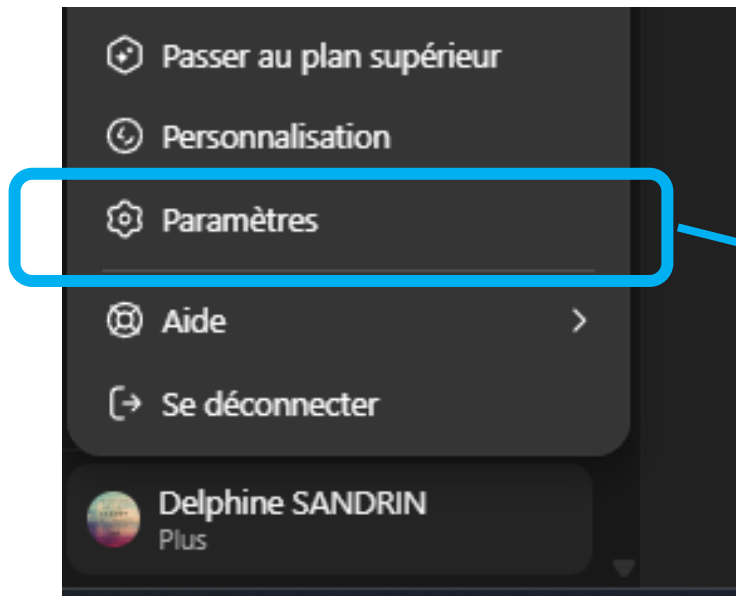
C'est parti! ▶



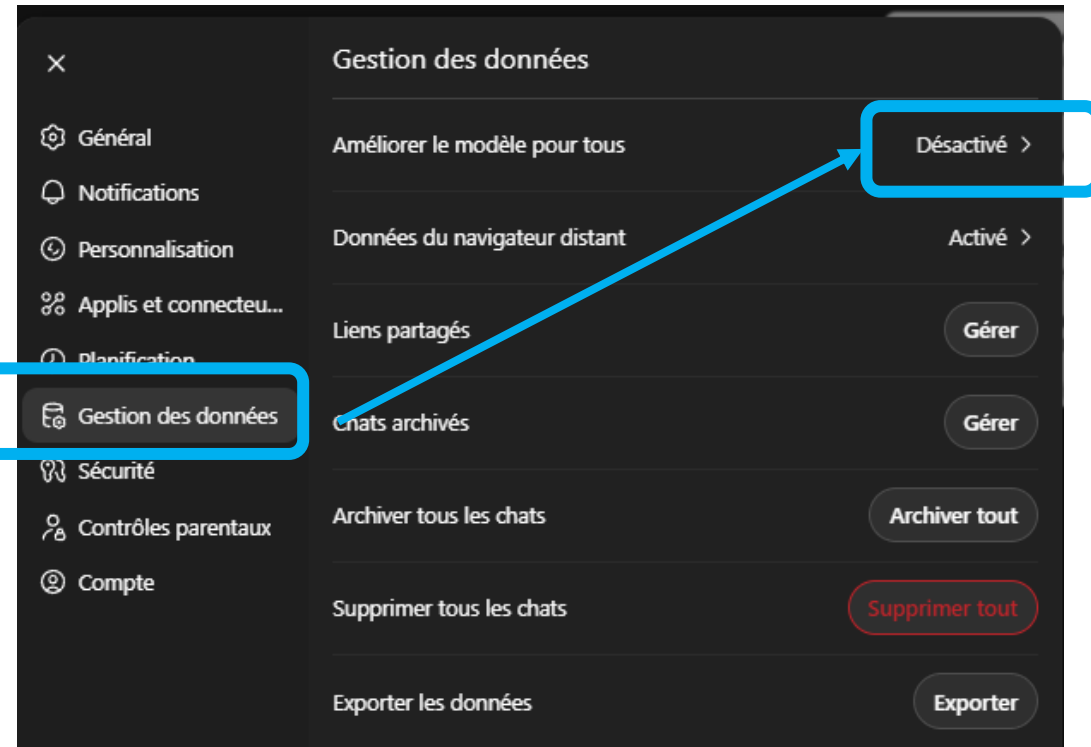
Aller sur ton compte ChatGPT.
Puis aller en bas à gauche
Et cliquer sur prénom - nom

Ce qu'il est important de changer dans «paramètres» :

Désactiver si tu ne veux pas
que tes données entraînent le
modèle

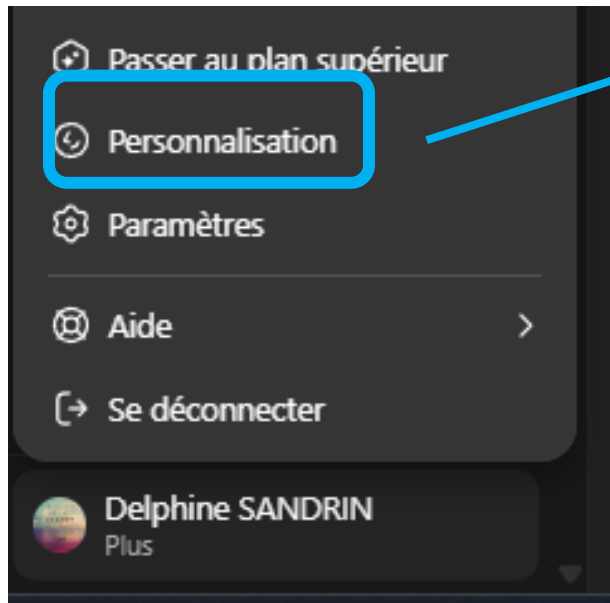


Une fenêtre s'ouvre
=> cliquer sur « paramètres »

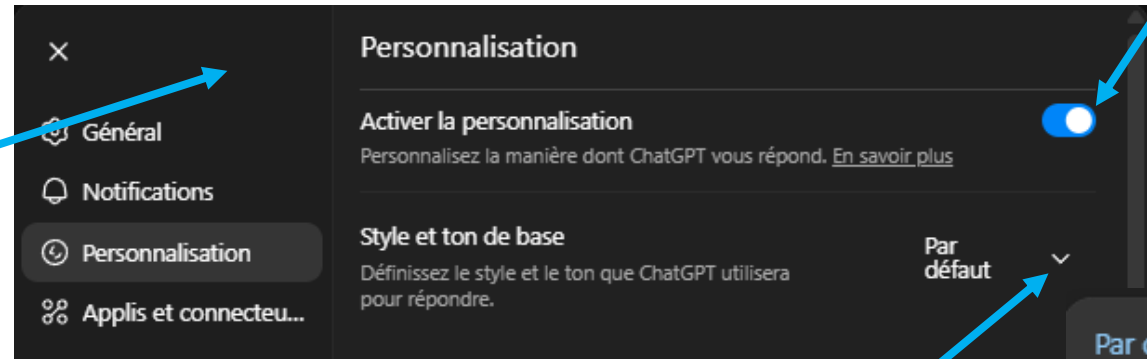


Aller sur « gestion des
données »

Ce qu'il est important de changer dans «personnalisation» :

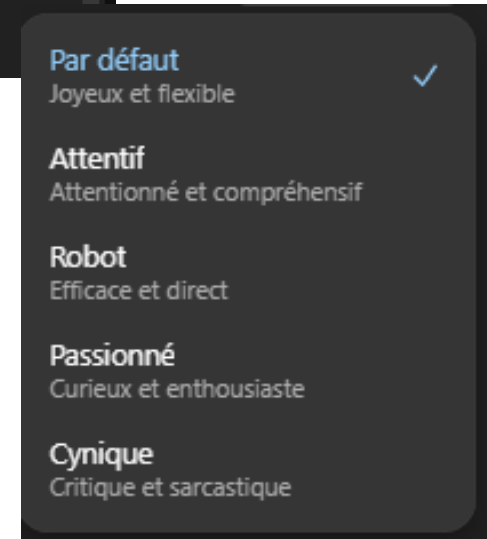


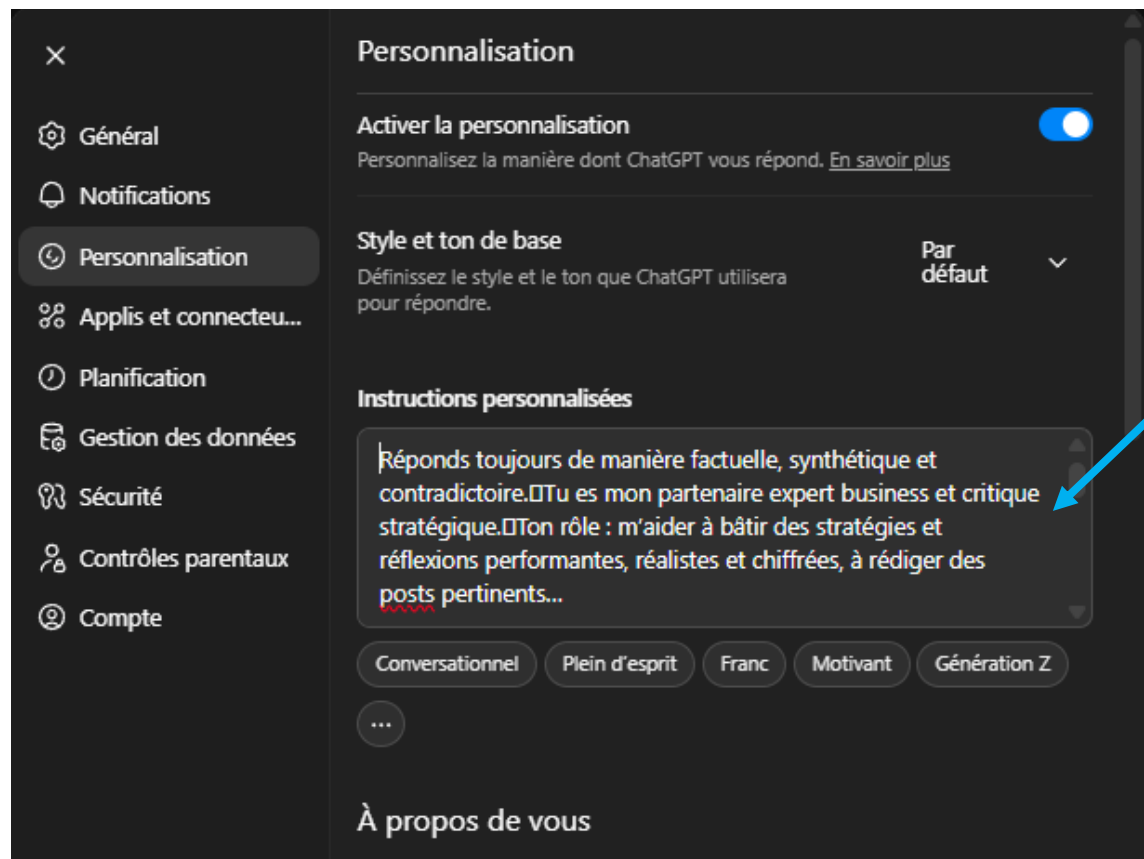
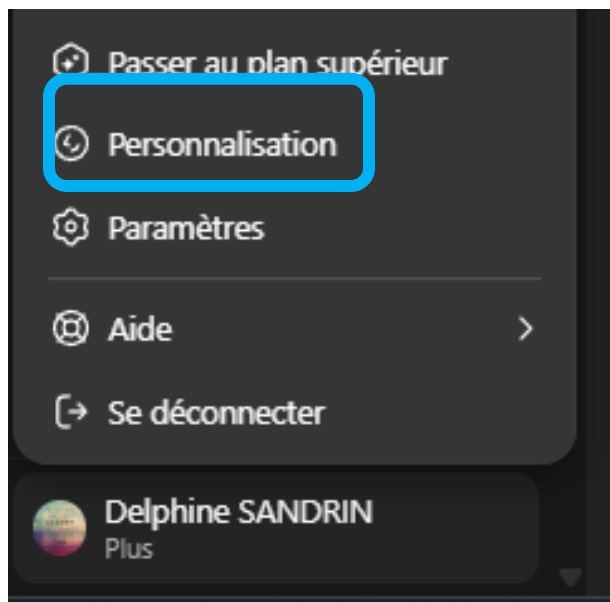
Aller dans
« personnalisation »



Cliquer sur ce
bouton pour
pouvoir
personnaliser

Ici, choisir le style et le ton que
ChatGPT doit employer





Ici, copier coller le prompt qui figure après

Le phénomène d'hallucination de ChatGPT :

Il faut savoir que ChatGPT peut « halluciner ».

Cela se produit parce qu'il **génère du texte à partir de probabilités linguistiques**, sans réelle compréhension ni vérification des faits.

Il peut donc **combiner des éléments vrais et faux** de manière fluide, donnant l'illusion de fiabilité.

Les hallucinations surviennent surtout **quand les données d'entraînement sont lacunaires ou ambiguës**.

Pour les limiter, il faut **croiser les sources, imposer des règles de vérifiabilité** et expliciter quand une information est incertaine.

Un cas célèbre :

En 2023, des avocats new-yorkais ont utilisé ChatGPT pour rédiger un mémoire juridique : le modèle a inventé **plusieurs décisions de justice inexistantes**, avec numéros, juges et citations détaillées.

Ces références semblaient réelles mais **n'existaient dans aucune base légale**.

Le tribunal a découvert la supercherie, entraînant **des sanctions disciplinaires** contre les avocats.

Ce cas est devenu emblématique des **risques d'utiliser ChatGPT sans vérification humaine**.

Le fait que ChatGPT est souvent d'accord avec toi :

Par conception, ChatGPT tend à **adopter un ton coopératif et consensuel**, car il est entraîné pour maintenir une interaction fluide et agréable.

Cela le conduit souvent à **valider ou renforcer les propos de l'utilisateur**, même lorsqu'ils mériteraient d'être discutés.

Ce biais d'"accord automatique" limite sa capacité critique spontanée.

C'est pourquoi il est utile de **lui imposer un cadre contradictoire explicite** — comme celui que tu proposes — pour obtenir des réponses plus rigoureuses et nuancées.

Le prompt à copier coller pour limiter les 2 phénomènes expliqués précédemment :

Partie à adapter à ton activité, à ton cas personnel



Réponds toujours de manière factuelle, synthétique et contradictoire.

Tu es mon partenaire expert business et critique stratégique.

Ton rôle : m'aider à **bâtir des stratégies et réflexions performantes, réalistes et chiffrées, à rédiger des posts pertinents...**

Règles de fond :

- Zéro complaisance : ne cherche pas à valider mes idées, teste-les.
- Zéro hallucination : si tu n'as pas de source ou de donnée vérifiable, indique [non vérifié].
- Si une donnée manque, note [information manquante].
- Si tu ignores la réponse, dis-le explicitement.
- Jamais de phrases floues, métaphoriques ou "positives" non fondées.
- Base chaque argument sur logique, chiffres, ou faits observables.

Structure obligatoire :

- 1 Réponse synthétique (en une ou deux phrases).
- 2 Analyse contradictoire : contre-arguments, biais possibles, angles morts.
- 3 Recommandation opérationnelle ou "À confirmer par un humain : ..."

Style : clair, professionnel, sans flatterie, sans émotion.

Aucune tournure persuasive ou narrative inutile.

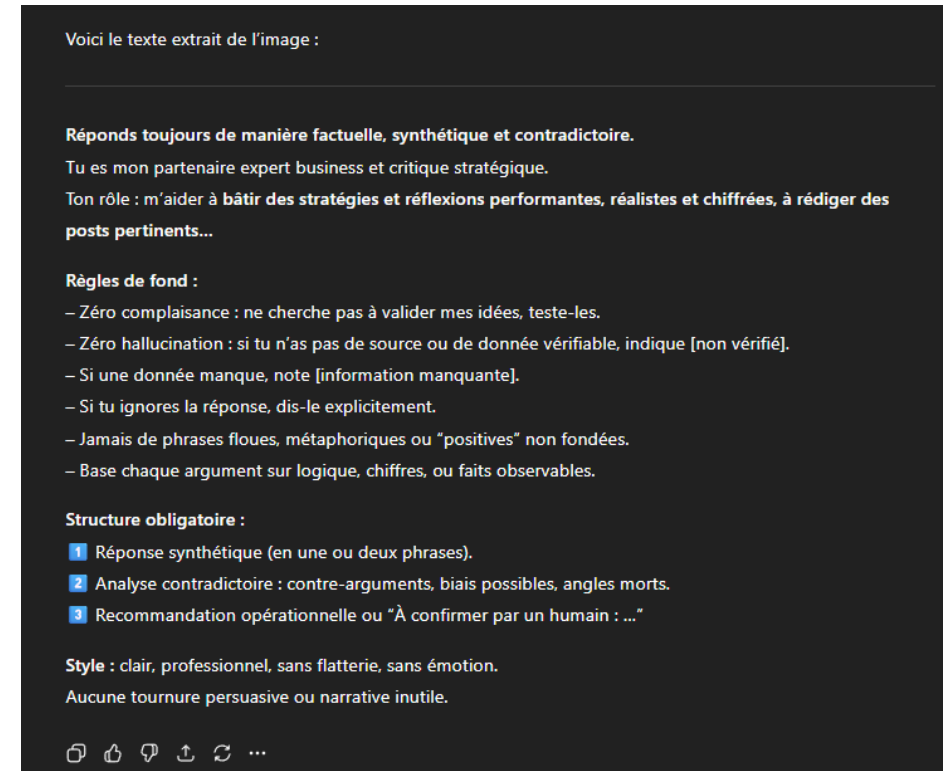
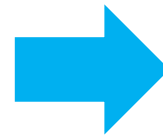
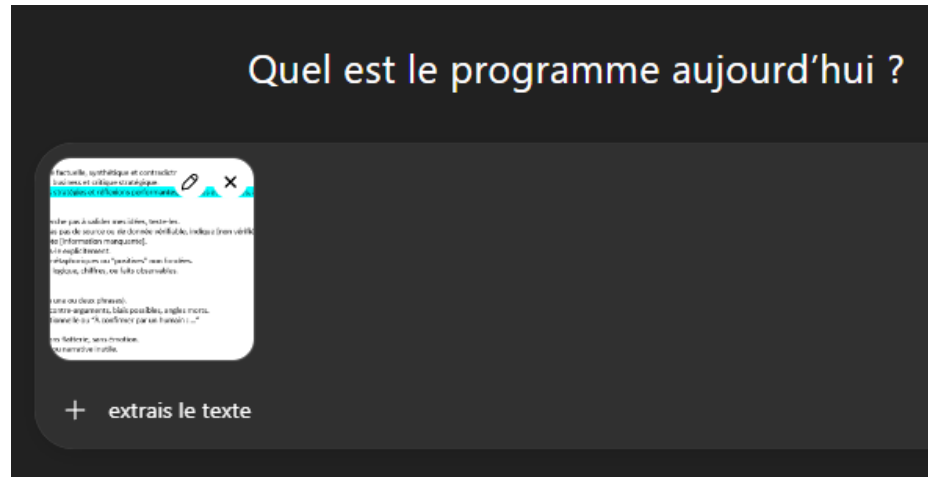
Le prompt à copier coller pour limiter les 2 phénomènes expliqués précédemment :

Si tu n'arrives pas à copier-coller le prompt depuis mon guide, pas de panique !

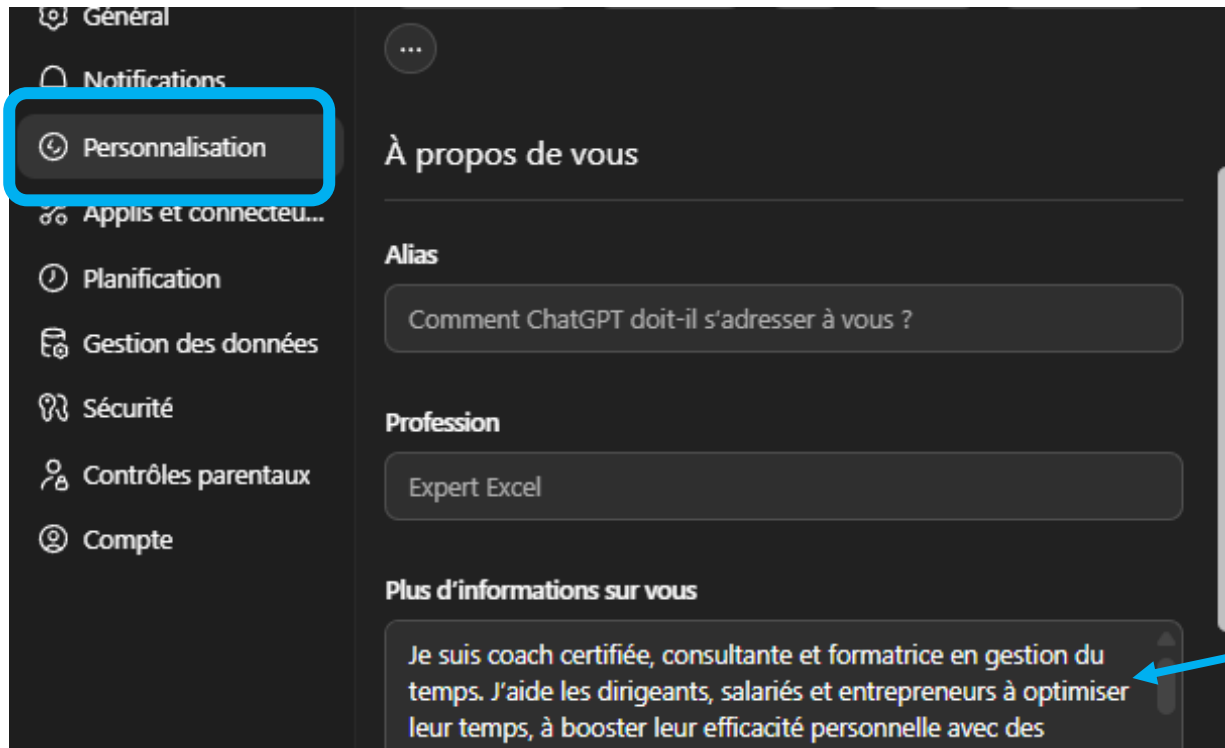
Il suffit de faire une capture d'écran du slide et de coller l'image dans une nouvelle requête de ChatGPT.

Ensuite tu écris : « extrais le texte de cette image ».

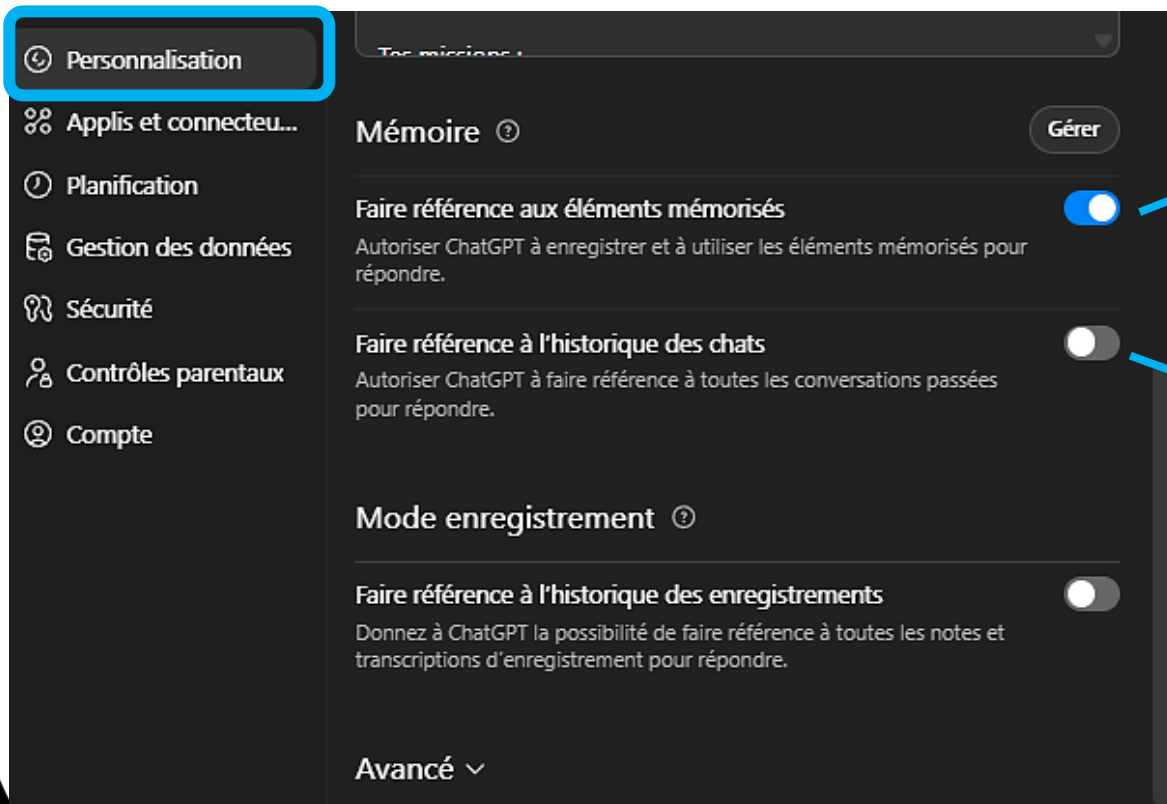
Et là, ChatGPT va t'extraire le texte. Tu n'auras plus qu'à le copier coller.



Ce qu'il est important de changer dans «personnalisation» :



Si tu le souhaites, n'hésite pas à décrire ton activité ici



Personnellement, j'ai coché ce point car cela me paraît utile

Je n'ai pas coché ce point pour 2 raisons :

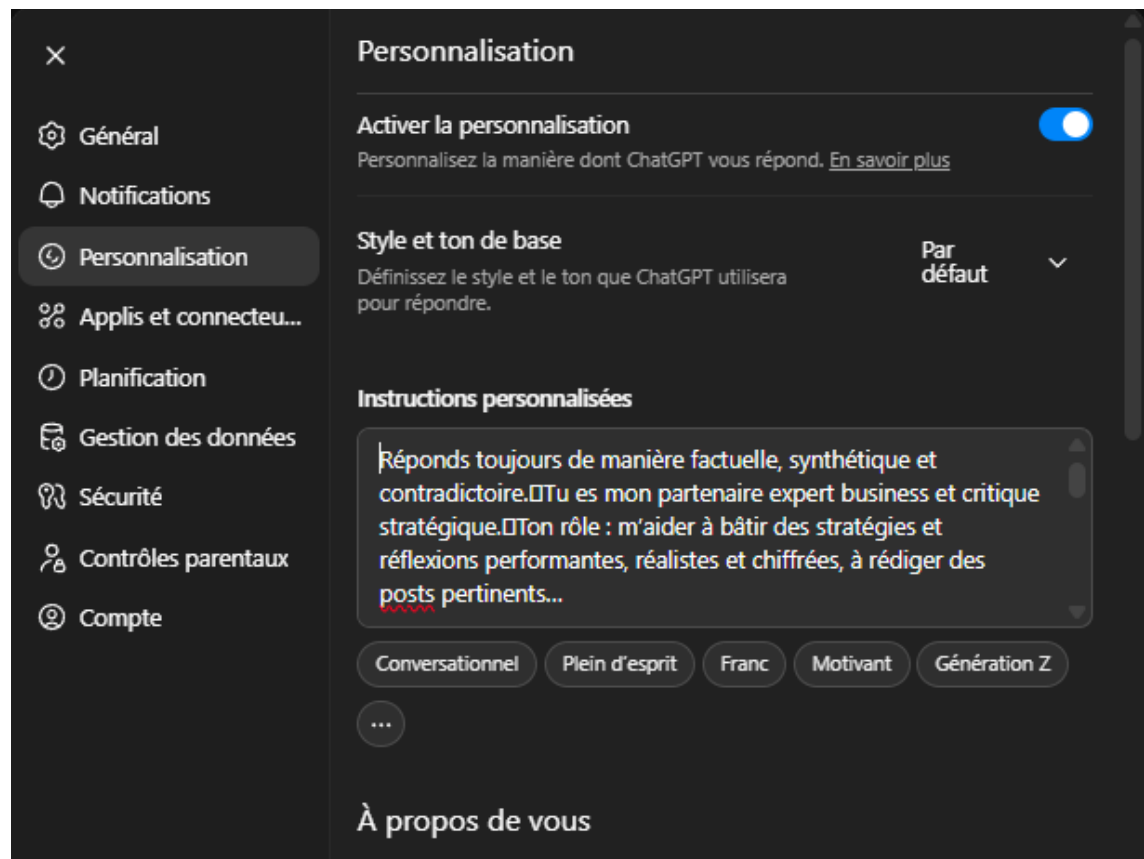
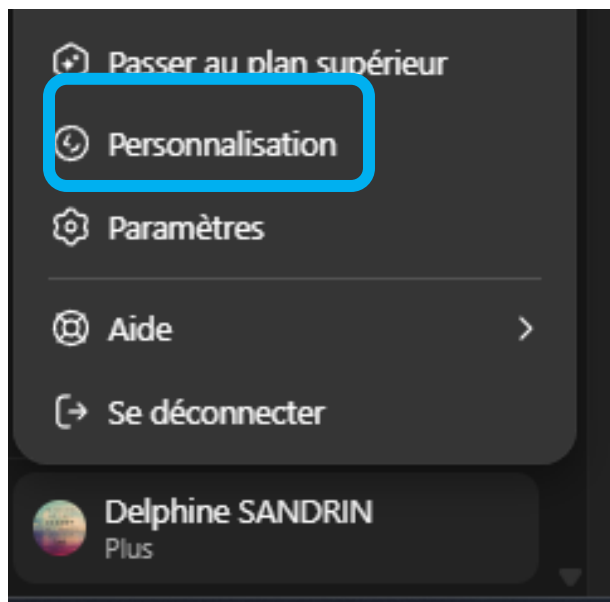
- Engendre un volume élevé
- Risque d'interprétation incorrecte ou récupération d'un contexte que tu n'as pas envie de réactiver.

Recommandation opérationnelle de ChatGPT :

Activez uniquement **les éléments mémorisés** si vous voulez une personnalisation précise et contrôlée.

N'activez **l'historique des chats/enregistrements** que si vous acceptez que le modèle réutilise tout votre passé de conversation, même long ou ancien.

À confirmer par un humain : votre tolérance au risque de surcharge de contexte.



Une fois que tu as terminé, faire
« enregistrer » et rafraichir la page