

Bien utiliser l'IA : les 6 principes essentiels

Les 6 réflexes qui font la différence entre un usage moyen et un usage maîtrisé de l'IA. À regarder avant de te lancer dans le reste de la formation : tu vas économiser des heures, payer beaucoup moins cher, et obtenir des réponses pertinentes dès le premier coup.

Source : <https://prisme.one/plateforme/ia-juillet-2026/semaine-1-parler-a-lia>

Objectif

Adopter les 6 réflexes qui font la différence entre un usage moyen et un usage maîtrisé de l'IA, économisant du temps et de l'argent au passage.

Output attendu

Réflexes plan mode + nouvelle conversation par objectif + MCP-check ancrés au quotidien. Au moins une session terminée par /done.

Bien utiliser l'IA : les 6 principes essentiels

Avant d'attaquer le contenu pédagogique (comprendre ton second cerveau, installer tes outils, créer tes skills), il y a une couche fondamentale à poser : **comment bien parler à l'IA**. Six principes simples qui vont changer du tout au tout la qualité de ce que tu obtiens, et qui te font économiser énormément de temps et d'argent au passage.

Bien utiliser l'IA : les fondamentaux

LOOM

Bien utiliser l'IA : les fondamentaux

<https://www.loom.com/share/7b753107da104f039b35d473e6a8fdee>

Clarté de pensée, plan mode, fenêtre de contexte, gestion des MCP, amélioration continue des skills, /done.

L'objectif de cette leçon : que tu partes de cette base et que tu évolues vers un usage qui t'est propre, en comprenant de plus en plus finement les limites et les capacités de l'IA. C'est ce qui va affiner ta façon de lui communiquer des informations et des demandes.

Principe 1 : Ne jamais déléguer la clarté de pensée à l'IA

C'est le premier réflexe à ancrer. Quand tu commences à prompter une IA pour qu'elle réalise un travail, **tu dois avoir une clarté de pensée absolue** sur ce que tu veux. Pas l'IA. Toi.

Anti-pattern

Laisser l'IA décider pour toi

Tu lances un prompt vague genre "fais-moi quelque chose de bien sur X" et tu acceptes ce qui sort. Résultat : l'IA hallucine ton goût, choisit à ta place, et tu te retrouves avec quelque chose de générique qui ne te ressemble pas.

Bon réflexe

Clarifier d'abord, prompter ensuite

Tu sais ce que tu veux, ton goût, ton intention, ta sensibilité. Tu le formules clairement, puis tu fais exécuter. Le résultat est juste, du premier coup.

Et si tu n'as pas encore les idées claires ?

Tu peux utiliser l'IA pour **t'aider à clarifier**, sans qu'elle décide à ta place. Le bootcamp inclut un skill dédié pour ça : **/thinking-partner**.

```
/thinking-partner
```

Tu lui dis l'objectif de ta session de réflexion, il devient un partenaire socratique qui pose les bonnes questions pour clarifier tes pensées, sans te donner de réponse prématurée. La clarté reste de toi, l'IA t'aide juste à l'extraire.

Règle simple : la clarté de pensée, le goût, l'attention, la sensibilité, ça repart toujours de toi. L'IA exécute, elle n'invente pas ton intention.

Principe 2 : Toujours commencer en Plan mode

Quand tu démarres une session de travail, **commence systématiquement en plan mode**. C'est le mode où l'IA a un fonctionnement restreint : elle peut lire des fichiers, appeler des outils pour enrichir son contexte, mais elle **ne peut pas écrire ni exécuter**.

Étape 1

Plan mode

L'IA lit, explore, appelle des outils. Elle ne touche à rien.



Étape 2

Elle propose un plan

Plan structuré : "voilà ce que je vais faire, dans cet ordre".



Étape 3 : STOP

Tu valides (ou tu ajustes)

C'est ton checkpoint. Tu alignes l'IA avant qu'elle exécute.



Étape 4

Build mode

L'IA exécute le plan validé. Liste de to-do qui apparaît au-dessus de la barre.

Pourquoi c'est puissant

Le plan mode te permet de faire des tâches **en one-shot**. C'est-à-dire que si tu clarifies bien avant, tu évites toutes les itérations inutiles qui viennent ensuite. Au lieu de "ah non, pas comme ça, refais", tu valides le plan, l'IA exécute, c'est juste du premier coup.

Concrètement, l'IA propose un plan du genre :

Plan proposé :

1. Comprendre la structure actuelle du site
2. Créer la nouvelle page
3. Rédiger le texte
4. Te faire valider le texte
5. Intégrer sur le site

Tu valides, et l'affichage passe en mode build avec ta to-do list visible. Tu vois exactement où l'IA en est, étape par étape.

Principe 3 : Comprendre la fenêtre de contexte

C'est probablement le concept le plus important à intégrer pour bien utiliser l'IA. Tous les modèles ont une **fenêtre de contexte** : la quantité maximale d'information qu'ils peuvent garder en mémoire dans une conversation.

Les tailles courantes en 2026

100 k

Petits modèles

~75 k mots

200 à 260 k

Standard 2026

~150 à 200 k mots

500 k à 1 M

Frontière haut de gamme

~1,3 M mots

+1 M

Expérimental

Très coûteux

Pourquoi c'est crucial : le coût computationnel et la perte de précision

Cette mémoire est **extrêmement coûteuse en termes de calcul**. Chaque nouveau token qui arrive déclenche deux multiplications avec **tous les tokens précédents**, et ce **60 fois d'affilée** (une fois par couche du modèle). On parle de **centaines de millions de multiplications par mot généré**.

Pour rendre ça gérable à grande échelle, les constructeurs d'IA utilisent des astuces. Par exemple, sur DeepSeek :

- Une couche sur deux ne lit que **1 000 tokens sur 1 million** (les plus pertinents)
- Les tokens sont **compactés par blocs de 128** (l'IA garde une idée globale de chaque bloc, pas les détails)

Résultat : **plus ta conversation est longue, plus l'IA devient imprécise**. Si tu es au token 900 000 dans une conversation très longue, l'IA va commencer à **se perdre**, oublier les instructions du début, halluciner sur les détails. C'est ce qui te frustre quand tu sens que "l'IA ne suit plus".

La solution : une conversation = un objectif

Règle d'or

Ne pas tout enchaîner dans une seule conversation longue. Ouvre une nouvelle conversation à chaque changement d'objectif.

Tentation classique : tu commences ta journée avec "tiens, j'aimerais clarifier ma journée", puis "ok, du coup commençons cette tâche", puis "ah, continuons là-dessus"... C'est agréable parce qu'on a l'impression que l'IA nous connaît. Mais ça pollue ta fenêtre de contexte avec tout le journaling de la veille, les détails du repas, etc. Au bout d'un moment, l'IA devient lente et hallucine.

Le pattern à adopter

1. Conversation A : Planification de la journée

Tu clarifies tes priorités, l'IA t'aide à structurer.

↓

2. Fin de la conversation A

Tu dis : "Génère-moi un brief court pour transmettre à la prochaine conversation."

↓

3. Conversation B : Tâche du jour (par exemple créer une présentation)

Tu colles le brief, tu pars en plan mode pour clarifier, puis build mode pour exécuter en one-shot.

Bonus : **les conversations courtes coûtent beaucoup moins cher** que les longues, parce que tu ne refais pas tourner la fenêtre saturée à chaque message. Ton portefeuille te dit merci.

Principe 4 : Attention aux MCP activés

Voici un point que la plupart des gens ignorent et qui peut **multiplier ton coût par 8** sans que tu t'en aperçoives.

L'expérience qui parle

Tu ouvres une nouvelle session, tu actives 4 MCP (par exemple Firecrawl, Playwright, Notion, Linear), et tu envoies un simple "test". Concrètement, "test" c'est **un seul mot**, donc à peu près **1 token**. Avec le prompt système, tu devrais être à **environ 1 000 tokens** au total.

Sauf que tu te retrouves à... **47 000 tokens d'entrée**.

Avec 4 MCP activés

47 000 tokens

Pour envoyer juste "test". Si tu es à la consommation, ça te coûte déjà 2 centimes pour rien.

Avec 0 MCP activé

6 000 tokens

Pour exactement la même requête. Moins de bruit, plus de pertinence, beaucoup moins cher.

Pourquoi cette différence

Quand tu actives un MCP, l'IA reçoit non seulement l'accès à l'outil, mais aussi **une description détaillée de chaque outil exposé par ce MCP**. Notion MCP par exemple, c'est 30 à 50 tools, chacun avec sa description, ses paramètres, ses exemples. Multiplié par 4 MCP, tu te retrouves avec des dizaines de milliers de tokens uniquement de description, **avant même que ta vraie question soit traitée**.

Les bonnes pratiques

- **Désactive tous les MCP au démarrage** d'une nouvelle conversation. Pars de zéro.
- **Active uniquement ce dont tu as strictement besoin** pour la conversation en cours.
Tu écris ta newsletter ? Active Firecrawl seul. Tu travailles sur du code ? Active Playwright seul.
- **Vérifie en haut à droite de l'écran** les MCP actifs avant chaque nouvelle session. Le réflexe "MCP-check" doit devenir aussi naturel que vérifier ton micro avant un call.
- **Bénéfice secondaire** : ton contexte de base (AGENTS.md, Moi.md) prend une part beaucoup plus grande dans le contexte total. Donc ton style, tes préférences, tes process pèsent plus dans la réponse de l'IA.

Principe 5 : Faire évoluer tes skills par métacognition

Tes skills (les commandes IA que tu crées en semaine 2) ne sont jamais parfaits du premier coup. La bonne nouvelle : **ils peuvent s'améliorer tout seuls**, à condition que tu leur en donnes l'occasion.

Le réflexe à ancrer

À la fin de chaque session où tu as utilisé un skill, et où tu as dû reprendre l'IA plusieurs fois ("non, pas ce titre", "reformule plus court", "manque ça"), demande-lui ceci :

```
Fais de la métacognition sur tous les points où je t'ai repris
dans cette
session. Mets à jour le skill [nom-du-skill] pour intégrer ces
consignes
dans ton fonctionnement de base. La prochaine fois, je ne devrais
plus
avoir à te les redire.
```

L'IA va lire son propre skill, repérer ce qui manquait, et le réécrire avec les nouvelles règles. À la session suivante, le skill sera meilleur. Au bout de 3 ou 4 itérations, il est fiable.

Et pour les usages pas encore en skill ?

Si tu as un usage que tu fais "à la main" pour la première fois (par exemple écrire une vidéo YouTube sans avoir de skill /youtube), tu peux le transformer en skill **à partir de la conversation où tu l'as fait** :

```
À partir de tout ce qu'on s'est dit dans cette conversation,
génère-moi
un skill pour qu'on puisse refaire ce type de tâche plus
facilement la
prochaine fois. Utilise /create-skill pour respecter les bonnes
normes
et la structure standard.
```

La commande /create-skill est déjà incluse dans ton coffre. Elle te guide pour créer un skill qui respecte le bon format (frontmatter, méthode, exemples, règles). Tu pars d'une vraie conversation, tu en extrais une méthode réutilisable.

Principe 6 : /done à la fin de chaque session

Le dernier réflexe, peut-être le plus simple. **À la fin de chaque session, tape /done.**

Ce que ça fait :

- L'IA fait un bilan court de ce qui s'est passé dans la session.
- Elle identifie ce qu'elle a appris sur toi, tes préférences, tes projets.
- Elle met à jour les contextes pertinents (Moi .md, fichier de phase de vie, note du projet en cours, AGENTS.md si besoin).

L'effet cumulé est puissant : **ton contexte holistique reste à jour en permanence**. Chaque conversation enrichit ton coffre. Au bout d'un mois, l'IA te connaît mieux que la plupart de tes collègues. Au bout de trois mois, c'est un vrai bras droit.

Sans /done, ton contexte stagne. Avec /done après chaque session, il s'enrichit en continu. C'est probablement le plus petit effort pour le plus gros effet long terme.

En résumé

1.
Ne délègue jamais ta clarté de pensée. Le goût et l'intention repartent toujours de toi.
2.
Commence toujours en plan mode. Valide le plan, puis passe en build mode pour exécuter en one-shot.
3.
Une conversation = un objectif. Brief de transition entre deux conversations, pas tout dans la même.
4.
Active uniquement les MCP nécessaires. Sinon tu paies 8x plus cher pour des descriptions d'outils que tu n'utilises pas.
5.
Améliore tes skills par métacognition. À chaque session, "améliore ton skill avec ce que je t'ai dit".
6.
/done à la fin de chaque session. Ton contexte holistique s'enrichit en continu.

C'est un excellent point de départ. Avec ces six réflexes, tu vas avoir des résultats radicalement meilleurs qu'avec un usage "ChatGPT classique". À partir de là, tu vas affiner ta façon de communiquer avec l'IA au fil du temps, et comprendre de mieux en mieux ses limites et ses capacités.

- ☐ J'ai compris pourquoi la clarté de pensée doit rester de moi (l'IA exécute, n'invente pas l'intention)
- ☐ Je sais utiliser /thinking-partner pour clarifier mes idées avant de prompter
- ☐ Je commence systématiquement mes sessions en plan mode et je valide le plan avant build mode
- ☐ Je comprends pourquoi une conversation longue dégrade la précision (compaction, oublis)
- ☐ Je crée une nouvelle conversation par objectif au lieu de tout enchaîner dans la même
- ☐ Je vérifie les MCP activés en haut à droite et je désactive ceux dont je n'ai pas besoin
- ☐ J'ai testé d'améliorer un skill via 'fais de la métacognition sur cette session'
- ☐ Je tape /done à la fin de chaque session pour enrichir mes contextes

Vidéos

LOOM

Bien utiliser l'IA : les fondamentaux

<https://www.loom.com/share/7b753107da104f039b35d473e6a8fdee>

Extraction verbatim depuis Prisme One · <https://prisme.one/plateforme/ia-juillet-2026/semaine-1-parler-a-lia>