



FORMATION IA

IA generative

LLM, tokens, prompts, RAG, agents,
multimodal, evaluation, couts et securite.

DanielCraft - Livre debutant

Sommaire

Clique un chapitre ou un sous-chapitre pour y aller.

1. Chapitre 1 - Salut, c'est quoi l'IA generative ?

- Ce que ce n'est pas
- Ce que tu vas savoir faire
- Comment lire ce livre
- Petite histoire
- Erreur classique
- En vrai
- A toi
- Zoom : generative vs predictive
- Petite scene DanielCraft
- Ce que "50 pages" doivent changer chez toi

2. 02-histoire-simple.md

- Pourquoi ca compte pour toi
- Quelques jalons (version poche).
- Ce qui a change dans le quotidien
- Ce qui n'a pas change
- Erreur classique
- En vrai
- A toi
- Du laboratoire au telephone
- Ce que l'histoire enseigne aux prudents

3. 03-types-ia.md

- 1) Prediction et classement
- 2) Generation de contenu
- 3) Multimodal
- 4) Agents et automatisations
- Ou ranger ce que tu utilises
- Comment choisir mentalement
- Erreur classique
- En vrai
- A toi
- Cas limites utiles
- Carte pour arbitrer un achat d'outil

4. Chapitre 4 - Les LLM : tokens, contexte, temperature

- Ce que fait vraiment un LLM
- Tokens : l' unite de lecture et d'ecriture
- Fenetre de contexte : ce que le modele "voit" maintenant
- Temperature : creatif ou strict
- ChatGPT et la famille 2026
- Forces et faiblesses typiques
- Erreur classique

- [En vrai](#)
- [A toi](#)
- [Tokens : exemples concrets](#)
- [Contexte long : pas une memoire magique](#)
- [Temperature et alternatives](#)

5. [05-prompts.md](#)

- [Le squelette du prompt utile](#)
- [System prompt : le cadre durable](#)
- [Iterer sans tourner en rond](#)
- [Techniques simples qui marchent](#)
- [Ce qu'il ne faut pas faire](#)
- [Erreur classique](#)
- [En vrai](#)
- [A toi](#)
- [Anatomie d'un system prompt solide](#)
- [Prompts de relecture](#)
- [Bibliotheque personnelle](#)
- [Scene de terrain \(developpee\)](#)
- [Pieges subtils](#)
- [Lien avec le reste du livre](#)
- [Pour aller plus loin sans te perdre](#)

6. [Chapitre 6 - Limites et hallucinations : l'IA peut inventer](#)

- [Pourquoi ca arrive](#)
- [Autres limites importantes](#)
- [Comment reduire les inventions](#)
- [Exemple concret \(invente\) : chiffre invente](#)
- [Evaluer une reponse \(idee simple\)](#)
- [Erreur classique](#)
- [En vrai](#)
- [A toi](#)
- [Typologie rapide des inventions](#)
- [Quand le modele refuse](#)
- [Culture d'equipe anti-hallucination](#)
- [Scene de terrain \(developpee\)](#)
- [Pieges subtils](#)
- [Lien avec le reste du livre](#)
- [Pour aller plus loin sans te perdre](#)

7. [07-outils-quotidien.md](#)

- [Ecrire et reformuler](#)
- [Resumer](#)
- [Ideer sans te noyer](#)
- [Apprendre et s'entrainer](#)
- [Integrer sans tout automatiser](#)
- [Erreur classique](#)

- En vrai
- A toi
- Semaine type d'un pilote
- Modeles de messages
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

8. Chapitre 8 - Multimodal : image, audio, video

- Image en entree : decrire, extraire, expliquer
- Image en sortie : generer un visuel
- Audio : dicter, transcrire, resumer
- Video : encore plus lourd
- Brief multimodal utile
- Erreur classique
- En vrai
- A toi
- Qualite d'entree = qualite de sortie
- Droits et deepfakes
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

9. Chapitre 9 - RAG et agents : documents et actions

- RAG : brancher le modele sur tes documents
- Agents : enchaîner des actions vers un but
- Comment cadrer un agent (checklist mentale).
- RAG + agents : la combo
- Erreur classique
- En vrai
- A toi
- Conception d'un agent en une page
- RAG : hygiene documentaire
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

10. 10-ethique-rgpd.md

- Minimiser
- Anonymiser et pseudonymiser
- Bases RGPD (idee debutant).
- Ethique au-dela de la loi
- Securite des donnees (ponte vers le chapitre securite).
- Erreur classique

- En vrai
- A toi
- Scenarios a refuser
- Transparence utile
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

11. Chapitre 11 - Travail et metiers : l'IA amplifie le jugement

- Ce qui monte en valeur
- Ce qui se commoditise
- Trois portraits
- Comment parler de l'IA a un client / collegue
- Se former sans paniquer
- Erreur classique
- En vrai
- A toi
- Competences a cultiver cette annee
- Management et IA
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

12. Chapitre 12 - Mini-projet : un usage IA utile et verifie

- Etape 1 - Choisir la tache
- Etape 2 - Preparer le brief
- Etape 3 - Produire
- Etape 4 - Evaluer
- Etape 5 - Industrialiser legerement
- Variante avancee (optionnelle)
- Erreur classique
- En vrai
- A toi
- Criteres de reussite du mini-projet
- Scene de terrain (developpee).
- Pieges subtils
- Lien avec le reste du livre
- Pour aller plus loin sans te perdre

13. Chapitre 13 - A retenir (l'essentiel)

- Les idees solides
- La boucle DanielCraft
- Ce que tu peux oublier
- Petite checklist de poche
- A toi

- Phrase talisman
- Note de rythme
- Pour aller plus loin sans te perdre

14. Chapitre 14 - Atelier : prompts et system prompts

- Exercice 1 - Mauvais vs bon (10 min)
- Exercice 2 - Mode strict vs creatif (10 min)
- Exercice 3 - System prompt (15 min)
- Exercice 4 - Anti-hallucination (10 min)
- Livrable
- Erreur a eviter
- Variante avantee
- Note de rythme
- Pour aller plus loin sans te perdre

15. Chapitre 15 - Atelier : workflow du quotidien

- Exercice 1 - Cartographier (10 min)
- Exercice 2 - Construire le workflow (15 min)
- Exercice 3 - Mini-RAG manuel (10 min)
- Exercice 4 - Point agent (5 min)
- Livrable
- Conseil DanielCraft
- Mesure sur 7 jours
- Note de rythme
- Pour aller plus loin sans te perdre

16. Chapitre 16 - Atelier : evaluer et verifier une reponse

- Pourquoi evaluer
- Exercice 1 - Grille a 6 criteres (10 min)
- Exercice 2 - Blind test (15 min)
- Exercice 3 - Chasse a l'hallucination (10 min)
- Livrable
- Erreur a eviter
- Revue a deux
- Note de rythme
- Pour aller plus loin sans te perdre

17. Chapitre 17 - Choisir un outil et comprendre les couts

- Criteres utiles
- L'idee des couts API
- Protocole de choix en une heure
- Quand payer
- Erreur classique
- En vrai
- A toi
- Lire une grille tarifaire API sans paniquer
- Scene de terrain (developpee)

- [Pieges subtils](#)
- [Lien avec le reste du livre](#)
- [Pour aller plus loin sans te perdre](#)

18. [Chapitre 18 - Securite : comptes, secrets, pouvoirs](#)

- [Comptes et acces](#)
- [Secrets](#)
- [Donnees sensibles](#)
- [Pouvoirs limites pour les agents](#)
- [Ingenierie sociale et arnaques](#)
- [Erreur classique](#)
- [En vrai](#)
- [A toi](#)
- [Incident : que faire](#)
- [Scene de terrain \(developpee\)](#)
- [Pieges subtils](#)
- [Lien avec le reste du livre](#)
- [Pour aller plus loin sans te perdre](#)

19. [Chapitre 19 - Bonnes pratiques : la routine du pilote](#)

- [Avant](#)
- [Pendant](#)
- [Apres](#)
- [Hebdomadaire](#)
- [Mensuel](#)
- [Equipe \(si tu n'es pas solo\)](#)
- [Erreur classique](#)
- [A toi](#)
- [Anti-liste des mauvaises habitudes](#)
- [Scene de terrain \(developpee\)](#)
- [Pieges subtils](#)
- [Lien avec le reste du livre](#)
- [Pour aller plus loin sans te perdre](#)

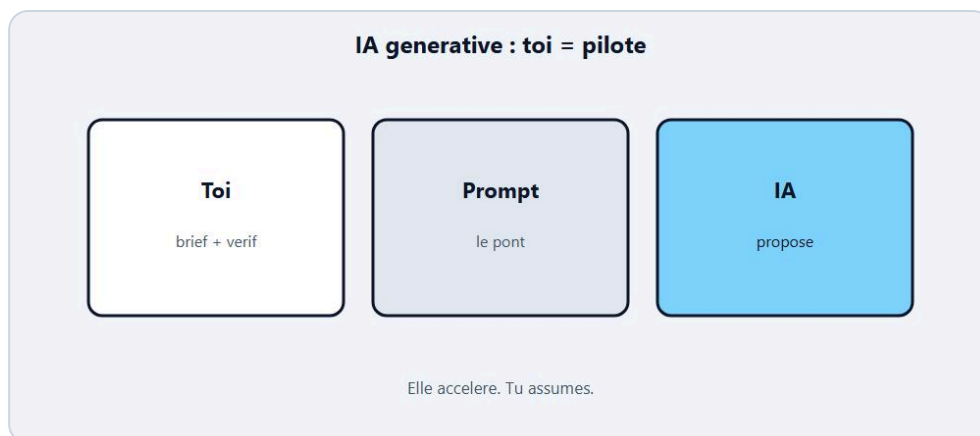
20. [Quiz final](#)

- [Questions](#)
- [Corriges](#)
- [Questions bonus \(auto-eval\)](#)
- [Note de rythme](#)
- [Pour aller plus loin sans te perdre](#)

21. [Chapitre 21 - Bravo](#)

- [La suite](#)
- [Un dernier rappel](#)
- [Engagement de 30 jours](#)
- [Note de rythme](#)
- [Pour aller plus loin sans te perdre](#)

Chapitre 1 - Salut, c'est quoi l'IA generative ?



L'IA generative : un outil. Toi : le pilote.

L'**IA generative**, ce n'est pas un robot qui "pense" comme toi. Ce n'est pas non plus de la magie noire cachee dans un nuage. C'est une famille d'outils informatiques capables de produire du texte, des images, du son, du code, parfois de la video, a partir d'une demande que tu formules. En 2026, quand quelqu'un dit "j'ai utilise l'IA", il parle le plus souvent d'un assistant conversationnel : tu ecris, il repond. Derriere, il y a des modeles entraines sur d'énormes quantites de donnees. Pour toi, le geste reste simple : poser une question utile, lire, verifier, garder ce qui sert.

Chez DanielCraft, on aime une image nette. L'IA generative, c'est un stagiaire ultra rapide qui a lu beaucoup de choses, qui n'a pas de jugement moral propre, qui invente parfois avec assurance, et qui travaille mieux quand tu lui donnes un brief clair. Tu restes le responsable. Lui, il accelere. Si tu confonds vitesse et verite, tu te fais avoir. Si tu gardes le pilotage, tu gagnes du temps sans perdre ta credibilite.

★ A RETENIR

Tu pilotes, elle propose. L'IA accelere le brouillon ; toi tu verifies et tu assumes.

Ce que ce n'est pas

Ce n'est pas une conscience. Ca ne "comprend" pas comme un humain, meme si les phrases sonnent humaines. Ce n'est pas une source de verite automatique. Ce n'est pas un remplacant de ton metier du jour au lendemain. Et ce n'est surtout pas une excuse pour coller un texte non relu a un client, a un eleve, a un juge, ou a un patient.

Ce n'est pas non plus "un seul produit". ChatGPT, Claude, Gemini, Copilot, Mistral, Midjourney, des assistants dans Word, dans ton telephone, dans ton editeur de code : meme famille large, usages differents. On va les demeler sans jargon opaque. Tu n'as pas besoin de devenir chercheur. Tu as besoin de devenir un bon pilote.

Tu as une tache. Tu as des contraintes : temps, ton, public, regles, budget. L'IA propose des pistes. Toi, tu choisis, tu corriges, tu assumes. Le pont, c'est le **prompt** : la facon dont tu formules la demande. Sans pont, tu obtiens du flou poli. Avec un pont, tu obtiens quelque chose de reutilisable. Plus loin dans le livre, on ajoutera des notions utiles : **tokens**, fenetre de **contexte**, temperature, system prompts, RAG,

agents, multimodal, evaluation, couts d'API, securite des donnees. Pas pour te faire peur. Pour que tu saches ce que tu manipules.

♦ ASTUCE

Commence par une tache repete et basse risque (mail, plan, reformulation). Monte ensuite vers les sujets plus techniques.

Lea, freelance web, utilise l'IA pour ebaucher un mail client et un plan de page. Max, artisan plombier, s'en sert pour reformuler un devis plus clair et une fiche d'entretien. Sam, enseignant, prepare un quiz et une explication plus simple pour ses eleves. Ines, qui decouvre l'IA au bureau, commence par anonymiser ses notes avant tout collage. Quatre profils, meme logique : gagner du temps sur le brouillon, garder le cerveau pour le jugement.

Ce que tu vas savoir faire

Dans ce livre, tu vas comprendre l'IA generative en mots simples, un peu d'histoire utile, les grands types d'outils, le fonctionnement pratique des **LLM** (tokens, contexte, temperature), l'art du prompt et des consignes systeme, les limites et les hallucinations, les usages du quotidien, le multimodal (image, audio, video), l'idee du RAG et des agents, l'ethique et la securite des donnees, l'evaluation des reponses, et l'idee des couts d'API. Puis un mini-projet, un recap, trois ateliers, comment choisir un outil, les bonnes pratiques, un quiz, et un bravo.

Niveau debutant solide. Pas besoin de coder pour comprendre. Pas besoin d'etre "tech". Besoin de curiosite et d'honnete : l'IA aide ; elle ne remplace pas ta verification.

Comment lire ce livre

Lis dans l'ordre au debut. Les premiers chapitres posent le sol. Les ateliers font faire. Le quiz verifie. Tu peux revenir ensuite a un chapitre precis (prompts, hallucinations, securite, couts) comme a une fiche. A chaque fin de chapitre, il y a un "A toi". Fais-le. Cinq minutes valent mieux qu'une lecture passive de quarante pages.



Exemple : Lea briefe, l'IA propose, Lea verifie.

Petite histoire

Lea devait écrire une proposition pour un fleuriste. Avant, elle passait deux heures à regarder le curseur clignoter. Maintenant, elle demande à l'IA un plan en cinq parties, un ton simple, et trois questions à poser au client. Elle relit. Elle coupe le blabla. Elle ajoute son prix et ses limites. Quarante minutes, proposition nette. L'IA n'a pas "fait le devis". Elle a déblocqué le début.

Max, lui, avait honte de ses mails. Trop courts, trop secs. Il demande une version plus claire, garde sa voix, envoie. Le client répond plus vite. Ce n'est pas de la triche. C'est de l'aide à la rédaction, comme un correcteur en plus fort - à condition de ne jamais coller des données clients sensibles sans réfléchir. On y revient longuement.

Erreur classique

Croire que "l'IA a dit" égal "c'est vrai". Ou croire que "je ne sais pas formuler" égal "l'IA ne marche pas pour moi". Souvent, le problème n'est pas le modèle. C'est la demande floue : "écris-moi quelque chose sur mon business". Quel business ? Pour qui ? Quel ton ? Quelle longueur ? Quelle action attendue ?

⚠ ATTENTION

"L'IA a dit" n'est pas "c'est vrai". Fluide n'est pas compétent ; tu restes le filtre.

Autre piège : tout automatiser trop tôt. Commence par une tâche répétée et basse risque (mail, plan, reformulation). Monte ensuite vers les sujets plus techniques du livre. Tu seras prêt.

En vrai

Ouvre l'outil auquel tu as déjà accès. Sans te former plus, pose une vraie question de ton quotidien : "aide-moi à expliquer mon offre en cinq lignes à un client pressé". Lis. Note ce qui est juste, ce qui est vague, ce qui est faux. Ce livre sert à transformer ce premier essai en habitude propre.

A toi

Écris en trois phrases : (1) une tâche que tu fais souvent et qui te fatigue, (2) ce que tu accepterais qu'une IA ébauche, (3) ce que tu ne lui confieras jamais sans contrôle humain. Garde ce papier. On y reviendra au mini-projet.

Zoom : generative vs predictive

Beaucoup de gens mélangent deux familles. L'IA predictive répond souvent à une question étroite : ce client va-t-il résilier ? Ce pixel appartient-il à un panneau ? L'IA generative produit un contenu nouveau : un paragraphe, une image, une piste audio. Les deux peuvent cohabiter dans le même produit (un assistant qui résume puis propose un mail). Pour piloter, tu dois savoir laquelle tu actives. Si tu demandes une génération quand tu as besoin d'une classification fiable, tu obtiens du blabla là où il fallait un score. Si tu demandes un score là où tu as besoin d'idées, tu obtiens une case trop sèche.

Dans la pratique 2026, le grand public touche surtout la génération via le chat. Les entreprises, elles, ont encore des modèles predictifs partout dans leurs logiciels métier, parfois sans le mot "IA" sur le bouton. Ton avantage, après ce livre, c'est de ne plus être impressionné par le label. Tu regardes le geste : générer, prédire, enchaîner des actions, ou chercher dans des documents.

Petite scene DanielCraft

Lea ouvre son assistant pour "ecrire une proposition". Elle a deja perdu vingt minutes a reformuler la meme phrase. Elle colle son brief signature, ajoute trois faits clients anonymises, demande un plan puis un texte. Elle coupe deux adjectifs marketing, remet son prix, envoie. Max, le meme matin, dicte une note de chantier dans le metro, demande une version claire, verifie qu'aucun delai n'a ete invente, envoie au client. Sam prepare un quiz : l'IA propose dix questions, il en jette six parce qu'elles sont ambiguës, il garde quatre, il assume. Trois usages, une meme posture de pilote.

Ce que "50 pages" doivent changer chez toi

A la fin, tu ne seras pas chercheur. Tu seras quelqu'un qui sait briefer, qui connait tokens et contexte assez pour ne pas noyer le modele, qui gere temperature mentale, system prompts, hallucinations, multimodal, RAG, agents, evaluation, couts, et securite des donnees. C'est deja beaucoup. C'est surtout actionnable demain matin.

Chapitre 2 - Une petite histoire utile (sans chronologie ennuyeuse)

Tu n'as pas besoin d'une liste de dates pour utiliser l'IA. Tu as besoin de comprendre d'où vient le comportement actuel des outils, pour ne pas être surpris. L'histoire de l'IA, c'est surtout l'histoire de machines qui apprennent à reconnaître des motifs : d'abord sur des règles écrites à la main, puis sur des données, puis sur des modèles de plus en plus grands capables de générer du contenu.

Dans les années où l'IA restait "dans les labos", on parlait surtout de systèmes experts, de reconnaissance de formes, de jeux (échecs, go), de traduction automatique encore raide. Puis le **machine learning** a pris le dessus : au lieu d'écrire toutes les règles, on montre des exemples et on laisse le modèle ajuster ses paramètres. Ensuite est arrivé le **deep learning** : des réseaux de neurones profonds, très bons sur l'image, la voix, puis le langage. Enfin, les grands modèles de langage (**LLM**) ont rendu l'IA conversationnelle accessible à tout le monde.

★ A RETENIR

L'IA générative est statistique, pas morale : suite plausible, pas encyclopédie officielle.

Pourquoi ça compte pour toi

Parce que les outils d'aujourd'hui ne "savent" pas au sens scolaire. Ils ont été formés à prédire des suites plausibles à partir de données. C'est pour ça qu'ils sont brillants pour reformuler, résumer, idéer... et dangereux si tu les traites comme une encyclopédie officielle. L'histoire explique le tempérament de la machine : elle est statistique, pas morale ; générative, pas omnisciente.

Chez DanielCraft, on résume souvent ainsi : avant, l'IA classait ou prédissait dans des cases étroites. Maintenant, elle écrit, dessine, parle, code, et enchaîne parfois des actions. Le passage à l'**IA générative** grand public a changé le rapport au travail quotidien. Ce n'est plus "un outil pour data scientists seulement". C'est un outil pour Lea, Max, Sam - et toi.

Quelques jalons (version poche)

Les années 2010 ont popularisé le deep learning sur l'image et la voix. Les années suivantes ont vu des modèles de langage de plus en plus grands. Puis les interfaces de chat ont rendu ça utilisable sans installation savante. En parallèle, les générateurs d'images, d'audio et de vidéo ont suivi. En 2025-2026, on parle aussi d'**agents** (des systèmes qui enchaînent des outils), de **RAG** (relier le modèle à tes documents), de multimodal (texte + image + son dans le même fil), et de coûts d'usage plus visibles via les API.

Tu n'as pas à retenir les noms de papier de recherche. Retiens le mouvement : plus de données, plus de calcul, meilleures interfaces, plus de risques de confiance aveugle. Chaque vague apporte des gains et de nouveaux pièges.

⚠ ATTENTION

Chaque vague a son accident typique : photo fake, voix clonée, agent qui envoie trop tôt, résumé qui invente une décision.

Ce qui a change dans le quotidien

Avant, pour obtenir un resume, tu lisais et tu ecrivais. Pour une traduction, tu payais ou tu utilisais un outil specialise. Pour une image, tu cherchais une banque ou tu dessinais. Aujourd'hui, tu peux demander une ebauche en trente secondes. Ca ne veut pas dire que le resultat est pret a publier. Ca veut dire que le cout du premier jet a chute. La valeur se deplace vers le brief, le gout, la verification, et la responsabilite.

Lea a vu ses clients demander "une version IA" de leurs textes. Max a vu des devis generes trop beaux et trop vagues. Sam a vu des eleves coller des copies fluides mais fausses. L'histoire recente, c'est aussi celle de l'adaptation humaine : apprendre a briefer, a verifier, a dire non.

Ce qui n'a pas change

Les gens veulent encore de la clarte, de la confiance, et des preuves. Un client lit encore un devis. Un eleve a encore besoin de comprendre. Un artisan a encore besoin de connaitre ses prix. L'IA n'efface pas ces realites. Elle change la vitesse a laquelle on produit des brouillons. Si tu oublies ca, tu produis plus de bruit, pas plus de valeur.

Erreur classique

Croire que "l'IA est nouvelle donc les anciennes methodes sont mortes". Non. Relire, tester, demander un avis humain, verifier une source : ca reste. Autre erreur : croire que "comme ca existe depuis longtemps dans les labos, je n'ai rien a apprendre". Les interfaces grand public ont change la donne. Savoir cliquer n'est pas savoir piloter.

En vrai

Demande a ton assistant : "Explique en dix lignes comment on est passe des systemes a regles aux LLM, pour un debutant, sans jargon." Lis. Souligne ce qui te parle. Ce chapitre doit te laisser une sensation simple : l'IA generative est une suite logique d'outils d'apprentissage, pas une espece magique tombee du ciel.

A toi

Ecris deux phrases : (1) ce que l'IA generative change deja dans ton quotidien, (2) ce qu'elle ne changera jamais dans ton metier (jugement, relation, responsabilite...).

Du laboratoire au telephone

Le basculement grand public n'est pas seulement une affaire de puissance de calcul. C'est une affaire d'interface. Tant que l'IA restait derriere des API obscures ou des logiciels specialises, elle concernait surtout des equipes techniques. Le chat a change la donne : tu ecris en francais, tu obtiens une reponse en francais. Le cout d'entree est tombe. Le risque de confiance aveugle a grimpe en meme temps, parce que la barriere technique qui forçait a ralentir a disparu.

♦ ASTUCE

Apprends les pieges (invention, biais, fuite, automatisation prematuree) : tu survivras aux renommages d'outils.

Entre temps, les modeles d'images ont rendu possible un moodboard en trente secondes. Les modeles de voix ont rendu possible une fausse presence orale. Les agents ont commence a enchaîner des outils. Chaque vague a son slogan marketing. Chaque vague a son accident typique : photo fake prise pour vraie, voix clonee pour une arnaque, agent qui envoie un mail trop tot, resume qui invente une decision de reunion.

Ce que l'histoire enseigne aux prudents

Les outils changent de nom. Les pieges se rhument rarement : invention, biais, fuite de donnees, automatisation prematuree, confusion entre fluidite et competence. Si tu apprends les pieges, tu survis aux renommages. C'est pour ca que ce livre insiste plus sur les gestes que sur les logos.

Chapitre 3 - Les grands types d'IA (carte simple)

"**IA**" est un mot parapluie. Sous le meme terme, on range des choses tres differentes : un filtre anti-spam, un systeme qui reconnait un visage, un moteur de recommandation, un chatbot, un generateur d'images, un agent qui reserve un billet. Si tu melanges tout, tu prends de mauvaises decisions : tu demandes a un **LLM** de faire ce qu'un classifieur simple ferait mieux, ou tu crois qu'un generateur d'images "comprend" ton contrat.

Chez DanielCraft, on propose une carte en quatre zones utiles pour debuter. Tu n'as pas besoin d'etre exhaustif. Tu as besoin de savoir ou tu te situes quand tu ouvres un outil.

♦ **ASTUCE**

Quand un commercial te vend "de l'IA", demande : ca predict, ca genere, ca voit/ecoute, ou ca agit ?

1) Prediction et classement

Beaucoup d'IA "classiques" repondent a une question etroite : ce mail est-il du spam ? Ce client va-t-il resilier ? Cette photo contient-elle un chat ? On parle souvent de **machine learning** : on entraine un modele sur des exemples labels, puis on predict sur de nouveaux cas. C'est puissant, parfois invisible (dans ta banque, ton reseau social, ton antivirus). Ce n'est pas la meme chose qu'un chat qui ecrit un roman.

2) Generation de contenu

L'**IA generative** produit quelque chose de nouveau : un texte, une image, un morceau de musique, une variante de logo, un bout de code. Les LLM sont le coeur de la generation de texte. Les modeles diffusion et cousins generent des images. D'autres modeles s'occupent de la voix ou de la video. Ici, la qualite se juge a l'usage : utile, coherent, fidele au brief - pas "vrai" au sens scientifique a chaque phrase.

3) Multimodal

Un systeme **multimodal** accepte plusieurs formes d'entree ou de sortie : tu colles une photo de tableau et tu demandes un resume ; tu dictes un message ; tu generes une image a partir d'un texte. En 2026, beaucoup d'assistants melangent deja texte et image. L'idee : un seul fil de conversation pour plusieurs types de media. Attention : decrire une image n'est pas "voir" comme un humain expert. Verifier reste obligatoire, surtout pour des chiffres, des signatures, des etiquettes medicales.

4) Agents et automatisations

Un **agent**, en idee simple, n'est pas seulement un chat qui repond une fois. C'est un systeme qui peut enchaîner des etapes vers un but : chercher, lire, ecrire un fichier, appeler une API, proposer une action. Plus de puissance, plus de risques. On y reviendra. Pour l'instant, retiens : chat = conversation ; agent = parcours d'actions a cadrer.

⚠ **ATTENTION**

Ce qui est grave, ce n'est pas de melanger les zones dans un produit : c'est de ne pas savoir laquelle est active quand tu cliques "autoriser l'envoi".

Où ranger ce que tu utilises

ChatGPT / Claude / Gemini en mode chat : generation de texte (LLM), parfois multimodal. Copilot dans Word : generation assistee dans un document. Midjourney / generateurs d'images : generation visuelle. Filtre spam Gmail : prediction / classement. Recommandations Netflix : prediction de preference. Un "agent" qui lit tes mails et repond seul : zone agents - a manier avec des garde-fous.

Lea utilise surtout la generation de texte pour ses propositions. Max utilise la reformulation et parfois la dictee (multimodal léger). Sam utilise la generation pour des exemples pedagogiques, puis filtre. Aucun des trois n'a besoin de tout maitriser. Ils ont besoin de savoir quel type d'outil ils tiennent.

Comment choisir mentalement

Si tu veux "ecrire / ideer / expliquer" : LLM. Si tu veux "trier / predire une case" : souvent un modele specialise (ou une fonction dans un logiciel metier). Si tu veux "a partir d'une image / d'un son" : multimodal. Si tu veux "faire une sequence d'actions" : agent, avec prudence. Cette carte te fera gagner des heures de frustration.

Erreur classique

Appeler tout "ChatGPT" et croire que le meme geste marche partout. Ou croire qu'un generateur d'images peut remplacer une photo reelle de ton chantier pour un devis legal. Autre piege : demander a un LLM des calculs comptables critiques sans verification - il peut ecrire joliment une erreur.

En vrai

Prends trois outils que tu as vus cette semaine (meme dans la pub). Classe-les : prediction, generation, multimodal, agent. Si tu hesites, note pourquoi. L'hésitation est deja de l'apprentissage.

A toi

Dessine (meme mal) quatre cases avec un exemple personnel dans chacune. Garde le schema. Tu t'en serviras quand on parlera de choisir un outil et d'evaluer un usage.

Cas limites utiles

Un correcteur orthographique "intelligent" : prediction / suggestion locale. Un outil qui genere un paragraphe entier dans ton ton : generation. Un assistant qui lit une capture d'ecran de tableau Excel et explique la tendance : multimodal. Un systeme qui lit tes mails, classe, brouillonne, et propose d'envoyer : agent. Tu verras que beaucoup de produits commerciaux melangent deux ou trois zones. Ce n'est pas grave. Ce qui est grave, c'est de ne pas savoir laquelle est active quand tu cliques "autoriser l'envoi".

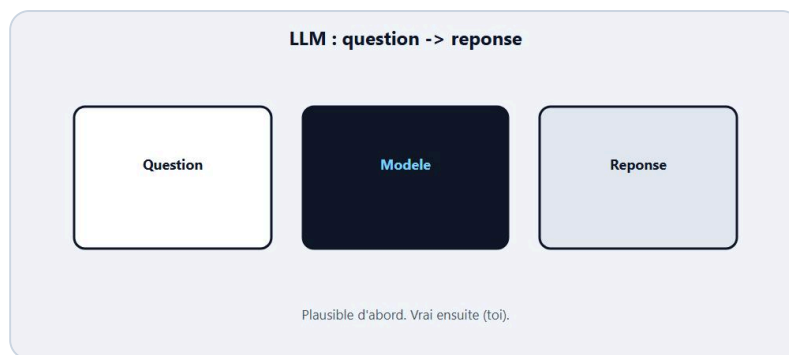
Carte pour arbitrer un achat d'outil

Quand un commercial te vend "de l'IA", demande : ca predit, ca genere, ca voit/ecoute, ou ca agit ? Sur quelles donnees ? Avec quelle validation humaine ? Quel journal d'actions ? Quel prix si le volume double ? Ces cinq questions filtrent deja beaucoup de fumee. Chez DanielCraft, on prefere un outil ennuyeux et clair a une demo spectaculaire sans freins.

★ A RETENIR

Quatre zones : predire, generer, multimodal, agent - savoir laquelle tu actives change tout.

Chapitre 4 - Les LLM : tokens, contexte, temperature



Question, modele, reponse.



Tokens et fenetre de contexte.

LLM veut dire Large Language Model : grand modele de langage. En francais simple : un systeme entraine a manipuler le langage en predisant la suite probable d'un texte, a une echelle enorme. ChatGPT est un produit construit autour de tels modeles (avec interface, memoire de conversation, regles de securite, options payantes). Claude, Gemini, Copilot, Le Chat, Mistral et d'autres jouent dans la meme cour, avec des differences de style, de limites, et d'integration.

Tu n'as pas besoin de comprendre les maths. Tu as besoin de comprendre le comportement - et trois idees pratiques qui changent tout : les **tokens**, la fenetre de **contexte**, et la **temperature** (ou l'esprit "creatif vs strict").

★ A RETENIR

Tokens + contexte + temperature mentale : trois leviers pour piloter un LLM sans maths.

Ce que fait vraiment un LLM

Il ne "cherche pas la verite dans une base officielle" par default (sauf si on le couple a une recherche ou a tes documents). Il produit une reponse coherente avec ta demande et avec ce qu'il a appris. Coherent n'est pas synonyme de vrai. Fluide n'est pas synonyme de competent.

Imagine un improvisateur genial qui a lu internet, des livres, des forums, du code. Il peut jouer le medecin, l'avocat, le marketeur. Parfois brillant. Parfois il invente une loi, une citation, un prix de piece detachee. D'ou l'importance des **prompts**, puis des hallucinations, puis de l'evaluation.

Tokens : l'unité de lecture et d'écriture

Le modèle ne lit pas vraiment "mot par mot" comme à l'école. Il découpe le texte en petits morceaux appelés tokens. Un token, ce peut être un mot court, un bout de mot, une ponctuation. "Bonjour" peut être un token ; un mot long ou rare peut en prendre plusieurs. Pourquoi tu t'en fiches... jusqu'au jour où tu colles un roman, ou tu paies une API au token, ou tu te demandes pourquoi le modèle "oublie" le début.

En pratique, plus tu envoies de texte, plus tu consommes de tokens en entrée. Plus il répond long, plus tu consommes en sortie. Les interfaces grand public masquent souvent ce compteur. Les API le montrent. Pour un débutant, retiens : sois précis, évite de coller dix PDF inutiles, demande la longueur dont tu as besoin. Tu gagnes en clarté et parfois en argent.

Ordre de grandeur pédagogique (pas une loi exacte) : en français, compte souvent ~0,7 à 1 token par mot. Un mail de 120 mots ~ 100 à 150 tokens. Un chapitre de 800 mots ~ 600 à 900 tokens. Si l'API facture 0,50 EUR / million de tokens en entrée (chiffre inventé pour l'exo), coller 10 fois le même PDF de 50 000 tokens coûte vite cher... pour peu de gain.

```
# Approximation pédagogique (pas un vrai tokenizer)
texte = "Bonjour Max, peux-tu confirmer le devis à 180 euros ?"
mots = texte.split()
tokens_approx = int(len(mots) * 0.9)
print(len(mots), "mots ~", tokens_approx, "tokens")
```

Avec une API, tu envoies souvent une liste de messages (système + utilisateur). Idée du format :

```
messages = [
    {"role": "system", "content": "Tu réponds en français simple. Signale les incertitudes."},
    {"role": "user", "content": "Reformule ce devis en 5 lignes. Prix exacts uniquement."},
]
# Puis: client.chat.completions.create(model="...", messages=messages, temperature=0.2)
```

Tu n'as pas à coder pour utiliser un chat. Mais voir ce squelette t'aide à comprendre température, rôles, et coût.

♦ ASTUCE

Impose "8 lignes", "pas d'introduction", "commence par la réponse" : tu pilotes tokens et patience d'un coup.

Fenêtre de contexte : ce que le modèle "voit" maintenant

La fenêtre de contexte, c'est la quantité maximale de conversation + documents que le modèle peut prendre en compte à un instant T. Ce n'est pas une mémoire humaine infinie. Si tu dépasses, le système tronque, résume, ou refuse. Si la conversation devient très longue et confuse, le début s'estompe ou se noie. Solution pratique : reformuler un bref propre, recommencer un fil, ou n'attacher que l'extrait utile.

Certains produits ajoutent une "mémoire" entre conversations (ils retiennent que tu es plombier à Lyon). Pratique. Aussi risque : ils retiennent des infos que tu n'aurais pas dû coller. Gère ça comme un carnet partagé, pas comme un coffre-fort.

Temperature : creatif ou strict

La temperature, dans beaucoup d'API, est un réglage qui influence le caractère aléatoire des choix du modèle. Temperature basse : réponses plus stables, plus "serrees", parfois répétitives. Temperature haute : plus de variété, plus de surprise, parfois plus d'invention. Dans les chats grand public, tu n'as pas toujours un bouton "temperature". Tu as l'équivalent mental : "sois créatif, propose 10 variantes" versus "sois strict, factuel, dis inconnu si tu ne sais pas".

Pour un devis, un règlement, une note aux parents : mode strict. Pour un slogan, un brainstorm, une métaphore pédagogique : mode créatif, puis filtre humain. Lea demande dix accroches (créatif) puis en garde deux (filtre). Max demande une reformulation fidèle de ses notes (strict). Sam demande trois analogies (créatif) puis choisit celle qui colle à sa classe.

⚠ ATTENTION

Un long contexte mal organisé noie le signal. Un court contexte bien choisi bat souvent un dump géant.

ChatGPT et la famille 2026

En pratique, tu ouvriras souvent : un chat web ou app ; un assistant dans Word / Docs / Outlook ; un assistant dans ton éditeur de code ; parfois un modèle local sur ta machine (plus avancé). Les noms changent. Les gestes restent : brief, iteration, verification. Chez DanielCraft, on recommande aux débutants de maîtriser un outil principal pendant un mois, plutôt que d'ouvrir cinq comptes. La compétence se transfère. La dispersion moins.

Forces et faiblesses typiques

Forces : rédiger, reformuler, résumer, traduire, brainstormer, expliquer simplement, générer des variantes, structurer un plan, ébaucher du code, jouer un rôle pour s'entraîner. Faiblesses : faits récents non vérifiés, chiffres inventés, sources bidon, raisonnement juridique / médical / fiscal dangereux si pris au pied de la lettre, flatterie, style générique, difficultés sur ce qui exige une vraie expérience terrain absente du prompt.

Erreur classique

Traiter le chat comme Google + expert + notaire. Ou jeter l'outil après une mauvaise réponse à une question floue. Autre erreur : enchaîner vingt messages confus sans jamais reposer le cadre. Souvent, un nouveau fil avec un bon brief bat une conversation pourrie. Et coller tout ton historique "au cas où" : tu saturer le contexte et tu mélanges les sujets.

En vrai

Pose la même question de deux façons. Version floue : "parle-moi du SEO". Version cadre : "Tu es conseiller pour un artisan fleuriste à Toulouse. Explique le SEO local en 8 lignes simples. Donne 3 actions concrètes cette semaine. Pas de jargon. Si tu n'es pas sûr, dis-le." Compare. Tu viens de toucher tokens (longueur), contexte (cadre), et temperature mentale (strict).

A toi

Choisis ton outil principal pour ce livre. Note son nom. Ecris une phrase d'identite a recoler en debut de prompt : qui tu es, pour qui tu ecris, ton ton. Exemple : "Je suis Max, plombier independant. Public : clients particuliers. Ton : clair, calme, sans blabla."

Tokens : exemples concrets

Une phrase courte consomme peu. Un collage de cinquante pages de PDF consomme beaucoup, parfois au-dela de ce que le modele peut vraiment "exploiter" utilement meme si ca rentre. Une reponse qui demarre par trois pages de preambule "Dans un monde en constante evolution" consomme ta patience et tes tokens de sortie pour rien. D'ou les consignes : "8 lignes", "pas d'introduction", "commence par la reponse". Tu ne micro-management pas par plaisir. Tu pilotes la ressource.

Certains outils affichent un compteur. D'autres non. En API, tu le verras sur la facture. Meme en forfait, un usage gaspille te rapproche des limites. Habitude saine : chaque piece jointe doit justifier sa presence en une phrase dans ton brief ("j'attache la page tarifs seulement").

Contexte long : pas une memoire magique

Les fenetres de contexte ont grandi. Ca n'a pas aboli le besoin de structure. Un long contexte mal organise noie le signal. Un court contexte bien choisi bat souvent un dump. Pareil pour la memoire produit : utile pour "je suis plombier", dangereuse pour "voici mon IBAN". Configure, nettoie, revois.

Temperature et alternatives

Meme sans bouton, tu influences l'aleas : "donne 1 reponse stricte" versus "donne 12 pistes tres differentes". Tu peux aussi demander "deux options : sage et audacieuse". Tu exteriorises ainsi le reglage. Pour les contenus critiques, baisse l'aleas et augmente la verification humaine. Pour la creativite, monte l'aleas et garde un filtre de gout.

Chapitre 5 - Prompts et system prompts : briefier comme un pro

Un **prompt**, c'est ta demande. Pas une formule magique. Pas une incantation. Un brief. Plus il est clair, plus la reponse a des chances d'etre utile. Chez DanielCraft, on enseigne un squelette simple que tu peux enrichir : role, tache, public, contraintes, format, et faits fournis. Ensuite, on ajoute l'idee du **system prompt** : la consigne durable qui cadre le comportement de l'assistant avant meme ta question du jour.

★ A RETENIR

Prompt = brief clair (role, tache, public, contraintes, format, faits). System prompt = politique de maison durable.

Le squelette du prompt utile

Role : qui doit "jouer" le modele (conseiller sober pour artisans, prof de 6e, relecteur exigeant). **Tache** : ce que tu veux exactement (plan, mail, checklist, reformulation). **Public** : pour qui c'est ecrit. **Contraintes** : longueur, ton, interdits, niveau de jargon, langue. **Format** : puces, tableau, email pret a envoyer, sections numerotees. **Faits** : ce que tu apportes (prix, dates, details vrais) pour eviter l'invention.

Exemple pour Max. Mauvais prompt : `ameliorer mon devis` . Bon prompt :

```

Role : relecteur pour un artisan plombier.
Tache : reecrire le devis ci-dessous.
Public : particulier presse, francais simple.
Contraintes :
- garde les prix exacts (ne rien inventer)
- 12 lignes max
- pas de promesse de delai si absente des notes
Format : texte pret a coller dans un email.
Notes brutes :
- debouchage evier cuisine
- deplacement 45 EUR
- main d oeuvre 1h30 a 55 EUR/h

```

Compare les deux sorties une fois. Tu verras clarte, prix fideles, moins de blabla.

Mini-appel API (schema pedagogique, adapte a ton fournisseur) :

```

prompt = ""Tu es relecteur pour un artisan.
Reecris le devis en francais simple, 12 lignes max.
Garde les prix exacts. Notes : debouchage, 45 EUR deplacement, 1h30 a 55 EUR/h.""

messages = [
    {"role": "system", "content": "Francais simple. N'invente aucun prix."},
    {"role": "user", "content": prompt},
]
# reponse = client.chat.completions.create(
#     model="ton-modele",
#     messages=messages,
#     temperature=0.2,
# )

```

♦ ASTUCE

"Tu es relecteur pour un artisan. Reecris ce devis en français simple, 12 lignes max, garde mes prix exacts." bat "améliorer mon devis".

System prompt : le cadre durable

Dans beaucoup d'outils (API, assistants custom, "projets", "GPTs", instructions persistantes), tu peux définir une consigne système : un texte qui dit comment l'assistant doit se comporter à chaque échange. Exemple : "Tu réponds toujours en français simple. Tu signales les incertitudes. Tu ne inventes pas de chiffres. Tu proposes d'abord un plan court puis le détail si on te le demande."

Le system prompt n'est pas une baguette. C'est une politique de maison. Il réduit les dérives de style et les oublis. Il ne remplace pas un bon prompt de tâche. Lea met dans ses instructions : "Ecris comme une freelance web française, ton clair, pas de marketing agressif, toujours proposer 2 options." Sam met : "Niveau collège, analogies du quotidien, jamais de contenu sensible non demandé."

Iterer sans tourner en rond

Premier jet : demande large mais cadrée. Deuxième : "garde la structure, coupe le blabla, renforce l'exemple 2". Troisième : "maintenant version plus courte pour un SMS". L'**iteration** intelligente bat le "regénérer" aveugle. Si ça part en vrille, nouveau fil + brief reconstitué. Ton **contexte** restera propre.

Techniques simples qui marchent

Donne un exemple de ton style ("écris comme ceci : ..."). Demande des variantes puis choisis. Impose "si information manquante, pose 3 questions avant de rédiger". Demande une auto-critique : "liste 5 faiblesses de ta réponse". Sépare clairement : "CONTEXTE :" puis "DEMANDE :". Pour le mode strict, ajoute "n'invente aucune source ; dis inconnu".

♦ ASTUCE

Sépare critique et réécriture : "liste objections d'abord" puis "corrige uniquement les points 2 et 5".

Ce qu'il ne faut pas faire

Prompt vide de sens. Prompt qui mélange cinq tâches sans priorité. Prompt qui demande "des sources académiques" sans intention de les ouvrir. Prompt qui colle des données sensibles "parce que ce sera plus perso". Prompt hostile ("tu es nul, recommence") : inutile. Mieux vaut préciser le critère manqué.



Exemple : Max precise role, public, format.

Erreur classique

Croire qu'un mega-prompt de quatre pages est toujours mieux. Parfois tu noies le modèle. Mieux vaut un brief net + pièces jointes utiles. Autre erreur : mettre toute la politique dans chaque message au lieu d'utiliser les instructions persistantes quand elles existent. Tu te fatigues, et tu oublies une règle sur deux.

En vrai

Prends une tâche réelle. Écris un mauvais prompt en une ligne. Puis réécris-le avec le squelette. Compare les deux réponses. Note trois différences concrètes (clarté, utilité, inventions).

A toi

Crée ton "prompt signature" en 8 à 12 lignes : identité, public, ton, 3 interdictions, format préféré. Colle-le dans les instructions de ton outil si possible. Sinon, garde-le dans un fichier texte à recopier.

Anatomie d'un système prompt solide

Un bon système prompt dit qui tu es (ou qui l'assistant doit être pour toi), pour qui tu écris, quel ton, quelles langues, quels interdictions, comment gérer l'inconnu, quel format par défaut, et parfois quelles questions poser avant de rédiger. Il évite les romans. Il évite aussi le vide. Dix à vingt lignes suffisent souvent. Au-delà, tu risques les contradictions ("sois concis" + "développe toujours").

Exemple compact pour une TPE : "Tu aides une petite entreprise française. Français simple. Pas de jargon. N'invente aucun prix ni délai. Si l'info manque, pose des questions. Structure : d'abord réponse courte, puis détails. Signale les incertitudes clairement."

Prompts de relecture

Souvent sous-estimes : "Relis ce texte comme un client méfiant. Liste objections, ambiguïtés, promesses dangereuses. Ne réécris pas encore." Puis seulement : "Corrige uniquement les points 2 et 5." Tu séparer critique et réécriture. Tu gardes le contrôle.

Bibliothèque personnelle

Rangé tes prompts par intention : Écrire, Résumer, Idées, Expliquer, Roleplay client, Anti-hallucination. Une bibliothèque de huit prompts battent un dossier de deux cents copies de "prompt miracle" trouvées sur les réseaux.

Scène de terrain (développée)

Imagine une matinée ordinaire. Tu ouvres l'outil, tu as une tâche, tu as dix minutes. Sans méthode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec méthode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu génères. Tu vérifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ça a marché. Le résultat n'est pas seulement "plus joli". Il est plus sûr, plus réutilisable, plus respectueux des gens dont les données pourraient trainer dans le fil.

Cette différence se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliothèque de huit prompts, une charte données, une grille d'évaluation, et zéro incident majeur. C'est ça que ce chapitre prépare : pas l'effet wow, l'effet fiable.

Pièges subtils

Le piège du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piège de la paresse : ne jamais retoucher. Le piège de la nouveauté : changer d'outil chaque semaine. Le piège de la peur : ne rien automatiser jamais, même le bas risque. Le juste milieu se construit en écrivant tes règles personnelles et en les testant. Ce livre te donne des règles candidates ; toi tu les adaptes à ton métier, ton risque, ton budget.

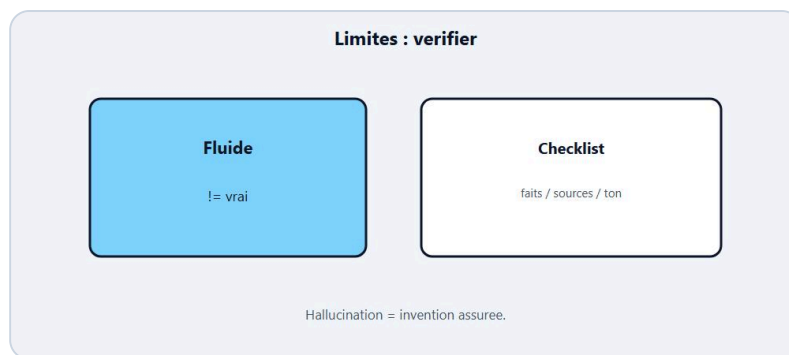
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la température (strict ou créatif), le system prompt (cadre durable), les hallucinations (vérifier), le multimodal (entrée propre), le RAG (documents rangés), les agents (freins), l'évaluation (grille), les coûts (ordre de grandeur), la sécurité (2FA, secrets). Tu n'as pas à tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Chapitre 6 - Limites et hallucinations : l'IA peut inventer



Verifier : l'IA peut inventer.

Une **hallucination**, en langage IA, ce n'est pas quand ton ecran scintille. C'est quand le modele invente des faits, des sources, des chiffres, des citations, des fonctionnalites logicielles - souvent avec un ton tres sur de lui. C'est le piege numero un pour les debutants, parce que la phrase sonne juste. Fluide. Professionnelle. Fausse.

Chez DanielCraft, on le dit cash : un **LLM** est un generateur de suite plausible, pas un notaire. Tant que tu traites ses reponses comme des hypotheses a verifier, tu es en securite relative. Des que tu signes sans lire, tu prends le risque a ta place.

⚠ ATTENTION

Fluide + assure ≠ vrai. Verifie chaque chiffre, source et promesse avant d'envoyer.

Pourquoi ca arrive

Le modele a ete entraine a continuer le texte de facon coherente. Quand tu demandes une reference precise, un prix de piece, une jurisprudence, un numero d'article de loi, il peut "completer" comme un improvisateur. Il n'a pas toujours un mecanisme interne qui dit "je ne sais pas" aussi fort qu'un humain prudent. Les produits ajoutent des garde-fous, de la recherche web, des refus - utiles, pas parfaits.

Les hallucinations augmentent quand tu demandes des details rares, des faits tres recents, des chiffres exacts, des listes de sources, ou quand tu pousse le mode creatif sur un sujet qui devrait etre factuel. Elles diminuent quand tu fournis le materiel, quand tu limites le perimetre, quand tu demandes l'aveu d'incertitude, et quand tu verifies ailleurs.

Autres limites importantes

Fenetre de **contexte** : oubli du debut, confusion entre deux sujets. Connaissances coupees a une date (selon l'outil). **Biais** : stereotypes herites des donnees. Flatterie : le modele veut souvent "aider" et te dire oui. Style generique. Faiblesse sur le raisonnement multi-etapes non verifie. Cout et latence si tu abuses. Et l'illusion de competence : plus c'est bien ecrit, plus tu baisses la garde.

Comment reduire les inventions

Fournis tes faits. Demande "si tu ne sais pas, ecris INCONNU". Interdis les fausses sources. Demande des hypotheses separees des faits. Pour tout chiffre critique, ouvre une source humaine ou officielle. Pour le code, execute et teste. Pour un mail client, relis les promesses. Pour un contenu medical, juridique, fiscal : expert humain, point.

♦ ASTUCE

Ajoute dans ton prompt : "n'invente aucune reference ; si tu n'en as pas, dis-le et propose comment chercher".

Lea demande un plan SEO, puis verifie les pretentions "Google adore X en 2026" sur des sources serieuses. Max ne laisse jamais l'IA inventer un prix de piece : il colle son tarif. Sam fait generer un exercice, puis le resout lui-meme avant de le donner aux eleves.

Exemple concret (invente) : chiffre invente

Demande floue : "Quel est le prix officiel de la piece X12 chez le fabricant Y ?" Reponse fluide possible : "La piece X12 coute 47,90 EUR HT, reference catalogue 2024."

Sans source fournie, ce 47,90 peut etre **entierement invente**. Max, lui, colle ses faits :

```
FAITS (ne pas modifier) :  
- piece : joint robinet cuisine  
- prix achat Max : 6,40 EUR  
- prix vente Max : 14,00 EUR  
DEMANDE : redige 3 lignes pour le client. Interdit d'ajouter d'autres prix.
```

Checklist rapide avant envoi (code mental, version Python) :

```
reponse = "..." # texte IA  
faits_a_verifier = ["14,00", "joint robinet"]  
for f in faits_a_verifier:  
    assert f in reponse, f"Manque ou altere: {f}"  
print("Controle basique OK - relis quand meme a l'oeil")
```

Ce n'est pas un antivirus magique. Ca force le reflexe : les chiffres critiques viennent de toi.

Evaluer une reponse (idee simple)

Avant d'envoyer, pose quatre questions : Est-ce fidele a mon brief ? Y a-t-il des faits inventes ? Le ton me ressemble-t-il ? Est-ce que je signerais ca demain matin devant un client difficile ? Si une reponse est non, tu iteres ou tu jettes. On approfondira l'**evaluation** plus loin, mais ce mini-filtre sauve deja des galeres.



Exemple : Sam croise un fait avant d'envoyer.

Erreur classique

Demander "cite tes sources" et se contenter de liens qui ont l'air vrais. Ou regénérer dix fois jusqu'à ce que ça sonne joli, sans vérifier. Ou croire que "le modèle premium n'hallucine jamais". Il hallucine moins souvent sur certaines tâches - pas jamais.

En vrai

Demande à ton outil une référence bibliographique précise sur un sujet niche que tu connais un peu. Vérifie chaque titre. Note le taux de solidité. Puis refais la même demande en mode : "n'invente aucune référence ; si tu n'en as pas, dis-le et propose comment chercher". Compare.

A toi

Choisis une tâche à risque dans ton métier (prix, délai, conseil sensible). Écris la règle personnelle : "l'IA peut ébaucher X, jamais signer Y sans vérification Z". Garde-la visible.

Typologie rapide des inventions

Invention de source (article, loi, page web). Invention de chiffre (prix, statistique, date). Invention de fonctionnalité logicielle ("dans Word tu cliques sur..."). Invention de procédure interne ("votre entreprise fait déjà..."). Confusion d'homonymes. Mélange de deux sujets proches. Extrapolation confiante hors données. Chacune a son antidote : fournir le matériel, interdire, vérifier, limiter le périmètre.

Quand le modèle refuse

Parfois l'outil refuse une demande sensible. Ce n'est pas toujours absurde. Parfois le refus est trop large, parfois trop étroit. Contourne éthiquement : reformule un usage légitime et précis, sans chercher à "jailbreaker" pour le sport. Si ton usage pro est légitime et bloqué, change d'outil ou de canal selon ta politique. Le refus n'est pas une preuve que ta demande était mauvaise, ni une preuve que tu dois ruser.

Culture d'équipe anti-hallucination

Interdiction de coller une réponse non lue dans un ticket client. Obligation de sourcer les faits critiques hors modèle. Rituel de double lecture sur les contenus à enjeu. C'est moins glamour qu'une démo. Ça évite les crises.

★ A RETENIR

L'IA ebauche ; toi tu verifies. Une regle personnelle "jamais signer Y sans Z" vaut plus qu'un modele premium.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 7 - L'IA dans le quotidien : écrire, resumer, ideer

Le meilleur usage debutant n'est pas "remplacer mon metier". C'est "debloquer les taches repetitives a bas risque". **Ecrire** un mail, **resumer** une reunion, clarifier une offre, preparer un plan, reformuler un message trop sec, **ideer** puis choisir. Ces gestes, faits proprement, changent une semaine de travail sans mettre en danger un client.

Chez DanielCraft, on recommande de construire une petite boite a outils personnelle : trois a cinq **prompts** types, ranges quelque part, reutilises, ameliores. Pas cinquante. Trois qui marchent battent cinquante oublies.

♦ ASTUCE

Trois prompts types nommes (Mail clair, Resume action, Idees classees) battent cinquante "prompts miracles" oublies.

Ecrire et reformuler

Tu as des notes brutes. Tu demandes une version claire pour un public precise. Tu fournis le ton. Tu interdis les ajouts de promesses. Tu relis. Exemples : mail de relance, reponse a une objection, page "a propos", script d'appel de deux minutes, message LinkedIn sans jargon. Lea garde un prompt "proposition client" et un prompt "mail gentil mais ferme". Max garde "devis en francais simple" et "compte-rendu de chantier".

Avant / apres concret (Max) :

```
AVANT (notes) :
ok client - fuite sous evier - passe mardi 10h - 90 EUR

APRES (demande IA + relecture Max) :
Bonjour,
Je confirme l'intervention mardi a 10h pour la fuite sous evier.
Montant convenu : 90 EUR.
Cordialement,
Max
```

Petit script pour enchaîner "resumer puis reformuler" (sans coller de secrets) :

```
notes = "fuite evier, mardi 10h, 90 EUR"
etape1 = f"Extrais les faits en puces, rien d'autre:\n{notes}"
# -> envoie etape1 au LLM, recupere faits
faits = "- fuite evier\n- mardi 10h\n- 90 EUR"
etape2 = f"Redige un mail court avec UNIQUEMENT ces faits:\n{faits}"
# -> envoie etape2, puis Max relit avant envoi
```

Deux appels battent souvent un seul prompt fourre-tout.

Resumer

Colle un compte-rendu long, une trame d'appel, un article. Demande : resume en 8 lignes, 5 decisions, 3 actions avec responsables. Ou : "explique ca a un collegue presse". Attention au **contexte** : ne colle pas un dossier confidentiel dans un outil non adapte. Anonymise. Coupe les noms si besoin. Le resume est un gain de temps enorme - et un risque **RGPD** si tu es negligent.

⚠ ATTENTION

Ne colle pas un dossier confidentiel pour "gagner deux minutes" : anonymise d'abord, ou change d'outil.

Ideer sans te noyer

Brainstorm utile : "donne 12 angles, puis classe-les par facilite / impact". Ou "3 idees sages, 3 idees audacieuses, 3 idees a eviter et pourquoi". Tu ne cherches pas la verite. Tu cherches des pistes. Ensuite tu filtres avec ton metier. Sam demande des analogies ; il en jette huit ; il en garde une.

Apprendre et s'entrainer

L'IA peut jouer l'eleve, le client difficile, le recruteur. Utile pour preparer un oral, une objection, une reunion. Cadre le role. Demande un feedback structure. Ne prends pas le score invente comme une science. C'est un simulateur, pas un jury officiel.

Integrer sans tout automatiser

Beaucoup d'outils du quotidien embarquent deja l'IA : messagerie, suite bureautique, CRM, IDE. Commence la ou tu travailles deja. Evite d'empiler cinq abonnements. Un assistant bien briefe dans ton flux reel vaut mieux qu'un compte "jouet" ouvert une fois.

Erreur classique

Utiliser l'IA pour tout, y compris ce qui demande une presence humaine (annonce de mauvaise nouvelle delicate, negociation sensible mal preparee). Ou a l'inverse, refuser toute aide par principe alors que tes mails te coutent deux heures par jour. Le juste milieu : bas risque d'abord, haut jugement toujours.

En vrai

Choisis une tache de cette semaine. Chronometre ta version sans IA. Puis avec IA (brief + relecture). Compare temps et qualite. Note une regle : "je garde l'IA pour X chaque semaine".

A toi

Cree trois prompts types nommes (ex. Mail clair, Resume action, Idees classees). Teste-les une fois. Range-les.

Semaine type d'un pilote

Lundi : resume de reunion -> actions. Mardi : deux mails difficiles. Mercredi : plan de contenu ou de seance. Jeudi : reformulation d'une offre. Vendredi : retrospective "qu'est-ce que l'IA a rate cette semaine ?". Ce rythme bat l'usage impulsif du dimanche soir sur un sujet critique.

Modeles de messages

Garde des squelettes : relance douce, relance ferme, excuse claire, demande d'information, confirmation de rendez-vous. L'IA remplit. Toi tu calibres le niveau de chaleur et les faits. Avec le temps, tu écriras parfois plus vite sans elle - et ce sera aussi une victoire : tu auras digéré des structures.

★ A RETENIR

Bas risque d'abord, haut jugement toujours : l'IA débloque le brouillon, pas la relation humaine.

Scene de terrain (developpee)

Imagine une matinée ordinaire. Tu ouvres l'outil, tu as une tâche, tu as dix minutes. Sans méthode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec méthode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu génères. Tu vérifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ça a marché. Le résultat n'est pas seulement "plus joli". Il est plus sûr, plus réutilisable, plus respectueux des gens dont les données pourraient trainer dans le fil.

Cette différence se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliothèque de huit prompts, une charte données, une grille d'évaluation, et zéro incident majeur. C'est ça que ce chapitre prépare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piège du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piège de la paresse : ne jamais retoucher. Le piège de la nouveauté : changer d'outil chaque semaine. Le piège de la peur : ne rien automatiser jamais, même le bas risque. Le juste milieu se construit en écrivant tes règles personnelles et en les testant. Ce livre te donne des règles candidates ; toi tu les adaptes à ton métier, ton risque, ton budget.

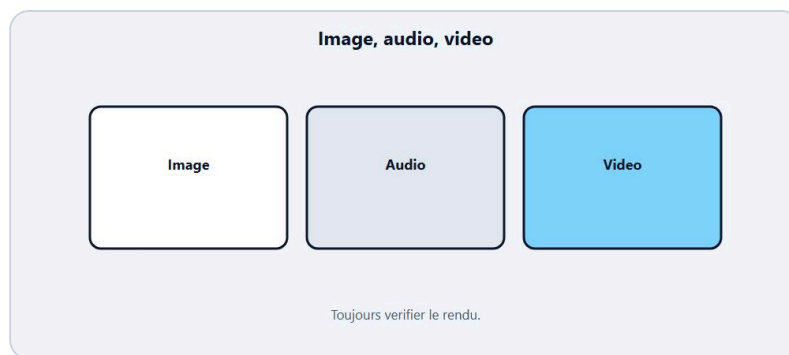
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la température (strict ou créatif), le system prompt (cadre durable), les hallucinations (vérifier), le multimodal (entrée propre), le RAG (documents rangés), les agents (freins), l'évaluation (grille), les coûts (ordre de grandeur), la sécurité (2FA, secrets). Tu n'as pas à tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Chapitre 8 - Multimodal : image, audio, video



Image, audio, video : toujours verifier.

Multimodal, en mots simples, veut dire : l'outil accepte ou produit plusieurs formes de media, pas seulement du texte. Tu colles une photo et tu demandes ce qu'elle montre. Tu dictes un message. Tu generes une image a partir d'une description. Tu transcris une reunion. Tu resumes une video. En 2026, beaucoup d'assistants melangent deja ces gestes dans le meme fil. C'est pratique. Ce n'est pas magique.

Chez DanielCraft, on traite le multimodal comme un accellerateur de comprehension et de brouillon visuel ou oral - jamais comme une preuve legale automatique, jamais comme un remplacement d'une photo reelle quand la realite compte.

⚠ ATTENTION

"L'IA a vu la photo" n'egal pas "c'est exact". Verifie nombres, noms propres et zones floues.

Image en entree : decrire, extraire, expliquer

Tu photographies un tableau blanc, une etiquette, un schema. Tu demandes une transcription, un resume, une checklist. Utile pour Max qui capture une note de chantier, pour Lea qui photographie un wireframe papier, pour Sam qui scanne un exercice. Limites : mauvaise photo, ecriture illisible, chiffres mal lus, confusion entre produits proches. Verifie toujours les nombres et les noms propres.

Ne colle pas une piece d'identite, un dossier medical complet, ou une fiche paie complete dans un outil grand public "pour voir". Anonymise. Recadre. Demande-toi si un humain non autorise pourrait nuire avec cette image.

Image en sortie : generer un visuel

Tu decris une ambiance, un style, un sujet. Le modele propose une image. Utile pour moodboards, illustrations de blog non critiques, idees de mise en page, variantes de logo tres en amont. Dangereux ou fragile pour : fausse photo de produit vendu, **deepfake**, imitation d'un artiste vivant sans cadre clair, visuel "temoignage" invente. Lea peut generer une direction artistique ; elle ne publie pas une fausse photo de boutique comme si c'etait la sienne.

Audio : dicter, transcrire, resumer

La **dictee** accelere la prise de notes. La **transcription** transforme une reunion en texte. Le resume tire des actions. Cadre : qui a consenti a l'enregistrement ? Ou vont les fichiers ? Combien de temps sont-ils gardes ? Pour une reunion d'equipe interne avec accord, souvent OK avec outil adapte. Pour un client qui n'a pas ete informe, pause. Sam peut transcrire un oral d'entrainement ; il ne diffuse pas la voix d'un eleve sans cadre.

Video : encore plus lourd

Generer ou editer de la video coute cher en calcul, en temps, en risques de fake. Pour un debutant, l'usage sain est souvent : resumer une video existante dont tu as le droit, extraire des idees, preparer un **script** - pas fabriquer une fausse interview. Les outils evoluent vite ; la prudence humaine doit rester stable.

◆ ASTUCE

Brief multimodal comme le texte : "liste les elements visibles, signale l'illisible, n'invente aucune dimension."

Brief multimodal utile

Comme pour le texte : sois precis. "Decris cette image pour un devis plomberie : liste les elements visibles, signale ce qui est illisible, n'invente aucune dimension." Ou : "Transcris cet audio, marque [inaudible] si besoin, puis donne 5 actions." Le meme squelette role / tache / contraintes / format s'applique.

Erreur classique

Croire que "l'IA a vu la photo donc c'est exact". Ou publier un visuel genere comme preuve sociale. Ou dicter des secrets clients dans le metro via un assistant cloud. Autre piege : juger un outil multimodal sur une image parfaite de demo, puis l'utiliser sur des photos floues reelles sans controle.

En vrai

Prends une image non sensible de ton travail (schema, capture d'ecran anonyme, objet). Demande une description stricte. Corrige les erreurs a la main. Note ce que le modele rate systematiquement chez toi.

A toi

Ecris ta regle multimodal : ce que tu acceptes (dictee, moodboard, transcription interne...) et ce que tu refuses (piece ID, fausse photo produit, voix d'autrui sans accord).

Qualite d'entree = qualite de sortie

Photo floue, contre-jour, document plie : la description souffrira. Audio avec tele et micro loin : la transcription souffrira. Avant d'accuser le modele, ameliore la capture. Puis demande explicitement de signaler les zones incertains. Un "je vois 12, peut-etre 18" est plus utile qu'un "18" invente.

Droits et deepfakes

Generer l'image d'une personne reelle, cloner une voix, fabriquer une video d'un evenement invente : zone legale et ethique sensible. Par default, ne le fais pas sans cadre clair et consentements. Pour une marque, prefere des visuels clairement illustres / fictifs quand ce n'est pas une photo reelle, et dis-le si besoin. La confiance se construit lentement et se perd vite.

★ A RETENIR

Multimodal = accélérateur de brouillon visuel ou oral, jamais preuve légale automatique.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 9 - RAG et agents : documents et actions



Agents : but, outils, boucle.

Deux idées changent le niveau au-dessus du simple chat : le **RAG** et les **agents**. Tu n'as pas besoin de les implémenter demain. Tu as besoin de comprendre ce qu'elles promettent, et ce qu'elles cassent si on les lâche sans cadre.

★ A RETENIR

RAG = répondre avec tes documents recherches. Agent = enchaîner des actions, toujours avec freins humains.

RAG : brancher le modèle sur tes documents

RAG veut dire Retrieval Augmented Generation. En français simple : avant de répondre, le système cherche des morceaux pertinents dans une base (tes PDF, ton wiki, ta FAQ), puis demande au **LLM** de répondre en s'appuyant sur ces extraits. L'idée : moins d'invention sur ton catalogue, tes procédures, tes tarifs - parce que la réponse est "augmentée" par du contenu que tu contrôles.

Ca ne rend pas le modèle infallible. Si la recherche ramène le mauvais paragraphe, la réponse sera belle et fautive. Si tes documents sont pourris ou obsolètes, le RAG amplifie le pourri. Si tu ne cites pas les extraits, tu ne peux pas vérifier. Un bon usage RAG montre ses sources internes : "d'après la fiche tarif 2026, page 2...".

Lea rêve d'un assistant qui connaît ses modèles de proposition. Max voudrait un outil qui lit sa bibliothèque de fiches d'entretien. Sam voudrait interroger son cours sans halluciner le programme. Dans les trois cas, le travail invisible est le même : ranger des documents propres, à jour, avec des droits d'accès clairs.

♦ ASTUCE

Sans outil spécial, simule un RAG : copie l'extrait utile, demande "réponds uniquement à partir de cet extrait, cite les phrases".

```
docs = {
    "tarif": "Deplacement 45 EUR. Main d oeuvre 55 EUR/h.",
    "secu": "Couper l'eau avant toute intervention.",
}
question = "Combien coute le deplacement ?"
extrait = docs["tarif"] # "recherche" naive
prompt = (
    "Reponds UNIQUEMENT avec cet extrait. Sinon dis INCONNU.\n"
    f"Extrait: {extrait}\nQuestion: {question}"
)
print(prompt)
# Envoie prompt au LLM (temperature basse), puis verifie la cite.
```

Agents : enchaîner des actions vers un but

Un agent, en idee simple, c'est un systeme qui ne se contente pas de repondre : il planifie, appelle des outils (recherche, calendrier, email, navigateur, base de donnees), observe le resultat, continue. Exemple : "prepare un dossier de veille sur X, resume 5 articles, propose un mail". Puissant. Aussi dangereux : un agent mal cadre peut envoyer un mail, modifier un fichier, acheter, publier.

Boucle pedagogique (pseudo-code) :

```
but = "Resumer 3 articles sur le plomb sans envoyer d'email"
outils_autorises = {"rechercher", "lire_page", "rediger_brouillon"}
# outils INTERDITS ici : envoyer_email, payer, supprimer_fichier

etat = {"notes": [], "brouillon": None}
for etape in range(5): # frein : max 5 tours
    action = llm_decide(but, etat, outils_autorises)
    if action["type"] not in outils_autorises:
        raise PermissionError("Outil non autorise")
    if action["type"] == "rediger_brouillon":
        etat["brouillon"] = action["texte"]
        break # humain relit avant tout envoi
    etat = executer(action, etat)

print("Brouillon pret - validation humaine obligatoire")
```

Sans liste blanche d'outils et sans plafond de tours, "agent" veut souvent dire "accident plus vite".

Chez DanielCraft, la regle debutant est dure : pas d'agent avec pouvoir d'envoi autonome tant que tu n'as pas de validations humaines sur les etapes critiques. L'agent propose ; toi tu valides ; ensuite seulement l'action part. Les demos marketing sautent souvent cette etape. La vraie vie ne devrait pas.

Comment cadrer un agent (checklist mentale)

But clair. Perimetre etroit. Outils autorises / interdits. Budget (temps, argent, **tokens**). Criteres de succes. Points de **validation humaine**. Journal de ce qu'il a fait. Plan de secours si ca part en vrille. Sans ca, tu as un stagiaire avec les cles de la boite et zero briefing.

⚠ ATTENTION

Un "GPT custom" sans documents a jour n'est pas un RAG solide. Et un agent "utile" sans liste d'interdits devient une faille.

RAG + agents : la combo

Beaucoup de produits melangent les deux : l'agent cherche dans tes docs (RAG), redige, puis propose une action. C'est souvent l'avenir des assistants metier. Ce n'est pas une raison de brancher ton CRM entier le premier jour. Commence lecture seule. Puis proposition. Puis action sur bac a sable. Puis production avec garde-fous.

Erreur classique

Confondre "j'ai un GPT custom" avec "j'ai un RAG solide". Un custom sans documents a jour reste un LLM generaliste deguise. Autre erreur : laisser un agent "etre utile" sans liste d'interdits. L'utilite sans frein devient une faille.

En vrai

Sans outil special, simule un RAG humain : ouvre un de tes PDF, copie l'extrait utile dans le chat, demande une reponse "uniquement a partir de cet extrait, cite les phrases". Tu viens de comprendre l'esprit du RAG. Puis ecris les 5 regles que tu imposerais a un agent avant de lui donner acces a ta boite mail.

A toi

Decris un usage RAG utile chez toi (quels documents ?) et un usage agent que tu refuses encore (quelle action ?). Une page max.

Conception d'un agent en une page

But : "preparer une veille hebdo sur X". Outils : recherche web + doc interne lecture seule. Interdits : envoi email, publication, achat. Budget : 20 appels outils max. Validation : humain lit le brouillon avant diffusion. Journal : conserver les sources. Critere de succes : 5 liens verifiees + resume fidele. Avec cette page, tu peux discuter avec un prestataire sans te faire vendre une boite noire.

RAG : hygiene documentaire

Documents dates, titres clairs, versions, droits d'accès, purge des doublons, separation du brouillon et de la procedure officielle. Le RAG n'est pas un sortilege sur un Drive bordelique. C'est une lampe torche : elle eclaire ce qui est la. Si ce qui est la est faux, tu eclaireras du faux.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 10 - Ethique, RGPD et securite des donnees

L'IA mange du **contexte**. Plus tu lui donnes d'infos, plus la reponse peut etre "perso". Plus tu lui donnes d'infos, plus tu exposes des personnes. Ethique et **RGPD** ne sont pas des chapitres ennuyeux reserves aux juristes. Ce sont des reflexes de respect : **minimiser**, anonymiser, informer, securiser, assumer.

Chez DanielCraft, on pose une regle simple : si tu n'aimerais pas voir cette donnee affichee sur un ecran de metro, ne la colle pas dans un outil grand public. Et si la donnee concerne quelqu'un d'autre, la barre monte encore.

★ A RETENIR

Si tu n'aimerais pas voir cette donnee sur un ecran de metro, ne la colle pas dans un outil grand public.

Minimiser

Tu n'as pas besoin du numero de securite sociale pour reformuler un mail. Tu n'as pas besoin de la date de naissance pour un plan de site. Tu n'as pas besoin du nom complet d'un eleve pour generer un exercice. Enleve ce qui ne sert pas. Remplace "Mme Dupont, 12 rue..." par "une cliente particuliere a Lyon". Souvent, la qualite du brief reste bonne.

Anonymiser et pseudonymiser

Anonymiser vraiment est plus dur qu'on croit (un petit village + un metier rare = reidentifiability). Fais de ton mieux : enleve noms, adresses exactes, telephones, numeros de dossier. Pour un usage interne serieux, choisis des outils et des contrats adaptes - pas seulement "c'est dans le cloud donc OK".

Bases RGPD (idee debutant)

Le RGPD encadre le traitement des donnees personnelles en Europe. Idees cles : finalite claire, minimisation, securite, droits des personnes, responsabilite. L'IA n'efface pas ca. Elle l'accentue, parce que coller un fichier dans un chat est devenu trop facile. Si tu es pro, parle a quelqu'un de competent pour ton cas. Ce livre n'est pas un conseil juridique personnalise ; c'est une boussole de prudence.

⚠ ATTENTION

Dire "c'est anonyme" des qu'on enleve le prenom ne suffit pas - surtout en petit village ou metier rare.

Ethique au-dela de la loi

Meme "legal", un usage peut etre moche : generer un faux avis client, cloner une voix sans accord, truquer une photo de resultat avant/apres, laisser un eleve croire qu'un devoir IA non assume est OK sans discussion. Ethique = ce que tu assumes le matin devant un miroir et devant les gens concernees. Transparence quand c'est utile. Respect quand c'est obligatoire.

Securite des donnees (ponte vers le chapitre securite)

Mots de passe uniques, double authentification, pas de compte partage "equipe" avec le meme login, pas de collages de cles API dans le chat, choix d'outils avec options pro / retention / training opt-out quand c'est critique. On detaille plus loin. Ici, retiens le lien : ethique sans **securite** technique, c'est un voeu pieux.

Erreur classique

Dire "c'est anonyme" des qu'on enleve le prenom. Ou "tout le monde fait ca". Ou "l'outil a une jolie page confiance donc je suis couvert". Autre piege : utiliser un compte perso gratuit pour des donnees pro sensibles parce que "c'est plus simple".

En vrai

Ouvre ton historique de chat recent. Identifie une donnee que tu n'aurais pas du coller (meme mineure). Note la lecon. Mets en place une regle d'anonymisation pour la semaine.

A toi

Ecris ta charte perso en 6 lignes : ce que je colle, ce que je n'y colle jamais, outils autorises, qui est informe, ou je stocke les prompts, qui valide les contenus sensibles.

Scenarios a refuser

Coller la liste complete d'eleves pour "personnaliser". Generer un faux avis 5 etoiles. Cloner la voix du dirigeant pour un message interne "drole". Utiliser des photos de clients sans base legale pour entrainer un truc. Laisser un stagiaire brancher un agent sur la boite mail partagee sans charte. Si un scenario te met mal a l'aise avant meme l'avis juridique, c'est deja un signal.

Transparence utile

Tu n'as pas a mettre un bandeau "IA" sur chaque virgule. Tu as a etre honnete quand ca compte : devoir evalue, contenu journalistique sensible, decision qui affecte quelqu'un, engagement contractuel. La transparence n'est pas un gadget ; c'est une maintenance de confiance.

♦ ASTUCE

Ecris ta charte perso en 6 lignes et colle-la la ou tu ouvres ton outil - avant le prochain collage.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 11 - Travail et metiers : l'IA amplifie le jugement



Idee RAG : docs + recherche + reponse.

L'IA ne "vole" pas tous les metiers d'un coup. Elle change le prix du brouillon. Ce qui devient rare, ce n'est pas ecrire une phrase fluide. C'est savoir quoi demander, quoi garder, quoi promettre, quoi refuser. Chez DanielCraft, on dit souvent : l'IA amplifie celui qui sait juger et briefer. Elle amplifie aussi les erreurs de celui qui colle sans lire.

Ce qui monte en valeur

Clarifier un besoin client. Poser un bon diagnostic. Choisir une strategie. Assumer une responsabilite legale ou morale. Creer une relation de confiance. Gouter un ton. Connaitre le terrain. Verifier un chiffre. Negocier. Trancher. Ces competences ne disparaissent pas parce qu'un modele ecrit vite. Elles deviennent le coeur du travail visible.

Ce qui se commoditise

Le premier jet generique. Le mail "dans le monde d'aujourd'hui". Le plan bateau. La reformulation moyenne. Si ton offre repose uniquement sur produire du texte moyen rapidement, tu es en concurrence avec un bouton. Si ton offre repose sur un resultat, une expertise, une presence, une garantie humaine, tu as encore un metier.

Trois portraits

Lea integre l'IA comme assistante de structure : plus de propositions envoyees, meme exigence de relecture, meilleurs questions clients. Son avantage n'est pas "j'ai ChatGPT". C'est "je livre clair et juste". Max gagne du temps sur l'administratif, pas sur le diagnostic plomberie. Sam utilise l'IA pour varier les exemples, pas pour corriger sans lire. Dans les trois cas, le metier reste le pilote.

Comment parler de l'IA a un client / collegue

Sans fanfaronnade, sans honte. "Je m'aide d'outils pour ebaucher, je verifie et j'assume." Si un client exige "zero IA", clarifie le perimetre. Si un client exige "tout full IA pour moins cher", explique ce que tu ne delegueras pas. La transparence utile bat le mystere.

Se former sans paniquer

Tu n'as pas a suivre chaque sortie de modele. Tu as a construire des gestes stables : brief, iteration, verification, securite donnees, evaluation. Ces gestes se transferent d'un outil a l'autre. La panique vient de la dispersion. La competence vient de la repetition sur tes vraies taches.

Erreur classique

Attendre que "l'IA soit stable" pour commencer. Elle bougera encore. Autre erreur : tout miser sur un outil unique comme identité de marque ("agence 100% GPT"). Les clients achètent un résultat, pas ton stack.

En vrai

Liste cinq tâches de ton métier. Classe-les : IA utile / IA dangereuse / IA inutile. Choisis une tâche "utile" à systématiser cette semaine avec un prompt type.

A toi

Ecris une phrase de positionnement : "Dans mon métier, l'IA m'aide à ... et je garde ...". Dis-la à voix haute. Si elle sonne juste, garde-la.

Compétences à cultiver cette année

Briefing. Esprit critique. Hygiène des données. Évaluation. Goût éditorial. Pédagogie (expliquer une offre). Négociation du périmètre avec les clients ("je m'aide d'outils, je vérifie"). Curiosité technique légère (assez pour lire une fiche produit sans paniquer). Ces compétences se transfèrent d'un modèle à l'autre. Les raccourcis "prompt pack" se périment.

Management et IA

Si tu managères, interdis la culture du collage non relu. Offre du temps pour construire des prompts types. Clarifie les outils autorisés. Célébrer les gains de temps vérifiés, pas les volumes de texte générés. Une équipe qui produit trois fois plus de brouillons non lus n'a pas gagné : elle a pollué.

Scène de terrain (développée)

Imagine une matinée ordinaire. Tu ouvres l'outil, tu as une tâche, tu as dix minutes. Sans méthode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec méthode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu génères. Tu vérifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ça a marché. Le résultat n'est pas seulement "plus joli". Il est plus sûr, plus réutilisable, plus respectueux des gens dont les données pourraient trainer dans le fil.

Cette différence se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliothèque de huit prompts, une charte données, une grille d'évaluation, et zéro incident majeur. C'est ça que ce chapitre prépare : pas l'effet wow, l'effet fiable.

Pièges subtils

Le piège du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piège de la paresse : ne jamais retoucher. Le piège de la nouveauté : changer d'outil chaque semaine. Le piège de la peur : ne rien automatiser jamais, même le bas risque. Le juste milieu se construit en écrivant tes règles personnelles et en les testant. Ce livre te donne des règles candidates ; toi tu les adaptes à ton métier, ton risque, ton budget.

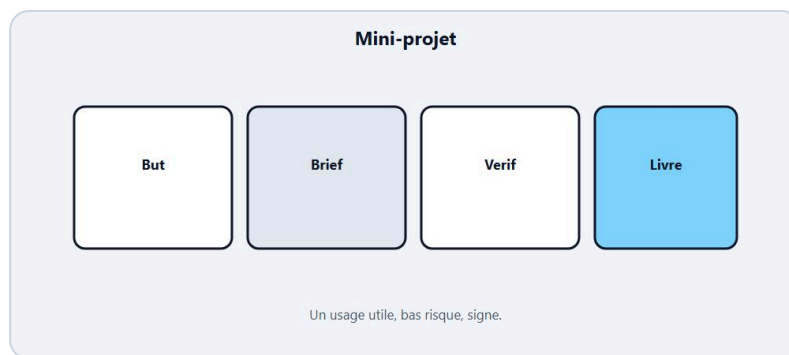
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 12 - Mini-projet : un usage IA utile et verifie



Mini-projet : usage utile et verifie.

Assez de theorie. On construit un usage complet, petit, reel, verifiable. Objectif : partir d'une tache fatigante, produire un livrable avec l'IA, appliquer les gardes (brief, limites, donnees, relecture), et pouvoir le reutiliser la semaine prochaine.

Etape 1 - Choisir la tache

Reprenons ta note du chapitre 1. Choisis une tache basse risque : mail type, page de presentation, fiche produit, plan de seance, checklist d'entretien, resume de process. Evite pour ce mini-projet : conseil medical, juridique, fiscal critique, envoi automatique a des clients, generation de fausse preuve.

Etape 2 - Preparer le brief

Ecris ton prompt signature + le brief de tache : role, public, contraintes, format, faits fournis, interdits. Anonymise les donnees. Decide du mode : strict ou creatif. Decide du critere de succes ("je peux envoyer apres 10 minutes de relecture").

Etape 3 - Produire

Lance le prompt. Si l'outil le permet, colle aussi tes instructions persistantes (system). Itere au plus trois fois avec des demandes precises. Si ca part en vrille, nouveau fil + brief reconstitue.

Etape 4 - Evaluer

Checklist : fidelite au brief, faits verifies, ton OK, pas de promesse inventee, pas de donnee sensible residuelle, longueur adaptee, tu signerais. Note les hallucinations eventuelles. Corrige a la main. Le livrable final est le tien, pas "celui de l'IA".

Etape 5 - Industrialiser legerement

Range le prompt, le livrable modele, et trois lecons dans un dossier. Donne un nom clair. Si c'est un mail type, laisse des zones [NOM], [DATE], [PRIX]. Demain, tu gagnes dix minutes. Dans un mois, des heures.

Variante avancee (optionnelle)

Si tu as des documents internes non sensibles, simule un mini-RAG : extrait utile colle dans le prompt, consigne "reponds uniquement a partir de cet extrait". Compare avec une reponse sans extrait. Tu verras souvent la difference de solidite.



Exemple : Ines livre un mini-projet signé.

Erreur classique

Choisir un projet trop large ("refondre tout mon business"). Ou s'arreter au premier jet "parce que c'est fluide". Ou ne pas ranger le prompt : tu recommenceras de zero a chaque fois.

En vrai

Fais le mini-projet aujourd'hui, meme imparfait. Un usage fini bat dix intentions.

A toi

Livre attendu : (1) prompt final, (2) livrable relu, (3) 5 lignes de retrospective : ce qui a marche, ce qui a trompe, ce que tu changes la prochaine fois.

Criteres de reussite du mini-projet

Le livrable a ete utilise (envoye, publie en interne, enseigne) ou est pret a l'etre. Le prompt est range et reutilisable. Aucune donnee sensible n'a fuite dans un outil inadapte. Tu peux expliquer en deux minutes ce que l'IA a fait et ce que tu as corrige. Tu as une lecon ecrite. Si ces cinq points sont verts, le mini-projet est reussi - meme si le texte n'est pas litteraire.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

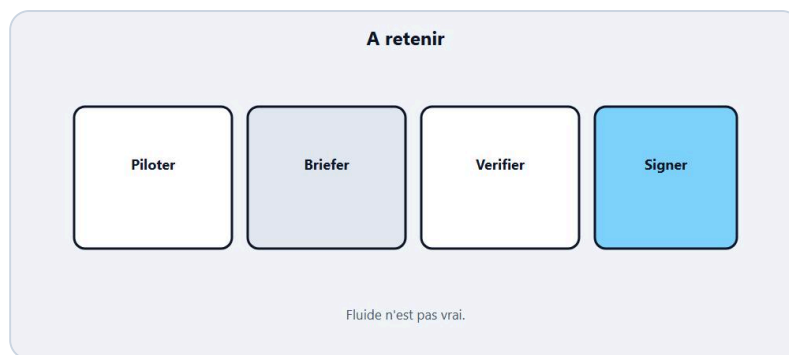
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 13 - A retenir (l'essentiel)



Piloter, briefer, verifier, signer.

Si tu ne devais garder qu'une page, ce serait celle-ci. L'IA generative est un accélérateur de brouillon, pas une source de verite automatique. Tu pilotes. Elle propose. Tu verifies. Tu assumes.

Les idées solides

LLM : modele de langage qui produit des suites plausibles. Tokens : unites de texte consommees en entree/sortie. Contexte : ce que le modele "voit" maintenant, limite. Temperature mentale : creatif vs strict. Prompt : brief clair (role, tache, public, contraintes, format, faits). System prompt : politique durable de comportement. Hallucinations : inventions assurees. Multimodal : texte + image + audio + video, avec verification. RAG : repondre avec tes documents recherches. Agents : enchaîner des actions, a cadrer. Ethique / RGPD : minimiser, anonymiser, respecter. Evaluation : fidelite, faits, ton, signature. Cout : le temps et parfois les tokens / API. Securite : comptes, donnees, pouvoirs limites.

La boucle DanielCraft

1) Clarifier le but. 2) Briefer. 3) Generer. 4) Verifier. 5) Corriger. 6) Ranger le prompt. 7) Ne brancher l'autonomie (agent) qu'avec freins.

Ce que tu peux oublier

La hype de la semaine. Le nom exact de chaque modele. La peur panique. Le perfectionnisme du mega-prompt de vingt pages. Garde les gestes. Change d'outil si besoin. Les gestes restent.

Petite checklist de poche

Est-ce bas risque ? Ai-je anonymise ? Mon brief est-il net ? Ai-je demande l'incertitude ? Ai-je verifie les faits critiques ? Est-ce que je signe ? Si oui partout, envoie. Sinon, itere.

A toi

Recopie cette carte sur papier en 10 lignes max, avec TES mots. Si tu peux l'expliquer a un ami sans regarder, c'est bon signe.

Phrase talisman

"Fluide n'est pas vrai ; utile n'est pas autonome ; rapide n'est pas responsable." Garde-la. Elle resume le livre mieux qu'un glossaire.

Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 14 - Atelier : prompts et system prompts



Atelier : du vague au brief net.

Objectif : construire et comparer des prompts, puis poser une consigne systeme durable. Duree : 30 a 45 minutes. Matériel : ton outil de chat + un editeur de texte.

Exercice 1 - Mauvais vs bon (10 min)

Choisis une tache réelle. Ecris un mauvais prompt d'une ligne. Lance-le. Puis écris un bon prompt avec le squelette (role, tache, public, contraintes, format, faits). Lance-le. Colle les deux reponses dans un fichier. Annote en rouge ce qui est vague, invente, utile.

Modele a copier-coller :

```
# mauvais
ameliorer ce texte

# bon
Role : relecteur exigeant
Tache : clarifier le texte ci-dessous
Public : client non technique
Contraintes : 8 lignes max, aucun chiffre invente
Format : paragraphes courts
Faits fournis : ...
Texte : ...
```

Exercice 2 - Mode strict vs creatif (10 min)

Sur le meme sujet, demande d'abord : "sois strict, factuel, dis INCONNU si besoin, 8 lignes". Puis : "sois creatif, 10 variantes audacieuses". Observe la difference de temperament. Note pour quels livrables tu utiliseras chaque mode.

Exercice 3 - System prompt (15 min)

Redige une consigne systeme de 8 a 15 lignes : identite, ton, interdits, gestion de l'incertitude, format par default. Colle-la dans les instructions persistantes si disponibles, sinon simule en la collant en tete. Pose ensuite trois taches differentes. Verifie si le comportement reste stable. Ajuste une seule regle a la fois.

Exercice 4 - Anti-hallucination (10 min)

Demande une liste de sources ou de chiffres sur un sujet que tu connais. Puis refais avec interdiction d'inventer. Compare. Ecris ta phrase magique anti-invention, celle que tu recoleras souvent.

Livrable

Un fichier "atelier-prompts.md" avec : mauvais prompt, bon prompt, system prompt, phrase anti-invention, 5 lignes de lecons.

Erreur a eviter

Changer dix variables a la fois. Tu ne sauras plus ce qui a ameliore quoi. Un changement, un test.

Variante avancee

Cree deux system prompts pour deux publics (ex. client particulier vs partenaire pro). Pose la meme tache. Compare. Tu verras a quel point le cadre durable change la voix sans retoucher le brief de tache. Documente lequel tu actives quand.

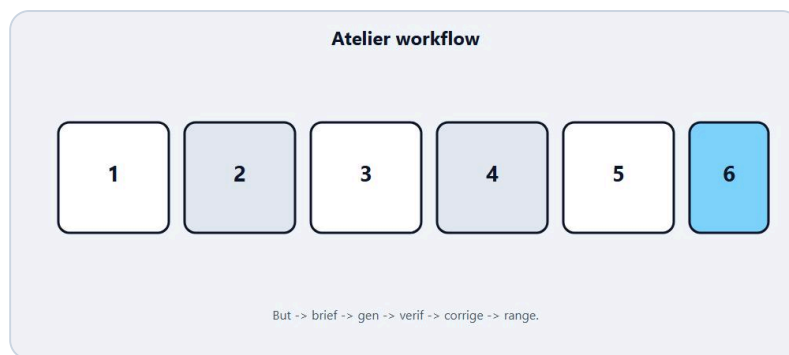
Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 15 - Atelier : workflow du quotidien



Atelier : la boucle en 6 etapes.

Objectif : monter une chaine simple "brief -> generation -> verification -> rangement" sur une vraie tache hebdomadaire. Duree : 40 minutes.

Exercice 1 - Cartographier (10 min)

Prends une tache que tu fais chaque semaine (mail, compte-rendu, prep de seance, proposition). Decoupe-la en etapes humaines actuelles. Marque ou l'IA peut aider, ou elle ne doit pas toucher, ou elle peut proposer seulement.

Exercice 2 - Construire le workflow (15 min)

Ecris ton workflow en 6 cases max : 1) Entree (notes brutes, faits) 2) Anonymisation 3) Prompt type 4) Generation 5) Checklist de verification 6) Sortie (modele range + envoi humain)

Teste-le une fois chronometre. Note le temps gagne ou perdu. Un workflow plus long qui evite une erreur client est un gain.

Exercice 3 - Mini-RAG manuel (10 min)

Si tu as une fiche tarif / FAQ / process non sensible, copie un extrait pertinent. Demande une reponse uniquement a partir de l'extrait, avec citation. Compare a une reponse "de tete" du modele. Decide si un vrai outil RAG vaudrait le coup plus tard.

Exercice 4 - Point agent (5 min)

Imagine une seule automatisation (ex. "preparer un brouillon de mail le lundi"). Ecris les validations humaines obligatoires avant envoi. Si tu ne trouves pas de validation, l'automation est trop dangereuse pour l'instant.

Livrable

Schema du workflow + prompt type + checklist + regle d'anonymisation + point de validation.

Conseil DanielCraft

Industrialise seulement ce qui a marche trois fois a la main. Pas l'inverse.

Mesure sur 7 jours

Applique le workflow sept jours. Note temps passe, incidents (hallucination, ton faux, donnee sensible evitee de justesse), gains. Ajuste une seule etape en fin de semaine. L'amelioration continue bat la refonte permanente.

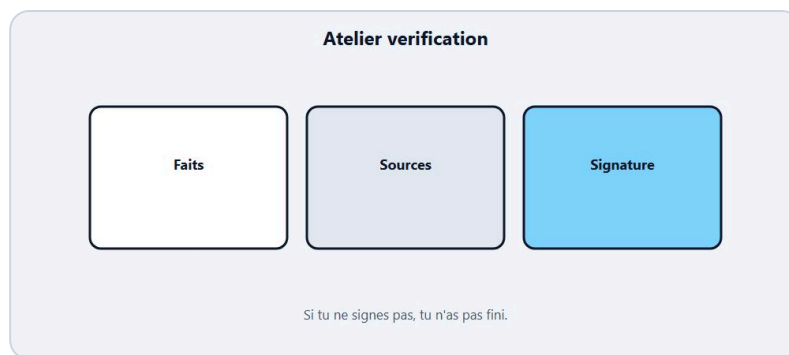
Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 16 - Atelier : evaluer et verifier une reponse



Atelier : faits, sources, signature.

Objectif : te fabriquer une grille d'evaluation simple, reutilisable, sans jargon data science. Duree : 35 minutes.

Pourquoi evaluer

Sans grille, tu juges au feeling : "ca sonne bien". Avec une grille, tu attrapes les inventions, les ton faux, les promesses dangereuses. L'evaluation n'est pas reservee aux labos. C'est un geste pro.

Exercice 1 - Grille a 6 criteres (10 min)

Recopie et adapte : 1) Fidelite au brief (0-2) 2) Exactitude des faits (0-2) 3) Ton et voix (0-2) 4) Utilite actionnable (0-2) 5) Risque (donnees / promesse) (0-2, 2 = peu de risque) 6) Pret a signer (oui/non)

Total sur 10 + veto signature. Un 9/10 avec "non" au veto = on n'envoie pas.

Exercice 2 - Blind test (15 min)

Genere deux reponses differentes a la meme tache (deux prompts ou deux modes). Melange-les. Note-les avec la grille sans regarder lequel vient d'ou. Garde la meilleure. Note pourquoi.

Exercice 3 - Chasse a l'hallucination (10 min)

Demande volontairement des details verifies (dates, prix, references). Surligne tout fait checkable. Verifie trois faits. Calcule un petit taux : faits OK / faits checkes. Garde ce reflexe pour les contenus critiques.

Livrable

Ta grille personnalisee + un exemple note + ta regle de veto.

Erreur a eviter

Noter seulement le style. Une reponse elegante et fausse doit perdre sur l'exactitude et prendre un veto. Toujours.

Revue a deux

Si tu peux, echange avec un pair : chacun note la reponse de l'autre avec la grille, sans voir le prompt. Les ecart de notation revelent des criteres implicites. Alignez votre definition de "pret a signer". En equipe, cet alignement vaut de l'or.

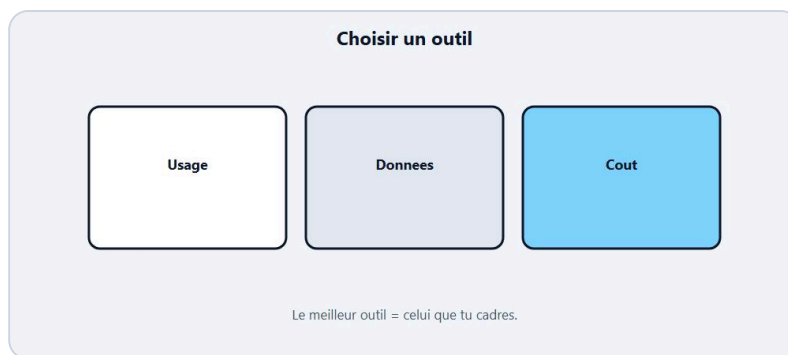
Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 17 - Choisir un outil et comprendre les couts



Choisir un outil selon usage et donnees.

Choisir un outil IA, ce n'est pas epouser une marque. C'est matcher un besoin, des contraintes de donnees, un budget, et une habitude de travail. Les noms changent. Tes criteres peuvent rester stables.

Criteres utiles

Qualite sur TES prompts (pas sur une demo marketing). Respect de tes contraintes de donnees (training, retention, options pro). Integration (mail, docs, IDE, mobile). Limites de contexte. Multimodal si besoin. Possibilite d'instructions persistantes. Prix clair. Support / statut (compte pro, SSO, facture). Pour un artisan solo, un bon chat suffit souvent. Pour une equipe, les questions admin et RGPD montent.

L'idée des couts API

En interface grand public, tu paies souvent un abonnement mensuel : un forfait d'usage. En API (quand un developpeur branche le modele dans une appli), on facture souvent a l'usage : tokens en entree, tokens en sortie, parfois le type de modele, parfois des outils (vision, recherche). Plus tu envoies de contexte, plus tu paies. Plus tu demandes des reponses longues, plus tu paies. Un agent qui boucle vingt fois peut couter cher sans que tu t'en rendes compte.

Meme si tu ne codes pas, comprendre ca t'aide : sois concis, attache l'extrait utile, evite les regenerations inutiles, choisis le modele "assez bon" pour la tache (pas toujours le plus gros pour reformuler un SMS). Lea a appris a couper ses PDF avant de les coller. Max a appris a ne pas demander un roman pour un mail. Sam a appris a limiter les variantes a un nombre utile.

Protocole de choix en une heure

1) Ecris 5 prompts types de ton metier. 2) Teste-les sur 1 a 2 outils max. 3) Note avec ta grille d'evaluation. 4) Verifie options donnees / prix. 5) Choisis pour un mois. 6) Reevalue. Pas besoin de cinq abonnements le jour 1.

Quand payer

Paye quand le free te freine vraiment (limites, confidentialite, qualite, integration), pas par peur de manquer. Un abonnement inutile est un cout. Un outil qui te fait gagner deux heures semaine a un prix raisonnable est un investissement. Fais le calcul simple : heures gagnees x ta valeur temps vs prix.

Erreur classique

Suivre chaque hype la semaine de sa sortie. Ou choisir uniquement sur "c'est gratuit". Gratuit peut être cher en données et en temps. Autre piège : rester sur un outil médiocre par habitude sans jamais retester tes 5 prompts ailleurs.

En vrai

Fais le protocole sur une heure cette semaine, même avec un seul concurrent. Écris la décision et la date de réévaluation (dans 30 jours).

A toi

Tableau simple : outil, prix, point fort, point faible, verdict 30 jours. Une page.

Lire une grille tarifaire API sans paniquer

Tu verras souvent un prix pour 1 million de tokens entrée, un autre pour la sortie, parfois un modèle "mini" moins cher, un modèle "large" plus cher, des options vision. Fais un calcul d'ordre de grandeur : un prompt de 1 000 tokens + réponse de 500 tokens, combien de fois par jour, fois 30. Compare à un forfait chat. Ajoute ton temps. Choisis. Recalcule dans trente jours avec du réel, pas de la science-fiction.

Scène de terrain (développée)

Imagine une matinée ordinaire. Tu ouvres l'outil, tu as une tâche, tu as dix minutes. Sans méthode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec méthode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu génères. Tu vérifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ça a marché. Le résultat n'est pas seulement "plus joli". Il est plus sûr, plus réutilisable, plus respectueux des gens dont les données pourraient trainer dans le fil.

Cette différence se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliothèque de huit prompts, une charte données, une grille d'évaluation, et zéro incident majeur. C'est ça que ce chapitre prépare : pas l'effet wow, l'effet fiable.

Pièges subtils

Le piège du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piège de la paresse : ne jamais retoucher. Le piège de la nouveauté : changer d'outil chaque semaine. Le piège de la peur : ne rien automatiser jamais, même le bas risque. Le juste milieu se construit en écrivant tes règles personnelles et en les testant. Ce livre te donne des règles candidates ; toi tu les adaptes à ton métier, ton risque, ton budget.

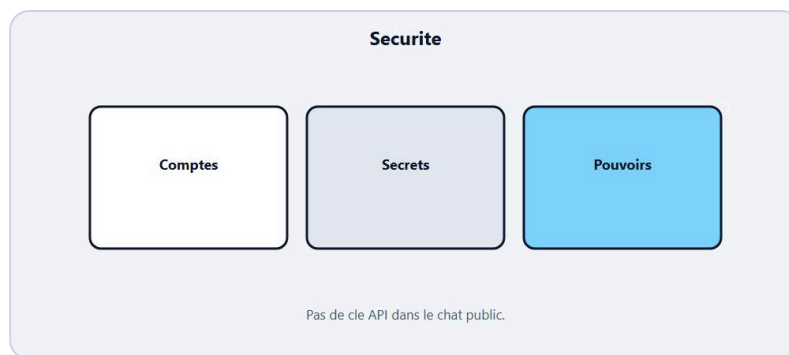
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la température (strict ou créatif), le system prompt (cadre durable), les hallucinations (vérifier), le multimodal (entrée propre), le RAG (documents rangés), les agents (freins), l'évaluation (grille), les coûts (ordre de grandeur), la sécurité (2FA, secrets). Tu n'as pas à tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Chapitre 18 - Securite : comptes, secrets, pouvoirs



Comptes, secrets, pouvoirs limites.

La securite IA debutant n'est pas de la paranoia. C'est de l'hygiene. Tu manipules des textes qui peuvent contenir des secrets, des acces, des pouvoirs d'action. Un bon brief ne suffit pas si ton compte est partage ou si un agent peut envoyer seul.

Comptes et acces

Mot de passe unique et long. Double authentication. Pas de compte "equipe" avec le meme login pour six personnes si tu peux eviter. Separe perso / pro quand c'est possible. Revoke les sessions perdues. Si tu quittes un outil, exporte ce qui compte, puis ferme proprement.

Secrets

Ne colle jamais de cle API, mots de passe, tokens d'accès, seeds crypto, ou extraits de bases clients dans le chat "pour debugger". Si tu developpes, utilise des variables d'environnement et des coffres. Si tu n'es pas tech : ne mets pas ces chaines dans un prompt, point. Un modele peut echo ce que tu as colle dans des logs ou des historiques.

Donnees sensibles

Pieces d'identite, dossiers sante, donnees scolaires nominatives, secrets d'affaires, conversations privees de tiers : regarde le chapitre ethique, puis ajoute la couche technique. Outil adapte, options de retention, desactivation de l'usage pour entraînement si disponible et necessaire, chiffrement et bonnes pratiques du poste (verrouillage ecran, pas de captures partagees a la legere).

Pouvoirs limites pour les agents

Lecture seule d'abord. Brouillons ensuite. Envoi / publication / paiement seulement avec validation humaine explicite. Journalise. Fixe un budget. Interdis des domaines d'action. Un agent sans perimetre est une faille autonome.

Ingenierie sociale et arnaques

L'IA ecrit de beaux mails de phishing. Mele a la voix clonee ou a un faux site, ca devient persuasif. Verifie les demandes d'argent et les changements d'IBAN par un second canal. Ne fais pas confiance a un style "parfaitement professionnel". Les attaquants aiment le parfait.

Erreur classique

Partager son écran en visio avec un chat ouvert plein de secrets. Ou laisser Copilot / assistant actif sur un depot qui contient des .env. Ou croire que "c'est interne donc safe" sans regarder qui a acces a l'espace workspace.

En vrai

Passe 15 minutes : active 2FA partout ou c'est possible, nettoie les sessions, cherche une donnee sensible dans tes historiques, supprime ou anonymise. Petite session, gros soulagement.

A toi

Checklist securite perso a 8 cases, cochee aujourd'hui. Date la prochaine revue (dans 90 jours).

Incident : que faire

Tu as colle un secret par erreur. Revoque la cle / change le mot de passe. Purge l'historique si l'outil le permet. Previens qui de droit selon ta politique. Note l'incident et le correctif (checklist). La honte retarde ; la revocation protege. Mieux vaut un signalement interne rapide qu'une fuite longue.

Scene de terrain (developpee)

Imagine une matinee ordinaire. Tu ouvres l'outil, tu as une tache, tu as dix minutes. Sans methode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec methode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu generes. Tu verifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ca a marche. Le resultat n'est pas seulement "plus joli". Il est plus sur, plus reutilisable, plus respectueux des gens dont les donnees pourraient trainer dans le fil.

Cette difference se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliotheque de huit prompts, une charte donnees, une grille d'evaluation, et zero incident majeur. C'est ca que ce chapitre prepare : pas l'effet wow, l'effet fiable.

Pieges subtils

Le piege du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piege de la paresse : ne jamais retoucher. Le piege de la nouveaute : changer d'outil chaque semaine. Le piege de la peur : ne rien automatiser jamais, meme le bas risque. Le juste milieu se construit en ecrivant tes regles personnelles et en les testant. Ce livre te donne des regles candidates ; toi tu les adaptes a ton metier, ton risque, ton budget.

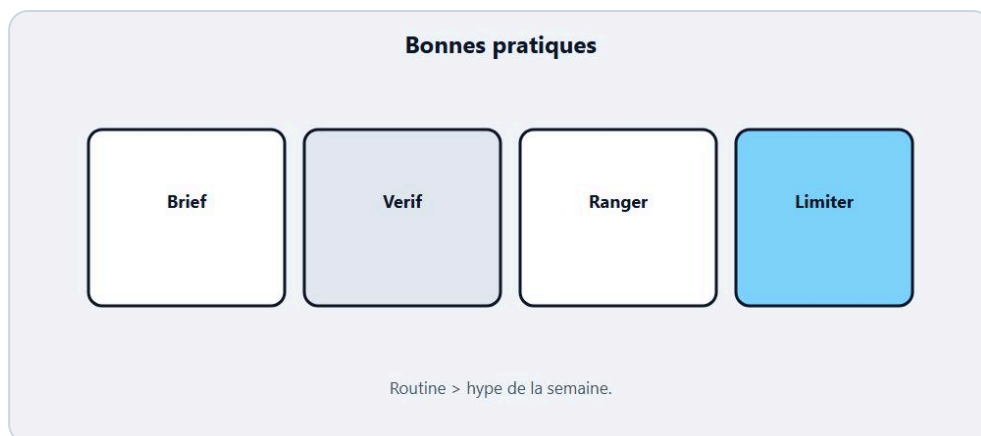
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la temperature (strict ou creatif), le system prompt (cadre durable), les hallucinations (verifier), le multimodal (entree propre), le RAG (documents ranges), les agents (freins), l'evaluation (grille), les couts (ordre de grandeur), la securite (2FA, secrets). Tu n'as pas a tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Chapitre 19 - Bonnes pratiques : la routine du pilote



Routine plutôt que hype.

Les bonnes pratiques, ce n'est pas une morale abstraite. C'est une routine qui t'évite les bêtises répétitives. Voici la version DanielCraft, courte assez pour être suivie, complète assez pour couvrir le livre.

Avant

Clarifie le but. Choisis le mode (strict / créatif). Anonymise. Ouvre le bon fil (propre). Préfère un extrait utile à un dump géant. Si tâche critique, prépare ta grille d'évaluation.

Pendant

Brief net. Une tâche principale par message important. Itère avec des critiques précises. Signale l'incertitude attendue. N'accorde pas de pouvoirs d'action sans frein. Surveille le volume (tokens / temps) si tu es en API ou sur un forfait limité.

Après

Vérifie les faits critiques. Corrige le ton. Enlève les résidus sensibles. Signe seulement si tu assumes. Range le prompt qui a marché. Note une leçon si tu t'es fait avoir.

Hebdomadaire

Relis un livrable envoyé. Améliore un prompt type. Vérifie que tes instructions système sont encore justes. Regarde si un abonnement sert encore. Purge un historique inutile.

Mensuel

Reteste tes 5 prompts sur ton outil (et éventuellement un concurrent). Révalue coût / gain. Relis ta charte données. Forme 15 minutes sur un seul sujet (RAG, multimodal, évaluation...) au lieu de scroller dix threads de hype.

Equipe (si tu n'es pas solo)

Même charte minimale. Même idée de ce qui ne se colle pas. Revue des outils autorisés. Pas d'agent partagé avec droits larges sans responsable. Culture du "je vérifie" plutôt que "l'IA a dit".

Erreur classique

Collectionner des "prompts miracles" sans routine. Ou écrire une charte de vingt pages que personne ne lit. Une page vivante bat un PDF mort.

A toi

Imprime ou épingle ta routine Avant / Pendant / Après en 12 lignes. Suis-la sept jours. Ajuste ensuite seulement.

Anti-liste des mauvaises habitudes

Régénérer vingt fois sans changer le brief. Coller tout le Drive. Faire confiance aux fausses sources. Laisser un agent envoyer. Partager le compte. Ignorer la facture tokens. Ne jamais ranger un prompt qui a marche. Chaque mauvaise habitude a son contraire dans ta routine Avant/Pendant/Après.

Scène de terrain (développée)

Imagine une matinée ordinaire. Tu ouvres l'outil, tu as une tâche, tu as dix minutes. Sans méthode, tu tapes une phrase vague, tu obtiens un texte poli, tu colles, tu regrettes. Avec méthode, tu prends deux minutes pour cadrer : but, public, contraintes, faits, interdits. Tu génères. Tu vérifies les faits critiques. Tu corriges le ton. Tu ranges le prompt si ça a marché. Le résultat n'est pas seulement "plus joli". Il est plus sûr, plus réutilisable, plus respectueux des gens dont les données pourraient trainer dans le fil.

Cette différence se voit peu le premier jour. Elle se voit au bout d'un mois, quand tu as une bibliothèque de huit prompts, une charte données, une grille d'évaluation, et zéro incident majeur. C'est ça que ce chapitre prépare : pas l'effet wow, l'effet fiable.

Pièges subtils

Le piège du perfectionnisme : retoucher le prompt une heure pour un mail de huit lignes. Le piège de la paresse : ne jamais retoucher. Le piège de la nouveauté : changer d'outil chaque semaine. Le piège de la peur : ne rien automatiser jamais, même le bas risque. Le juste milieu se construit en écrivant tes règles personnelles et en les testant. Ce livre te donne des règles candidates ; toi tu les adaptes à ton métier, ton risque, ton budget.

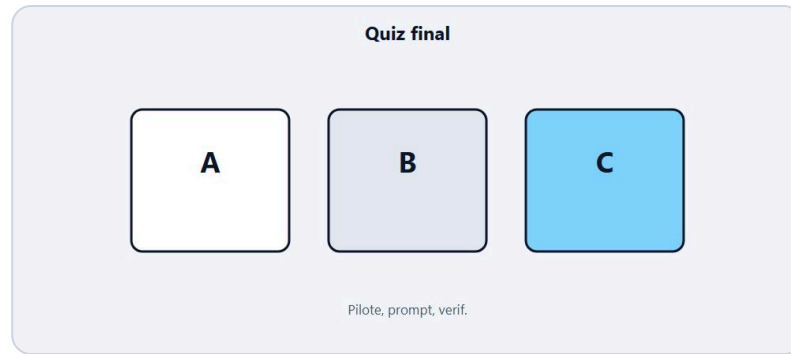
Lien avec le reste du livre

Ce que tu lis ici se branche sur les tokens (ne noie pas), le contexte (un fil propre), la température (strict ou créatif), le system prompt (cadre durable), les hallucinations (vérifier), le multimodal (entrée propre), le RAG (documents rangés), les agents (freins), l'évaluation (grille), les coûts (ordre de grandeur), la sécurité (2FA, secrets). Tu n'as pas à tout activer d'un coup. Active une brique, solidifie, ajoute.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Quiz final



Quiz final.

Pas de piège. L'idée : vérifier que tu pilotas l'IA générative avec méthode.

Questions

1. 1. Un LLM, c'est surtout :

- A) Un robot conscient qui détient la vérité
- B) Un modèle qui produit des suites de langage plausibles, pas toujours vraies
- C) Un moteur de recherche officiel

1. 2. Un token, c'est :

- A) Une pièce de monnaie virtuelle obligatoire
- B) Un petit morceau de texte lu ou écrit par le modèle
- C) Un mot de passe

1. 3. La fenêtre de contexte, c'est :

- A) La taille de ton écran
- B) La quantité de conversation / documents que le modèle peut prendre en compte à un instant T
- C) Le nombre d'utilisateurs simultanés

1. 4. "Temperature" mentale haute veut souvent dire :

- A) Plus de créativité / variété, parfois plus d'invention
- B) Un serveur en surchauffe
- C) Une réponse toujours plus vraie

1. 5. Un système prompt sert surtout à :

- A) Remplacer toute vérification
- B) Cadre durable le comportement de l'assistant
- C) Payer moins cher automatiquement

1. 6. Une hallucination, c'est :

- A) Quand le modèle invente des faits / sources avec assurance
- B) Quand l'écran scintille

- C) Quand l'outil refuse de répondre

1. 7. Le RAG, en idée simple :

- A) Un antivirus
- B) Rechercher des extraits dans tes documents puis générer une réponse appuyée dessus
- C) Générer des images aléatoires

1. 8. Un agent, c'est surtout :

- A) Un chat a une seule réponse
- B) Un système qui peut enchaîner des actions vers un but, à cadrer
- C) Un câble réseau

1. 9. Face aux données personnelles, le réflexe de base est :

- A) Tout coller pour personnaliser
- B) Minimiser / anonymiser, choisir des canaux adaptés
- C) Dire "c'est anonyme" dès que le prénom est coupé, même en village

1. 10. Évaluer une réponse, c'est surtout :

- A) Regarder si c'est fluide
- B) Grille : fidélité, faits, ton, utilité, risque, veto signature
- C) Compter les emojis

1. 11. Les coûts API dépendent souvent :

- A) Uniquement de la couleur du logo
- B) Des tokens entrée/sortie (et du modèle / options)
- C) Du signe astrologique

1. 12. Bonne pratique sécurité débutant :

- A) Partager le compte et coller les clés API dans le chat
- B) 2FA, minimiser les secrets, pas d'agent à envoi autonome sans validation
- C) Faire confiance à toute facture bien rédigée

Corrigés

1-B, 2-B, 3-B, 4-A, 5-B, 6-A, 7-B, 8-B, 9-B, 10-B, 11-B, 12-B.

Si tu as 9/12 ou plus : tu es prêt à pratiquer proprement. Sinon, relis les chapitres liés (LLM, prompts, limites, RAG/agents, éthique, évaluation, coûts, sécurité), puis refais le quiz après une semaine d'usage réel.

Questions bonus (auto-évaluation)

1. 13. Cite trois leviers anti-hallucination. 14. Différence RAG / simplement coller un PDF entier au hasard. 15. Pourquoi un système prompt n'exempte pas de vérification. Réponses libres : compare à tes chapitres 5, 6, 9.

Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas a tout maitriser d'un coup. Reviens a ce chapitre quand tu butes sur un cas reel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on prefere trois relectures actives a une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le a voix haute avec ton propre exemple. Des que ca passe a l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente cree une competence rapide sur la duree - le contraire du binge de tutos oublies le lendemain.

Chapitre 21 - Bravo



Bravo. Tu pilotes l'IA generative.

Tu as parcouru un chemin solide : comprendre l'IA generative, les LLM, les tokens et le contexte, la temperature, les prompts et system prompts, les limites, le multimodal, le RAG et les agents, l'ethique et la securite des donnees, l'evaluation, les couts, le choix d'outil, et une routine de pilote.

Ce n'est pas un diplome magique. C'est mieux : une tete claire et des gestes. Tu sais pourquoi une reponse fluide peut etre fausse. Tu sais briefer. Tu sais anonymiser. Tu sais dire non a un agent trop libre. Tu sais signer seulement ce que tu assumes.

Chez DanielCraft, on considere que c'est ca, la vraie competence IA debutant en 2026. Pas connaitre vingt logos. Piloter.

La suite

Reutilise tes prompts types. Refais le mini-projet sur une nouvelle tache. Si tu veux aller plus loin techniquement, les livres Machine Learning et Deep Learning du meme parcours t'expliquent l'envers du decor sans jargon opaque. Si tu restes sur l'usage quotidien, approfondis surtout evaluation, RAG propre, et securite equipe.

Un dernier rappel

L'IA accelere. Toi, tu orientes. Garde le gouvernail. Et quand un outil te promet la magie sans verification, souviens-toi de ce livre : beau texte, mauvais maitre - si tu abduques.

Bravo. Maintenant, va faire un usage utile, verifie, et range le prompt.

Engagement de 30 jours

Chaque jour ouvert : un usage bas risque avec brief + verification. Chaque semaine : ameliorer un prompt type. J+30 : refaire le quiz et le protocole de choix d'outil. Tu transformeras ce livre en muscle. C'est le vrai bravo.

Note de rythme

Prends le temps. Un atelier fait a fond vaut mieux que trois ateliers survolés. Si tu es presse, fais la moitie aujourd'hui et l'autre demain - mais ecris le livrable. Sans livrable, le cerveau classe ca comme "lu", pas comme "su". DanielCraft forme des gens qui livrent, meme petit.

Pour aller plus loin sans te perdre

Tu n'as pas à tout maîtriser d'un coup. Reviens à ce chapitre quand tu butes sur un cas réel. Relis l'erreur classique. Refais le "A toi". Chez DanielCraft, on préfère trois relectures actives à une lecture passive de vingt pages. Si un paragraphe te semble encore flou, reformule-le à voix haute avec ton propre exemple. Dès que ça passe à l'oral, c'est que c'est entre. Ensuite seulement, passe au chapitre suivant. Cette discipline lente crée une compétence rapide sur la durée - le contraire du binge de tutos oubliés le lendemain.

Tu as les briques. Maintenant, construis.