

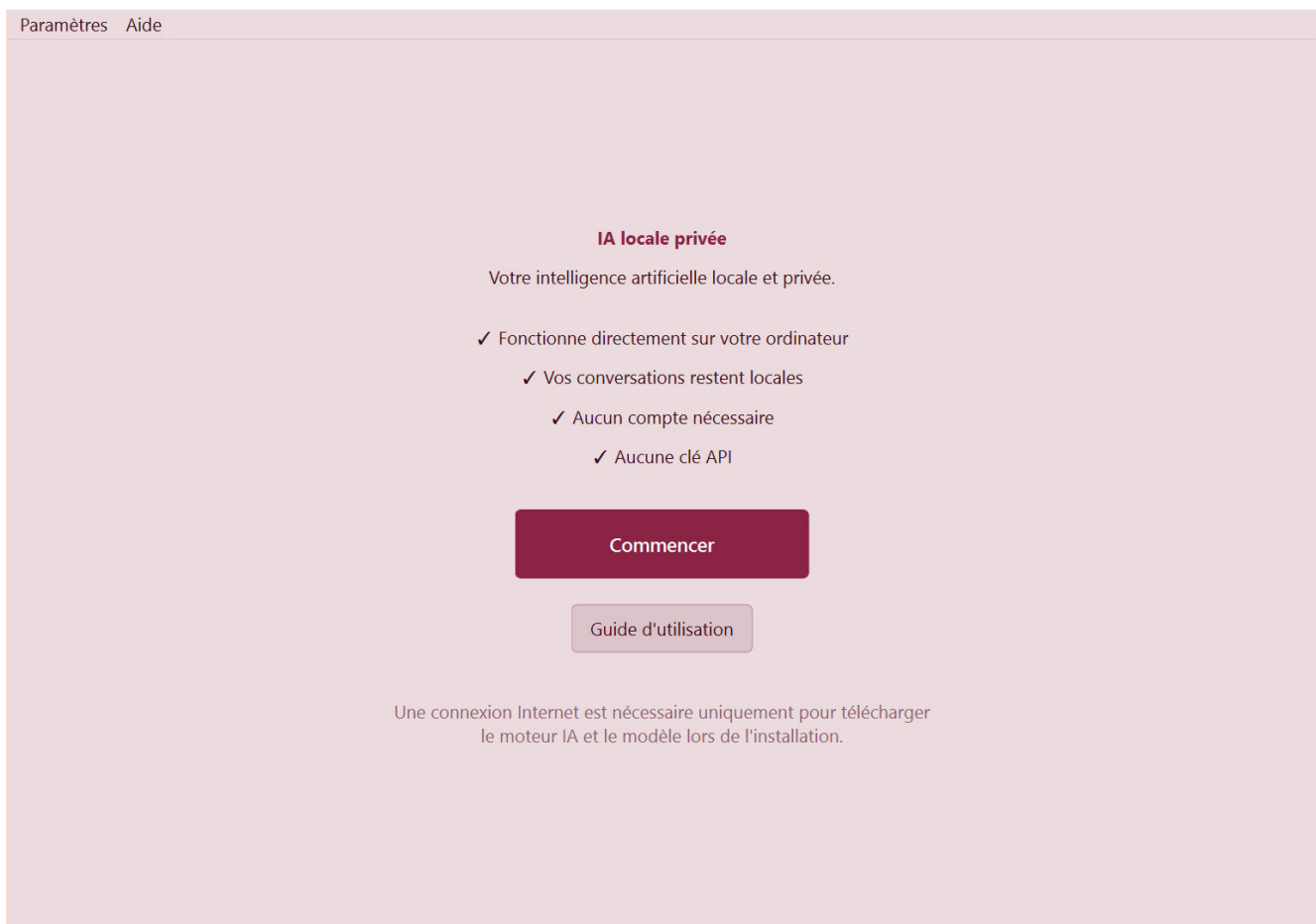
GUIDE UTILISATEUR

IA locale privée

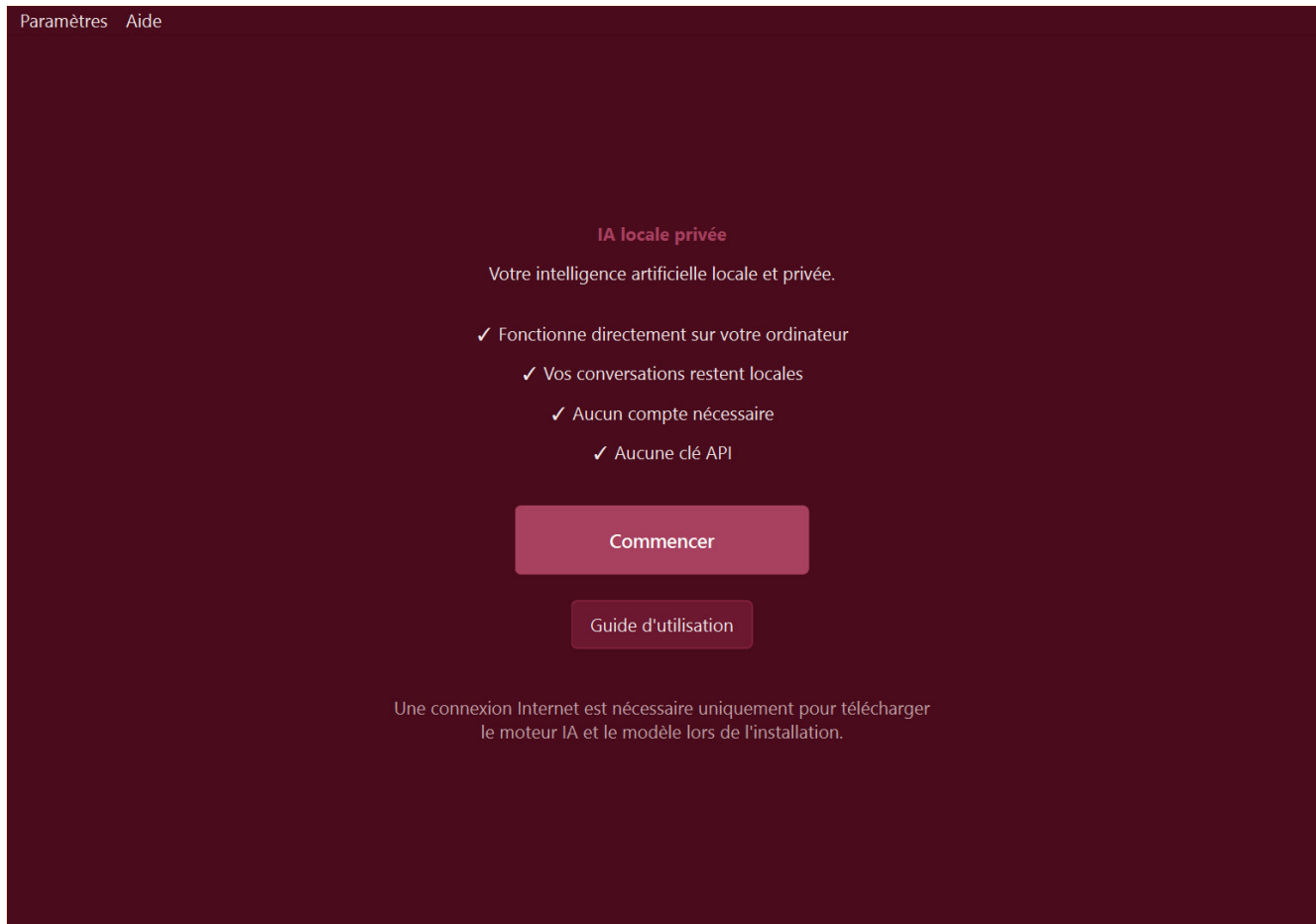
Installer le logiciel, démarrer l'IA locale, discuter, joindre un fichier, créer des images (Studio) et gérer les modèles — sans connaissance technique. Les captures ci-dessous sont les écrans réels du logiciel.

1. À quoi sert ce logiciel ?

- IA locale privée fait tourner une intelligence artificielle sur votre PC, sans compte, sans abonnement et sans jargon technique.
- Vous posez une question, joignez un fichier ou demandez une image : le logiciel installe et choisit un modèle adapté à votre machine.
- Chat, analyse de documents, programmation, vision, et Studio image optionnel restent sur cet ordinateur.
- Vos conversations et documents ne quittent pas le PC ; Internet sert seulement aux téléchargements que vous demandez (moteur, modèles, modules image).
- Livré en un seul fichier Windows : LocalAIPrivate-V3.exe (aucune installation de Python).



Accueil — cliquez sur Commencer. Pas de compte, pas de clé API.



Même accueil, thème Bordeaux sombre.

2. Installation (Windows)

- 1) Copiez LocalAIPrivate-V3.exe dans un dossier de votre choix.
- 2) Conservez ce fichier d'aide (guide_utilisation.pdf) dans le même dossier que LocalAIPrivate-V3.exe.
- 3) Double-cliquez sur LocalAIPrivate-V3.exe pour démarrer.
- Au premier lancement, un assistant analyse votre ordinateur et propose un modèle adapté.
- Le menu Aide → Guide d'utilisation ouvre ce document à tout moment.

3. Premiers pas

- Étape A — Cliquez sur Commencer à l'écran d'accueil.
- Étape B — Laissez le logiciel analyser le PC (processeur, mémoire, carte graphique), installer le moteur IA si besoin, puis proposer un modèle RECOMMANDÉ (ou Plus rapide / Meilleure qualité). Les ☐ (1 à 5) indiquent l'adéquation à votre machine : plus il y en a, mieux le modèle convient à votre PC.
- Étape C — Attendez la fin du téléchargement et du test de vitesse : vous arrivez ensuite à l'écran de discussion.
- Si le disque est trop petit, choisissez un autre dossier ou un autre disque quand le logiciel le propose.
- Une mise à jour du moteur IA peut être proposée si Internet est disponible : elle n'est jamais imposée.





[Paramètres](#) [Aide](#)

Votre ordinateur

Processeur : Intel Core i7

Mémoire : 16 Go

Graphique : NVIDIA GeForce RTX 3060

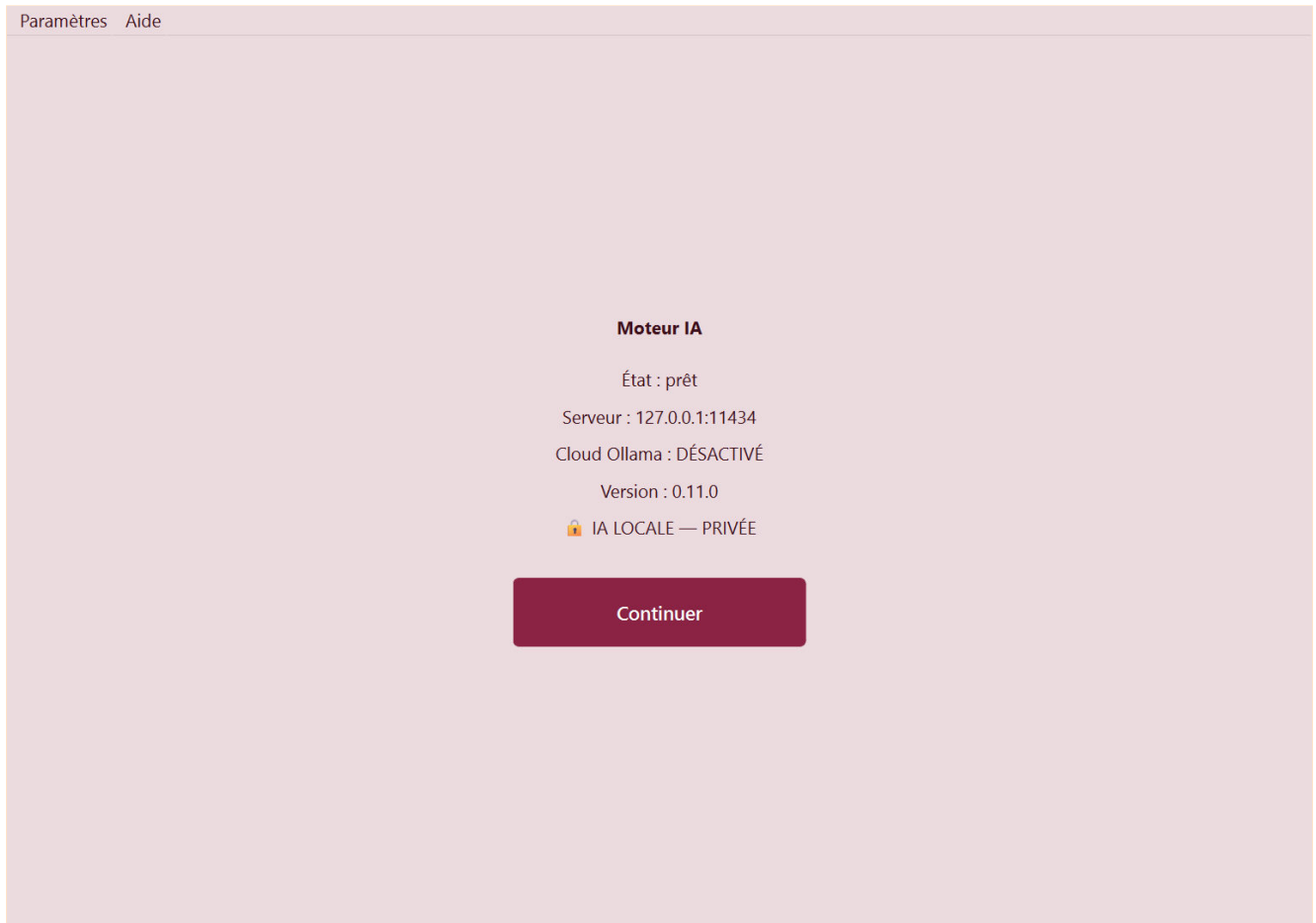
Disque : 180 Go disponibles

Capacité IA locale : TRÈS BONNE

Compatibilité estimée

[Voir les détails techniques](#) [Continuer](#)

Résumé de votre ordinateur.



Moteur IA local prêt.



Paramètres

Aide

Modèles adaptés à votre ordinateur

Tous les modèles

Nom	Taille	Vitesse	Qualité	Choix	État
Gemma 4 12B (non censuré) ★★★★★	7,6 Go	30,8 tokens/s	Excellente	RECOMMANDÉ	À installer
Qwen 2.5 Coder 14B ★★★★★	9,0 Go	26,0 tokens/s	Excellente	Meilleure qualité	À installer
Dolphin 3.0 8B (non censuré) ★★★	4,9 Go	47,8 tokens/s	Très bonne		À installer
Qwen 3 8B ★★	5,2 Go	45,0 tokens/s	Excellente		À installer
DeepSeek R1 8B ★	5,2 Go	45,0 tokens/s	Excellente		À installer
Mistral 7B v0.3 (non censuré) ★	4,4 Go	53,2 tokens/s	Très bonne		À installer
Qwen 2.5 Coder 7B ★	4,7 Go	49,8 tokens/s	Excellente		À installer
DeepSeek Coder 6,7B ★	3,8 Go	61,6 tokens/s	Très bonne		À installer
Gemma 4 E4B (non censuré) ★	6,3 Go	111,4 tokens/s	Très bonne		À installer
Qwen 3 4B ★	2,6 Go	90,0 tokens/s	Très bonne		À installer
Phi-4 Mini ★	2,5 Go	93,6 tokens/s	Très bonne		À installer
Qwen 2.5 VL 3B ★	3,2 Go	73,1 tokens/s	Bonne		À installer
Gemma 4 E2B (non censuré) ★	4,4 Go	212,7 tokens/s	Bonne		À installer
DeepSeek Coder V2 16B ★	8,9 Go	175,3 tokens/s	Très bonne		À installer ↴

Gemma 4 12B (non censuré)

Équilibre :
Qualité :
Filtrage :
RAM requise estimée :

Très bon équilibre

Excellente

Non censuré

1,5 Go

Vitesse :
Taille :
Mode estimé :
VRAM détectée :

Correct sur votre machine (~30,8 tokens/s estimés)

7,6 Go

GPU

12 Go

Utiliser

Tester

Installer

Détails

Modèle recommandé pour cette machine (à installer).



Paramètres

Aide

Modèles adaptés à votre ordinateur

Tous les modèles

Nom	Taille	Vitesse	Qualité	Choix	État
Gemma 4 12B (non censuré) ★★★★★	7,6 Go	30,8 tokens/s	Excellente	RECOMMANDÉ	Déjà installé
Qwen 2.5 Coder 14B ★★★★★	9,0 Go	26,0 tokens/s	Excellente	Meilleure qualité	À installer
Dolphin 3.0 8B (non censuré) ★★★	4,9 Go	47,8 tokens/s	Très bonne		À installer
Qwen 3 8B ★★	5,2 Go	45,0 tokens/s	Excellente		À installer
DeepSeek R1 8B ★	5,2 Go	45,0 tokens/s	Excellente		À installer
Mistral 7B v0.3 (non censuré) ★	4,4 Go	53,2 tokens/s	Très bonne		À installer
Qwen 2.5 Coder 7B ★	4,7 Go	49,8 tokens/s	Excellente		À installer
DeepSeek Coder 6,7B ★	3,8 Go	61,6 tokens/s	Très bonne		À installer
Gemma 4 E4B (non censuré) ★	6,3 Go	111,4 tokens/s	Très bonne		À installer
Qwen 3 4B ★	2,6 Go	90,0 tokens/s	Très bonne		À installer
Phi-4 Mini ★	2,5 Go	93,6 tokens/s	Très bonne		À installer
Qwen 2.5 VL 3B ★	3,2 Go	73,1 tokens/s	Bonne		À installer
Gemma 4 E2B (non censuré) ★	4,4 Go	212,7 tokens/s	Bonne		À installer
DeepSeek Coder V2 16B ★	8.9 Go	175.3 tokens/s	Très bonne		À installer

Gemma 4 12B (non censuré)

RECOMMANDÉ ★★★★★

Équilibre :

Très bon équilibre

Vitesse :

Correct sur votre machine (~30,8 tokens/s estimés)

Qualité :

Excellente

Taille :

7,6 Go

Filtrage :

Non censuré

Mode estimé :

GPU

RAM requise estimée :

1,5 Go

VRAM détectée :

12 Go

Retour

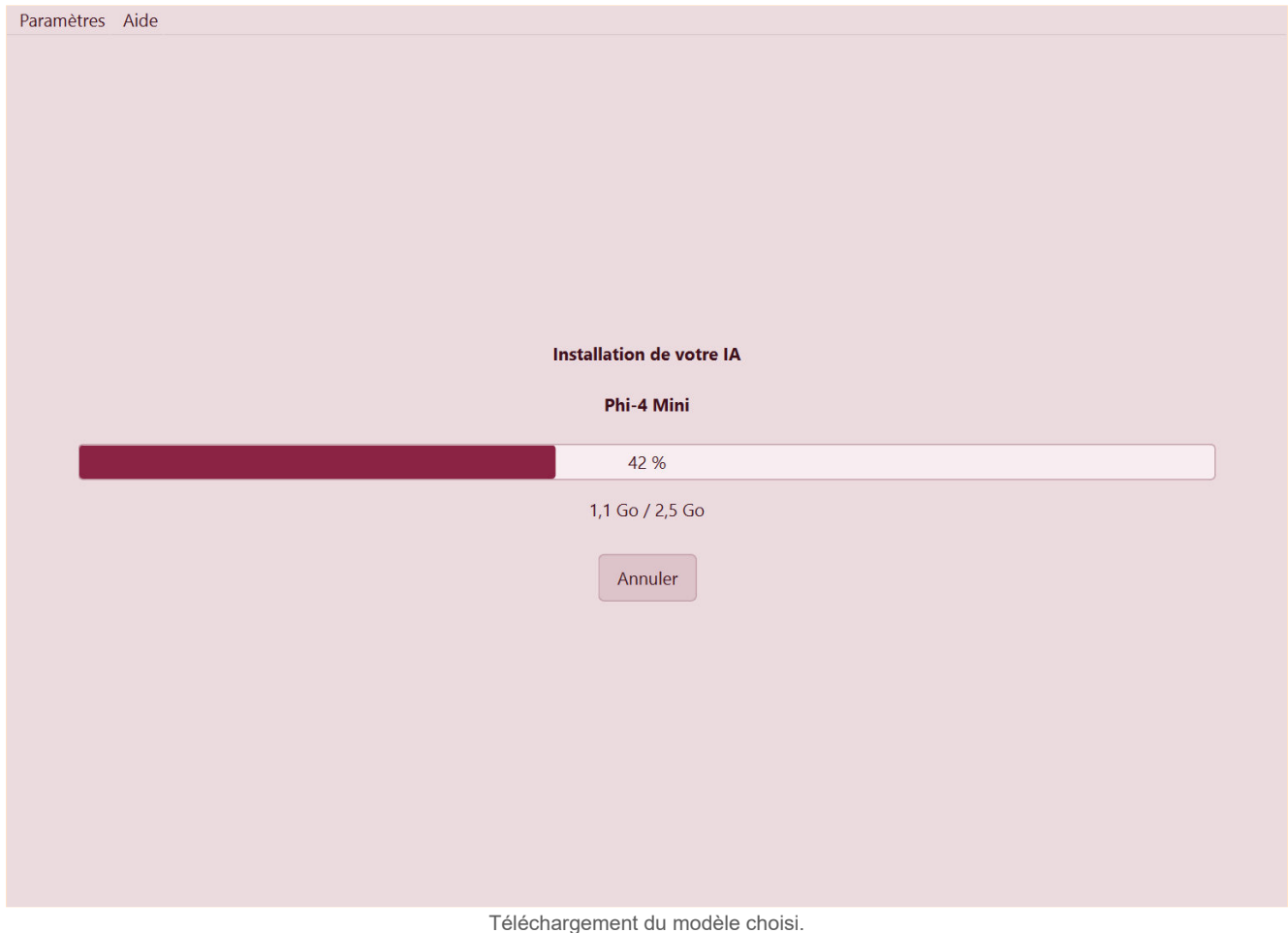
Utiliser

Tester

Déjà installé

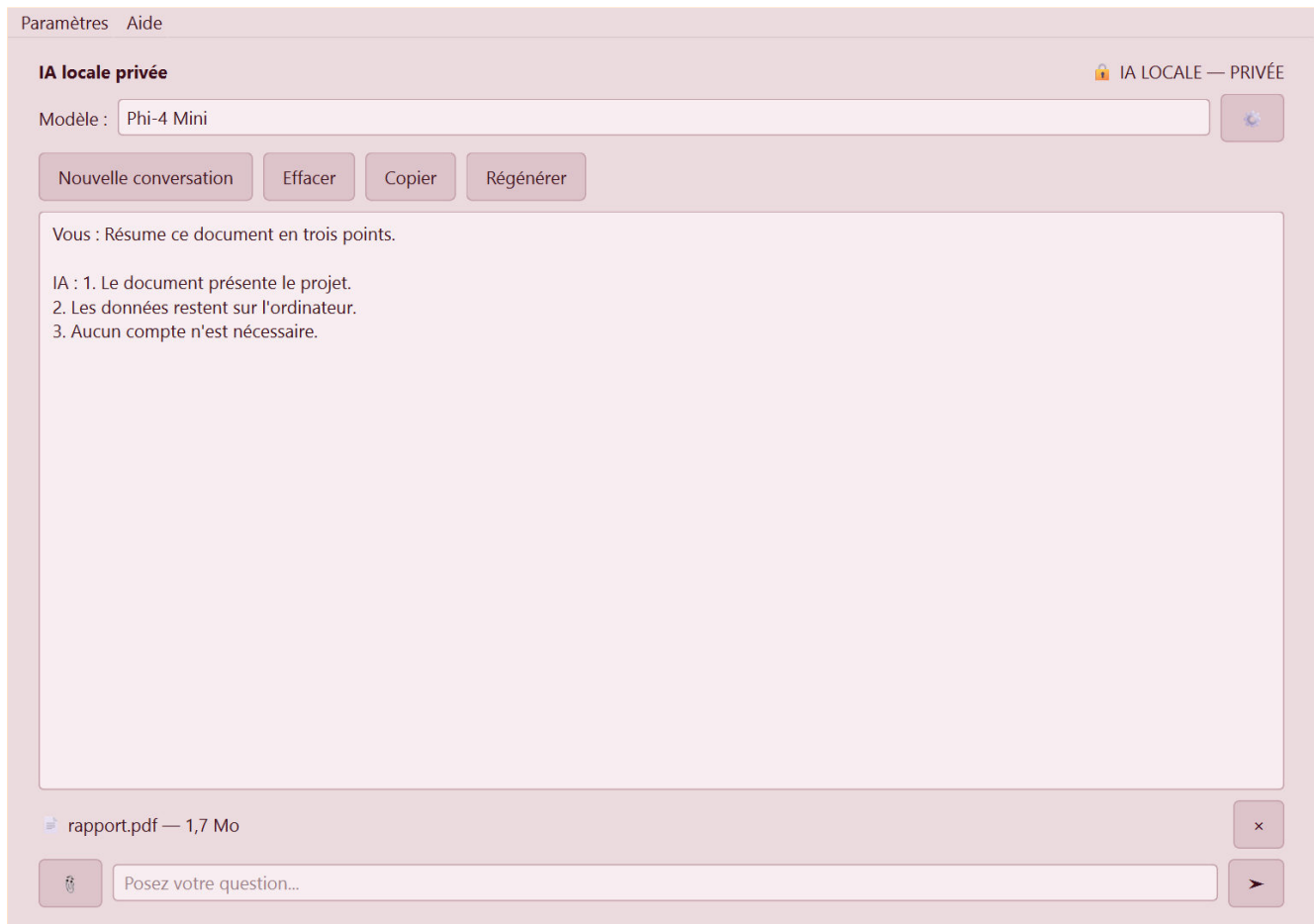
Détails

Modèle déjà présent : Utiliser (ouvrir le chat) ou Tester (mesure de vitesse).



4. Discuter avec l'IA

- Tapez votre question en bas (« Posez votre question... ») puis Entrée ou le bouton d'envoi. La réponse s'écrit au fur et à mesure.
- Arrêter interrompt la réponse en cours. Régénérer reformule la dernière réponse.
- Nouvelle conversation démarre un autre fil. Effacer vide le fil actuel (après confirmation). Copier recopie toute la conversation affichée.
- Le bandeau IA LOCALE — PRIVÉE confirme que la discussion reste sur cet ordinateur.
- À côté de « Modèle : », une liste permet de changer rapidement : elle contient les modèles déjà installés et testés. L'ancien modèle est déchargé de la mémoire pour libérer la carte graphique.
- Si l'IA a basculé temporairement (image, document, code), le nom affiche « · temporaire ». Le choisir dans la liste le rend permanent.
- La liste Conversations (si l'historique est activé) rouvre un fil précédent. Les titres restent locaux (« Sans titre » par défaut).
- Vous pouvez écrire « génère une image de... » : le Studio s'ouvre si le module est prêt, sinon un message vous renvoie vers Paramètres.



Discussion : liste Modèle (installés et testés) et fichier joint.

5. Analyser un fichier

- Joignez un fichier avec le bouton trombone ou par glisser-déposer ; une seule pièce à la fois. Le fichier n'est ni modifié ni envoyé sur Internet.
- Formats : TXT, MD, PDF (texte), DOCX, CSV, JSON, images (JPG, PNG, WEBP) et fichiers de code courants (Python, JavaScript, HTML, Java, Kotlin...).
- Si besoin, le logiciel bascule temporairement vers un modèle déjà installé (vision pour une image, document long, programmation, raisonnement) puis revient au modèle précédent et le décharge.
- Confirmation seulement s'il faut encore télécharger un modèle et que l'ordinateur est en ligne. Le modèle par défaut ne change pas tout seul.
- Joindre une image + question = analyse (description, texte visible). Ce n'est pas créer une image, ni la modifier.
- Si le Studio est prêt et qu'une image est jointe : « Améliorer l'image... » agrandit (Real-ESRGAN) ; « Créer une variation... » ouvre le Studio.
- Un PDF scanné (image seule, sans texte) peut donner peu de contenu : le logiciel n'a pas d'OCR.

Votre modèle actuel ne peut pas analyser cette image.
Votre ordinateur peut cependant utiliser Moondream (1,7 Go).
Ce modèle n'est pas encore installé.

L'installer et l'utiliser temporairement pour cette analyse ?

Oui

Non, garder le modèle actuel

Confirmation seulement s'il faut encore télécharger un modèle vision.

Paramètres Aide

IA locale privée IA LOCALE — PRIVÉE

Modèle : Phi-4 Mini

Nouvelle conversation Effacer Copier Régénérer

IA : Bonjour. Comment puis-je vous aider ?

_sample_lighthouse.png — 2,9 Ko

Améliorer l'image... Créer une variation... x

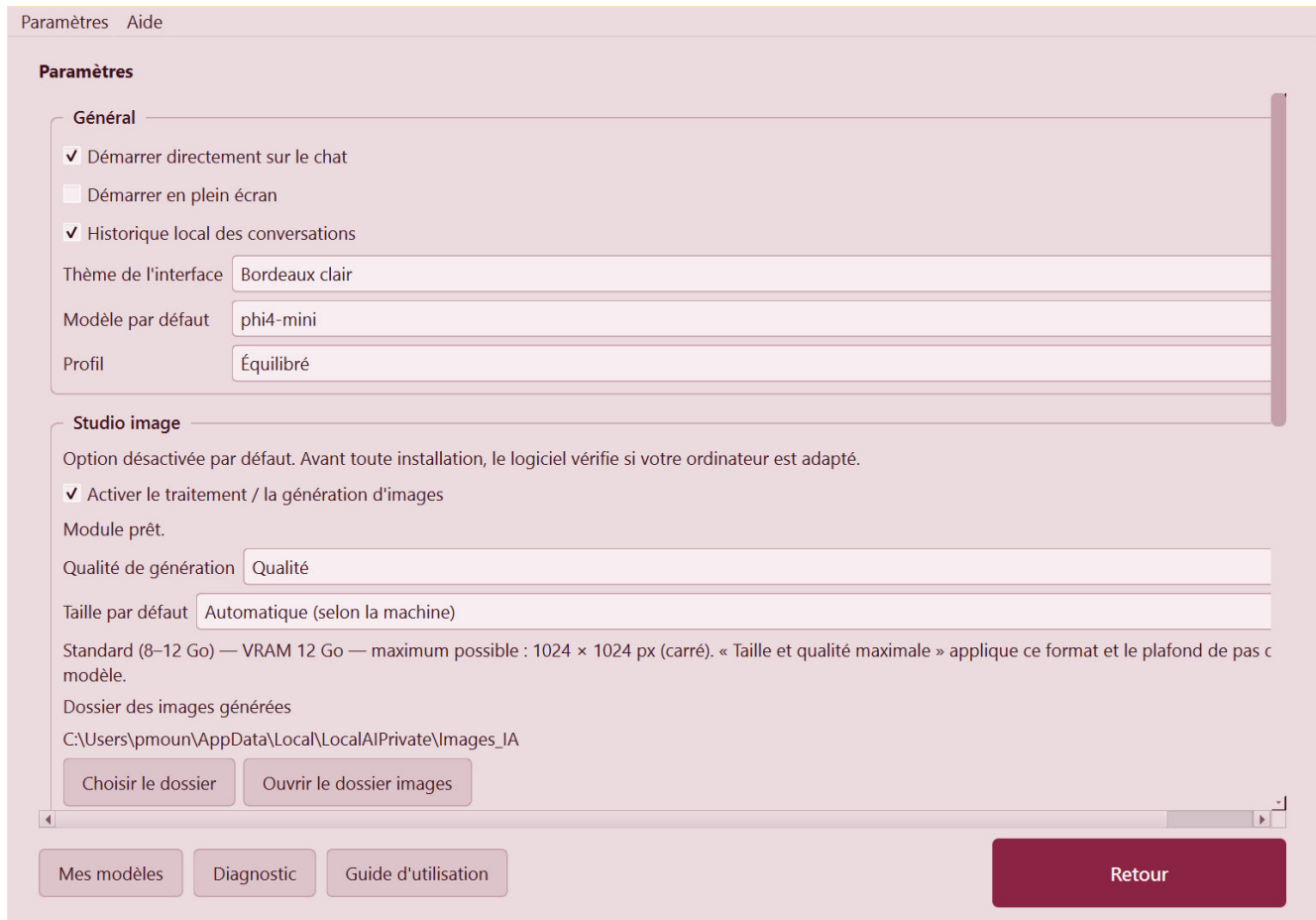
Créer une image... remplacer le ciel par des aurores boréales >

Image jointe : Améliorer l'image... ou Créer une variation...

6. Activer le Studio image

- La génération d'images est désactivée par défaut et n'utilise pas le moteur de chat (Ollama).
- Paramètres → Studio image → cochez « Activer le traitement / la génération d'images ».

- Le logiciel évalue votre PC : Possible, Possible mais limité, ou Impossible (selon GPU, VRAM, RAM et espace disque).
- Après confirmation, cliquez sur « Installer les modules complémentaires » (Internet requis, environ 4 Go : torch, diffuseurs, etc.).
- Quand le statut affiche « Module prêt », ouvrez le Studio ou utilisez « Créer une image... » dans le chat.
- Vous pouvez aussi écrire « génère une image de... » : le Studio s'ouvre si le module est prêt, sinon un message vous renvoie vers Paramètres.
- Qualité Rapide / Qualité, taille par défaut et dossier des images générées se règlent dans Paramètres → Studio image.



Paramètres Aide

Paramètres

Général

- ☒ Démarrer directement sur le chat
- ☐ Démarrer en plein écran
- ☒ Historique local des conversations

Thème de l'interface : Bordeaux clair

Modèle par défaut : phi4-mini

Profil : Équilibré

Studio image

Option désactivée par défaut. Avant toute installation, le logiciel vérifie si votre ordinateur est adapté.

- ☒ Activer le traitement / la génération d'images

Module prêt.

Qualité de génération : Qualité

Taille par défaut : Automatique (selon la machine)

Standard (8–12 Go) — VRAM 12 Go — maximum possible : 1024 × 1024 px (carré). « Taille et qualité maximale » applique ce format et le plafond de pas c modèle.

Dossier des images générées

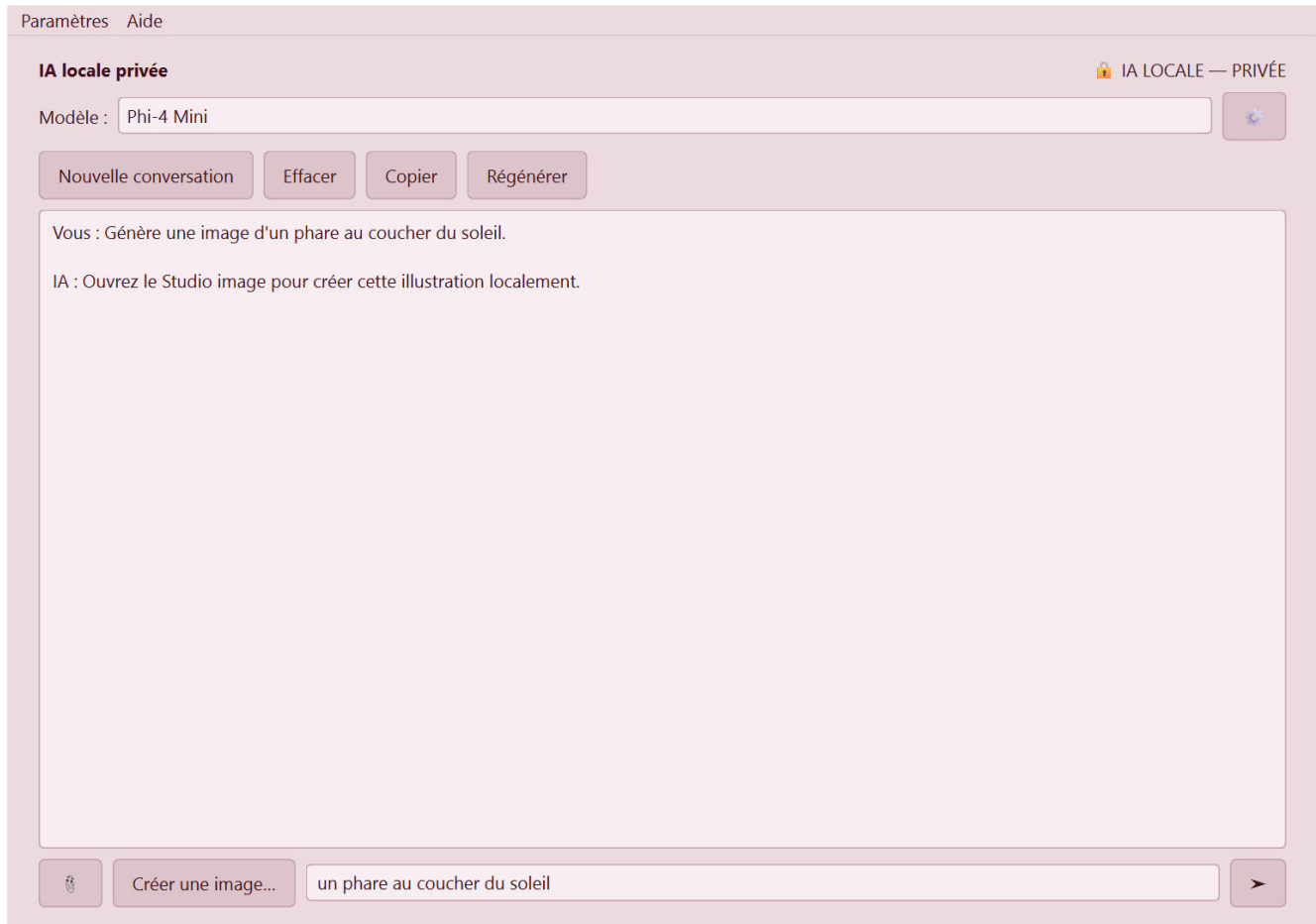
C:\Users\pmoun\AppData\Local\LocalAIPrivate\Images_IA

Choisir le dossier Ouvrir le dossier images

Mes modèles Diagnostic Guide d'utilisation

Retour

Paramètres — activation du Studio image et autres réglages.



Chat : bouton Créer une image... (module prêt).

7. Utiliser le Studio image

- Haut de page — ligne Modèle : choisissez le modèle, puis utilisez les boutons regroupés :
 - • Télécharger : premier téléchargement des poids (Internet requis).
 - • Reprendre : continue un téléchargement interrompu (fichiers partiels conservés).
 - • Mettre à jour : vérifie / rafraîchit un modèle déjà installé.
 - • Supprimer : efface uniquement les poids de ce modèle (pas le chat Ollama).
- Pas de case « Non censurés uniquement » : la liste montre tous les modèles actifs ; des ☐ (1 à 5, classement dégressif) indiquent l'adéquation à votre carte. Un modèle déjà téléchargé a toujours au moins 1 ☐.
- Taille de l'image : Automatique = taille médiane parmi les formats possibles ; « Vitesse maximale » utilise le plus petit format et le plancher de pas du modèle (brouillons rapides) ; « Taille et qualité maximale » affiche le plus grand format possible pour votre machine (ex. 1280×720) et applique le plafond de pas du modèle ; ou choix manuel (640×480 / 640×360 jusqu'à 1280×720 / 1024×768, plus carrés).
- Description : décrivez l'image (2 lignes maximum à l'écran). « Prompt de test » propose 5 scènes du plus simple au plus complexe pour comparer les modèles avec le même texte.
- Boutons Générer / Enregistrer... : à côté de la description. Pendant la génération, Générer devient Annuler.
- Bas de page — Comparatif des générations à gauche : tableau étiré (Modèle, Temps, Taille, Type, Pas, Statut). Chaque génération (prompt de test ou texte libre) ajoute une ligne, sans effacer les résultats précédents tant que le Studio reste ouvert. Cliquez une ligne pour revoir l'image. En bas du panneau : « Ouvrir le dossier des images » (dossier Images_IA à côté de l'exécutable, ou le chemin choisi dans Paramètres).
- Aperçu à droite : dernière image générée ou sélectionnée dans le tableau.

Paramètres Aide

Studio image Retour

Décrivez l'image souhaitée. Le modèle est choisi selon votre carte graphique (priorité aux modèles non censurés).

Modèle Juggernaut XL Lightning 4S (non censuré) — SDXL Lightning rapide (SDXL) ★★★★★ Télécharger Mettre à jour Supprimer

Type SDXL Spécialité SDXL Lightning rapide

Taille de l'image Automatique (selon la machine)

Standard (8–12 Go) — VRAM 12 Go — maximum possible : 1024 × 1024 px (carré). « Taille et qualité maximale » applique ce format et le plafond de pas du modèle.

Image source Aucune — génération texte → image Choisir une image... Retirer

Description Prompt de test — Choisir — Générer Enregistrer...

un phare au coucher du soleil, style photo réaliste

Prêt.

Comparatif des générations

Modèle	Temps	Taille	Type	Pas	Statut

Ouvrir le dossier des images

L'image générée apparaîtra ici.

Studio image à l'ouverture (description et modèle).

Paramètres
Aide

Studio image
Retour

Décrivez l'image souhaitée. Le modèle est choisi selon votre carte graphique (priorité aux modèles non censurés).

Modèle
Juggernaut XL Lightning 4S (non censuré) — SDXL Lightning rapide (SDXL) ★★★★★
Télécharger
Mettre à jour
Supprimer

Type
SDXL
Spécialité
SDXL Lightning rapide

Taille de l'image
Automatique (selon la machine)

Standard (8–12 Go) — VRAM 12 Go — maximum possible : 1024 × 1024 px (carré). « Taille et qualité maximale » applique ce format et le plafond de pas du modèle.

Image source
Aucune — génération texte → image
Choisir une image...
Retirer

Description
Prompt de test
— Choisir —
Générer
Enregistrer...

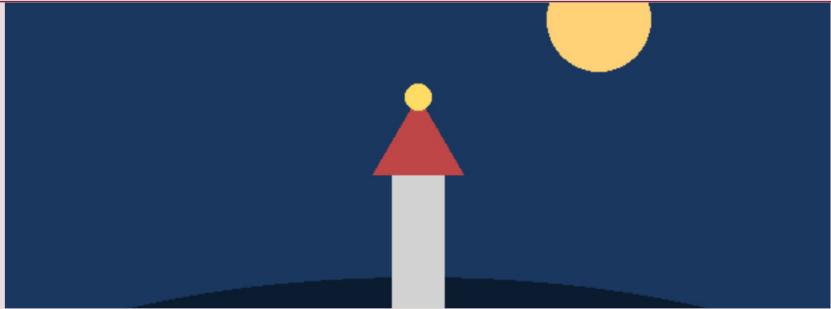
un phare au coucher du soleil, style photo réaliste

Prêt.

Comparatif des générations

Modèle
Photon v1 (non censuré)
DreamShaper 8 (non censuré)
Juggernaut XL Lightning 4S (non censuré)
Z-Image Turbo (non censuré)

Ouvrir le dossier des images



Comparatif des générations : chaque essai ajoute une ligne, sans tout effacer.

8. Modèles image et matériel

- Catalogue détaillé (développement) : DOCUMENTATION_CURSOR/15_MODELES_IMAGE.md.
- Photon v1 — photo légère, dès ~4 Go VRAM (18–28 pas, CFG 4,5 ; SD 1.5) — cible basse.
- DreamShaper 8 — 4–6 Go VRAM, PC modeste (25 pas, CFG 7) — cible basse.
- Z-Image Turbo — DiT few-step (9 pas, CFG 0). Conçu pour ≥8 Go VRAM ; très lent en offload sur 4–6 Go (préférez Photon / DreamShaper sur la cible basse).
- CyberRealistic XL v5.7 — photo / portraits / cinéma (25–30 pas, CFG 7,5 ; checkpoint converti en Diffusers).
- Juggernaut XL Lightning 4S — ~8 Go, génération rapide (8 pas, CFG 2 — profil Lightning, pas de 20–30 pas).
- RealVisXL v5 — ~8 Go, photoréalisme (32 pas, CFG 7).
- SANA 1.5 1.6B — généraliste moderne / réalisme (18–24 pas, CFG 4,5 ; DiT / SanaPipeline ; modèle standard Gemma, hors priorité non censuré ; pas de variation).
- Juggernaut XI v11 — ~10 Go, qualité SDXL maximale (28 pas, CFG 7 ; téléchargement Hugging Face parfois soumis à acceptation de conditions).
- Les modèles adaptés à votre VRAM (ou déjà installés) portent des □ dégressives (5 pour le meilleur match, puis 4, 3... jusqu'à 1). Après des générations dans le Studio, les □ tiennent compte des temps réels du comparatif : le plus rapide sur votre PC remonte, le plus lent descend. Le meilleur match compatible est sélectionné par défaut.
- Sans GPU CUDA utilisable, la génération peut basculer en CPU (beaucoup plus lente, plusieurs minutes).
- Chaque modèle a son profil d'inférence (plage de pas + CFG). La taille et le mode mémoire (offload) restent adaptés à la VRAM.
- Emplacement des poids : %LOCALAPPDATA%\LocalAIPrivate\images_models\
- Modules Python image : %LOCALAPPDATA%\LocalAIPrivate\image_runtime\

- Images générées : dossier Images_IA à côté de l'exécutable (repli %LOCALAPPDATA%\LocalAIPrivate\Images_IA\ si non inscriptible, ou le dossier choisi dans Paramètres). L'ancien dossier images_generated est migré automatiquement puis supprimé.

9. Commandes sur une image chargée

- Chargez une image (Choisir une image...) ou utilisez « Créer une variation... » depuis le chat après avoir joint une image.

- Mode « Commande (image entière) » en premier : décrivez en français ce que vous voulez — fond, retirer un élément, couleurs, météo / heure, cadrage, style photo ou illustration, retouché (netteté, bruit, peau). Pas de pinceau : l'IA applique la commande automatiquement.

- Menu Commandes : presets prêts (fond plage / studio / saison, noir et blanc, cinéma, vintage, ciel plus beau, cadre serré / large, adoucir la peau, etc.).

- Après une édition : curseur Avant/Après sous l'aperçu (glisser pour comparer), ou « Voir l'original » ; « Annuler l'édition » pour revenir à l'image précédente.

- « Améliorer pour partage » : enchaîne une retouche lumière puis l'agrandissement (Real-ESRGAN), sans jargon technique. Annuler interrompt aussi l'agrandissement en cours.

- Modes Zone (centre / haut / bas) : régénèrent seulement une bande approximative selon votre texte (sans dessin manuel).

- « Améliorer l'image... » dans le chat reste l'agrandissement seul (Real-ESRGAN), distinct des commandes d'édition.

- Le bouton Envoyer du chat sert toujours à poser une question sur l'image (analyse vision), pas à l'éditer.

- Portrait / peau : le preset « adoucir peau » reste subtil ; l'identité de la personne doit être préservée — résultat indicatif, pas une retouche professionnelle.

Paramètres
Aide

Studio image
Retour

Décrivez l'image souhaitée. Le modèle est choisi selon votre carte graphique (priorité aux modèles non censurés).

Modèle
Juggernaut XL Lightning 4S (non censuré) — SDXL Lightning rapide (SDXL) ★★★★★
Télécharger
Mettre à jour
Supprimer

Type
SDXL
Spécialité
SDXL Lightning rapide

Taille de l'image
Automatique (selon la machine)

Standard (8–12 Go) — VRAM 12 Go — maximum possible : 1024 × 1024 px (carré). « Taille et qualité maximale » applique ce format et le plafond de pas du modèle.

Image source
_sample_lighthouse.png
Choisir une image...
Retirer

Image source : tapez une commande (fond, retirer, couleurs, météo, cadrage, style, retouche). Pas de pinceau — l'IA applique la commande sur toute l'image. Comparez avec le curseur Avant/Après ou « Voir l'original » ; zones approximatives via Mode → Zone.

Mode
Zone — centre
Voir l'original
Annuler l'édition
Améliorer pour partage

Zone approximative : le centre, le haut ou le bas de l'image sera régénéré selon votre commande (sans dessin manuel).

Commande
Commandes
— Choisir —
Générer
Enregistrer...

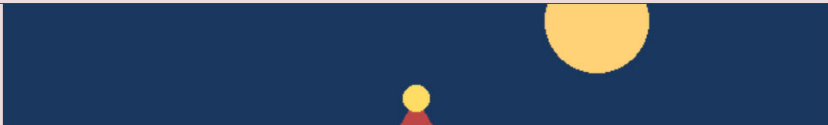
des aurores boréales dans le ciel

Prêt.

Comparatif des générations

Modèle	Temps	Taille	Vuée	Pas	St
4					

Ouvrir le dossier des images



Studio : image source et mode Zone (centre / haut / bas).

10. Modèles de chat et réglages

- Modèles adaptés à votre ordinateur : tableau filtrable (Tous, Généraliste, Programmation, Vision, Raisonnement). Installer télécharger puis teste. Si le modèle est déjà là : Utiliser (en faire le modèle du chat), Tester (mesure de vitesse), Détails.
- Mes modèles : colonnes Nom (avec ☐) , Taille, Actif, Vitesse mesurée, Usage et capacités, Filtrage (Non censuré / Standard). Actions : Utiliser, Tester, Détails, Supprimer, Installer un autre modèle.
- Les ☐ (1 à 5) classent les modèles pour VOTRE PC, pas une note universelle. Un modèle déjà téléchargé a toujours au moins 1 ☐.
- La recommandation vise 10 à 35 mots-jetons par seconde : trop lent est écarté, trop rapide (souvent un tout petit modèle) est moins bien classé.
- Profils Paramètres : Rapide, Équilibré, Qualité — ils changent l'ordre proposé, pas le fonctionnement du chat.
- Thème : plusieurs palettes (Bordeaux clair par défaut, sombre, crème, etc.).
- Autres réglages : démarrer sur le chat, plein écran, historique, dossier des modèles, contexte avancé, keep-alive, ouvrir les journaux.
- Une mise à jour du moteur IA peut être proposée au démarrage si Internet est disponible (uniquement si vous cliquez sur Mettre à jour).
- Diagnostic fournit un rapport simple à copier, sans le texte de vos conversations (y compris l'état du Studio image).
- Pour libérer de l'espace image : Paramètres → Supprimer les modules et modèles image (le chat Ollama n'est pas touché).

[Paramètres](#) [Aide](#)

Mes modèles

1 modèle(s) compatible(s) à installer [Voir les modèles](#)

Nom	Taille	Actif	Vitesse	Usage et capacités	Filtrage
phi4-mini ★	2,5 Go	oui	42,7 t/s	Qualité, Raisonnement, Programmation, Outils, Contexte 128K	Standard
qwen3:4b ★	2,6 Go	non	64,3 t/s	Français, Documents, Raisonnement, Outils, Programmation, Contexte 40K	Standard
moondream ★	1,7 Go	non	inconnu	Vision, Vitesse, Contexte 2K	Standard

[Utiliser](#) [Tester](#) [Détails](#) [Supprimer](#) [Installer un autre modèle](#) [Retour](#)

Mes modèles : liste installée, ☐ et actions.

[Paramètres](#) [Aide](#)

DIAGNOSTIC

DIAGNOSTIC

Windows

✓ Windows 11

CPU

✓ Intel Core i7

RAM

✓ 16 Go

RAM disponible

✓ 9,0 Go

GPU

✓ NVIDIA GeForce RTX 3060

VRAM

✓ 12 Go

Ollama

✓ 0.11.0

Serveur local

✓ 127.0.0.1:11434

Cloud Ollama

✓ Désactivé

Modèle

✓ phi4-mini

Mode réel

✓ GPU

Vitesse

✓ 22,4 tokens/s

Espace modèles

✓ 6,8 Go

Studio image

✓ Prêt

État de l'IA : Très bon

Copier

Enregistrer diagnostic.txt

Ouvrir les logs

Retour

Diagnostic — rapport simple à copier, sans vos conversations.

11. Problèmes fréquents

- Message de sécurité Windows : Informations complémentaires → Exécuter quand même.
- Échec de téléchargement (chat ou image) : vérifiez Internet et l'espace disque, puis Réessayer / Reprendre.
- Réponses trop lentes : choisissez un modèle Plus rapide dans Mes modèles.
- Studio image — modules absents : Paramètres → Installer les modules complémentaires (y compris après une mise à jour du logiciel).
- Studio image — « tiktoken / sentencepiece / protobuf / scipy manquant » : réinstallez les modules complémentaires depuis Paramètres.
- Studio image — modèle inutilisable / model_index.json : Supprimer le modèle puis Télécharger à nouveau (ancien format incompatible).
- Studio image — échec immédiat ou mémoire insuffisante : choisissez un modèle plus léger ou une taille plus petite.
- Téléchargement image très long : normal pour 7–10 Go (voir l'estimation et la progression Mo / temps restant).
- Guide introuvable : remplacez guide_utilisation.pdf à côté de LocalAIPrivate-V3.exe.
- Détails techniques : Paramètres → Ouvrir les logs (%LOCALAPPDATA%\LocalAIPrivate\logs\app.log).

12. Mentions

- Produit : IA locale privée V3.0a
- Auteur : Pascal Mouneyrat
- Éditeur : PM Software — (c) PM Software
- Données utilisateur : %LOCALAPPDATA%\LocalAIPrivate\



- Chat / Ollama : models\ et runtime\ollama\
- Studio image (optionnel) : image_runtime\, images_models\, Images_IA\
- Journaux : logs\ (ou dossier logs à côté de l'exécutable si accessible).

IA locale privée V3

IA locale privée V3 est un assistant Windows tout-en-un qui analyse automatiquement votre ordinateur, sélectionne et teste les modèles adaptés, change d'IA selon la tâche et intègre un Studio de génération et de retouche d'images locales optimisé automatiquement selon votre matériel.

- ✓ Chat privé
- ✓ Documents
- ✓ Programmation
- ✓ Analyse d'images
- ✓ Choix automatique du modèle
- ✓ Génération d'images
- ✓ Retouche et variations
- ✓ Aucun abonnement
- ✓ Fonctionnement local



Site PM Software

<http://pascal.mouneyrat.free.fr>

PM Software — Pascal Mouneyrat
Logiciels indépendants & outils spécialisés