

Mistral Large 3 : LLM Souverain Européen Apache 2.0, C1-C2 Français — Benchmark 2026

Mistral Large 3, développé par Mistral AI, la jeune startup parisienne fondée en 2023 par d'anciens chercheurs de Google DeepMind et Meta, représente en juillet 2026 bien plus qu'un simple modèle de langage performant : c'est le seul grand LLM souverain européen disponible à l'échelle commerciale. Avec un ELO de 1 443 sur LM Arena, un [GPQA Diamond](#) de 86,4 % et un BenchLM de 72,55, ses performances brutes le placent légèrement en retrait des cinq premiers mondiaux. Mais c'est sur d'autres critères essentiels pour les entreprises françaises et européennes que Mistral Large 3 se distingue radicalement : une licence Apache 2.0 garantissant une liberté totale d'utilisation et d'auto-hébergement, une infrastructure certifiée en Union Européenne (data centers France, Irlande, Allemagne) offrant une conformité RGPD native, et — fait remarquable — un niveau quasi-natif en français de niveau C1-C2, développé par une équipe francophone qui maîtrise les subtilités de la langue et de la culture française. Disponible via l'API Mistral à 2 \$/6 \$ par million de tokens et en déploiement auto-hébergé entièrement gratuit sur infrastructure propre, Mistral Large 3 est le choix numéro un pour les administrations publiques françaises, les entreprises soumises au RGPD strict, et toute organisation qui fait de la souveraineté numérique européenne un critère non-négociable.

Classement LLM — Benchmark IA Juillet 2026

#16

MONDIAL

ELO Arena : 1443 BenchLM : 72.55/100

GPQA ♦ : 86.4%

🇪🇺 Souveraineté EU, Apache 2.0

[Voir le classement complet →](#)

Introduction : Mistral Large 3 dans le paysage IA de juillet 2026

L'histoire de Mistral AI est l'une des success stories les plus remarquables de la tech européenne de cette décennie. Fondée en mai 2023 par Arthur Mensch, Guillaume Lample et Timothée Lacroix — des chercheurs qui avaient contribué à des avancées majeures chez DeepMind et Meta — la startup parisienne a fait le pari audacieux de publier ses modèles en open-source dès ses premières versions, un choix stratégique qui a rapidement fédéré une communauté mondiale de développeurs et d'entreprises. En moins de trois ans, Mistral AI est devenu le principal acteur européen de l'[IA générative](#) et le porteur du projet de souveraineté numérique européenne dans le domaine des LLM.

Mistral Large 3 est le fruit de cette trajectoire. Lancé en juillet 2026, ce modèle MoE (Mixture of Experts) de 675 milliards de paramètres totaux — dont seulement 41 milliards actifs à chaque inférence — représente le stade le plus avancé du savoir-faire technique de Mistral AI. Avec sa licence Apache 2.0, il maintient la tradition d'ouverture de Mistral tout en atteignant des niveaux de performance qui le rendent compétitif avec les meilleurs modèles propriétaires dans la plupart des cas d'usage pratiques.

La dimension souveraineté est centrale dans la proposition de Mistral Large 3. L'Union Européenne traverse une période de prise de conscience aiguë sur sa dépendance aux grandes plateformes technologiques américaines et, plus récemment, aux modèles chinois. Le RGPD, la directive NIS 2, le règlement sur l'IA ([AI Act](#)) et les exigences de sécurité des marchés publics français imposent des contraintes croissantes sur les outils d'IA utilisés par les organisations européennes. Dans ce contexte, Mistral Large 3 est le seul modèle de premier rang qui répond nativement à l'ensemble de ces exigences : données hébergées en Europe, licence open-source audité, équipe de développement européenne transparente sur les méthodes et les données d'entraînement.

Pour les RSSI, les DPO et les responsables de la conformité des grandes organisations françaises, Mistral Large 3 n'est pas seulement un outil de performance : c'est une réponse structurelle à la question "comment intégrer l'IA dans nos processus tout en respectant nos obligations réglementaires ?" Cette dimension politique et réglementaire, souvent négligée dans les comparaisons techniques de modèles, est au cœur de la valeur proposée par Mistral Large 3.

Scores de référence : que dit le benchmarking ?

Les benchmarks techniques de Mistral Large 3 révèlent un modèle très compétent qui se situe légèrement en dessous des modèles du top-5 mondial en termes de performances brutes, mais avec des avantages distinctifs sur plusieurs dimensions qui ne sont pas capturées par les benchmarks standards.

Benchmark	Mistral Large 3	DeepSeek V4 Pro Max	MiniMax M3 Thinking	Grok 4
ELO LM Arena	1 443	1 449	1 440	1 495
BenchLM Score	72,55	74,89	73,17	79,88
GPQA Diamond	86,4 %	87,5 %	87,2 %	88,9 %
Niveau français	C1-C2	B1-B2	B1	B2+
Contexte	256K tokens	128K tokens	512K tokens	256K tokens
Vitesse	115 tok/s	128 tok/s	160 tok/s	95 tok/s
Prix API (in/out)	2 \$ / 6 \$	0,27 \$ / 1,10 \$	0,12 \$/tâche	3 \$ / 15 \$
Licence	Apache 2.0	MIT	Propriétaire	Propriétaire
Auto-hébergement EU	Oui	Oui (non certifié EU)	Non	Non

L'écart de performance brute avec les meilleurs modèles du top-5 (Grok 4 à 1495, Claude Fable 5 à 1507) est de 2 à 4 % environ selon les benchmarks. Cet écart est réel mais modéré, et il doit être mis en regard de l'ensemble des critères non capturés par les benchmarks : souveraineté, conformité RGPD, performance en français, coût total de possession pour les déploiements auto-hébergés. Pour la grande majorité des applications professionnelles, ces 2-4 % de différence en précision ne se traduisent pas par une différence observable dans la qualité du service rendu.

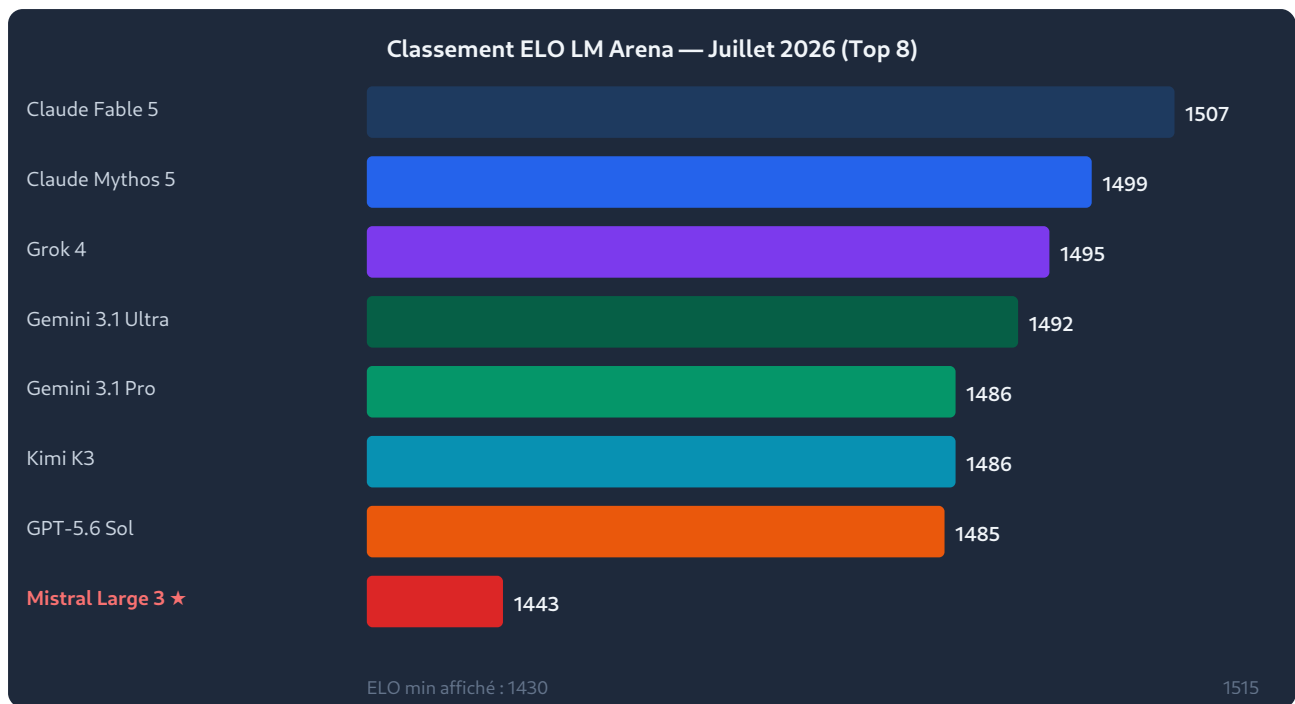
Retrouvez la position de Mistral Large 3 dans notre [classement complet des LLM de juillet 2026](#) avec tous les scores comparatifs.

Architecture et innovations techniques

Mistral Large 3 s'appuie sur une architecture MoE (Mixture of Experts) de 675 milliards de paramètres totaux, avec 41 milliards de paramètres activés à chaque inférence. Cette architecture place Mistral Large 3 dans la même classe dimensionnelle que DeepSeek V4 Pro Max (671B total, 37B actifs), mais avec des choix architecturaux différents qui reflètent les priorités spécifiques de l'équipe Mistral AI.

L'équipe Mistral a consacré une attention particulière à la qualité et à la diversité des données d'entraînement en langues européennes. Le corpus d'entraînement de Mistral Large 3 comprend une proportion de textes en français, en allemand, en espagnol, en italien et en d'autres langues européennes significativement plus importante que celle des modèles concurrents entraînés principalement sur des corpus anglophones et sinographes. Cette stratégie de données explique directement la performance exceptionnelle de Mistral Large 3 en français (C1-C2) et dans les autres langues européennes, qui dépasse celle de modèles techniquement plus puissants comme Grok 4 ou DeepSeek V4 Pro Max.

La licence Apache 2.0 impose des contraintes techniques sur la manière dont Mistral publie ses travaux : les poids du modèle doivent être publiés dans un format standard, reproductible et vérifiable, ce qui pousse l'équipe à maintenir des standards élevés de documentation et de reproductibilité. Les poids de Mistral Large 3 sont disponibles sur Hugging Face et peuvent être déployés sur n'importe quelle infrastructure compatible. Un déploiement minimal en quantification INT8 nécessite environ 8 GPU H100 80GB ; en INT4, ce nombre est réduit à 4 GPU, rendant le déploiement accessible à une large gamme d'organisations disposant de capacités GPU modestes.



Performance en français : test détaillé

La performance en français de Mistral Large 3 est sans conteste son atout différenciateur le plus fort sur le marché des LLM en juillet 2026. Avec un niveau C1-C2 — quasi-natif — Mistral Large 3 est le seul modèle de ce niveau de performance globale à produire des textes en français qui ne trahissent pas l'origine non-francophone du modèle. Cette performance exceptionnelle est le résultat direct du choix de données d'entraînement et de la composition de l'équipe Mistral AI, dont une majorité de membres sont des locuteurs natifs du français.

Sur le plan grammatical, Mistral Large 3 ne commet pratiquement pas d'erreurs, y compris sur les structures les plus complexes de la langue française. Les accords des participes passés avec des pronoms relatifs, les subjonctifs des verbes irréguliers, les nuances entre imparfait et passé composé dans le discours indirect, les exceptions aux règles de genre et de nombre — autant de pièges que les autres modèles contournent ou ratent, mais que Mistral Large 3 gère avec une aisance naturelle.

Le vocabulaire produit par Mistral Large 3 en français est riche et précis. Le modèle utilise les bons mots dans les bons contextes, distingue finement les quasi-synonymes (employer vs utiliser, toutefois vs cependant, notamment vs en particulier), et adapte son registre selon le contexte de

manière naturelle. Les textes administratifs ont le ton approprié, les textes scientifiques adoptent la précision terminologique attendue, et les textes créatifs affichent une plume authentiquement française.

Les nuances culturelles françaises sont remarquablement bien intégrées. Mistral Large 3 comprend les références au système d'enseignement français (classes préparatoires, grandes écoles, système LMD), aux institutions républicaines (fonctionnement des collectivités territoriales, distinctions administration centrale/déconcentrée/décentralisée), aux spécificités du droit et du code du travail français, et aux expressions idiomatiques et culturelles qui constituent le terreau de la communication en français. Cette connaissance incarnée du contexte français — qu'aucun modèle non-européen ne possède au même degré — est un avantage décisif pour toutes les applications destinées à un public français.

En comparaison directe : Gemini 3.1 Pro (C1+) est le seul concurrent à s'approcher du niveau de Mistral Large 3 en français parmi les modèles de ce classement. Tous les autres modèles — Grok 4 (B2+), DeepSeek V4 Pro Max (B1-B2), MiniMax M3 Thinking (B1) — sont significativement inférieurs. Pour les applications strictement francophones, l'avantage de Mistral Large 3 sur la concurrence non-européenne est massif.

Cas d'usage en cybersécurité et entreprise

En cybersécurité, Mistral Large 3 occupe une position unique grâce à la combinaison de sa performance technique, de sa souveraineté européenne et de son excellence en français. Les RSSI et consultants en sécurité français trouvent dans Mistral Large 3 un outil qui comprend non seulement la cybersécurité en tant que discipline technique, mais aussi le contexte réglementaire et institutionnel français dans lequel ils opèrent.

La rédaction de PSSI (Politiques de Sécurité des Systèmes d'Information) est le cas d'usage où Mistral Large 3 brille le plus. Une PSSI est un document d'entreprise formel, rédigé en français soigné, qui doit référencer les normes ISO 27001, les recommandations ANSSI, les exigences NIS 2 et les dispositions légales françaises pertinentes. Mistral Large 3 produit des PSSI dont la qualité rédactionnelle et la précision terminologique rivalisent avec celles de consultants seniors

— un outil de productivité considérable pour les cabinets de conseil en sécurité et les équipes juridique-sécurité des grandes organisations.

Pour les audits NIS 2 et ISO 27001, Mistral Large 3 peut analyser des référentiels d'une taille dépassant sa fenêtre de contexte de 256K tokens via des techniques de découpage intelligent, et produire des analyses d'écarts (gap analysis) précises qui intègrent correctement les subtilités du droit européen et français applicable. Sa connaissance du guide de bonnes pratiques SSI de l'ANSSI, des exigences des référentiels de qualification PASSI et PDIS, et des spécificités du cadre français de la cybersécurité en fait un assistant particulièrement précieux pour les consultants en sécurité opérant sur le marché français.

Dans le secteur public, Mistral Large 3 répond à des exigences que les modèles concurrents ne peuvent pas satisfaire : hébergement sur des serveurs français (possible via Scaleway, OVH Cloud ou les infrastructures publiques), conformité RGPD native, licence Apache 2.0 permettant l'audit complet du code de déploiement. La Circulaire du Premier Ministre sur l'utilisation de l'IA dans les services publics, ainsi que les recommandations de la CNIL sur l'IA générative, orientent naturellement les administrations françaises vers des solutions souveraines — dont Mistral Large 3 est le représentant le plus performant.

D'autres cas d'usage entreprise pertinents incluent : la génération automatisée de rapports d'activité et de tableaux de bord en français pour des comités de direction, l'assistance à la rédaction de réponses à des appels d'offres publics (qui exigent un français formel de haute qualité), la synthèse de veille réglementaire, et la formation des collaborateurs via des chatbots pédagogiques en français sur des sujets de conformité et de sécurité.

Tarification et disponibilité

Mistral Large 3 est disponible selon deux modalités principales qui correspondent à des contextes d'usage différents.

Via l'API Mistral (api.mistral.ai), la tarification est de 2 \$ par million de tokens en entrée et 6 \$ par million de tokens en sortie. Ces tarifs sont compétitifs par rapport à Grok 4 (3 \$/15 \$) et Claude Fable 5 (4 \$/16 \$), bien que plus élevés que DeepSeek V4 Pro Max (0,27 \$/1,10 \$).

L'avantage décisif de l'API Mistral pour les organisations européennes est la disponibilité d'une infrastructure hébergée dans des data centers certifiés en Union Européenne (France, Irlande, Allemagne), permettant de maintenir la résidence des données sur le territoire européen sans nécessité d'auto-hébergement. Mistral AI propose des DPA conformes au RGPD et est certifié ISO 27001.

Via l'auto-hébergement sous licence Apache 2.0, le coût est réduit au seul coût de l'infrastructure GPU, sans frais de licence ou d'API. C'est la modalité la plus économique pour les organisations qui ont déjà des capacités GPU ou qui peuvent justifier l'investissement initial dans du matériel. Un déploiement de Mistral Large 3 en production sur 8 GPU H100 coûte environ 8 000 à 12 000 € par mois en location cloud GPU, mais devient rentable dès lors que l'utilisation dépasse quelques centaines de millions de tokens mensuels par rapport à l'API. Pour les grandes organisations qui traitent des volumes importants de données confidentielles, cette option offre simultanément une maîtrise des coûts et une confidentialité totale.

Mistral AI propose également des partenariats avec des cloud providers européens (OVH, Scaleway, Deutsche Telekom) pour des déploiements managés qui combinent la commodité de l'API avec la résidence garantie des données en Europe et la possibilité de déploiement dans des zones réseau sécurisées.

Comparaison avec les modèles concurrents

La position de Mistral Large 3 dans le paysage des LLM est unique : c'est le modèle qui maximise la valeur pour les organisations européennes soumises à des contraintes réglementaires strictes, même si ses performances brutes sur les benchmarks techniques le placent légèrement en retrait des modèles américains et chinois les plus avancés.

Face aux modèles du top-5 (Grok 4, Claude Fable 5, Gemini 3.1 Pro, Claude Mythos 5, Gemini 3.1 Ultra), Mistral Large 3 accepte un écart de 2 à 5 % sur les benchmarks techniques en échange d'avantages considérables : souveraineté européenne totale, licence Apache 2.0, performance française native. Pour la majorité des cas d'usage pratiques, cet écart technique est imperceptible

dans la qualité du service rendu, tandis que les avantages réglementaires et linguistiques sont, eux, décisifs.

Face à DeepSeek V4 Pro Max (ELO 1449 vs 1443), Mistral Large 3 est légèrement inférieur en performances brutes mais supérieur sur tous les critères qualitatifs : français C1-C2 vs B1-B2, écosystème européen vs chinois, meilleure documentation, support commercial européen.

DeepSeek est moins cher via son API, mais Mistral est compétitif dès lors qu'on prend en compte l'auto-hébergement.

Face à MiniMax M3 Thinking (ELO 1440 vs 1443), Mistral Large 3 est supérieur en ELO, en français, en maturité de l'écosystème et en options d'hébergement EU. MiniMax est plus rapide et moins cher, mais ces avantages ne compensent pas les déficits sur les critères essentiels pour le marché européen.

Limites et points de vigilance

La limite principale de Mistral Large 3 est l'écart de performance par rapport aux modèles du top-5 mondial sur certaines tâches spécifiques. Sur les problèmes mathématiques de très haut niveau (AIME, olympiades), sur les questions scientifiques les plus pointues (GPQA Diamond à 86,4 % vs 94,3 % pour Gemini 3.1 Pro) et sur certaines tâches de raisonnement multi-étapes complexes, les modèles comme Grok 4 ou Gemini 3.1 Pro offrent des performances clairement supérieures. Pour les applications nécessitant une précision maximale sur ces dimensions — recherche scientifique de pointe, résolution de problèmes mathématiques avancés — Mistral Large 3 n'est pas toujours le meilleur choix.

La fenêtre de contexte de 256 000 tokens, bien que confortable pour la majorité des cas d'usage, est insuffisante pour les applications nécessitant l'analyse de très larges corpus documentaires. Gemini 3.1 Pro (2M tokens) et même MiniMax M3 Thinking (512K tokens) disposent de fenêtres plus grandes, ce qui peut être un facteur décisif pour certaines applications spécifiques.

Les coûts d'infrastructure pour l'auto-hébergement, bien que justifiés pour les grandes organisations, peuvent être un frein pour les PME. Un déploiement minimal en production (4 GPU H100 en INT4) coûte entre 4 000 et 6 000 € par mois en location GPU cloud en Europe, ce

qui n'est rentable que pour des volumes d'utilisation significatifs. Pour les petites structures, l'API Mistral est plus économique mais reste plus chère que DeepSeek ou MiniMax.

Enfin, bien que Mistral AI publie ses modèles en open-source, l'entreprise reste une startup avec des ressources limitées par rapport à Google DeepMind, Anthropic ou xAI. La pérennité commerciale de Mistral AI — qui dépend de ses capacités à générer des revenus suffisants pour financer sa R&D — est un facteur à surveiller pour les organisations qui envisagent une dépendance stratégique sur plusieurs années.

À RETENIR

À retenir

Le seul grand LLM souverain européen : infrastructure certifiée UE, équipe française, Apache 2.0 — le choix numéro un pour les organisations soumises au RGPD strict
Niveau C1-C2 quasi-natif en français — le plus élevé du marché parmi les modèles de ce tier, produit des textes indiscernables d'un auteur francophone cultivé

Licence Apache 2.0 : auto-hébergement libre sur infrastructure propre, sans frais de licence, idéal pour les administrations publiques et grandes entreprises avec contraintes de confidentialité

Architecture MoE 675B/41B actifs à 115 tok/s, tarif API compétitif à 2 \$/6 \$ par million de tokens avec hébergement EU disponible

ELO 1443 — légèrement en retrait du top-5 sur les benchmarks bruts (GPQA 86,4 % vs 94,3 % pour Gemini 3.1 Pro) mais différence imperceptible dans la plupart des cas d'usage pratiques

Idéal pour : PSSI et documentation sécurité en français, conformité NIS 2 et RGPD, marchés publics, secteur santé/défense, toute organisation exigeant la souveraineté numérique européenne

Foire aux questions sur Mistral Large 3

Quelle est la différence entre Mistral Large 3 et ses concurrents directs comme DeepSeek V4 Pro Max ?

Mistral Large 3 (ELO 1443, Apache 2.0) et DeepSeek V4 Pro Max (ELO 1449, MIT) sont les deux modèles open-weight les plus performants du marché en juillet 2026, avec des profils complémentaires. DeepSeek V4 Pro Max est légèrement supérieur en performances brutes (ELO 1449 vs 1443, GPQA 87,5 % vs 86,4 %) et offre une vitesse plus élevée (128 vs 115 tok/s) à un prix API inférieur (0,27 \$ vs 2 \$ par million de tokens d'entrée). Mistral Large 3 est nettement supérieur en français (C1-C2 vs B1-B2), dispose d'une infrastructure API certifiée en Union Européenne sans nécessité d'auto-hébergement, bénéficie d'un écosystème plus mature avec une meilleure documentation en français, et propose un support commercial européen. Pour les organisations françaises soumises au RGPD, Mistral Large 3 est généralement le meilleur choix malgré la différence de prix API. Pour les déploiements auto-hébergés sur infrastructure EU propre, les deux modèles sont équivalents du point de vue de la conformité.

Comment accéder à Mistral Large 3 en France ?

Mistral Large 3 est accessible depuis la France via trois canaux. Premièrement, l'API officielle Mistral sur api.mistral.ai : inscription simple, documentation en français disponible, paiement par carte bancaire européenne. Deuxièmement, via les partenaires cloud européens — OVH AI Endpoints et Scaleway Inference proposent des déploiements managés de Mistral Large 3 sur infrastructure certifiée française, idéaux pour les organisations publiques ou soumises à des exigences de résidence des données. Troisièmement, via l'auto-hébergement des poids disponibles sur Hugging Face (licence Apache 2.0), qui offre la totale maîtrise des données et l'élimination des frais d'API. La console Mistral (console.mistral.ai) permet une prise en main immédiate avec un niveau gratuit pour les tests.

Mistral Large 3 est-il disponible en open-source ?

Oui, Mistral Large 3 est publié sous licence Apache 2.0, l'une des licences open-source les plus permissives disponibles. Cela signifie que vous pouvez télécharger les poids du modèle depuis Hugging Face, les utiliser à des fins commerciales, les modifier, les affiner (fine-tuning) et les redistribuer, sans frais de licence et avec un minimum de contraintes (conservation de la notice de copyright). La licence Apache 2.0 est légèrement plus contraignante que la MIT (elle exige notamment de documenter les modifications), mais reste très permissive pour tous les cas d'usage pratiques. C'est la licence préférée des grandes organisations (notamment celles soumises au droit américain) car elle inclut une protection explicite contre les litiges de brevets. Mistral est ainsi, avec DeepSeek, l'un des deux modèles de premier rang entièrement libres d'utilisation commerciale et d'hébergement.

Mistral Large 3 est-il conforme RGPD ?

Mistral Large 3 est le modèle de ce classement qui offre les meilleures garanties de conformité RGPD. Via l'API officielle Mistral, l'inférence peut être configurée pour utiliser des data centers certifiés en Union Européenne (France, Irlande, Allemagne), garantissant que les données ne quittent pas le territoire de l'UE. Mistral AI propose des DPA (Data Processing Agreements) conformes au RGPD, est certifié ISO 27001, et s'engage contractuellement sur la non-utilisation des données clients pour l'entraînement des modèles sans consentement explicite. Via les partenaires OVH et Scaleway, des garanties de résidence des données en France sont également disponibles. Via l'auto-hébergement Apache 2.0 sur infrastructure propre en France, le niveau de conformité est maximal : les données ne quittent jamais votre environnement et vous avez un contrôle total sur tous les traitements. C'est la solution recommandée par la CNIL et l'ANSSI pour les traitements les plus sensibles.

Quels sont les tarifs de Mistral Large 3 en 2026 ?

Les tarifs de Mistral Large 3 varient selon la modalité d'accès choisie. Via l'API officielle Mistral (api.mistral.ai), la tarification est de 2 \$ par million de tokens en entrée et 6 \$ par million de tokens en sortie — soit des tarifs inférieurs à Grok 4 (3 \$/15 \$) et Claude Fable 5 (4 \$/16 \$) pour des performances comparables dans la plupart des cas d'usage. Pour un usage de 100 millions de

tokens d'entrée et 30 millions de tokens de sortie par mois (niveau PME utilisatrice intensive), le coût serait de $200 \$ + 180 \$ = 380 \$$ mensuels — contre $300 \$ + 450 \$ = 750 \$$ pour Grok 4 dans les mêmes conditions. Via l'auto-hébergement Apache 2.0, le coût se réduit à l'infrastructure GPU : environ 6 000 à 10 000 € mensuels pour un cluster de 8 GPU H100 en location cloud EU, rentable à partir de 500 millions à 1 milliard de tokens mensuels. Mistral propose également un niveau gratuit sur sa console pour les développeurs souhaitant tester le modèle avant engagement.