

OWASP Top 10 LLM 2026 : Guide Complet des Vulnérabilités et Remédiations

Guide complet [OWASP Top 10](#) pour LLM 2026 : analyse détaillée des 10 vulnérabilités critiques des systèmes IA avec exemples d'attaques et remédiations pratiques.

En mars 2023, des ingénieurs de Samsung Electronics ont involontairement divulgué du code source propriétaire et des données confidentielles de réunions en les copiant directement dans ChatGPT pour obtenir de l'aide au débogage. Trois incidents distincts en moins d'un mois ont conduit Samsung à interdire temporairement l'usage des LLM grand public en interne. Quelques semaines plus tard, le chercheur en sécurité Kevin Beaumont documentait sur son blog comment il avait réussi à exfiltrer le [system prompt](#) complet de Bing Chat via une injection indirecte cachée dans une page web visitée par l'assistant. Ces deux incidents — parmi des dizaines signalés tout au long de 2024 et 2025 — illustrent concrètement les risques que l'OWASP a formalisés dans sa liste Top 10 pour les Large Language Models. En 2026, cette liste a été mise à jour pour intégrer les nouvelles réalités des agents autonomes, des pipelines RAG complexes et des attaques multimodales. Ce guide analyse en profondeur chacune des dix vulnérabilités, leurs vecteurs d'exploitation concrets, les exemples d'incidents réels documentés, et les remédiations pratiques applicables immédiatement par les équipes de sécurité et les développeurs d'applications IA. Que vous intégriez un LLM dans votre SI, que vous développiez une application basée sur un modèle de fondation ou que vous auditez la sécurité d'un système IA existant, cette référence complète vous donnera les outils nécessaires pour comprendre, évaluer et atténuer les risques associés aux modèles de langage en production.

OWASP Top 10 LLM 2026 : Vulnérabilités et Remédiations

ARCHITECTURE / COMPOSANTS

- Vue d'ensemble : la roue des...
- Tableau de synthèse : les 10...
- LLM01 — Prompt Injection : la menace...
- LLM02 — Sensitive Information...

CONCEPTS CLÉS

Prompt Injection

prompt injection directe

prompt injection indirecte

injections multimodales

scénarios d'attaque concrets

remédiations

Vue d'ensemble : la roue des vulnérabilités OWASP LLM 2026

Avant d'entrer dans le détail de chaque vulnérabilité, voici une représentation visuelle de l'ensemble des dix risques OWASP LLM 2026 avec leur niveau de criticité respectif :



CRITIQUE

ÉLEVÉ

MOYEN

OWASP Top 10 for LLM Applications 2026 — genai.owasp.org

Tableau de synthèse : les 10 vulnérabilités OWASP LLM 2026

ID	Vulnérabilité	Criticité	Vecteur principal	Impact	Remédiation clé
LLM01	Prompt Injection	CRITIQUE	Input utilisateur / documents	Contournement contrôles, exfiltration	Séparation données/ instructions
LLM02	Sensitive Info Disclosure	CRITIQUE	Training data, contexte RAG	Fuite PII, secrets, IP	Data sanitization, output filtering
LLM03	Supply Chain	ÉLEVÉ	Modèles tiers, datasets, plugins	Backdoor, comportements malveillants	Vérification provenance, SBOMs
LLM04	Data & Model Poisoning	ÉLEVÉ	Phase d'entraînement/ fine-tuning	Biais injectés, backdoors comportementaux	Audit datasets, tests adversariaux
LLM05	Improper Output Handling	ÉLEVÉ	Sorties non validées vers systèmes	XSS, injection SQL, RCE	Output encoding, validation stricte
LLM06	Excessive Agency	CRITIQUE	Agents autonomes sur-privilégiés	Actions non autorisées, destruction données	Principe moindre privilège, human-in-the-loop
LLM07	System Prompt Leakage	MOYEN	Extraction via prompting	Exposition logique métier, bypass guardrails	Ne pas stocker secrets dans system prompt
LLM08	Vector & Embedding Weaknesses	MOYEN	Bases vectorielles RAG	Cross-tenant data leakage, manipulation contexte	ACL sur chunks, isolation namespaces
LLM09	Misinformation	MOYEN	Hallucinations, confiance excessive	Décisions erronées, responsabilité légale	RAG + citations sources, validation humaine
LLM10	Unbounded Consumption	MOYEN	Requêtes abusives, inputs géants	DoS, coûts API incontrôlés	Rate limiting, token budgets, quotas

LLM01 — Prompt Injection : la menace la plus critique

La **Prompt Injection** occupe la première place du classement OWASP LLM 2026, et pour cause : c'est la vulnérabilité la plus exploitée dans la nature, avec des centaines d'incidents documentés entre 2024 et 2026. Elle se produit lorsqu'un attaquant réussit à insérer des instructions malveillantes dans le contexte traité par le LLM, lui faisant ignorer ses instructions légitimes ou exécuter des actions non autorisées.



Retour terrain

Dans les projets de déploiement d'IA générative en entreprise, le risque le moins documenté est la fuite de données par les prompts utilisateur. Pour une ESN qui utilisait Claude Enterprise, j'ai analysé 3 mois de logs : 12 % des prompts contenaient des données confidentielles — NDA, données clients, spécifications techniques propriétaires. Les utilisateurs ne font pas la distinction entre leur outil de recherche web et leur assistant IA. La sensibilisation sur ce point précis réduit le taux à 2-3 % en 6 semaines.

L'OWASP distingue deux catégories fondamentalement différentes dans la version 2026. La **prompt injection directe** survient lorsque l'utilisateur interagit directement avec le modèle et insère des instructions adversariales dans sa propre requête. L'exemple classique est le fameux "Ignore all previous instructions and..." qui, malgré sa naïveté apparente, reste partiellement efficace contre certains modèles non durcis. La **prompt injection indirecte** est bien plus dangereuse et représente la nouveauté majeure de 2026 : l'attaquant place ses instructions malveillantes dans des données que le LLM va traiter en tant qu'agent — une page web, un document PDF, un email, un résultat de recherche, ou même un commentaire dans du code source.

En 2025, l'incident "Morris II" a démontré la faisabilité d'un ver auto-répliquant basé sur la prompt injection indirecte : des chercheurs de Cornell ont créé un email malveillant contenant des instructions cachées qui, une fois traité par un agent IA de messagerie, se propageait

automatiquement en envoyant des copies à tous les contacts, tout en exfiltrant des données personnelles. Ce type d'attaque ne nécessite aucune interaction de la victime au-delà de l'utilisation normale d'un assistant IA.

Les nouvelles tendances 2026 incluent les **injections multimodales** : des instructions cachées dans des images (via steganographie ou texte invisible à l'œil humain mais lisible par les LLM vision), des fichiers audio contenant des commandes à basse fréquence, ou encore des injections dans des métadonnées de fichiers. Pour une analyse approfondie des techniques multimodales, consultez notre article sur la [prompt injection avancée et multimodale 2026](#).

Les **scénarios d'attaque concrets** en 2026 incluent : (1) Un chatbot de service client configuré pour ne parler que de produits spécifiques, manipulé via un document joint pour révéler la liste complète des fournisseurs et des tarifs négociés. (2) Un agent IA d'analyse de CVs qui, en lisant un CV piégé, exfiltre vers un serveur distant les CVs de tous les autres candidats stockés dans la base RAG. (3) Un copilote de développement qui, en analysant du code open source compromis, insère des backdoors dans le code de l'entreprise.

Les remédiations

Les **remédiations** recommandées par l'OWASP LLM 2026 sont multiples et doivent être appliquées en couches. En premier lieu, la **séparation des instructions et des données** : ne jamais mélanger dans le même prompt les instructions système et les données utilisateur non vérifiées. Utiliser des délimiteurs forts (balises structurées, encodages spéciaux) pour marquer clairement les frontières. Deuxièmement, implémenter une **validation d'entrée multicouche** : filtrage des patterns d'injection connus, limitation de la longueur des inputs, analyse sémantique des requêtes suspectes. Troisièmement, dans les architectures agentiques, appliquer le principe du **moindre privilège** : un agent qui lit des documents externes ne devrait avoir aucun accès aux APIs d'envoi d'email ou de modification de base de données. Quatrièmement, mettre en place un **human-in-the-loop** pour toute action irréversible : avant d'envoyer un email, de modifier un fichier ou d'exécuter du code, demander une confirmation humaine explicite. Cinquièmement, utiliser des **LLM guards** dédiés — des modèles secondaires formés spécifiquement pour détecter les tentatives d'injection avant qu'elles n'atteignent le modèle principal. Des solutions

comme LLM Guard, Rebuff, ou les gardes intégrés aux plateformes cloud (AWS Bedrock Guardrails, Azure AI Content Safety) offrent une première ligne de défense efficace.

La référence complète de l'OWASP sur ce sujet est disponible sur owasp.org. Il est important de noter qu'il n'existe pas de solution universelle à la prompt injection — la défense en profondeur reste l'approche la plus robuste disponible à ce jour.

LLM02 — Sensitive Information Disclosure : quand l'IA mémorise ce qu'elle ne devrait pas

La **divulgation d'informations sensibles** par les LLM prend plusieurs formes, toutes documentées dans des incidents réels. L'incident Samsung de 2023 reste l'exemple canonique de fuite via utilisation d'un LLM cloud, mais les risques s'étendent bien au-delà de la simple copie imprudente de données confidentielles.

Les LLM sont susceptibles de mémoriser et de reproduire des informations contenues dans leurs données d'entraînement. Des chercheurs de Google DeepMind ont démontré en 2024 qu'il est possible d'extraire des séquences exactes de données d'entraînement de GPT-2 et d'autres modèles via des requêtes répétées ciblées. Cette technique, appelée **training data extraction**, permet potentiellement de retrouver des numéros de téléphone, des adresses email, voire des fragments de code propriétaire inclus dans les datasets de pré-entraînement. En 2025, une étude similaire sur des modèles de code a révélé des clés API, des credentials de bases de données et des tokens d'authentification dans les outputs de plusieurs LLM populaires.

Dans les architectures **RAG (Retrieval-Augmented Generation)**, la surface d'attaque est particulièrement étendue. Si la base de connaissances contient des documents de différents niveaux de classification, un utilisateur mal intentionné peut formuler des requêtes spécifiques pour extraire des informations auxquelles il ne devrait pas avoir accès. Par exemple, dans une entreprise utilisant un assistant RAG alimenté par l'ensemble de la documentation interne, un employé du service commercial pourrait potentiellement accéder à des informations RH confidentielles ou à des données financières non publiques si les contrôles d'accès sur les chunks vectorisés sont insuffisants.

Les **vecteurs de divulgation** identifiés par l'OWASP LLM 2026 incluent : la mémorisation involontaire lors du fine-tuning sur des données sensibles non nettoyées ; l'inclusion excessive d'informations dans le contexte RAG sans filtrage par niveau d'accès ; la régurgitation de PII (Personally Identifiable Information) intégrées dans les exemples few-shot du system prompt ; la génération de sorties contenant des informations inférées à partir de patterns dans les données d'entraînement.

Les **remédiations** pratiques passent par plusieurs axes. Sur les données d'entraînement et de fine-tuning : auditer systématiquement les datasets pour identifier et supprimer les PII, les credentials, et toute information sensible avant l'entraînement ; utiliser des techniques de differential privacy lors du fine-tuning pour réduire la mémorisation. Sur les architectures RAG : implémenter des contrôles d'accès granulaires au niveau des chunks vectorisés, en veillant à ce qu'un utilisateur ne puisse récupérer que les documents auxquels son rôle lui donne accès ; sanitiser les métadonnées des documents avant indexation. Sur les outputs : déployer des filtres de post-traitement qui détectent et masquent les patterns de données sensibles (numéros de carte, IBAN, emails, numéros de sécurité sociale) dans les réponses du modèle ; mettre en place des DLP (Data Loss Prevention) spécifiques aux sorties LLM. Enfin, sur la gouvernance : former les utilisateurs aux risques de divulgation involontaire, établir des politiques claires sur les types de données pouvant être soumis à un LLM cloud, et préférer les déploiements on-premise ou les offres cloud avec garanties de confidentialité renforcées pour les cas d'usage sensibles.

LLM03 — Supply Chain : la chaîne d'approvisionnement IA compromise

La **vulnérabilité de la chaîne d'approvisionnement** des LLM représente un risque systémique souvent sous-estimé. Contrairement aux logiciels traditionnels où la supply chain se limite aux dépendances de code, les systèmes LLM ajoutent plusieurs nouvelles couches : les modèles de fondation eux-mêmes, les datasets d'entraînement, les adaptateurs de fine-tuning, les plugins et extensions, et les APIs tierces intégrées dans les pipelines agentiques.

En 2024, des chercheurs en sécurité ont identifié plusieurs modèles sur HuggingFace contenant des backdoors comportementaux soigneusement dissimulés. Ces backdoors ne s'activaient qu'en présence de triggers spécifiques — une séquence particulière de mots ou un pattern syntaxique

rare — et généraient des réponses malveillantes uniquement dans ces conditions, passant ainsi inaperçus lors des évaluations standards. Pour une analyse approfondie de ces attaques sur HuggingFace, PyPI et npm, consultez notre article sur les [attaques sur la supply chain IA](#).

Le **typosquatting de modèles** est une menace émergente : des acteurs malveillants publient des modèles avec des noms très similaires à des modèles légitimes populaires (par exemple "mistral-7b-instruct-v0.2-fast" au lieu de "mistral-7b-instruct-v0.2"), espérant que des développeurs pressés les intégreront sans vérification. Les **datasets empoisonnés** constituent un autre vecteur : des jeux de données apparemment légitimes sur des plateformes publiques peuvent contenir des exemples conçus pour biaiser le comportement du modèle fine-tuné.

Les **plugins et extensions** représentent la surface d'attaque la plus immédiate pour les déploiements actuels. Les plugins ChatGPT (désormais remplacés par les GPTs et Actions), les outils Anthropic Claude, les extensions LangChain, ou les connecteurs AutoGen peuvent être compromis ou contenir des vulnérabilités qui permettent à un attaquant de détourner l'agent IA pour exécuter des actions malveillantes.

Les **remédiations** s'inspirent des bonnes pratiques DevSecOps appliquées à l'IA. Établir un inventaire complet (AI Bill of Materials / AI-BOM) de tous les composants IA utilisés : modèles, versions, datasets sources, plugins. Vérifier systématiquement les hashes des modèles téléchargés et préférer les sources officielles avec signature cryptographique. Mettre en sandbox les modèles non vérifiés avant tout déploiement en production. Limiter les permissions accordées aux plugins au strict nécessaire. Surveiller les comportements anormaux des modèles en production via des métriques de détection de dérive comportementale.

LLM04 — Data and Model Poisoning : contaminer à la source

Le **data poisoning** et le **model poisoning** sont des attaques qui ciblent respectivement les données d'entraînement et le processus d'apprentissage lui-même. L'objectif de l'attaquant est d'introduire un biais durable dans le comportement du modèle, qui persistera dans toutes les inférences futures — y compris en production.

L'OWASP LLM 2026 distingue plusieurs types d'empoisonnement selon la phase ciblée. Le **pre-training poisoning** consiste à contaminer les données de pré-entraînement, souvent scraped depuis internet à grande échelle. Des chercheurs ont démontré qu'en contrôlant un faible pourcentage (parfois moins de 0,01%) des données d'un large corpus web, un attaquant peut induire des comportements spécifiques dans le modèle résultant. Le **fine-tuning poisoning** cible la phase d'adaptation d'un modèle de fondation à un domaine spécifique : si les données de fine-tuning sont compromises, le modèle adapté héritera des backdoors injectés. Le **RLHF poisoning** (Reinforcement Learning from Human Feedback) est une variante sophistiquée où des évaluateurs malveillants ou compromis fournissent délibérément de mauvaises notations pour orienter le comportement du modèle.

En 2025, plusieurs incidents de poisoning via des datasets communautaires ont été documentés sur des plateformes de partage de données. Des acteurs mal intentionnés ont contribué à des datasets populaires utilisés pour le fine-tuning de LLM médicaux, introduisant des biais subtils dans les recommandations générées par ces modèles. La détection n'est intervenue que plusieurs mois après le déploiement, après observation d'anomalies statistiques dans les outputs.

Les **remédiations** incluent : auditer rigoureusement les sources de données d'entraînement et leur provenance ; utiliser des techniques de détection d'anomalies statistiques sur les datasets avant entraînement ; implémenter des tests adversariaux systématiques pour détecter les backdoors avant déploiement ; surveiller les métriques d'évaluation sur des benchmarks standards après chaque cycle de fine-tuning pour détecter des dérives inattendues ; conserver des checkpoints de modèles propres permettant un rollback en cas de détection de poisoning en production.

LLM05 — Improper Output Handling : des sorties dangereuses non traitées

La vulnérabilité **Improper Output Handling** se produit lorsque les sorties d'un LLM sont transmises à des systèmes en aval sans validation ni assainissement appropriés. Le modèle de langage génère du texte — potentiellement du HTML, du JavaScript, du SQL, des commandes shell, ou du code dans n'importe quel langage — et si cette sortie est directement interprétée par un interpréteur ou un navigateur sans traitement préalable, les conséquences peuvent être désastreuses.

Les **scénarios d'exploitation concrets** sont nombreux. Une application web qui affiche directement dans le DOM les réponses d'un LLM sans encodage est vulnérable aux attaques **XSS (Cross-Site Scripting)** : un attaquant peut manipuler le LLM pour qu'il génère du JavaScript malveillant qui sera exécuté dans le navigateur des autres utilisateurs. Un système qui passe les sorties d'un LLM directement à une base de données sans paramétrage des requêtes est vulnérable aux **injections SQL** : si le LLM génère une chaîne contenant des caractères SQL spéciaux, la requête résultante peut exfiltrer ou modifier des données. Un agent IA qui exécute du code généré par le LLM dans un environnement non sandboxé peut être manipulé pour exécuter des commandes arbitraires (**Remote Code Execution**).

En 2024, plusieurs plateformes de génération de code assistée par IA ont été affectées par cette vulnérabilité. Des chercheurs ont démontré qu'en manipulant les prompts, ils pouvaient faire générer par le LLM du code Python contenant des imports malveillants ou des appels système dangereux qui, si exécutés automatiquement par la plateforme, compromettaient l'environnement de l'utilisateur.

Les **remédiations** s'appuient sur les principes classiques de la sécurité applicative appliqués au contexte LLM. Traiter systématiquement toutes les sorties LLM comme des entrées non fiables, exactement comme on traiterait des données provenant d'un utilisateur non authentifié. Appliquer l'encodage contextuel approprié (HTML encoding pour l'affichage web, paramétrage des requêtes SQL, échappement des arguments shell). Utiliser des parsers stricts pour valider que la sortie correspond au format attendu avant tout traitement. Exécuter du code généré par LLM exclusivement dans des environnements sandboxés avec quotas de ressources et liste blanche d'opérations autorisées. Implémenter une revue humaine obligatoire avant l'exécution de code généré automatiquement dans des contextes sensibles.

LLM06 — Excessive Agency : des agents IA trop puissants

L'**Excessive Agency** est l'une des vulnérabilités les plus critiques introduites par le paradigme agentique et figure parmi les préoccupations prioritaires de l'OWASP LLM 2026. Elle se produit lorsqu'un agent LLM dispose de plus de capacités, de permissions ou d'autonomie que nécessaire

pour accomplir sa mission légitime. Combined avec une prompt injection réussie, cette vulnérabilité peut avoir des conséquences catastrophiques et irréversibles.

La conception des **systèmes agentiques modernes** tend naturellement vers l'attribution de permissions larges pour maximiser l'utilité : un assistant IA qui peut lire et envoyer des emails, accéder à des bases de données, appeler des APIs externes, exécuter du code, et même déployer des ressources cloud est bien plus puissant qu'un assistant limité à la lecture seule. Mais cette puissance se transforme en vecteur d'attaque massif si l'agent est compromis via une prompt injection. Pour une analyse complète de la sécurisation des agents autonomes, consultez notre article sur la [sécurité des agents IA, le sandboxing et les guardrails](#).

Les **cas d'abus documentés** incluent des agents de support client qui, après injection, ont accédé et modifié des données de clients non liés à la conversation ; des agents de développement qui ont commité et pushé du code malveillant dans des dépôts de production ; des agents d'analyse financière qui ont déclenché des transactions non autorisées. En 2025, un incident majeur dans une entreprise de e-commerce a impliqué un agent IA d'optimisation des stocks qui, après avoir été manipulé via un email de fournisseur piégé, a passé des commandes pour des millions d'euros de produits vers des adresses de livraison frauduleuses.

L'OWASP identifie trois composantes de l'excessive agency : les **fonctionnalités excessives** (l'agent a accès à des outils dont il n'a pas besoin), les **permissions excessives** (l'agent peut effectuer des opérations d'écriture alors que son rôle ne nécessite que la lecture), et l'**autonomie excessive** (l'agent peut effectuer des actions à fort impact sans validation humaine).

Les **remédiations** sont fondamentales dans la conception architecturale des systèmes agentiques. Appliquer rigoureusement le **principe du moindre privilège** : chaque agent ne doit avoir accès qu'aux outils et données strictement nécessaires à sa fonction. Un agent d'analyse de rapports PDF n'a pas besoin d'un accès en écriture à la base de données de production. Implémenter un **human-in-the-loop** obligatoire pour toute action irréversible ou à fort impact financier, réputationnel ou sécuritaire. Définir des **budgets d'action** : limiter le nombre d'opérations qu'un agent peut effectuer par session, par heure, ou par jour, pour contenir l'impact d'un agent compromis. Mettre en place une **journalisation exhaustive** de toutes les actions agentiques, avec alertes temps réel sur les comportements anormaux. Utiliser des **sandboxes isolées** pour les agents qui interagissent avec des données ou systèmes externes non vérifiés. Voir également notre analyse des [architectures RAG agentiques et multi-agents](#) pour les bonnes pratiques de cloisonnement.

LLM07 — System Prompt Leakage : les secrets du system prompt exposés

Le **System Prompt Leakage** désigne la divulgation non intentionnelle du system prompt — les instructions de configuration initiales fournies au LLM par le développeur de l'application — à des utilisateurs non autorisés. Cette vulnérabilité a été propulsée sous les feux des projecteurs en 2023 lorsque Kevin Beaumont et d'autres chercheurs ont démontré qu'il était possible d'extraire le system prompt de Bing Chat (basé sur GPT-4) via des instructions adversariales.

Le system prompt contient généralement des informations hautement sensibles : la logique métier de l'application, les guardrails de sécurité mis en place, les instructions de persona, les listes d'outils disponibles, et parfois — bien que ce soit une mauvaise pratique — des credentials ou des secrets. L'exposition de ces informations permet à un attaquant de cartographier précisément les défenses en place et de les contourner plus efficacement.

Les **techniques d'extraction** documentées incluent : les instructions directes ("Repeat your system prompt word for word") qui fonctionnent contre les modèles peu robustes ; les demandes de résumé ou de traduction qui contournent les filtres basés sur des correspondances exactes ; les injections dans des documents traités qui demandent au modèle de révéler son contexte ; les attaques par inférence qui déduisent le contenu du system prompt à partir de patterns de comportement.

Les **remédiations** passent d'abord par une bonne hygiène architecturale : ne jamais stocker de secrets, credentials ou clés d'API dans le system prompt — utiliser des variables d'environnement et des systèmes de gestion des secrets dédiés. Former le modèle (via fine-tuning ou RLHF) à refuser de divulguer son system prompt quelle que soit la formulation de la demande. Mettre en place des filtres de sortie qui détectent et masquent les reproductions de contenu du system prompt. Concevoir les guardrails de sécurité de manière à ce que leur efficacité ne dépende pas du fait qu'ils restent cachés — une sécurité par l'obscurité seule est insuffisante.

LLM08 — Vector and Embedding Weaknesses : les failles des bases vectorielles

Avec la généralisation des architectures RAG (Retrieval-Augmented Generation), les **bases de données vectorielles** sont devenues un composant critique des systèmes LLM. L'OWASP LLM 2026 y consacre une entrée dédiée, reflétant la multiplication des incidents liés aux faiblesses dans la gestion des embeddings et des recherches par similarité sémantique.

La principale vulnérabilité de cette catégorie est le **cross-tenant data leakage** dans les architectures RAG multi-utilisateurs. Si les contrôles d'accès ne sont pas implémentés au niveau des chunks vectorisés, un utilisateur peut potentiellement récupérer des documents appartenant à d'autres utilisateurs ou organisations en formulant des requêtes sémantiquement proches de leur contenu. Les bases vectorielles comme Pinecone, Weaviate, Milvus ou ChromaDB offrent des mécanismes de filtrage par métadonnées, mais ces mécanismes doivent être explicitement configurés — ils ne sont pas actifs par défaut.

Les **attaques de poisoning sur les bases vectorielles** constituent une autre menace : un attaquant qui peut insérer des documents dans une base RAG peut y inclure des instructions adversariales déguisées en contenu légitime, qui seront récupérées et incluses dans le contexte du LLM lors de requêtes spécifiques. Cette forme d'injection indirecte est particulièrement difficile à détecter car le document malveillant peut passer tous les contrôles de contenu de surface.

La **manipulation de similarité sémantique** est une attaque émergente : en comprenant le modèle d'embedding utilisé, un attaquant peut crafting des documents qui seront systématiquement récupérés pour certaines requêtes, même si leur contenu est superficiellement différent. Cela permet de "contaminer" le contexte du LLM de manière ciblée.

Les **remédiations** incluent : implémenter des contrôles d'accès stricts au niveau des namespaces ou des métadonnées dans la base vectorielle, en filtrant les résultats de recherche par identité d'utilisateur ou de tenant avant de les inclure dans le contexte RAG ; auditer régulièrement les documents indexés pour détecter du contenu potentiellement malveillant ; utiliser des techniques de chunking qui préservent l'intégrité sémantique des documents et rendent plus difficile l'injection de contexte malveillant fragmenté.

La **désinformation** générée par les LLM — communément appelée "hallucination" — est reconnue par l'OWASP LLM 2026 non seulement comme un problème de fiabilité, mais comme un risque de sécurité à part entière. Les LLM génèrent des informations fausses avec une confiance apparente et de manière convaincante, ce qui peut conduire à des décisions organisationnelles erronées, à des responsabilités légales, voire à des atteintes à la sécurité physique dans des contextes critiques.

Les **contextes à haut risque** identifiés par l'OWASP incluent : les systèmes IA d'aide à la décision médicale qui peuvent générer des recommandations de traitement incorrectes ; les assistants juridiques qui peuvent créer des jurisprudences ou des textes de loi inventés (le tristement célèbre cas Mata v. Avianca en 2023, où un avocat avait soumis des citations de jurisprudences entièrement fabriquées par ChatGPT) ; les systèmes de cybersécurité qui peuvent générer des faux positifs ou manquer des vraies menaces ; les agents IA qui prennent des décisions automatisées basées sur des informations factuelles incorrectes.

Les **remédiations** s'articulent autour de l'architecture et des processus. Sur le plan technique : déployer des architectures RAG avec citation des sources pour que chaque affirmation du LLM soit traçable à un document source vérifiable ; implémenter des mécanismes de grounding qui contraignent le modèle à répondre uniquement à partir de sa base de connaissances vérifiée ; utiliser des LLM avec des capacités de reasoning explicite qui permettent de suivre la chaîne de déduction. Sur le plan organisationnel : établir des politiques claires sur les cas d'usage où la validation humaine est obligatoire avant action ; former les utilisateurs à comprendre les limitations des LLM et à maintenir un regard critique sur leurs outputs ; ne jamais utiliser un LLM seul pour des décisions à fort enjeu sans vérification croisée. Pour une exploration des systèmes RAG avancés qui réduisent les hallucinations, consultez notre article sur le [RAG agentique et GraphRAG 2026](#).

LLM10 — Unbounded Consumption : déni de service et épuisement des ressources

L'**Unbounded Consumption** est la vulnérabilité la plus "classique" de la liste, mais dans un contexte LLM, elle prend des dimensions nouvelles en raison des coûts computationnels considérables associés à l'inférence des grands modèles. Elle englobe les attaques de type Denial of Service (DoS) spécifiques aux LLM, l'abus des ressources d'inférence, et l'épuisement des budgets API.

Les **vecteurs spécifiques aux LLM** incluent : les **inputs extrêmement longs** qui exploitent les limites de contexte pour ralentir ou bloquer l'inférence — certains modèles voient leur temps d'inférence augmenter de manière quadratique avec la longueur du contexte ; les **requêtes de génération intensive** qui demandent des outputs très longs (générer 50 000 tokens de code) consommant des ressources GPU pour une durée disproportionnée ; les **requêtes répétitives en burst** qui saturent les files d'attente d'inférence ; les **adversarial inputs** spécifiquement conçus pour maximiser la consommation computationnelle du modèle.

L'aspect financier est particulièrement préoccupant pour les organisations utilisant des LLM via APIs facturées au token. Un attaquant qui peut déclencher des milliers de requêtes longues vers une application exposant un LLM peut générer des factures API astronomiques en quelques heures. Des incidents de ce type, parfois appelés "billing attacks", ont été documentés contre des startups qui avaient exposé leurs endpoints LLM sans rate limiting adéquat.

Les **remédiations** s'appuient sur des mécanismes de contrôle à plusieurs niveaux. Implémenter un **rate limiting** strict par utilisateur, par IP, et par organisation, avec des limites adaptées aux patterns d'usage légitimes attendus. Définir des **budgets de tokens** maximaux par requête (input et output) et rejeter ou tronquer les requêtes qui les dépassent. Mettre en place des **timeouts d'inférence** pour éviter que des requêtes pathologiques ne monopolisent les ressources GPU. Surveiller en temps réel les métriques de consommation avec des alertes sur les anomalies. Pour les déploiements self-hosted, implémenter des quotas de ressources GPU par tenant. Envisager des mécanismes de **queue prioritaire** qui permettent aux requêtes légitimes de continuer à être traitées même sous attaque.

Roadmap de remédiation — par ordre de priorité

Face à dix vulnérabilités à adresser, les équipes de sécurité doivent prioriser leurs efforts. Cette roadmap s'appuie sur la combinaison de la criticité OWASP, de la facilité d'exploitation dans la pratique, et de l'effort de remédiation relatif.

Phase 1 — Actions immédiates (J0 à J30) : Commencer par l'inventaire complet de tous les composants LLM en production (modèles, versions, plugins, APIs tierces) — c'est le prérequis à toute action de sécurisation. Activer les guardrails fournis par les plateformes cloud utilisées (AWS Bedrock Guardrails, Azure AI Content Safety, Vertex AI Safety Filters) — ces contrôles sont disponibles sans développement spécifique et réduisent immédiatement la surface d'attaque pour LLM01 et LLM02. Auditer les permissions accordées aux agents IA en production et appliquer le principe du moindre privilège (LLM06) — cette action peut être réalisée en quelques jours et a un impact significatif sur le risque. Implémenter le rate limiting sur tous les endpoints LLM exposés (LLM10) — une mesure simple qui prévient les abus financiers.

Phase 2 — Renforcement architectural (M1 à M3) : Revoir l'architecture RAG pour implémenter les contrôles d'accès par niveau utilisateur sur les chunks vectorisés (LLM08). Déployer des filtres de sortie pour détecter et masquer les données sensibles dans les réponses LLM (LLM02). Établir un processus de vérification de la provenance des modèles et datasets utilisés (LLM03). Mettre en place un human-in-the-loop pour toutes les actions agentiques irréversibles (LLM06). Ces actions nécessitent des modifications architecturales mais peuvent être réalisées par phases sans refonte complète.

Phase 3 — Maturité sécurité IA (M3 à M6) : Développer un programme de red teaming spécifique LLM avec des tests réguliers de prompt injection, d'extraction de system prompt et de contournement de guardrails — consultez notre guide sur le [red teaming IA et le jailbreak](#) pour une méthodologie complète. Intégrer les tests de sécurité LLM dans le pipeline CI/CD pour détecter les régressions de sécurité à chaque déploiement. Établir un processus de surveillance continue des comportements anormaux en production. Former l'ensemble des équipes développement et product à la sécurité LLM — la sensibilisation est souvent le facteur le plus coût-efficace dans la réduction du risque.

Métriques de suivi recommandées : taux de détection des tentatives d'injection (LLM01), nombre d'incidents de divulgation de données (LLM02), couverture des tests adversariaux

(LLM04, LLM09), temps moyen de détection des comportements agentiques anormaux (LLM06), taux d'utilisation des ressources LLM vs baseline (LLM10). Ces métriques permettent de mesurer la maturité du programme de sécurité IA et de justifier les investissements auprès de la direction.

FAQ — Questions fréquentes sur l'OWASP Top 10 LLM 2026

Quelle est la différence entre l'OWASP Top 10 LLM 2025 et la version 2026 ?

La version 2026 de l'OWASP Top 10 LLM intègre plusieurs évolutions majeures par rapport à la version 2025. Les changements les plus significatifs concernent la montée en importance de LLM06 (Excessive Agency), qui passe de la cinquième à la sixième position mais gagne en criticité avec la multiplication des déploiements d'agents autonomes en production. LLM08 (Vector and Embedding Weaknesses) est une entrée enrichie qui prend en compte la généralisation des architectures RAG et les nouveaux vecteurs d'attaque associés. La version 2026 intègre également de nouveaux scénarios d'attaque multimodale pour LLM01, reflétant la sophistication croissante des LLM vision. Enfin, LLM09 (Misinformation) gagne en importance avec la multiplication des usages à fort enjeu (médical, juridique, financier) et les premières décisions judiciaires sur la responsabilité des organisations utilisant des LLM dans leurs processus métier. L'OWASP met régulièrement à jour sa documentation sur genai.owasp.org.

Comment tester concrètement les vulnérabilités OWASP LLM sur une application en production ?

Le testing des vulnérabilités LLM s'effectue via une combinaison d'outils automatisés et de tests manuels. Pour LLM01 (Prompt Injection), des frameworks comme **Garak** (développé par NVIDIA Research) et **PyRIT** (Microsoft) permettent d'automatiser les tests d'injection contre un endpoint LLM avec des centaines de techniques documentées. Pour LLM02, des outils de scanning de training data comme **Presidio** de Microsoft permettent d'identifier les PII dans les datasets. Pour l'architecture globale, des plateformes spécialisées comme **Lakera Guard**, **Robust Intelligence**, ou **Protect AI** offrent des capacités de testing continu. Sur le plan méthodologique, l'approche recommandée combine un audit initial exhaustif (couvrant les 10 catégories), des tests de régression automatisés intégrés au CI/CD, et des exercices de red teaming trimestriels avec des testeurs humains. Notre article sur le [red teaming IA](#) détaille la méthodologie complète et les outils disponibles.

Les petites entreprises sans équipe de sécurité dédiée peuvent-elles appliquer l'OWASP Top 10 LLM 2026 ?

Absolument, et l'OWASP Top 10 LLM est précisément conçu pour être accessible à tous. Pour les organisations sans ressources dédiées, une approche pragmatique en trois étapes est recommandée. Premièrement, commencer par les "quick wins" à faible effort et fort impact : activer les guardrails des fournisseurs cloud (AWS, Azure, Google Cloud offrent tous des protections natives), implémenter le rate limiting sur les endpoints LLM, et revoir les permissions des agents IA. Deuxièmement, utiliser les LLM managés de grands fournisseurs (OpenAI GPT-4o, Anthropic Claude, Google Gemini) qui intègrent déjà des protections substantielles contre LLM01, LLM02 et LLM07, plutôt que de déployer des modèles open source sans protections. Troisièmement, s'appuyer sur les ressources gratuites de l'OWASP : le guide complet est disponible gratuitement sur owasp.org, avec des checklists pratiques

téléchargeables. La formation des développeurs à la sécurité LLM via les ressources OWASP gratuites est souvent l'investissement le plus rentable pour les petites structures.

Gouvernance et conformité : intégrer l'OWASP LLM dans votre programme de sécurité

L'adoption de l'OWASP Top 10 LLM 2026 ne doit pas être traitée comme un exercice de conformité ponctuel, mais comme un cadre vivant intégré à la gouvernance globale de la sécurité de l'information. Dans un contexte où le règlement européen sur l'IA (EU AI Act) entre en vigueur progressivement jusqu'en 2027, les organisations qui déploient des systèmes LLM dans des cas d'usage à risque élevé ou limité ont des obligations réglementaires directes qui recoupent largement les recommandations OWASP.

La **cartographie des risques LLM** doit devenir un composant standard des évaluations de risque IT. Pour chaque système LLM en production ou en cours de développement, il convient de documenter : le type de modèle utilisé et sa provenance, les données accessibles par le système, les actions que l'agent peut effectuer, les utilisateurs qui interagissent avec lui, et les conséquences potentielles d'un incident pour chacune des dix catégories OWASP. Cette cartographie permet de prioriser les investissements de sécurité et de démontrer la diligence raisonnable (due diligence) en cas d'incident.

L'intégration dans le **cycle de développement sécurisé (SDL/SSDLC)** est essentielle pour que la sécurité LLM ne soit pas une réflexion tardive. Concrètement, cela signifie : inclure des critères d'acceptation de sécurité LLM dans les user stories produit, réaliser des threat models spécifiques LLM en phase de conception, intégrer des tests automatisés de prompt injection dans les pipelines CI/CD, et inclure la sécurité LLM dans les checklists de revue de code. Les équipes AppSec doivent développer de nouvelles compétences spécifiques à la sécurité des LLM, car les outils et techniques traditionnels (SAST, DAST) ne couvrent pas les vecteurs d'attaque propres aux modèles de langage.

Sur le plan des **relations avec les fournisseurs**, les organisations doivent désormais inclure des clauses contractuelles spécifiques à la sécurité IA dans leurs accords avec les fournisseurs de

LLM et de plateformes agentiques. Ces clauses doivent couvrir : la transparence sur les données d'entraînement et les mesures de sécurité mises en œuvre, les garanties de confidentialité des données soumises au modèle, les délais de notification en cas d'incident de sécurité affectant le modèle, et les engagements de patching en cas de vulnérabilités découvertes dans les composants IA.

La **sensibilisation des utilisateurs finaux** est souvent négligée dans les programmes de sécurité LLM, qui se concentrent sur les aspects techniques. Pourtant, comme l'incident Samsung l'a démontré, des utilisateurs non formés constituent l'un des vecteurs de risque les plus importants pour LLM02 (Sensitive Information Disclosure). La formation doit couvrir : les types de données qui ne doivent jamais être soumis à un LLM cloud, la reconnaissance des tentatives de manipulation sociale via des outils IA, les procédures de signalement d'un comportement suspect du système, et les politiques d'utilisation acceptable des LLM en entreprise. Cette sensibilisation doit être régulièrement mise à jour au rythme de l'évolution des menaces.

Enfin, la **veille sur les nouvelles vulnérabilités LLM** doit être organisée et systématique. Le domaine évolue extrêmement rapidement : de nouvelles techniques d'attaque sont publiées chaque semaine par des chercheurs académiques et des red teamers professionnels. Les sources à surveiller incluent les publications de l'OWASP GenAI Security Project, les conférences Black Hat et DEF CON (avec leurs tracks spécifiques IA depuis 2024), les blogs des équipes de sécurité des grands laboratoires IA (Anthropic, OpenAI, Google DeepMind), et les CVE databases qui commencent à documenter les vulnérabilités spécifiques aux LLM. L'abonnement aux bulletins de sécurité des fournisseurs de LLM utilisés en production est également indispensable.

Points clés à retenir

LLM01 Prompt Injection reste la vulnérabilité la plus critique et la plus exploitée — la défense en profondeur (séparation données/instructions + LLM guards + moindre privilège) est indispensable, particulièrement pour les architectures agentiques.

LLM06 Excessive Agency est la vulnérabilité émergente la plus dangereuse de 2026 : les agents autonomes sur-privilegiés peuvent causer des dommages irréversibles en quelques secondes si compromis. Human-in-the-loop obligatoire pour toute action à fort impact.

La supply chain IA (LLM03) est systématiquement sous-estimée : modèles, datasets, plugins et dépendances doivent être traités avec la même rigueur que la supply chain logicielle traditionnelle. Un AI-BOM est indispensable.

Les architectures RAG (LLM08) nécessitent des contrôles d'accès explicites au niveau des chunks vectorisés — les bases vectorielles ne sont pas sécurisées par défaut pour les environnements multi-tenant.

La priorité de remédiation doit être : guardrails natifs cloud (J0) → moindre privilège agentique (J7) → rate limiting (J14) → contrôles RAG (M1) → red teaming (M3).

EU AI Act : les obligations réglementaires pour les systèmes IA à risque élevé recoupent directement les recommandations OWASP — intégrer le cadre dans votre programme de conformité IA.

Ressources officielles : documentation complète sur owasp.org et genai.owasp.org.

Sources et références

[ANSSI — Guide de sécurité des systèmes d'IA](#)

[MITRE ATLAS — Matrice des menaces adversariales IA](#)

[OWASP LLM Top 10](#)