

Shadow Agents : Le Nouveau Risque Insider Causé par l'IA Agentique

Les **shadow agents** représentent l'évolution la plus dangereuse du **shadow IT** : des agents IA autonomes déployés hors de tout contrôle SSI, capables d'agir sur des systèmes d'entreprise sans supervision humaine. Contrairement au shadow IT classique qui se contentait de consulter, un shadow agent exécute, modifie, envoie et compromet. Ce guide analyse la nature du risque, les vecteurs de déploiement, et les mesures de détection applicables dès maintenant.

Les shadow agents constituent le **risque insider de nouvelle génération** que la plupart des équipes SSI n'ont pas encore cartographié. Un shadow agent, c'est un agent IA autonome — construit sur [LangChain](#), AutoGPT, [CrewAI](#) ou les API Claude/GPT-4 — déployé par un collaborateur sans validation IT, sans revue de sécurité, sans contrôle des permissions. Selon les analyses publiées par l'[OWASP LLM Top 10 2025](#), le risque d'Excessive Agency (LLM06) est précisément ce scénario : un agent qui obtient des droits et des outils disproportionnés par rapport à la tâche définie. En 2026, avec la démocratisation des frameworks agentiques et la disponibilité d'API LLM à moins de 10 euros par million de tokens, le risque n'est plus théorique. Des collaborateurs dans des fonctions finance, RH, marketing et IT déploient des agents locaux ou [SaaS](#) qui se connectent à leur messagerie, leurs outils de collaboration, leurs accès cloud — sans que la DSI en soit informée.

À retenir

Shadow agents ≠ shadow IT : Un outil non-autorisé consulte. Un agent non-autorisé agit — il envoie des emails, modifie des fichiers, appelle des APIs, exfiltre des données.

OWASP LLM06 — Excessive Agency : C'est la taxonomie officielle de ce risque. Un agent qui reçoit plus de permissions qu'il n'en a besoin devient un vecteur d'attaque interne.

Vecteur RGPD : Un shadow agent qui accède à des données personnelles crée une obligation de notification CNIL en cas d'incident — même si l'action était "bénigne".

Détection par les logs API : Les appels aux APIs LLM externes (OpenAI, Anthropic, Mistral) laissent des traces dans les proxies et les CASB — c'est le premier signal de détection.

Politique d'urgence : Sans politique formelle sur les agents IA d'ici 6 mois, toute entreprise de plus de 50 collaborateurs est statistiquement exposée à un incident shadow agent.

Shadow Agents vs Shadow IT : pourquoi la différence est fondamentale ?

Le shadow IT classique, c'était Dropbox pour stocker des fichiers ou WhatsApp pour coordonner des équipes. Gênant, parfois coûteux en termes de conformité, mais passif. Le collaborateur consultait, copiait, partageait. Un **shadow agent** change de paradigme : il est doté d'outils (*tools*) qui lui permettent d'agir de manière autonome sur les systèmes de l'entreprise. La différence n'est pas de degré, c'est de nature.

Prenons un exemple concret. Un chef de projet déploie un agent CrewAI sur son poste qui se connecte à son compte Outlook via l'API Graph, à Salesforce via le connecteur natif, et à son espace SharePoint. Il lui demande de "gérer ses relances commerciales automatiquement". En pratique, l'agent envoie des emails en son nom, crée des opportunités Salesforce, lit et modifie des fichiers SharePoint partagés. Sans autorisation. Sans audit trail. Sans politique de rétention. C'est un incident RGPD potentiel qui court en silence.

Sur un audit de maturité IA réalisé début 2026 chez une ETI de 300 personnes, nous avons identifié 7 agents "maison" actifs dans l'organisation — aucun référencé dans l'inventaire IT. Deux d'entre eux avaient accès aux boîtes email de leurs créateurs via OAuth, avec des tokens jamais révoqués. L'un envoyait des emails automatiques depuis le compte d'un commercial parti 3 mois plus tôt, dont le compte n'avait pas été désactivé.

— *Retour de mission, audit IA, mars 2026*

Comment les shadow agents se déploient sans validation SSI

Les vecteurs de déploiement sont multiples et difficiles à bloquer sans une politique claire. Le plus courant reste l'installation locale : un développeur ou un power user installe Python, pip-installe LangChain ou AutoGPT, configure des clés API personnelles ou d'entreprise, et lance son agent sur son poste de travail. Aucune trace dans le CMDB. Aucun ticketing. Aucun contrôle.

Le deuxième vecteur, en forte croissance en 2026, passe par les **plateformes SaaS agentiques** : Zapier AI, Make.com agents, n8n hébergé, ou des outils verticaux comme Harvey AI (juridique) ou Jasper (marketing). Ces plateformes proposent des agents préconfigurés avec des connecteurs OAuth vers les outils d'entreprise. Un clic suffit pour autoriser l'accès. Côté DSI, le seul signal visible est un nouveau consentement OAuth dans les logs Azure AD — si quelqu'un regarde.

Agents locaux Python (LangChain, AutoGPT, CrewAI) avec clés API personnelles

Plateformes SaaS agentiques (Zapier AI, Make.com, n8n, Voiceflow)

Extensions VS Code ou Cursor AI avec accès aux repos Git d'entreprise

Agents MCP connectés à des outils internes via serveurs MCP non-contrôlés

Anatomie d'un incident shadow agent en entreprise

La chronologie type d'un incident shadow agent suit un schéma prévisible. Tout commence par une intention bénigne : un collaborateur veut automatiser une tâche répétitive. Il construit ou configure un agent, lui donne accès à ses outils professionnels, et le met en production sur son poste ou sur une plateforme cloud. L'agent fonctionne — peut-être très bien. Jusqu'au moment où quelque chose déraile.

Le déraillement peut prendre plusieurs formes. Une **prompt injection** via un email malveillant que l'agent traite et qui modifie son comportement (voir [notre analyse des attaques par prompt injection indirecte](#)). Un bug dans la logique de l'agent qui entraîne une boucle d'actions non désirées. Une révocation OAuth qui échoue et laisse l'agent actif après le départ du collaborateur. Ou simplement un agent qui accède à plus de données qu'il ne le devrait parce que les permissions OAuth accordées étaient trop larges.

Les actions que réalisent les shadow agents sans supervision

La question que posent souvent les DSI est : "Qu'est-ce qu'un agent non-supervisé peut faire concrètement ?" La réponse dépend des *tools* qui lui sont donnés. Avec un accès Gmail/Outlook : envoyer, supprimer, archiver, répondre en masse. Avec un accès SharePoint/Drive : lire tous les fichiers partagés, créer, modifier, télécharger des documents sensibles. Avec un accès Salesforce : exporter des données clients, modifier des opportunités, créer des contacts. Avec un accès GitHub : pusher du code, ouvrir des issues, accéder aux secrets stockés dans les repos.

Ce qui rend le risque particulièrement sérieux, c'est la capacité des agents à **chaîner ces actions**. Un agent qui a accès à l'email ET à SharePoint ET à une base de données clients peut construire une chaîne d'actions qui extrait des données personnelles, les formate et les envoie par email à

une adresse externe — le tout en réponse à une instruction ambiguë de son créateur, sans aucune intention malveillante initiale.

Pourquoi le risque insider est amplifié par les agents IA ?

Le risque insider traditionnel — un collaborateur qui exfiltre des données — est limité par la capacité humaine. Un humain peut copier quelques gigaoctets, pas des téraoctets. Il peut envoyer quelques centaines d'emails par jour. Il est contraint par la vitesse de traitement humaine. Un agent IA n'a pas ces contraintes. Il peut traiter des milliers de documents en quelques minutes, envoyer des milliers d'emails en quelques secondes, et coordonner des actions sur plusieurs systèmes simultanément.

Le [framework MITRE ATLAS](#), qui classe les tactiques d'attaque contre les systèmes IA, identifie précisément ce vecteur sous la tactique AML.T0040 — Exfiltration via ML Model. Un agent IA mal configuré ou compromis peut devenir un vecteur d'exfiltration bien plus efficace qu'un humain malveillant. Et contrairement à l'insider humain, il ne présente aucun des signaux comportementaux classiques : pas de connexion à des heures inhabituelles, pas de stress visible, pas de comportement suspect au bureau.

Comment détecter les shadow agents dans votre organisation ?

La détection des shadow agents repose sur cinq couches d'observabilité que peu d'organisations ont combinées en 2026. La première : les **logs OAuth et consentements API**. Chaque agent qui se connecte à un service Microsoft 365, Google Workspace ou Salesforce laisse une empreinte OAuth dans les logs d'identité. Un audit mensuel des applications tierces autorisées détecte les nouveaux agents en quelques minutes.

La deuxième couche : les **proxies réseau et CASB**. Les appels aux APIs LLM externes (api.openai.com, api.anthropic.com, api.mistral.ai) sont visibles dans les logs proxy. Un CASB

bien configuré peut détecter ces appels et les corrélérer avec les postes source. La troisième couche : les **logs de téléchargement massif**. Un agent qui exfiltre des données déclenche des patterns de téléchargement inhabituels dans SharePoint, OneDrive ou les SIEM.

Audit OAuth mensuel : lister toutes les apps tierces avec accès délégué (Azure AD → Enterprise Applications)

Blocage proxy des domaines API LLM non-approuvés (api.openai.com, api.anthropic.com non-canalises)

Alertes CASB sur volumes de téléchargement > seuil par utilisateur

Inventaire des installations Python/npm avec packages LangChain, AutoGPT, CrewAI

Revue des Copilot Studio workspaces créés sans ticket IT

Pour aller plus loin sur la détection des usages IA non-autorisés, notre guide sur le [Shadow AI en entreprise](#) couvre les techniques d'audit complémentaires.

Cas concrets : shadow agents et violations RGPD

L'angle RGPD des shadow agents est souvent sous-estimé. Un agent qui accède à des données personnelles — même en lecture seule — constitue un traitement de données au sens du RGPD. Si ce traitement n'est pas documenté dans le [registre des activités de traitement](#), l'organisation est déjà en infraction. Si l'agent est hébergé sur une plateforme SaaS américaine sans DPA adapté, c'est un transfert de données hors UE potentiellement illégal.

Cas documenté publiquement : en 2025, plusieurs entreprises européennes ont découvert que leurs collaborateurs utilisaient des agents IA tiers qui envoyaient des données clients vers des serveurs américains pour traitement. La CNIL a commencé à examiner ce type de cas dans le cadre de ses contrôles sur l'IA en 2026. La [loi sur l'IA \(EU AI Act\)](#), entrée en vigueur progressivement depuis 2024, impose des obligations supplémentaires sur certains usages agentiques classés à risque limité ou élevé.

Tableau : comparaison shadow IT vs shadow agents par surface de risque

Dimension	Shadow IT classique	Shadow Agents IA
Mode d'action	Passif (consultation, stockage)	Actif (création, modification, envoi, exécution)
Vitesse d'impact	Humaine (limitée)	Machine (milliers d'actions/minute)
Détection comportementale	Possible (logs accès, VPN, print)	Difficile (l'agent n'a pas de comportement "humain")
Vecteur RGPD	Transfert/stockage non-autorisé	Traitement + transfert + action sur données personnelles
Risque de cascade	Limité (1 système)	Élevé (chaînage multi-systèmes)
Réversibilité	Souvent réversible	Parfois irréversible (emails envoyés, données modifiées)
Facilité de déploiement	Téléchargement d'une app	pip install + clé API (< 10 min)

Quels outils de détection pour les shadow agents ?

L'arsenal de détection disponible en 2026 couvre plusieurs dimensions. Du côté des **CASB** (Cloud Access Security Broker), Microsoft Defender for Cloud Apps et Netskope permettent de surveiller les connexions OAuth et les appels APIs vers des services IA externes. Ces outils peuvent bloquer ou alerter sur les nouvelles applications qui demandent des permissions sur les données Microsoft 365 ou Google Workspace.

Du côté des **outils de gouvernance IA spécifiques**, des solutions comme Securiti.ai AI Governance, Credo AI, ou Fairly.ai commencent à proposer des inventaires d'agents IA et des frameworks de contrôle. Ce marché est encore jeune — la plupart de ces outils en sont à leur version 1.0 ou 2.0 — mais la direction est claire. Pour les équipes qui veulent une approche plus artisanale, un audit OAuth trimestriel combiné à des logs proxy bien analysés reste la méthode la plus accessible. Notre article sur la [gouvernance des shadow agents IA](#) détaille ces approches.

La responsabilité juridique de l'entreprise face aux shadow agents

Voici une réalité que peu de juristes ont encore intégrée : **l'entreprise est responsable des actions de ses collaborateurs, y compris celles réalisées par des agents IA déployés hors cadre**. Si un agent shadow envoie un email en violation du droit de la concurrence, c'est l'entreprise qui répond. Si un agent exfiltre des données personnelles, c'est l'organisation qui notifie la CNIL. La responsabilité ne se transfère pas au collaborateur sous prétexte qu'il a agi "de sa propre initiative".

Cette asymétrie juridique crée une urgence réglementaire pour les RSSI et DPO. Sans politique écrite sur les agents IA, sans clause dans le règlement intérieur, sans formation des collaborateurs, l'organisation ne dispose d'aucun levier pour démontrer la diligence raisonnable en cas d'incident. La gouvernance des shadow agents n'est pas qu'une question technique — c'est une question de gestion du risque juridique. Notre guide [Gouvernance IA en entreprise](#) couvre la construction de ces politiques.

Tableau : Shadow IT vs Shadow Agents — comparaison de la surface de risque

Dimension	Shadow IT classique	Shadow Agents (IA agentique)	Niveau de risque
Nature de l'action	Consultation, stockage	Exécution autonome d'actions	+++ Shadow Agents
Données exposées	Données copiées	Données accédées + traitées + envoyées	+++ Shadow Agents
Visibilité SSI	DéTECTABLE par DLP/CASB	Actions via API légitimes, peu visibles	+++ Shadow Agents
Impact RGPD	Article 28 (sous-traitant)	Art. 22 (décision automatisée) + Art. 28	+++ Shadow Agents
Remédiation	Bloquer l'accès	Identifier toutes les actions passées	+++ Shadow Agents
Responsabilité DPO	Violation de politique	Possible violation légale (AI Act article 13)	+++ Shadow Agents

Détection réseau : monitorer les connexions sortantes vers les APIs d'OpenAI, Anthropic, Mistral, Cohere — un shadow agent communique nécessairement avec un LLM externe ou local

Audit des tokens API : inventorier tous les tokens API de LLM créés dans l'organisation — c'est la trace la plus fiable de l'existence de shadow agents

Review des intégrations Make/Zapier/n8n : ces plateformes d'automatisation sont le vecteur de déploiement le plus courant des shadow agents non surveillés

Formation et politique : publier une politique claire sur les agents IA autorisés, avec un processus de validation rapide (< 5 jours) pour éviter que le shadow soit le chemin de moindre résistance

Notre guide sur le [Shadow AI en entreprise](#) couvre les méthodes de détection que j'applique en mission. La [gouvernance IA en entreprise](#) doit explicitement adresser le risque shadow agent dans sa politique. Et si vous cherchez à mettre en place une détection opérationnelle, notre service [RSSI externalisé](#) intègre systématiquement un audit shadow IA dans ses missions de démarrage.

Quelle est la responsabilité juridique de l'entreprise face aux shadow agents ?

La question juridique est tranchée par deux textes européens. L'[AI Act européen \(Règlement UE 2024/1689\)](#) classe certains usages d'agents IA dans la catégorie "haut risque" dès qu'ils touchent des décisions affectant des personnes. L'article 13 impose une documentation explicite et un contrôle humain. Les shadow agents, par définition non documentés, violent mécaniquement ces exigences. La CNIL, dans ses recommandations de 2026, a précisé que l'entreprise est responsable en tant que responsable de traitement même si l'agent a été déployé à l'insu de la SSI.

Le RGPD ajoute une couche supplémentaire : l'article 22 interdit les décisions entièrement automatisées affectant significativement des personnes, sauf consentement explicite. Un shadow agent qui traite automatiquement des candidatures, des évaluations ou des accès constitue potentiellement une violation de l'article 22, et la DPO peut être mise en cause. La CNIL peut infliger jusqu'à 4 % du chiffre d'affaires mondial pour de telles infractions.

Sur le plan pénal, la jurisprudence commence à se former. En Allemagne, une décision de 2025 a engagé la responsabilité du RSSI d'une entreprise ayant ignoré des signaux d'alerte sur des agents IA non autorisés. En France, la doctrine de la CNIL évolue vers une obligation de vigilance

active — pas seulement une obligation de réaction. **Ne pas savoir** qu'un shadow agent existe n'est plus une défense acceptable si les outils de détection n'ont pas été déployés.

Obligation de documentation (AI Act art. 13) : tout système IA à haut risque doit être documenté — les shadow agents y échappent par définition, créant une violation automatique

Responsabilité solidaire : en cas d'incident causé par un shadow agent (fuite de données, décision discriminatoire), l'entreprise et potentiellement le manager ayant toléré l'usage non autorisé sont co-responsables

Devoir de surveillance RSSI : la jurisprudence européenne tend à reconnaître un devoir de surveillance actif, pas seulement réactif, pour les responsables sécurité

Dans une mission d'audit de conformité AI Act pour un acteur de l'assurance (2026), nous avons identifié 7 workflows automatisés de traitement des sinistres développés par des équipes métier, sans aucune validation juridique ni SSI. Certains prenaient des décisions d'indemnisation de manière entièrement automatique. Qualification : article 22 RGPD + article 13 AI Act. La remédiation a nécessité 4 mois de refonte complète, avec suspension temporaire des workflows concernés.

— *Retour de mission, conformité AI Act, T2 2026*

Pour les entreprises soumises à NIS 2 ou à la réglementation financière (DORA), le risque shadow agent est doublement encadré. Notre service de [détection du shadow hacking par outils IA non autorisés](#) couvre ces dimensions réglementaires, et notre guide sur la [gouvernance de l'IA agentique](#) détaille les exigences AI Act article par article.

Questions fréquentes sur les shadow agents et le risque insider

Quelle est la différence entre un shadow agent et un usage non-autorisé de ChatGPT ?

Un usage non-autorisé de ChatGPT implique un humain qui envoie des données et reçoit une réponse — c'est passif et ponctuel. Un **shadow agent** est un programme autonome qui s'exécute en continu, connecté à plusieurs systèmes via des APIs, capable de prendre des décisions et d'agir

sans intervention humaine à chaque étape. L'agent peut fonctionner la nuit, le weekend, après le départ du collaborateur qui l'a créé. L'échelle et l'autonomie sont radicalement différentes.

Comment savoir si notre organisation a déjà des shadow agents actifs ?

Trois vérifications immédiates : auditez les applications tierces autorisées dans Azure AD ou Google Workspace (cherchez les apps avec permissions Mail.ReadWrite, Files.ReadWrite.All, Calendars.ReadWrite créées ces 12 derniers mois) ; analysez les logs proxy pour des appels répétés vers api.openai.com, api.anthropic.com ou api.mistral.ai depuis des postes hors politique ; et interrogez vos équipes IT sur les installations Python récentes avec des packages langchain, autogpt ou crewai. Dans la grande majorité des organisations, cette vérification révèle au moins 2 ou 3 agents non-répertoriés.

Les agents IA intégrés dans Microsoft 365 Copilot sont-ils des shadow agents ?

Non — **Microsoft 365 Copilot** et les agents Copilot Studio déployés via le canal IT officiel sont des agents "autorisés". Ils sont couverts par le DPA Microsoft, les permissions sont gérées via Entra ID, et les logs d'activité sont disponibles dans le portail admin. Le risque shadow agent naît quand des collaborateurs créent leurs propres agents Copilot Studio sans revue SSI, ou utilisent des plateformes tierces non-approuvées. La frontière est la validation IT, pas la technologie.

Quelles sanctions encourt une entreprise en cas d'incident shadow agent sous le RGPD ?

Les sanctions RGPD applicables à un incident causé par un shadow agent sont les mêmes que pour tout incident de sécurité impliquant des données personnelles : notification CNIL dans les 72 heures si le risque pour les personnes concernées est élevé, notification des personnes si le risque est très élevé, et potentiellement une amende jusqu'à 4% du chiffre d'affaires mondial annuel. Le fait que l'incident ait été causé par un agent IA non-autorisé plutôt qu'une attaque externe n'est pas une circonstance atténuante — c'est au contraire une circonstance aggravante en termes de défaut de mesures organisationnelles.

Par où commencer pour sécuriser les shadow agents dans une PME sans équipe SSI dédiée ?

Trois actions prioritaires, dans l'ordre : premièrement, auditez les consentements OAuth dans Azure AD ou Google Admin Console et révoquez les accès des applications non-reconnues. Deuxièmement, rédigez une politique d'usage des agents IA (une page suffit) et communiquez-la aux équipes avec un canal de validation IT pour les nouveaux projets. Troisièmement, bloquez au niveau proxy les appels directs vers les APIs LLM grand public non-canalises et orientez les usages légitimes vers une API centralisée avec logging. Si vous manquez de ressources internes, notre [service RSSI externalisé](#) peut accompagner cette mise en place.

Architecture de contrôle recommandée pour bloquer les shadow agents

La réponse architecturale au risque shadow agent s'organise en trois couches complémentaires. La première couche est la **gouvernance des identités** : chaque agent IA autorisé doit disposer d'une identité de service dédiée (service account ou managed identity), distincte de l'identité humaine du créateur. Cela permet de révoquer les accès de l'agent sans impacter le compte utilisateur, et d'auditer les actions de l'agent séparément des actions humaines.

La deuxième couche est la **politique de consentement OAuth**. Dans Azure AD, il est possible de désactiver le consentement utilisateur pour les applications tierces et d'imposer un processus d'approbation admin. En pratique, cette mesure bloque immédiatement les nouveaux shadow agents SaaS sans bloquer les usages légitimes approuvés. Google Workspace propose une configuration équivalente via les contrôles d'accès des API tierces.

La troisième couche est le **monitoring comportemental**. Un agent IA légitime a des patterns d'appel API prévisibles : même endpoint, même volume, horaires réguliers. Un agent shadow ou compromis présente souvent des anomalies : pics d'appels inhabituels, accès à de nouveaux endpoints, téléchargement de volumes anormaux. Un SIEM correctement configuré peut détecter ces patterns en quelques heures. Notre analyse du [SOC augmenté par IA](#) couvre ces cas d'usage de détection.

Le rôle du MCP (Model Context Protocol) dans la propagation des shadow agents

Le **Model Context Protocol (MCP)**, standardisé par Anthropic en 2024, permet aux agents IA de se connecter à des serveurs de ressources et d'outils de manière structurée. C'est une avancée réelle pour les agents légitimes — mais aussi un nouveau vecteur pour les shadow agents. Un collaborateur peut déployer un serveur MCP local qui expose des ressources d'entreprise (fichiers, bases de données, APIs internes) à n'importe quel client MCP, sans passer par la DSI.

Le risque MCP est encore peu documenté dans les référentiels officiels, mais il est réel et croissant. Notre analyse des [attaques par injection MCP](#) montre comment un serveur MCP mal configuré peut devenir un point de pivot pour des actions non-autorisées. Pour les équipes SSI, l'inventaire des serveurs MCP actifs dans l'organisation doit devenir une priorité au même titre que l'inventaire des applications OAuth.

Former les collaborateurs : la défense la plus efficace contre les shadow agents

La formation reste la mesure de prévention la plus rentable. Un collaborateur qui comprend pourquoi il doit déclarer son projet d'agent IA au service IT — et qui dispose d'un canal simple pour le faire — est bien moins susceptible de déployer un shadow agent. La formation ne doit pas être culpabilisante : l'intention derrière la création d'un agent est presque toujours bénigne. L'objectif est d'outiller les collaborateurs pour qu'ils agissent correctement, pas de les dissuader d'innover.

Les éléments clés d'une formation efficace sur les shadow agents : expliquer concrètement ce qu'un agent peut faire et pourquoi ça pose problème ; montrer le processus de validation IT (qui contacter, délai, critères d'approbation) ; donner accès à un catalogue d'agents approuvés que les collaborateurs peuvent utiliser sans risque. Cette approche "carotte et bâton" — règles claires + alternatives légitimes accessibles — est celle qui fonctionne le mieux selon les retours des programmes de sensibilisation IA en 2025-2026.

Shadow agents et chaîne d'approvisionnement logicielle

Un risque moins évident mais potentiellement très sérieux : les shadow agents peuvent introduire des vulnérabilités dans la **chaîne d'approvisionnement logicielle**. Un développeur qui utilise un agent IA pour générer du code et l'intégrer directement dans un repo sans revue humaine introduit un vecteur de supply chain attack. Si l'agent a été compromis par une prompt injection (via un document malveillant dans son contexte), le code généré peut contenir des backdoors ou des vulnérabilités délibérées.

Ce scénario n'est pas hypothétique. Les recherches publiées en 2025 sur les attaques de type [ML Supply Chain](#) montrent que les modèles IA utilisés dans les pipelines de développement peuvent être compromis à plusieurs niveaux. Un agent développeur shadow qui utilise un modèle tiers non-vérifié, connecté au repo Git de l'entreprise, représente un risque de supply chain que les outils de sécurité classiques (SAST, DAST) ne détecteront pas forcément.

Les shadow agents sont probablement déjà actifs dans votre organisation. Un audit ciblé de 2 jours peut cartographier l'existant et produire un plan d'action immédiat. [Notre service RSSI externalisé](#) propose ce type d'intervention sans engagement sur la durée.