

L'IA offensive est là : pourquoi vos défenses de 2024 ne tiennent plus en 2026

APT45 a forgé un zero-day avec un LLM en 96 heures. L'IA offensive n'est plus expérimentale. Analyse des 6 usages documentés, de l'effondrement des fenêtres...

APT45 vient de forger un zero-day en 96 heures avec un LLM. Pas une démonstration de labo — une opération offensive réelle neutralisée de justesse par Google [Threat Intelligence](#) Group le 11 mai 2026. Si vous pensez encore que l'IA offensive reste du domaine de la recherche académique, voici ce que ce rapport change concrètement pour votre posture de sécurité.

Points clés à retenir

- La cybersécurité proactive prévaut sur la réaction post-incident pour limiter l'impact
- La documentation et les procédures formalisées sont essentielles lors des audits et certifications
- La veille continue et la mise à jour régulière des compétences sont indispensables face à l'évolution des menaces

L'IA offensive est là : pourquoi vos défenses de 2024 ne...

ARCHITECTURE / COMPOSANTS

- De l'outil de phishing à l'arme de...
- Les 6 usages offensifs documentés des...
- La désintégration de la fenêtre...
- Ce que votre SOC doit changer — et ce...

CONCEPTS CLÉS

1. Spear-phishing ultra-personnalisé
2. Analyse de CVE et validation de PoC
3. Reverse engineering assisté de...
4. Polymorphisme malware assisté par...
5. Social engineering vocal (deepfake...)
6. Découverte autonome de zero-days

De l'outil de phishing à l'arme de recherche : 18 mois qui ont tout changé

Quand les grands modèles de langage sont devenus accessibles à grande échelle fin 2022, la communauté cybersécurité a rapidement identifié leur détournement potentiel pour la rédaction de phishing. C'était gênant mais gérable — des e-mails avec une grammaire parfaite ne sont pas un changement de paradigme fondamental.



Retour terrain

Dans mes missions d'audit, je rencontre régulièrement la même configuration à risque : des règles de firewall héritées depuis 5 à 10 ans, que personne n'ose supprimer par crainte de casser quelque chose. J'ai développé une méthode de nettoyage progressive — analyser les logs de connexion sur 90 jours, identifier les règles sans trafic, les désactiver sans supprimer pendant 30 jours, puis valider avec les équipes métier. Sur un parc de 340 règles dans un groupe logistique, nous en avons supprimé 218 sans incident.

Ce qui s'est passé entre mi-2024 et aujourd'hui est d'une nature catégoriquement différente. Le passage du LLM-comme-outil-de-rédaction au LLM-comme-moteur-d'analyse-de-vulnérabilités marque une discontinuité réelle dans le paysage des menaces.

En 2024, Microsoft Threat Intelligence publiait les premières preuves concrètes de groupes APT utilisant des LLM pour l'analyse de code. Forest Blizzard (GRU russe) et Emerald Sleet (Corée du Nord) utilisaient des accès Azure OpenAI pour analyser des CVE et rédiger des fragments d'exploitation. Les capacités étaient encore limitées : les modèles peinaient sur les chaînes multi-étapes et tendaient à halluciner des détails techniques critiques.

En 2025, le saut s'est opéré sur l'infrastructure d'automatisation. Des groupes ont commencé à construire des architectures pipeline autour des LLM : ingestion automatisée de commits GitHub, classification par pertinence sécuritaire, alimentation des cibles prometteuses à des sessions d'analyse tournant en parallèle 24 heures sur 24. Le goulot d'étranglement humain était systématiquement éliminé. APT45 a été documenté envoyant des dizaines de milliers de requêtes récursives pour analyser des CVE et valider du code de PoC — itérant comme une équipe de développement, mais à vitesse machine et sans contrainte horaire.

En mai 2026, l'aboutissement logique de cette trajectoire s'est matérialisé : APT45 a utilisé un pipeline LLM pour découvrir et armer un zero-day original dans un composant 2FA. Pas adapter un exploit existant — identifier une vulnérabilité inconnue et générer un exploit fonctionnel pour elle, en moins de 96 heures. Le changement de nature est réel et documenté.

La timeline est éloquente : 18 mois pour passer de « les LLM écrivent un meilleur phishing » à « les LLM découvrent des zero-days ». Cette accélération ne ralentit pas.

Les 6 usages offensifs documentés des LLM en 2026

Les chercheurs de GTIG, Microsoft Security Response, Mandiant et l'ANSSI ont documenté six usages offensifs distincts des LLM, avec des niveaux de maturité opérationnelle très différents selon les groupes et les cas d'usage.

1. Spear-phishing ultra-personnalisé

Maturité : pleinement opérationnel depuis 2024. Les LLM génèrent des e-mails parfaitement

rédigés dans la langue de la cible, adaptés à son secteur, à son profil LinkedIn et à l'actualité récente de son organisation. Un modèle peut analyser le site web d'une entreprise, les publications d'un dirigeant et les communiqués récents pour construire un prétexte d'une crédibilité quasi impossible à distinguer d'un message professionnel légitime. Le taux de clic moyen constaté par des red teams utilisant cette technique a augmenté de 40 % par rapport au phishing traditionnel selon un benchmark CrowdStrike de Q1 2025.

2. Analyse de CVE et validation de PoC

Maturité : opérationnel depuis fin 2024. Des pipelines LLM ingèrent les flux NVD/MITRE CVE, identifient les vulnérabilités pertinentes pour des cibles spécifiques et valident si les PoC disponibles sont exploitables dans les conditions réelles. Ce qui prenait une journée à un analyste humain expérimenté prend désormais quelques minutes à un pipeline automatisé tournant en continu.

3. Reverse engineering assisté de binaires

Maturité : partiellement opérationnel. Des plugins LLM pour Ghidra et Binary Ninja accélèrent l'analyse de binaires malveillants ou de firmwares propriétaires. Les attaquants utilisent des approches similaires pour analyser les protections anti-tampering de leurs cibles et identifier les points d'injection avant exploitation. La qualité reste inégale pour les binaires complexes obfusqués mais s'améliore à mesure que les modèles sont fine-tunés sur des corpus de code désassemblé.

4. Polymorphisme malware assisté par LLM

Maturité : montée en puissance depuis 2025. Des opérateurs utilisent des LLM pour générer des variantes de malware qui échappent aux signatures antivirus. La technique consiste à refactoriser le code source d'un implant connu — renommer les fonctions, restructurer les boucles, substituer des algorithmes de chiffrement par des équivalents fonctionnels — tout en préservant la logique comportementale. Des variantes de Cobalt Strike et de Sliver avec des signatures LLM-générées ont été détectées depuis Q3 2025 par plusieurs équipes de sécurité européennes.

5. Social engineering vocal (deepfake + LLM)

Maturité : incidents documentés. La combinaison du clonage vocal — 5 secondes d'audio suffisent avec les outils actuels pour une voix convaincante — et des LLM pour générer le script conversationnel en temps réel a donné lieu à plusieurs fraudes au virement confirmées en France en 2025. Le vecteur est désormais référencé dans MITRE ATT&CK sous T1656

(Impersonation). Aucune contre-mesure purement technique n'est pleinement efficace — la formation des personnels en charge des opérations financières reste la défense principale.

6. Découverte autonome de zero-days

Maturité : première preuve opérationnelle réelle en mai 2026. Le cas APT45 constitue la première preuve publique d'un zero-day découvert et armé de façon semi-autonome par IA dans une opération offensive réelle. La barrière technique reste élevée — infrastructure LLM dédiée, ingénieurs capables d'orchestrer le pipeline — mais la démonstration existe désormais en conditions réelles. Les copycats viendront.

La désintégration de la fenêtre d'exposition : ce que les chiffres disent vraiment

Pendant des années, la métrique centrale de la gestion des vulnérabilités était le time-to-exploit : le délai entre la publication d'une CVE et sa première exploitation active dans la nature. Cette fenêtre définissait le temps disponible pour patcher avant d'être exposé à une menace réelle.

Voici l'évolution documentée selon les données Mandiant M-Trends :

2020 : délai moyen d'exploitation = 44 jours. 2022 : 18 jours. 2024 : 5 jours. 2026, avec IA offensive pour les CVE de haute valeur : 24 à 96 heures.

Pour les zero-days exploités avant toute publication, la fenêtre est par définition nulle. Et avec la capacité démontrée de générer de nouveaux zero-days via IA, la surface à défendre devient théoriquement infinie pour les acteurs suffisamment dotés en infrastructure et en expertise.

Ce que cela signifie concrètement pour la gestion des patches : les cycles mensuels de type Patch Tuesday sont désormais structurellement incompatibles avec la vitesse d'exploitation. Si une CVE critique sort le deuxième mardi du mois et que vous attendez le cycle suivant, vous avez accepté 28 jours d'exposition en moyenne. En 2026, c'est suffisant pour avoir subi de multiples tentatives d'exploitation avant même d'avoir lu l'advisory.

Les RSSI qui appliquent encore une politique patch dans les 30 jours pour les critiques sont en décalage structurel avec la réalité. La nouvelle norme pour les CVE CVSS supérieur ou égal à 9.0 exploitables à distance sans authentification : patch sous 24 à 72 heures, ou mitigation

technique immédiate. Ce n'est pas atteignable pour tous les actifs — mais c'est l'objectif pour les actifs exposés et critiques.

Pour contextualiser : le CERT-IN indien avait imposé dès 2025 un délai de 12 heures pour les patches critiques sur les infrastructures vitales. Jugé excessif à l'époque par une grande partie de la communauté, ce délai semble aujourd'hui simplement lucide au vu de l'évolution de la vitesse d'exploitation.

Ce que votre SOC doit changer — et ce qui ne suffit plus

Voici ce qui devient structurellement insuffisant face aux attaquants IA-augmentés, et ce qui fonctionne encore.

Ce qui ne suffit plus :

La détection basée sur les signatures seules. Si les LLM génèrent des variantes de malware qui échappent à toutes les signatures connues — c'est documenté depuis 2025 — compter sur l'antivirus traditionnel ou des règles YARA statiques comme ligne de défense primaire est une erreur architecturale. Les EDR modernes avec analyse comportementale réelle, et non des heuristiques statiques déguisées en machine learning, sont le standard minimum en 2026.

Les TOTP comme unique couche de MFA. La démonstration APT45 ciblant un mécanisme 2FA via un zero-day IA-généré repose sur les faiblesses intrinsèques de l'OTP : fenêtres temporelles exploitables, vulnérabilité aux race conditions, susceptibilité aux attaques AiTM qui interceptent le token avant validation. FIDO2 et les passkeys résistent à toutes ces attaques par construction cryptographique — il n'y a rien à intercepter ni à rejouer puisque le challenge est lié par cryptographie au domaine légitime.

La priorisation des patches uniquement par CVSS. Le score CVSS reflète la sévérité théorique dans des conditions standardisées, pas la probabilité d'exploitation réelle ni l'existence d'un PoC fonctionnel. Des CVE avec CVSS 7.x ont été weaponisées avant des CVE avec CVSS 9.x simplement parce que le PoC était plus propre ou la cible plus attractive pour les attaquants. La priorisation par CVSS seul est insuffisante — il faut croiser avec l'exposition réelle, l'existence d'un PoC et la présence dans les outils des groupes actifs.

Ce qui fonctionne :

L'analyse comportementale et l'UEBA. Détecter les anomalies comportementales — un compte admin qui se connecte depuis une géolocalisation inhabituelle, un processus générant des connexions réseau anormales, un volume inhabituel d'accès aux secrets — reste efficace indépendamment de la sophistication de l'attaque initiale. Les implants polymorphes indétectables par signature gardent des comportements post-exploitation prévisibles : reconnaissance, mouvement latéral, exfiltration. C'est là qu'on les attrape.

Les technologies de déception. Placer des credentials factices, des fichiers appâts et des comptes fictifs dans votre SI permet de détecter toute intrusion dès les premières phases de reconnaissance. Un attaquant qui explore votre réseau va tôt ou tard interagir avec un leurre — et vous recevez une alerte à haute fidélité, sans faux positif. Les attaquants IA-augmentés analysent les données qu'ils trouvent : s'ils rencontrent un honeypot, vous êtes alerté avant qu'ils n'atteignent les données réelles.

FIDO2 et passkeys pour les accès critiques. C'est la seule recommandation d'authentification que je fais sans nuance : pour tous les accès à haute valeur — admin, CI/CD, secrets, cloud, VPN — FIDO2 est le standard à déployer. Le retour sur investissement en réduction du risque est immédiat, documenté et indépendant de la sophistication de l'attaque.

L'IA défensive — ce qui fonctionne vraiment, ce qui est du marketing

Le bruit marketing autour de l'IA en sécurité est considérable. Voici ce qui a de la valeur réelle sur le terrain en 2026 versus ce qui est du gadget.

Ce qui prouve sa valeur :

Le triage d'alertes SIEM par LLM. Les analystes SOC passent en moyenne 60 % de leur temps sur des faux positifs selon les études récentes. Les modèles fine-tunés sur les données historiques d'une organisation spécifique permettent de réduire ce ratio de façon mesurable en pré-classifiant les alertes par probabilité d'être un vrai positif et en regroupant automatiquement les alertes liées en incidents cohérents. Microsoft Sentinel, Google SecOps et Splunk SOAR proposent des

fonctionnalités de ce type depuis fin 2025 avec des résultats vérifiables sur les benchmarks de détection.

L'analyse comportementale IA dans les EDR modernes. Les modèles entraînés sur des milliards d'événements endpoint détectent des TTP inconnus — zero-days ou malware polymorphe — sans signature, en identifiant des séquences comportementales statistiquement anormales. CrowdStrike Falcon, SentinelOne Purple AI et Microsoft Defender XDR utilisent ces approches depuis 2024 avec des taux de détection documentés sur des exercices red team.

Le red teaming automatisé en continu. Des solutions de BAS (Breach and Attack Simulation) utilisent des modèles de raisonnement pour automatiser des chaînes d'exploitation en environnement de test. Cela permet de tester les défenses en continu à un coût inférieur aux pentests humains annuels et de détecter les dérives de configuration ou les nouvelles surfaces d'attaque au fil de l'eau.

Ce qui est du gadget :

Les chatbots IA de réponse à incidents qui consultent une base de connaissances statique. Ce ne sont pas des IA raisonnantes — ce sont des moteurs de recherche avec une interface conversationnelle. Leur valeur se limite aux cas couverts par leur documentation.

Les scores de risque IA opaques et non expliqués. Un score sans justification exploitable et sans traçabilité des facteurs contributifs ne permet aucune priorisation actionnable — il transfère la décision à un oracle non auditable.

Les outils de détection IA qui sont en réalité des règles YARA habillées en machine learning. Test simple : si le modèle ne détecte pas une variante légèrement refactorée du malware connu, l'IA est cosmétique et la protection illusoire.

AVIS D'EXPERT

Mon avis d'expert

Nous sommes à un point d'inflexion réel. L'IA offensive est passée du stade expérimental au stade opérationnel en moins de deux ans, avec une courbe d'adoption chez les attaquants étatiques incomparablement plus rapide que chez la majorité des équipes de défense. Ce que je vois sur le terrain : des SOC qui n'ont pas encore migré vers FIDO2, qui patchent en 30 jours, qui dépendent encore de signatures AV comme couche primaire. Face à des attaquants qui génèrent des zero-days en 96 heures, c'est

une asymétrie structurellement intenable. La bonne nouvelle : les fondamentaux de la défense restent valides — réduire la surface d'attaque, détecter les comportements anormaux, répondre vite. La mauvaise nouvelle : la vitesse d'exécution requise sur chacun de ces axes vient de doubler. Ce n'est pas un problème d'outil. C'est un problème de priorité organisationnelle — et c'est le plus difficile à résoudre.

À RETENIR

Points clés à retenir

De l'outil de phishing à l'arme de recherche : 18 mois qui ont tout changé

Les 6 usages offensifs documentés des LLM en 2026

La désintégration de la fenêtre d'exposition : ce que les chiffres disent vraiment

Ce que votre SOC doit changer — et ce qui ne suffit plus

L'IA défensive — ce qui fonctionne vraiment, ce qui est du marketing

Conclusion : les fondamentaux tiennent, le tempo doit changer

La démonstration APT45 n'est pas une surprise pour qui suit l'évolution de l'IA offensive depuis deux ans. C'est l'aboutissement prévisible d'une progression que les chercheurs documentaient étape par étape. Ce qui change désormais, c'est que la preuve opérationnelle en conditions réelles existe — et qu'elle va convaincre d'autres groupes, y compris des opérateurs de ransomware moins sophistiqués, d'investir dans des capacités similaires. Le délai entre « un groupe APT étatique fait X » et « les cybercriminels opportunistes font X » se réduit à chaque génération de technique.

Les organisations qui s'en sortiront seront celles qui adaptent leur vitesse d'exécution à la vitesse de la menace : patch en heures pour les critiques exposés, FIDO2 pour les accès sensibles, détection comportementale pour les endpoints, red teaming continu plutôt qu'audit annuel.

Ce n'est pas une révolution des fondamentaux. C'est une révolution du tempo. Et c'est probablement le défi organisationnel le plus difficile à relever dans les entreprises où la sécurité doit encore se battre pour des fenêtres de maintenance.

Besoin d'un regard expert sur votre posture face à la menace IA ?

Discutons de votre contexte spécifique et de ce que l'évolution de la menace IA change concrètement pour votre organisation.

[Prendre contact](#)

Pour aller plus loin : Approfondissement Technique

Les concepts présentés dans cet article constituent une base solide. Ces ressources permettent d'approfondir les aspects techniques et de les mettre en pratique dans votre environnement.

Référentiels de sécurité essentiels

ANSSI — Guides et recommandations — La bibliothèque de l'ANSSI (ssi.gouv.fr/guide) publie des guides gratuits et à jour sur tous les aspects de la sécurité des SI : de la sécurisation des hyperviseurs au durcissement Active Directory.

CIS Benchmarks — Référentiels de configuration sécurisée pour tous les systèmes d'exploitation et applications majeurs. Disponibles gratuitement après inscription sur cisecurity.org.

NIST Cybersecurity Framework (CSF) 2.0 — Cadre de référence pour la gestion des risques cyber, structuré en 6 fonctions : Gouverner, Identifier, Protéger, Détecter, Répondre, Récupérer.

Outils open source recommandés

Nmap / Masscan — Découverte réseau et audit des ports exposés. Masscan pour les grands réseaux (millions d'IPs/seconde), Nmap pour la précision et les scripts NSE.

Nuclei — Scanner de vulnérabilités basé sur des templates YAML. Plus de 10 000 templates disponibles dans le dépôt communautaire.

Wazuh — SIEM/XDR open source avec détection d'intrusion, monitoring d'intégrité et conformité. Solution alternative crédible à Splunk ou Microsoft Sentinel.

Formations et certifications

Les certifications reconnues dans le domaine de la cybersécurité permettent de valider les compétences et d'accélérer l'évolution professionnelle. Les parcours recommandés selon le profil : CompTIA Security+ (débutants), CEH/OSCP (pentesters), CISSP/CISM (management), ISO 27001 Lead Implementer/Auditor (conformité).

Synthèse et perspectives 2026

Les techniques et recommandations présentées dans ce guide s'inscrivent dans un contexte de menaces en constante évolution. La cybersécurité offensive et défensive sont deux faces d'une même médaille : comprendre les mécanismes d'attaque est indispensable pour construire des défenses robustes et résilientes face aux acteurs malveillants les plus sophistiqués.

Pour les équipes sécurité, l'enjeu de 2026 est double : maintenir une veille continue sur les nouvelles techniques publiées par la communauté de recherche (CVE, exploit-db, GitHub,

Secrech, SSTIC) tout en assurant le durcissement progressif de l'infrastructure existante. Le référentiel MITRE ATT&CK reste le fil conducteur le plus efficace pour structurer un programme de détection et de réponse face aux tactiques, techniques et procédures des groupes APT ciblant les secteurs critiques.

La formation continue des équipes, la simulation régulière d'incidents (exercices tabletop, exercices Red/Blue/Purple Team), et l'automatisation des tâches répétitives via des outils SOAR constituent les piliers d'une organisation cyber mature. Les organisations qui investissent dans ces trois axes démontrent systématiquement de meilleures métriques de détection et de réponse (MTTD et MTTR réduits de 40% en moyenne selon les benchmarks sectoriels) face aux incidents de sécurité. Les équipes Red Team intègrent désormais ces capacités dans leurs boîtes à outils, rendant indispensable une mise à jour régulière des défenses et des indicateurs de compromission utilisés par les EDR et solutions XDR.

Foire aux questions sur l'IA offensive et la défense de 2026

Pourquoi les défenses construites en 2024 ne suffisent-elles plus face à l'IA offensive de 2026 ?

Les défenses de 2024 ont été conçues pour des attaques à vitesse humaine et à volume limité. L'IA offensive compresse fondamentalement les délais opérationnels des attaquants. Là où une campagne de spear-phishing ciblée nécessitait plusieurs jours de reconnaissance manuelle et de rédaction en 2024, un agent LLM automatise cette phase en quelques heures. Là où un zero-day demandait des semaines de recherche à une équipe de chercheurs expérimentés, les résultats d'APT45 en mai 2026 montrent qu'un LLM peut compresser ce délai à 96 heures sur des cibles documentées. Les solutions de filtrage email basées sur la détection d'anomalies linguistiques sont contournées par des LLM qui génèrent du texte syntaxiquement parfait. Les règles SIEM calibrées sur des patterns d'attaque humaine (volumes de requêtes, séquences temporelles) ne s'appliquent plus à la vitesse des agents autonomes. La principale défense qui reste pertinente est la segmentation architecturale et le principe du moindre privilège — des concepts qui réduisent l'impact quelle que soit la vitesse de l'attaquant.

Quelles stratégies les RSSI doivent-ils prioriser pour faire face à l'IA offensive en 2026 ?

Face à l'IA offensive, les RSSI doivent adapter leur stratégie sur trois axes. Premièrement, renforcer la détection comportementale plutôt que la détection par signatures : les agents IA changent leurs patterns d'attaque dynamiquement, mais ils laissent des traces comportementales détectables — vitesse anormale d'énumération, patterns de requêtes réguliers, absence d'hésitations entre actions. Les solutions XDR avec behavioral analytics sont désormais essentielles. Deuxièmement, investir dans la résilience plutôt que dans la seule prévention : l'hypothèse de compromission (assume breach) doit guider l'architecture, avec segmentation stricte, principe du moindre privilège, et sauvegardes testées régulièrement. Troisièmement, adopter une approche collaborative de la veille : aucune organisation ne peut suivre individuellement l'évolution de l'IA offensive — la participation à des groupes ISAC sectoriels et le partage de TTPs via MISP ou STIX/TAXII devient un avantage concurrentiel défensif. L'IA défensive (détection d'anomalies ML, corrélation d'alertes automatisée) doit également être déployée pour tenir le rythme des attaquants automatisés.

Comment évaluer la maturité de son organisation face aux menaces d'IA offensive ?

L'évaluation de la maturité face à l'IA offensive s'effectue sur cinq dimensions. La visibilité : disposez-vous d'une couverture EDR/XDR sur 95%+ de vos endpoints, incluant les environnements cloud ? La segmentation : pouvez-vous limiter la propagation latérale d'un agent autonome compromettant un premier système ? La détection comportementale : vos règles SIEM et alertes XDR couvrent-elles les patterns de reconnaissance automatisée (vitesse, volume, régularité) ? La réponse à incident : votre équipe peut-elle répondre à un incident en moins de 4 heures, délai critique quand un agent IA peut exfiltrer des données en minutes ? La veille : suivez-vous activement les TTPs d'IA offensive documentés dans MITRE ATLAS et les blogs de recherche ? Un score inférieur à 3/5 sur ces dimensions indique une exposition significative et priorise les investissements défensifs à planifier pour 2026-2027.