

Attaques AiTM 2026 : Evilginx, Modlishka et MFA Bypass

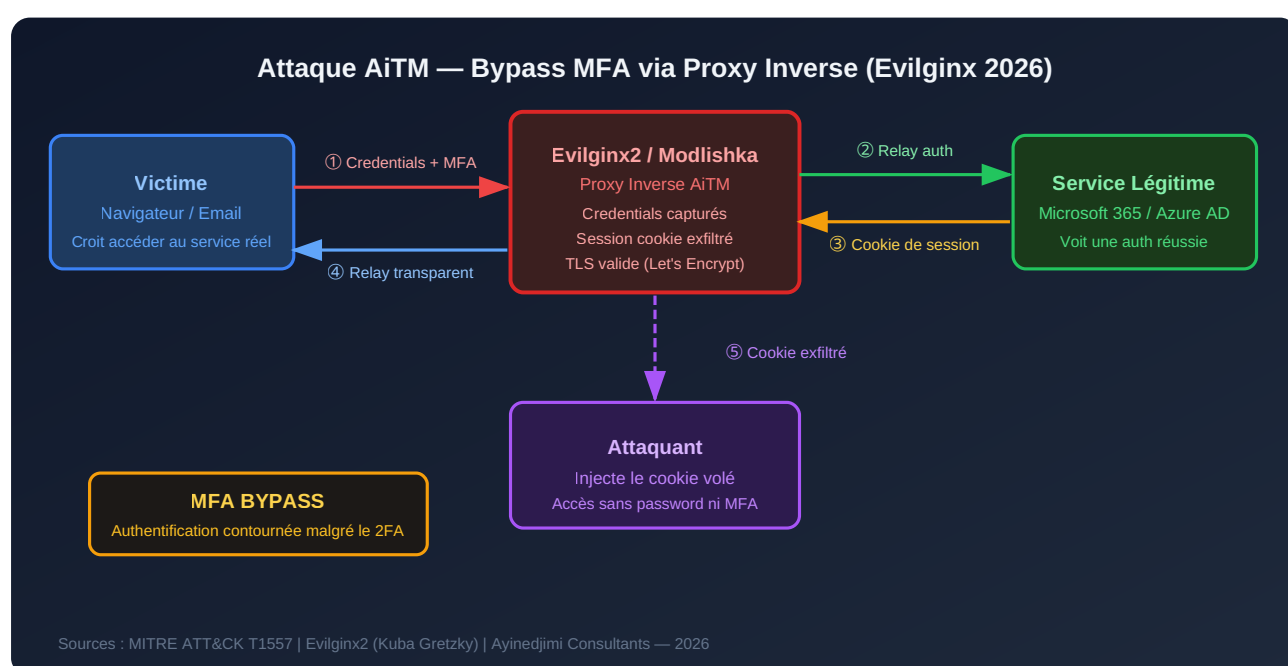
Maîtrisez les attaques AiTM avec Evilginx MFA bypass en 2026 : techniques, IoC et défenses pour sécuriser vos accès Microsoft 365 et Azure AD.

Résumé exécutif

Les attaques **AiTM (Adversary-in-the-Middle)** représentent en 2026 la menace la plus sophistiquée contre les architectures d'authentification multi-facteurs. En positionnant un proxy inverse entre la victime et le service cible, des outils comme **Evilginx2** et **Modlishka** capturent en temps réel les cookies de session post-authentification, rendant obsolètes les méthodes MFA classiques (SMS, TOTP, push). Ce guide technique analyse les mécanismes d'attaque, les configurations d'infrastructure malveillante, les indicateurs de compromission observés en 2026 et les stratégies de défense adaptées aux environnements Microsoft 365, Azure AD et Active Directory on-premises.

Les **attaques AiTM avec Evilginx MFA bypass** constituent en 2026 le vecteur de compromission initial le plus redouté des équipes de sécurité offensive et défensive à l'échelle mondiale. Contrairement aux attaques de [phishing](#) classiques qui cherchent à dérober des identifiants statiques, les attaques AiTM interceptent dynamiquement les cookies de session après une authentification MFA réussie, contournant ainsi intégralement les mécanismes de protection par SMS, TOTP ou notification push. Les outils open source comme Evilginx2, développé par Kuba Gretzky, et Modlishka, ont considérablement abaissé la barrière technique pour réaliser ce type d'attaque, permettant à des acteurs malveillants de toutes envergures de

cibler des environnements Azure Active Directory, Microsoft 365, GitHub ou des applications d'entreprise basées sur [SAML](#). Face à cette menace évolutive répertoriée sous MITRE ATT&CK T1557, les équipes [blue team](#) doivent désormais déployer des contrôles de détection avancés reposant sur l'analyse comportementale post-authentification, le [Conditional Access](#) avec Token Protection, et les technologies FIDO2 ou passkeys — seules véritablement résistantes aux attaques AiTM en 2026. Cet article vous fournit une analyse technique exhaustive des mécanismes d'attaque, des configurations d'infrastructure, des indicateurs de compromission documentés cette année, et des stratégies de mitigation adaptées aux organisations cherchant à se protéger durablement contre cette menace.



Architecture d'une attaque AiTM : Evilginx2 agit comme proxy inverse transparent, capturant credentials et cookies de session post-MFA (2026)

Qu'est-ce qu'une attaque AiTM (Adversary-in-the-Middle) en 2026 ?

Une attaque *AiTM* (*Adversary-in-the-Middle*) est une technique d'interception active dans laquelle un adversaire se positionne entre l'utilisateur et le service d'authentification légitime, jouant le rôle de proxy transparent en temps réel. MITRE ATT&CK répertorie cette technique sous la référence [T1557 — Adversary-in-the-Middle](#), distinguant les variantes réseau classiques

des attaques modernes par proxy HTTP inverse, qui dominent en 2026 par leur efficacité contre les défenses MFA traditionnelles.



Retour terrain

Dans mes missions de red team, la phase de post-exploitation révèle systématiquement des données que le client pensait protégées. Le cas le plus fréquent : des fichiers Excel de comptabilité ou RH stockés sur des partages réseau accessibles à tous les utilisateurs du domaine, sans restriction. Sur les 30 dernières missions, 27 avaient des partages réseau avec des données sensibles accessibles à n'importe quel utilisateur authentifié. La sensibilité des données n'est pas corrélée à leur niveau de protection réel.

Contrairement aux attaques Man-in-the-Middle traditionnelles basées sur l'empoisonnement ARP ou le détournement BGP, les attaques AiTM modernes exploitent des serveurs de **reverse proxy** qui imitent parfaitement l'interface des services cibles en proxifiant leur contenu en temps réel. Du point de vue de la victime, elle interagit avec une interface identique au service légitime, hébergée sur un domaine HTTPS avec un certificat valide. Du point de vue du service cible, l'authentification s'est déroulée normalement. C'est précisément cette transparence qui rend les attaques AiTM si redoutables en 2026.

La distinction fondamentale avec le phishing classique réside dans la gestion du MFA : là où un faux formulaire de connexion classique capture des credentials statiques mais échoue face à un second facteur, le proxy AiTM **relaie le flux d'authentification complet en temps réel**, y compris le challenge et la validation MFA. Le proxy récupère ensuite le cookie de session émis par le service légitime, indépendamment du mécanisme de second facteur utilisé par la victime.

Comment fonctionnent les attaques AiTM pour bypasser la MFA en 2026 ?

Le mécanisme de bypass MFA exploite une réalité fondamentale des protocoles web : une fois authentifié, un utilisateur est représenté par un cookie de session, pas par son identité continue. Ce cookie est émis par le serveur après validation de tous les facteurs d'authentification, y compris le MFA. Si un proxy capture ce cookie avant de le transmettre à la victime, l'attaquant dispose d'un accès authentifié complet sans jamais avoir eu les credentials ni le second facteur.

Voici le déroulé technique précis d'une attaque AiTM réussie contre un compte Microsoft 365 :

Reconnaissance et préparation : L'attaquant identifie la cible via OSINT ou des campagnes de phishing d'énumération, puis enregistre un domaine typosquatté avec certificat TLS Let's Encrypt auto-généré.

Déploiement du proxy AiTM : Evilginx2 ou Modlishka est configuré avec un phishlet décrivant le service cible (URLs, cookies à capturer, règles de réécriture d'hostname). Le proxy écoute sur le domaine malveillant en HTTPS.

Leurre de la victime : Un email de spearphishing dirige la victime vers le domaine de phishing. Le contenu de la page légitime est proxifié et affiché, avec tous les hostnames remplacés dynamiquement.

Capture des credentials : La victime saisit ses identifiants sur la fausse page. Le proxy les capture et les relaie immédiatement au service légitime, déclenchant la procédure MFA.

Relay du MFA : La victime reçoit et valide son code MFA (SMS, TOTP, push). Le proxy relaie cette validation au service légitime qui répond avec un cookie de session valide.

Exfiltration du cookie : Le proxy capture le cookie avant de le transmettre à la victime. L'attaquant dispose maintenant d'un token de session utilisable indépendamment.

Exploitation post-compromission : L'attaquant injecte le cookie dans son navigateur et accède au compte de la victime sans aucune friction d'authentification supplémentaire.

Avertissement légal et éthique : Les techniques décrites dans cet article sont présentées exclusivement à des fins éducatives, de sensibilisation et de défense en profondeur. La mise en œuvre de ces techniques sans autorisation écrite explicite du propriétaire du

système cible constitue une infraction pénale grave au titre de la loi Godfrain (article 323-1 et suivants du Code pénal) et du règlement européen sur la cybercriminalité. Tout test d'intrusion doit être réalisé dans le cadre d'un engagement contractuel signé.

Evilginx2 : architecture technique et AiTM Evilginx MFA bypass

Evilginx2 est le framework AiTM le plus répandu dans les arsenaux red team et chez les attaquants en 2026. Développé par **Kuba Gretzky** et disponible sur [GitHub \(evilginx2\)](#), il combine un serveur DNS autoritaire intégré, un moteur de reverse proxy HTTP/HTTPS et un gestionnaire de sessions dans un binaire Go autonome. Sa force principale réside dans son système de *phishlets* — des fichiers YAML décrivant la topologie des services cibles et les cookies à intercepter.

L'architecture d'Evilginx2 repose sur trois composants critiques. Le serveur DNS intégré résout les sous-domaines du domaine de phishing vers l'IP du proxy. Le moteur de proxy HTTP(S) intercepte, modifie à la volée et relaie le trafic en remplaçant dynamiquement tous les hostnames du service légitime par ceux du domaine malveillant, y compris dans le JavaScript et les en-têtes HTTP. Le gestionnaire de sessions stocke localement les credentials et cookies capturés, accessibles via l'interface CLI interactive en temps réel.

Environnement de laboratoire — Structure d'un phishlet Evilginx2

Illustration de la structure d'un phishlet YAML ciblant un service d'authentification entreprise (à des fins de compréhension technique uniquement) :

```
# Exemple structurel d'un phishlet (abstrait)
name: 'enterprise-sso'
proxy_hosts:
  - {phish_sub: 'login', orig_sub: 'login', domain: 'service-cible.com',
      session: true, is_landing: true}
auth_tokens:
  - domain: '.login.service-cible.com'
    keys: ['SESSION_TOKEN', 'AUTH_PERSISTENT', 'REFRESH_TOKEN']
credentials:
  username:
    key: 'username'
    search: '(.*)'
    type: 'post'
  password:
    key: 'password'
    search: '(.*)'
    type: 'post'
login:
  domain: login.service-cible.com
  path: '/oauth2/authorize'
```

Pour les environnements Microsoft Azure AD, les tokens critiques sont

ESTSAUTHPERSISTENT et **ESTSAUTH**. Une fois capturés, ils permettent d'établir une session complète sans re-authentification ni nouveau challenge MFA.

En 2026, les variantes d'Evilginx3 intègrent des capacités d'injection JavaScript avancées pour contourner certaines protections de **Token Binding** et détecter les environnements de sandbox analytiques. Les équipes red team documentent également l'émergence de variantes "headless" combinant un vrai navigateur automatisé (Playwright, Puppeteer) avec les techniques AiTM pour contourner les défenses basées sur les empreintes TLS et les challenges JavaScript anti-bot déployés par certains services.

Modlishka : l'alternative open source aux capacités étendues

Modlishka est un outil AiTM open source développé en Go, apprécié pour sa flexibilité et sa compatibilité avec les environnements utilisant des *Content Security Policy (CSP)* strictes. Là où

Evilginx2 repose sur des phishlets prédéfinis, Modlishka fonctionne comme un proxy inverse générique configurable via des règles de réécriture d'URL, d'injection de contenu et de manipulation des en-têtes HTTP, sans nécessiter de serveur DNS intégré.

Modlishka se distingue par sa gestion native des flux OAuth2/OIDC multi-étapes et des authentifications progressives (step-up authentication), ainsi que par sa capacité à injecter des scripts de tracking personnalisés dans les pages proxifiées pour capturer des métadonnées supplémentaires sur la victime. Sa modularité en fait un choix privilégié pour les campagnes ciblant des services SSO d'entreprise sur mesure ou des portails non couverts par les phishlets communautaires Evilginx2.

Critère	Evilginx2	Modlishka
Langage	Go	Go
Configuration cibles	Phishlets YAML	Flags CLI + config JSON
Serveur DNS intégré	Oui	Non (externe requis)
Injection JavaScript	Limitée (phishlets)	Avancée (règles flexibles)
Gestion certificats TLS	Let's Encrypt auto	Let's Encrypt auto
Évasion CSP stricte	Partielle	Avancée
Services documentés	+40 phishlets communauté	Générique (tout service)
Utilisation APT documentée 2026	Storm-0558, Scattered Spider	APT29, groupes FIN7/FIN8
Détection par solutions EDR/XDR	Moyenne	Faible à moyenne

Infrastructure de phishing AiTM : montage, obfuscation et résilience en 2026

La mise en place d'une infrastructure AiTM professionnelle en 2026 repose sur plusieurs couches d'obfuscation pour résister aux tentatives de takedown et contourner les systèmes de détection

automatisés. Les attaquants utilisent systématiquement des services cloud légitimes comme *redirecteurs initiaux* — Azure Static Web Apps, Cloudflare Workers, AWS Lambda@Edge — avant de proxifier vers le vrai serveur Evilginx2. Cette technique de multi-hop redirection exploite la réputation des domaines cloud pour passer les filtres de messagerie et les solutions DNS sécurisées des passerelles email.

Microsoft Security Operations a documenté dans son centre de ressources [Microsoft Security Operations](#) que plus de 70% des campagnes AiTM actives en 2025-2026 utilisent au moins un service cloud légitime comme premier hop. Le flux typique d'une campagne moderne est : email → lien Azure/AWS/Cloudflare légitime → redirection 302 vers domaine typosquatté → proxy Evilginx2 avec certificat TLS valide. Chaque étape ajoute une couche de légitimité perçue et complique considérablement l'investigation forensique post-incident.

Hébergement offshore : Serveurs dans des juridictions peu coopératives, enregistreurs de domaines permissifs, paiement en cryptomonnaie pour l'anonymat

Rotation d'infrastructure : Scripts Terraform automatisés pour changer d'IP et de serveur toutes les 24-72h, limitant l'impact des blocages réactifs

Certificats TLS automatiques : Let's Encrypt fournit des certificats HTTPS valides en minutes, supprimant l'avertissement navigateur qui alerterait la victime

Typosquatting Unicode : Utilisation de caractères homoglyphes Cyrillic ou IDN indétectables visuellement dans les URL de phishing

Géofencing anti-sandbox : Blocage automatique des requêtes provenant des plages IP de VirusTotal, Symantec, Palo Alto et des ranges datacenter connus

TTL court sur DNS : Enregistrements DNS avec TTL de 60 secondes pour faciliter la rotation rapide d'IP en cas de détection et blocage

À RETENIR

À retenir — Indicateurs d'infrastructure AiTM malveillante

Domaine enregistré il y a moins de 7 jours ressemblant à un service corporate ou partenaire

Certificat TLS récent (< 48h) sur un domaine ressemblant à un service d'entreprise connu

Redirection d'un lien hébergé sur Azure/AWS/Cloudflare vers un domaine tiers récent

Serveur retournant des erreurs de géolocalisation pour les IPs de sandbox analytiques connues

Présence de sous-domaines CNAME pointant vers des IPs bulletproof identifiées

Headers HTTP inhabituels révélant un framework Go (Server: Evilginx, via: nginx/go)

Campagnes AiTM documentées en 2026 : acteurs APT et secteurs ciblés

L'année 2026 a enregistré une recrudescence significative des campagnes AiTM contre les environnements Microsoft 365 et les plateformes de collaboration d'entreprise. Le groupe **Storm-0558**, attribué à la Chine, a perfectionné ses techniques AiTM pour cibler des agences gouvernementales et sous-traitants de défense européens, en développant des phishlets personnalisés pour Outlook Web Access et les portails ADFS. **Scattered Spider** (UNC3944) continue d'associer vishing et AiTM pour compromettre des Fortune 500 via des campagnes hautement ciblées contre les services IT help desk.

Ces campagnes illustrent une convergence inquiétante entre AiTM et d'autres vecteurs d'attaque avancés. Les compromissions initiales via AiTM sont fréquemment suivies de mouvements latéraux exploitant Active Directory, un sujet analysé en détail dans notre article sur les techniques [Purple Team 2026 pour AD et le cloud](#). La chaîne d'attaque typique observée en 2026 suit le schéma : **AiTM** → **accès cloud Microsoft 365** → **création de règles de forwarding** → **exfiltration d'emails** → **pivotement vers l'AD on-premises** → **déploiement ransomware ou fraude BEC**.

Un pattern particulièrement documenté en 2026 est l'utilisation d'AiTM comme premier vecteur dans des campagnes de **Business Email Compromise (BEC) de nouvelle génération**. Une fois le compte email d'un DAF ou d'un responsable financier compromis, les attaquants créent des règles silencieuses interceptant les conversations financières en cours, puis injectent leurs coordonnées bancaires au moment critique d'une transaction. Pour une analyse approfondie de ces vecteurs combinés, consultez notre article sur les [attaques BEC et vishing alimentées par l'IA en 2026](#).

Quels types de MFA sont vulnérables aux attaques AiTM bypass ?

L'idée reçue la plus dangereuse en matière d'authentification est que "tout MFA protège contre le phishing". Les attaques AiTM démontrent concrètement que la grande majorité des méthodes MFA courantes restent vulnérables à l'interception par proxy. Comprendre cette vulnérabilité est la première étape vers une architecture d'identité défensive efficace en 2026 :

SMS OTP : Vulnérable. Le code est saisi sur la page de phishing et relayé en temps réel au service légitime avant expiration (30-90s).

TOTP (Google Authenticator, Authy, Microsoft Authenticator code) : Vulnérable. Le code TOTP valide 30 secondes est largement suffisant pour l'interception et le relay AiTM instantané.

Notification push (Microsoft Authenticator, Duo Mobile) : Vulnérable à la *MFA fatigue*. Des demandes répétées finissent par être acceptées par inadvertance. Également relayable si la victime approuve normalement.

Email OTP : Vulnérable. Même mécanisme que le SMS OTP, avec une fenêtre d'expiration généralement plus longue.

Hardware token OATH HOTP/TOTP : Vulnérable. Le code est intercepté et relayé avant son expiration naturelle.

FIDO2 / WebAuthn / Passkeys : **Résistant**. Le challenge est cryptographiquement lié à l'origin du domaine légitime (rpId). Un proxy AiTM ne peut pas relayer ce challenge car l'authenticateur vérifie que l'origine correspond au domaine enregistré lors de l'enrôlement.

Certificats client (Smart Card, PIV, YubiKey PIV) : Résistant. L'authentification est liée à l'identité TLS du serveur cible, rendant tout relay structurellement impossible.

Cette distinction est capitale pour les équipes sécurité : seules les méthodes d'authentification *phishing-resistant* selon NIST SP 800-63B — FIDO2/WebAuthn et certificats client — offrent une protection réelle contre les attaques AiTM. La migration vers ces technologies représente la priorité défensive absolue pour toute organisation exposée à des acteurs sophistiqués en 2026.

Détection des attaques AiTM : SIEM, XDR et analyse comportementale en 2026

La détection des attaques AiTM est particulièrement délicate car, du point de vue du service cible, l'authentification s'est déroulée normalement. Les signaux de détection doivent se concentrer sur les **anomalies comportementales post-authentification** plutôt que sur le processus d'authentification lui-même. Les solutions XDR modernes, dont nous détaillons les capacités et les techniques de bypass dans notre article sur le [bypass EDR/XDR en 2026](#), intègrent désormais des modules spécifiques de détection AiTM basés sur la corrélation multi-sources.

Microsoft Defender XDR intègre depuis 2025 une règle de détection AiTM basée sur la corrélation entre les logs d'authentification Azure AD et les activités applicatives post-login. La règle s'active lorsqu'une lecture massive d'emails, une création de règle de forwarding ou des appels API inhabituels suivent une authentification depuis une IP et un User-Agent jamais vus. La technologie de **Continuous Access Evaluation (CAE)** permet en outre de révoquer les sessions suspectes en temps quasi-réel lors de la détection d'anomalies comportementales.

Source de log	Signal de détection AiTM	Sévérité	Taux de faux positifs
Azure AD Sign-in logs	Auth réussie + IP jamais vue + User-Agent nouveau simultanément	Haute	Faible (VPN, voyage)
Exchange Online	Règle inbox forwarding externe créée dans les 5 min post-login	Critique	Très faible
Microsoft Graph API	Appels massifs /me/messages ou /me/mailFolders dans les 10 min	Haute	Modéré
DNS / Proxy web	Résolution domaine ressemblant au corporate enregistré < 7 jours	Moyenne	Modéré
Email gateway (Defender for O365)	Lien Azure/AWS redirigeant vers domaine < 30 jours lors du clic	Haute	Faible
Azure AD Conditional Access	Token refresh depuis IP différente de l'IP d'émission (sans CAE)	Critique	Faible (NAT, proxies)
UEBA / Comportemental	Impossible travel < 5 min entre deux authentications distantes	Critique	Très faible

Pour les environnements disposant de ressources Active Directory on-premises, des techniques complémentaires comme le [Pass-the-Hash](#) sont fréquemment utilisées en post-exploitation après une compromission AiTM initiale. Une détection holistique doit couvrir l'ensemble de la chaîne d'attaque, du phishing initial jusqu'aux mouvements latéraux AD. De même, des vulnérabilités SSRF dans des applications web internes peuvent être exploitées comme relais dans une chaîne d'attaque complexe, comme détaillé dans notre article sur l'[exploitation SSRF](#).

Stratégies de défense contre les attaques AiTM bypass MFA en 2026

La défense efficace contre les attaques AiTM nécessite une approche en défense en profondeur, combinant des contrôles techniques, des politiques d'accès adaptatives et une sensibilisation ciblée des utilisateurs à risque. La mesure la plus impactante à court terme est l'activation du

Conditional Access avec Continuous Access Evaluation (CAE) dans Azure AD, qui permet de révoquer instantanément les sessions suspectes lors de la détection d'anomalies (changement d'IP, désactivation de compte, modification des groupes), éliminant la fenêtre d'exploitation des cookies interceptés par Evilginx2.

Le déploiement du **Token Protection** (anciennement Token Binding) dans Azure AD lie cryptographiquement les access tokens et refresh tokens à l'instance TLS spécifique de l'appareil ayant effectué l'authentification. Un cookie ou token volé via AiTM et rejoué depuis un autre appareil sera automatiquement rejeté. Cette fonctionnalité, disponible pour les applications Office 365 depuis 2024, représente l'une des défenses anti-AiTM les plus efficaces sans migration immédiate vers FIDO2.

Déployer FIDO2/Passkeys en priorité pour tous les comptes administrateurs, accès financiers et direction — seule protection cryptographiquement phishing-resistant disponible en 2026

Activer Conditional Access avec CAE et Token Protection pour tous les utilisateurs Microsoft 365 sans exception

Configurer des politiques basées sur la conformité de l'appareil (Intune compliance) pour n'autoriser que les appareils gérés, à jour et conformes

Activer Microsoft Defender for Office 365 Plan 2 avec analyse des liens à l'heure du clic et protection contre les URL de redirection suspectes

Déployer DMARC/DKIM/SPF en mode strict (p=reject) sur tous les domaines de l'organisation, y compris les domaines de marque secondaires et partenaires

Mettre en place une surveillance DNS proactive avec alertes sur les nouveaux domaines ressemblant au domaine corporate via les certificate transparency logs

Former les utilisateurs à risque élevé (DAF, RH, équipes IT) à identifier les signaux d'un proxy AiTM : URL légèrement différente, timing inhabituel, éléments d'interface manquants

Token Protection et authentification phishing-résistant : l'avenir de l'IAM en 2026

La réponse structurelle et pérenne aux attaques AiTM passe par une refonte des architectures d'identité vers des mécanismes *intrinsèquement phishing-résistant*. **FIDO2/WebAuthn** constitue le standard de référence : lors de l'authentification, le challenge cryptographique est signé avec la clé privée de l'authenticateur ET lié au **rpId** (relying party identifier) du service légitime. Si un proxy AiTM tente de relayer ce challenge, l'origine du proxy — domaine de phishing — ne correspond pas à l'rpId enregistré, et l'authenticateur refuse de signer, rendant l'attaque structurellement impossible quelle que soit la sophistication du proxy.

En 2026, les organisations les plus avancées déploient une stratégie d'identité en trois horizons : **immédiat** — Token Protection + CAE pour tous les utilisateurs Microsoft 365 — ; **moyen terme 6-12 mois** — FIDO2 pour les comptes à hauts privilèges et les accès sensibles — ; **long terme 12-24 mois** — migration complète vers les passkeys pour l'ensemble du parc utilisateur. Cette approche progressive réduit significativement la surface d'attaque AiTM dès les premières semaines, tout en planifiant la transformation complète de l'architecture d'identité.

Les équipes sécurité utilisant des approches [Purple Team en 2026](#) incluent systématiquement des simulations AiTM dans leurs exercices pour valider l'efficacité de ces contrôles en conditions réelles. Ces simulations permettent notamment de vérifier que les règles SIEM sont bien calibrées, que les procédures de réponse à incident couvrent le scénario AiTM, et que les utilisateurs à risque sont suffisamment sensibilisés pour signaler des URLs suspectes avant de saisir leurs credentials.

À RETENIR

À retenir — Priorités défensives anti-AiTM 2026

Priorité absolue : Migrer les comptes administrateurs et accès sensibles vers FIDO2/passkeys — seule protection structurellement résistante aux proxies AiTM

Court terme : Activer Azure AD Conditional Access avec CAE + Token Protection pour tous les utilisateurs Microsoft 365

Détection : Déployer des règles SIEM corrélant les anomalies post-authentification (impossible travel, inbox rules, appels API massifs) plutôt que surveiller uniquement l'authentification

Surveillance proactive : Monitorer les certificate transparency logs et les enregistrements DNS pour les domaines ressemblant aux domaines corporate

Validation : Exercices Purple Team avec simulations AiTM pour vérifier l'efficacité réelle des détections avant qu'un incident réel ne survienne

Email : DMARC p=reject + Defender for O365 avec analyse URL au clic + formation ciblée des profils à risque élevé

FAQ — Questions fréquentes sur les attaques AiTM et le bypass MFA Evilginx

Pourquoi les attaques AiTM avec Evilginx contournent-elles même la MFA la plus forte ?

Les attaques AiTM ne cassent pas la MFA à proprement parler : elles exploitent le fait que la session authentifiée est matérialisée par un cookie, indépendant de tout facteur d'authentification. Une fois la MFA validée par la victime, le service émet un cookie de session valide que le proxy Evilginx capture et exfiltre avant de le transmettre. L'attaquant réutilise ce cookie directement, sans avoir besoin du mot de passe ni du second facteur. Cette vulnérabilité est inhérente aux protocoles d'authentification web basés sur des cookies de session — TOTP, SMS OTP et notifications push sont tous vulnérables car le code éphémère est capturé et relayé en temps réel avant expiration. Seule une authentification FIDO2 ou par certificat client, cryptographiquement liée au domaine légitime via l'rpId, empêche structurellement ce relay en 2026, car l'authenticateur refuse de signer pour un domaine de phishing différent du service enregistré.

Comment détecter qu'un compte a été compromis via une attaque AiTM ?

Les indicateurs de compromission post-AiTM les plus fiables se distinguent par leur caractère post-authentification : une connexion réussie depuis une adresse IP géographiquement distante de la localisation habituelle, particulièrement si elle survient dans les secondes ou minutes suivant une authentification légitime (impossible travel) ; la création immédiate de règles de redirection d'emails vers une adresse externe non reconnue ; des appels API massifs à Microsoft Graph ou à l'API Exchange pour lire des emails en masse ; l'apparition d'appareils non enregistrés dans la liste des sessions Azure AD actives ; et des modifications de paramètres de compte (numéro de téléphone MFA, email de récupération) dans les minutes suivant la connexion. Les équipes SOC doivent créer des alertes SIEM corrélant ces événements plutôt que surveiller l'authentification seule, qui apparaît parfaitement légitime lors d'une attaque AiTM réussie depuis les logs du service cible.

Qu'est-ce qui différencie les attaques AiTM en 2026 des campagnes de phishing classiques ?

La différence fondamentale réside dans la sophistication technique, l'efficacité contre les défenses modernes et le niveau d'automatisation post-compromission. Un phishing classique ne vole que des credentials statiques, inefficaces si la MFA est activée. Les attaques AiTM de 2026 opèrent entièrement en temps réel, relayant l'intégralité du flux d'authentification incluant les tokens MFA éphémères, et déclenchent automatiquement l'exploitation post-compromission (extraction d'emails, création de règles de forwarding, initiation de transactions BEC) sans intervention manuelle immédiate de l'attaquant. En 2026, ces campagnes se caractérisent également par l'utilisation systématique de services cloud légitimes comme redirecteurs pour passer les filtres de messagerie, par une infrastructure hautement automatisée (déploiement, rotation, gestion des victimes via API), et par leur intégration dans des chaînes d'attaque multi-étapes visant aussi bien les environnements cloud que l'Active Directory on-premises.

Comment les équipes red team utilisent-elles légalement les outils AiTM en 2026 ?

Dans le cadre d'engagements red team contractuels dûment autorisés, Evilginx2 et Modlishka servent à valider la résistance réelle de l'organisation face au phishing ciblé et à tester la chaîne de détection complète du SOC. Le red team déploie une infrastructure AiTM réaliste, cible des comptes représentatifs identifiés dans le scope d'engagement, et mesure précisément le délai de détection par l'équipe blue team. L'objectif est triple : vérifier si la compromission initiale est possible malgré les contrôles déployés, valider que les alertes SIEM/XDR se déclenchent dans des délais acceptables (MTTD cible < 1h), et confirmer que les procédures de réponse à incident couvrent effectivement le scénario AiTM. Ces exercices, idéalement conduits en mode Purple Team avec une collaboration red/blue, permettent d'identifier précisément les lacunes dans la chaîne de détection et de les corriger avant qu'un attaquant réel ne les exploite dans un contexte malveillant.

Conclusion

Les attaques **AiTM avec Evilginx MFA bypass** représentent en 2026 la menace la plus directe et la plus efficace contre les architectures d'authentification basées sur MFA classique. La maturité des outils open source (Evilginx2, Modlishka), l'automatisation des infrastructures d'attaque, l'intégration dans des chaînes d'attaque multi-étapes (BEC, ransomware, mouvement latéral AD) et l'adoption par des groupes APT étatiques comme Storm-0558 placent cette technique au cœur des préoccupations des équipes de sécurité enterprise. La réponse ne peut être uniquement réactive : elle doit passer par une modernisation proactive des architectures d'identité vers FIDO2/passkeys, le déploiement du Conditional Access avec CAE et Token Protection, et des exercices Purple Team réguliers validant l'efficacité des contrôles en place. En 2026, toute organisation n'ayant pas démarré sa migration vers l'authentification phishing-resistant reste exposée à un vecteur d'attaque dont la barrière d'entrée continue de baisser, tandis que l'impact d'une compromission réussie — accès aux emails sensibles, fraude BEC, pivotement vers l'AD — demeure critique pour la continuité d'activité.

Évaluez et renforcez votre résistance aux attaques AiTM

Votre organisation utilise encore des méthodes MFA vulnérables aux attaques AiTM bypass ? Nos experts certifiés OSCP/CRTO réalisent des audits complets de votre architecture d'identité, des simulations d'attaque AiTM en conditions contrôlées et des exercices Purple Team pour identifier et corriger vos lacunes avant les attaquants réels.

[Demander un audit AiTM et MFA →](#)