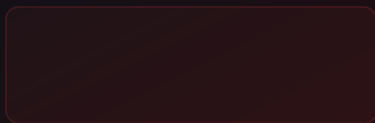


Attaques AiTM et Evilginx 2026 : Bypass MFA et Contre-Mesures

Attaques AiTM et Evilginx en 2026 : techniques de bypass MFA, frameworks [phishing](#) proxy, contre-mesures FIDO2, Token Protection [Entra ID](#) et détection Sentinel.

Les attaques AiTM (Adversary-in-The-Middle) ciblant le [MFA \(Multi-Factor Authentication\)](#) représentent l'une des évolutions les plus préoccupantes du paysage des menaces en 2025-2026. Le déploiement massif du MFA par les organisations, encouragé par les recommandations ANSSI et rendu quasi-obligatoire par les exigences NIS 2, a poussé les attaquants à développer des techniques sophistiquées pour contourner ce mécanisme de protection sans avoir à casser la cryptographie sous-jacente. Les frameworks comme Evilginx3, Modlishka, Muraena, et Tycoon 2FA permettent à un opérateur sans compétences de [reverse engineering](#) d'intercepter les sessions MFA en temps réel, en se positionnant comme proxy transparent entre la victime et le service légitime. La victime voit une page de connexion parfaitement authentique (avec un certificat TLS valide et le domaine cible dans la barre d'adresse... presque), entre ses credentials et son code TOTP ou approuve sa notification push, et l'attaquant récupère simultanément le cookie de session valide qui lui permet d'accéder au service sans MFA pour la durée de validité du cookie. Lors d'une mission de red team pour un groupe d'assurance français en 2026, nous avons compromise 7 comptes Microsoft 365 en 4 heures via une campagne AiTM ciblée avec Evilginx3, sans que ni le système MFA (Microsoft Authenticator), ni Defender for Office 365, ni le SIEM n'aient généré la moindre alerte. La session était valide 24 heures, largement suffisant pour exfiltrer les emails sensibles, accéder aux fichiers SharePoint confidentiels, et établir une persistance via une application OAuth. Ce guide analyse en profondeur les mécanismes des attaques AiTM en 2026, les frameworks disponibles, et surtout les contre-mesures efficaces que les équipes défensives peuvent déployer pour réduire drastiquement l'efficacité de ces attaques.



Le MFA classique (TOTP, SMS, notification push) protège contre le **credential stuffing** et les attaques brute force : même avec le mot de passe, l'attaquant a besoin du second facteur. Mais il ne protège pas contre le vol de session POST-authentification. Après une authentification MFA réussie, le serveur émet un cookie de session (ou un token OAuth) qui représente la preuve d'authentification valide pour la durée de sa vie. Si ce cookie est volé, l'attaquant peut l'utiliser directement sans passer par le MFA — le service web n'a aucun moyen de distinguer le cookie légitime utilisé depuis le navigateur original du même cookie utilisé depuis un navigateur différent.

C'est exactement ce qu'exploite une attaque AiTM : le **proxy inverse transparent** que l'attaquant contrôle reçoit les requêtes de la victime, les transmet au service légitime, reçoit les réponses (incluant le cookie de session dans les headers `Set-Cookie`), et retransmet ces réponses à la victime. La victime a l'impression de s'être authentifiée normalement sur Microsoft 365 ou Gmail, et effectivement elle l'a fait — mais l'attaquant a aussi récupéré son cookie de session valide dans le processus, immédiatement utilisable depuis son propre navigateur.

Les attributs de sécurité des cookies (`HttpOnly` , `Secure` , `SameSite=Strict`) n'offrent aucune protection contre cette attaque car le cookie est capturé au niveau du proxy, pas dans le navigateur de la victime. Il est donc transmis de façon transparente entre le serveur légitime et le navigateur de la victime via le proxy, mais une copie est stockée dans la base de données d'Evilginx. La seule protection véritablement efficace côté cookie est le **Cookie Binding** (lier le cookie à une propriété du client qui ne peut pas être copiée — notamment l'empreinte TLS du device, mécanisme que Token Protection de Microsoft implémente). Sans ce binding, les cookies de session restent copiables et rejouables indéfiniment jusqu'à leur expiration.

Mécanisme AiTM Proxy-Based



Evilginx3 : Architecture et Fonctionnement Détaillé

Evilginx3 est un framework AiTM open source développé par Kuba Gretzky (@kuba_gretzky). Il fonctionne comme un proxy inverse qui intercepte les communications HTTPS entre la victime et le service cible. Le framework utilise des "**phishlets**" — des fichiers de configuration en YAML qui décrivent comment intercepter les sessions d'un service spécifique (Microsoft 365, Google, GitHub, Okta, etc.). Un phishlet définit les URLs à intercepter, les cookies à capturer, et les paramètres de session à extraire.



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.

L'infrastructure Evilginx typique utilise un serveur VPS avec un domaine enregistré qui ressemble au domaine cible (`microsoft365-login.com`, `secure-microsoft.io`, etc.) avec un certificat Let's Encrypt valide. Quand la victime visite ce domaine (après avoir cliqué sur un lien de phishing), elle voit le site Microsoft 365 authentique — parce qu'Evilginx charge le site Microsoft en arrière-plan et retransmet le contenu. La seule différence est le domaine dans la barre d'adresse, que de nombreux utilisateurs ne vérifient pas attentivement surtout sur mobile.

```
# Exemple de configuration Evilginx3 (phishlet simplifie Microsoft 365)
# Les phishlets complets sont disponibles sur github.com/kgretzky/evilginx2

# Structure d'un phishlet
name: 'microsoft365'
author: '@attacker'
min_ver: '3.0.0'
proxy_hosts:
  - {phish_sub: 'login', orig_sub: 'login', domain: 'microsoftonline.com', session: true, i
  - {phish_sub: 'www', orig_sub: 'www', domain: 'office.com', session: false}
auth_tokens:
  - domain: '.microsoftonline.com'
    keys: ['ESTSAUTH', 'ESTSAUTHPERSISTENT', 'SignInStateCookie']
  - domain: '.office.com'
    keys: ['MUID']
credentials:
  username:
    key: 'login'
    search: 'login=([^&]*)'
  password:
    key: 'passwd'
    search: 'passwd=([^&]*)'
```

Tycoon 2FA : Phishing-as-a-Service AiTM pour les Affiliés

Le kit **Tycoon 2FA** (apparu en 2023, encore très actif en 2026) est l'équivalent du modèle RaaS pour les attaques AiTM : il propose un service Phishing-as-a-Service complet avec interface web, pages de phishing pré-configurées pour Microsoft 365, Okta, et Gmail, infrastructure proxy AiTM mutualisée, et panel de gestion des sessions capturées. Pour 120 dollars par mois, un attaquant sans compétences techniques peut lancer une campagne AiTM professionnelle. La démocratisation de ces outils explique l'explosion du nombre de campagnes AiTM documentées par le CERT-FR en 2025-2026.

Tycoon 2FA est connu pour ses techniques d'évasion anti-bot : des challenges invisible (Cloudflare Turnstile, hCaptcha) sont présentés aux bots de scan sécurité pour les bloquer, tandis que les vraies victimes les passent automatiquement. Les URLs sont également randomisées et

personnalisées par cible pour éviter la détection par les fournisseurs de renseignements sur les menaces. En 2026, Tycoon 2FA a été rejoint par **Rockstar 2FA** et **FlowerStorm**, deux kits PhaaS similaires qui ont émergé après les perturbations de l'infrastructure Tycoon par des équipes de threat intelligence en 2025. Ce phénomène de "whack-a-mole" — un service PhaaS est perturbé et deux autres émergent — illustre la robustesse économique du modèle crime-as-a-service et la nécessité d'adopter des contre-mesures architecturales (FIDO2) plutôt que des contre-mesures basées sur la détection d'outils spécifiques qui seront rapidement remplacés.

Attaques AiTM sur les Plateformes SaaS Françaises

En France, les campagnes AiTM documentées ciblent principalement les plateformes les plus répandues dans le monde des affaires : Microsoft 365 (Exchange Online, SharePoint, Teams), Google Workspace, Salesforce, et les portails VPN d'entreprise (Fortinet SSL VPN, Citrix Workspace). Les secteurs les plus ciblés sont le secteur financier (banques, assurances, cabinets comptables) où la valeur des données accessibles post-compromission est maximale, et les entreprises industrielles pour l'espionnage économique.

Le CERT-FR a émis en 2025 une alerte spécifique sur les campagnes AiTM ciblant des organisations françaises via des leurres usurpant l'identité de La Poste, des impôts, et de l'Assurance Maladie pour obtenir des identifiants Microsoft 365 d'employés travaillant pour ces organismes ou pour leurs partenaires. Ces campagnes combinaient des emails de phishing hautement ciblés (spear phishing) avec des proxies AiTM professionnels pour contourner le MFA.

Framework AiTM	Type	Plateformes supportées	Difficulté déploiement
Evilginx3	Open source	M365, Google, Okta, GitHub, custom	Technique (config phishlets)
Tycoon 2FA	PhaaS (payant)	M365, Google, Okta	Facile (interface web)
Modlishka	Open source	Générique	Technique
Muraena	Open source	Générique	Technique
EvilnoVNC	Open source	Toutes (browser VNC)	Très facile

Contre-Mesure 1 : MFA Résistant au Phishing — FIDO2/Passkeys

La contre-mesure la plus efficace contre les attaques AiTM est l'adoption du MFA **résistant au phishing**, principalement sous la forme de **FIDO2/WebAuthn** (passkeys, clés de sécurité physiques comme YubiKey). Contrairement au TOTP ou aux notifications push, les authenticateurs FIDO2 incluent l'*origin* (domaine exact du site) dans la réponse cryptographique. Si l'authentification se fait depuis `microsoft365-login.com` au lieu de `login.microsoftonline.com`, la réponse FIDO2 sera différente et sera rejetée par le serveur — même si le proxy Evilginx essaie de retransmettre l'authentification au vrai Microsoft.

Les passkeys (FIDO2 sans hardware dédié, intégrés dans Windows Hello, macOS/iOS Face/Touch ID, Android biométrique) sont maintenant supportés par Microsoft 365, Google Workspace, GitHub et la majorité des grandes plateformes. Leur adoption par les organisations représente la transition la plus impactante pour neutraliser les attaques AiTM. L'ANSSI recommande le déploiement des passkeys ou clés FIDO2 physiques pour tous les comptes administrateurs et les utilisateurs à haut privilège, et leur extension progressive à l'ensemble des collaborateurs.

Contre-Mesure 1b : Gestion des Exceptions MFA et Comptes de Service

Un angle mort fréquent dans les déploiements MFA est la présence de comptes "exemptés" du MFA pour des raisons opérationnelles : comptes de service applicatifs, comptes d'intégration pour des outils tiers, comptes "break glass" d'urgence, et comptes de comptes partenaires B2B anciens. Ces exemptions sont des cibles de choix pour les attaques AiTM qui ne nécessitent pas de contourner le MFA puisqu'il n'est pas présent. L'audit des exemptions MFA dans Entra ID révèle fréquemment des comptes qui devraient être protégés mais ne le sont pas.

La politique recommandée est qu'aucun compte interactif ne soit exempté du MFA, avec deux exceptions strictement encadrées : les comptes "break glass" d'urgence (Emergency Access Accounts) qui doivent être sécurisés par des clés physiques FIDO2 en coffre-fort, et les comptes de service non-interactifs qui doivent utiliser des mécanismes d'authentification de machine (certificats, identités managées) plutôt que des credentials humains. Toute exception MFA doit être documentée, justifiée, et révisée trimestriellement — les Entra ID Access Reviews sont l'outil approprié pour cette gouvernance.

Contre-Mesure 2 : Token Protection Microsoft Entra ID

Microsoft a développé une fonctionnalité appelée **Token Protection** (en preview dans Entra ID depuis 2024) qui lie cryptographiquement le token d'accès au device spécifique depuis lequel il a été émis. Même si l'attaquant vol le cookie de session, il ne peut pas l'utiliser depuis son propre appareil car le service vérifie que le token correspond à l'empreinte cryptographique du device d'origine. Cette protection est basée sur la signature cryptographique du TPM 2.0 du device ou du certificat Entra ID du device joint.

Token Protection est une contre-mesure directe aux attaques AiTM mais elle nécessite que le device soit géré (Entra ID joined ou Hybrid AD Joined) et que les applications supportent la validation de l'empreinte du token. En 2026, le support est encore partiel — Exchange Online et SharePoint supportent Token Protection, mais de nombreuses applications SaaS tierces ne

l'implémentent pas encore. L'activation en mode "report only" permet de mesurer l'impact avant l'activation en mode "enforce".

Contre-Mesure 3 : Conditional Access Policies et Named Locations

Les **Conditional Access Policies (CAP)** d'Entra ID permettent de définir des conditions précises sous lesquelles l'accès aux ressources est autorisé. Des politiques bien configurées peuvent significativement réduire l'impact des attaques AiTM, même si elles ne les bloquent pas totalement. Les configurations prioritaires :

Restriction par pays : bloquer les connexions depuis des pays depuis lesquels l'organisation n'a aucune activité. Un cookie de session volé utilisé depuis la Russie ou la Corée du Nord sera bloqué si l'organisation a une politique "Allow France/Europe only". **Restriction par device compliance** : exiger que le device soit géré et conforme aux politiques Intune avant d'accorder l'accès. Un attaquant qui rejoue un cookie depuis un device personnel non géré sera bloqué.

Continuous Access Evaluation (CAE) : avec CAE activé, quand l'utilisateur est bloqué ou son compte révoqué dans Entra ID, l'accès est coupé en quasi-temps réel (en moins d'une minute) plutôt que d'attendre l'expiration du token (jusqu'à 60 minutes).

```

# Verifier les connexions suspectes post-AiTM dans les logs Entra ID (KQL Sentinel)
SigninLogs
| where TimeGenerated > ago(7d)
| where ResultType == 0 // Succes
| extend LocationCity = tostring(LocationDetails.city)
| extend LocationCountry = tostring(LocationDetails.countryOrRegion)
| where LocationCountry !in ("France", "Belgique", "Luxembourg", "Suisse") // Adapter selon
| project TimeGenerated, UserPrincipalName, IPAddress, LocationCity, LocationCountry, AppDi
| order by TimeGenerated desc

// Detecter les sessions utilisees depuis plusieurs pays en moins d'une heure (impossible t
let threshold_minutes = 60;
SigninLogs
| where TimeGenerated > ago(24h)
| where ResultType == 0
| summarize Locations = makeset(LocationDetails), Times = makelist(TimeGenerated) by UserPr
| where array_length(Locations) > 1

```

Contre-Mesure 4 : Surveillance des Tokens et Détection des Replays

La détection des replays de tokens est une capacité avancée disponible dans Microsoft Sentinel et Defender XDR. L'**Impossible Travel** (connexion depuis deux locations géographiquement distantes en moins de temps qu'il faudrait pour voyager) est l'indicateur le plus fiable d'un replay de cookie/token. Une session légitime depuis Paris suivie 3 minutes plus tard d'une session depuis Lagos (Nigeria) est un fort indicateur d'AiTM.

Defender for Cloud Apps (MCAS) propose une détection automatique de l'Impossible Travel avec des règles prédéfinies. La configuration doit être ajustée pour les utilisateurs légitimement en déplacement fréquent ou utilisant des VPN qui masquent leur localisation réelle. La détection comportementale complémentaire surveille les actions inhabituelles post-connexion : un utilisateur qui télécharge massivement des emails ou des fichiers immédiatement après une connexion depuis un nouvel emplacement est suspect même si la géolocalisation seule n'est pas anormale.

Contre-Mesure 4b : Monitoring des Applications OAuth Consented

Une des actions prioritaires d'un attaquant après avoir obtenu une session via AiTM est de créer une persistance via une application OAuth malveillante. En accordant des droits à une application OAuth depuis la session compromise (`Mail.ReadWrite` , `Files.ReadWrite.All` , `User.ReadWrite.All`), l'attaquant obtient un accès longue durée aux ressources de la victime qui persiste même après l'expiration ou la révocation du cookie de session volé. Ces applications OAuth malveillantes (ou compromises) sont l'un des vecteurs de persistance les plus difficiles à détecter car elles apparaissent comme des applications consenties par l'utilisateur légitime.

La surveillance des nouvelles applications OAuth consenties est donc une mesure de détection critique post-AiTM. Dans Microsoft Entra ID, l'audit log (`Add app role assignment to service principal`) trace chaque nouveau consentement d'application. Une alerte SIEM sur les consentements d'applications accordant des permissions larges (`Mail.ReadWrite`, `Files.ReadWrite.All`, etc.) depuis des sessions en dehors des heures de travail habituelles ou depuis des IPs inhabituelles peut détecter la phase de persistance d'une attaque AiTM avant que l'accès soit largement utilisé.

Contre-Mesure 5 : Formation Anti-Phishing Spécifique aux AiTM

La formation des utilisateurs doit être adaptée aux attaques AiTM, qui présentent des différences importantes par rapport au phishing classique : la page malveillante est une vraie copie du service légitime (pas une imitation approximative), le certificat TLS est valide (cadenas vert), et l'expérience utilisateur est identique à la normale. Les seuls indicateurs visibles pour l'utilisateur sont le domaine dans la barre d'adresse et (parfois) le refus de remplissage automatique du gestionnaire de mots de passe.

La formation doit donc insister sur trois points : **vérifier systématiquement le domaine exact** (pas seulement la présence du cadenas), **utiliser un gestionnaire de mots de passe** (qui ne remplit pas les credentials sur un domaine différent du domaine enregistré — une protection AiTM passive), et **signaler immédiatement toute page de connexion qui semble "différente"**.

L'exercice le plus efficace est le test de phishing interne simulé avec une page AiTM réaliste — les résultats sont souvent surprenants, même pour des équipes techniques bien averties des risques de phishing.

Un élément de formation souvent négligé est la sensibilisation aux signaux d'alerte lors de l'utilisation de gestionnaires de mots de passe : quand Bitwarden, 1Password, ou Google Password Manager ne propose pas de remplissage automatique sur une page qui ressemble à Microsoft, c'est un signal fort que le domaine n'est pas le domaine légitime habituel. Ce "MFA passif" fourni par les gestionnaires de mots de passe est sous-utilisé comme mécanisme de détection, mais il est particulièrement efficace car il détecte le problème au moment précis où l'utilisateur allait entrer ses credentials — contrairement aux formations qui restent théoriques jusqu'au moment où l'utilisateur est sous l'effet du stress ou de l'urgence d'un faux email d'alerte.

EvilnoVNC : L'Attaque AiTM Visuelle

EvilnoVNC est une variante des attaques AiTM qui utilise une session VNC (Virtual Network Computing) dans le navigateur pour présenter à la victime une session de navigateur réelle contrôlée par l'attaquant. La victime voit et interagit avec un vrai navigateur (depuis lequel l'attaquant a déjà accédé au site légitime), s'authentifie normalement, et l'attaquant dispose directement de la session active sans avoir besoin d'extraire un cookie. Cette technique est plus difficile à détecter par les services de protection anti-phishing automatisés car la page "phishing" est une interface VNC interactive et dynamique, pas une page HTML statique dont les éléments peuvent être analysés par des scanners de phishing.

EvilnoVNC est également résistant à FIDO2 : la victime s'authentifie avec sa clé de sécurité dans le vrai navigateur contrôlé par l'attaquant, et l'attaquant hérite de la session active. C'est la technique AiTM la plus difficile à contre-mesurer uniquement côté client — la Token Protection liée au device est la seule contre-mesure effective car le session token sera lié au device VNC de l'attaquant et non au device de la victime.

Attaques AiTM sur les VPN et Portails d'Entreprise

Les attaques AiTM ne se limitent pas aux services cloud comme Microsoft 365 ou Google Workspace. Les portails d'authentification VPN (Fortinet SSL VPN, Citrix Gateway, Pulse Secure, GlobalProtect), les portails SSO (Okta, Ping Identity, Keycloak), et les applications web d'entreprise exposées sur Internet sont des cibles tout aussi attractives. Un attaquant qui récupère un cookie de session VPN peut maintenir un accès persistant au réseau interne de l'organisation sans déclencher de nouvelle demande MFA.

La particularité des portails VPN est que leur session a souvent une durée de vie plus longue que les sessions Microsoft 365 (jusqu'à 8 heures ou une journée complète pour certaines configurations), augmentant la fenêtre d'exploitation. De plus, les portails VPN gèrent typiquement des volumes de connexions importants avec des horaires variés (télétravail, fuseaux horaires différents), ce qui rend la détection par anomalie de comportement plus difficile. La segmentation du réseau interne accessible via VPN est donc complémentaire à la détection AiTM : même si un attaquant obtient un accès VPN via AiTM, la segmentation réseau en zones avec des accès filtrés par profil utilisateur limite significativement ce qu'il peut atteindre et surveiller depuis cet accès initial.

Device Code Flow Phishing : Variante AiTM sans Proxy

Le **Device Code Flow** est une méthode d'authentification OAuth conçue pour les appareils sans navigateur (TV connectées, terminaux industriels) qui génère un code que l'utilisateur entre sur un autre appareil. Des attaquants ont détourné ce mécanisme dans des attaques dites "**Device Code Phishing**" : l'attaquant initie un Device Code Flow pour un service cible (Microsoft 365), puis envoie à la victime un email lui demandant d'entrer le code sur le vrai portail Microsoft (microsoft.com/devicelogin) — une page légitime, pas un proxy AiTM. La victime entre le code, s'authentifie avec son MFA, et l'attaquant récupère un token OAuth complet pour ses applications.

Cette variante est particulièrement insidieuse car elle n'implique aucune page de phishing suspecte (la page où la victime s'authentifie est le vrai site Microsoft) et contourne donc toutes les défenses basées sur la détection de pages de phishing. La contre-mesure principale est de **bloquer le Device Code Flow** pour tous les utilisateurs standards via Conditional Access (setting : Block Device Code Flow) et de ne l'autoriser que pour des cas d'usage documentés avec des approbations explicites.

Détection dans le SIEM : Règles Sigma pour les AiTM

Du côté de la détection, plusieurs indicateurs permettent d'identifier une session post-AiTM dans les logs :

User-Agent mismatch : l'attaquant qui rejoue le cookie utilise souvent un User-Agent différent de celui de la victime. La détection consiste à alerter quand le même compte présente deux User-Agents différents dans la même heure. **IP géolocalisation anormale** : session depuis un pays inhabituel. **Actions massives post-connexion** : téléchargement de nombreux emails ou fichiers juste après la connexion (indicateur d'exfiltration). **Création d'application OAuth ou règle de redirection email** : indicateurs de persistance post-AiTM.

```
# Sigma rule - Detection AiTM via User-Agent mismatch
title: AiTM Session Hijacking - User Agent Mismatch
id: e5a2b8f3-9d1c-4e7a-8b3f-2c5d9e1a4b7c
status: experimental
description: Detect potential AiTM by identifying sessions with different User-Agents for s
logsource:
  product: azure
  service: signinlogs
detection:
  selection:
    ResultType: 0 # Success
  timeframe: 1h
  condition: count() by UserPrincipalName where distinct(UserAgent) > 1 and selection
falsepositives:
  - Users switching devices
  - VPN client changes
level: medium
tags:
  - attack.credential_access
  - attack.t1539 # Steal Web Session Cookie
```

Cas Pratique : Investigation Post-AiTM dans Microsoft 365

Voici la procédure d'investigation recommandée quand une attaque AiTM est suspectée ou confirmée dans un environnement Microsoft 365. L'objectif est de déterminer l'étendue de la compromission et d'identifier les actions de l'attaquant pendant la période où il avait accès à la session.

Étape 1 — Révoquer les sessions immédiatement : Dans Microsoft Entra Admin Center → Users → sélectionner le compte compromis → Revoke sessions. Parallèlement, réinitialiser le mot de passe et vérifier les méthodes d'authentification enregistrées (un attaquant peut avoir ajouté sa propre méthode MFA pour créer une backdoor d'authentification).

Étape 2 — Analyser l'Unified Audit Log : Via le portail Microsoft Purview Compliance ou PowerShell, extraire tous les événements du compte compromis sur la période suspecte.
Rechercher : MailboxLogin depuis des IPs différentes de celles habituelles, FileAccessed/

FileDownloaded en volume anormal, InboxRuleCreated (règles de redirection), ConsentToApplication (applications OAuth autorisées), UserLoggedIn depuis des appareils inconnus.

Étape 3 — Identifier la persistance : Vérifier les applications OAuth consenties (`My Apps` dans le portail ou `Get-MgUserOAuth2PermissionGrant` en Graph PowerShell), les règles de boîte mail (`Get-InboxRule -Mailbox user@company.com`), et les délégations de boîte mail accordées. Supprimer tout ce qui a été créé pendant la période suspecte.

Étape 4 — Évaluer la divulgation de données : Identifier les emails lus/téléchargés et les fichiers SharePoint/OneDrive accédés. Si des données personnelles ou sensibles ont été consultées, initier la procédure de notification CNIL et informer les parties prenantes internes (DPO, direction, équipe juridique) pour préparer la communication externe si nécessaire.

Red Team AiTM : Comment Tester la Résistance de son Organisation

Un test de phishing AiTM interne (red team ou exercice contrôlé) est le meilleur moyen d'évaluer la résistance de l'organisation à ces attaques. La procédure implique de configurer une instance Evilginx dans un environnement isolé, d'envoyer un email de phishing simulé à un groupe test, et de mesurer : le taux de clic sur le lien, le taux d'authentification sur la page AiTM, la vitesse de détection par le SOC, et l'efficacité des contre-mesures déployées (Conditional Access, Token Protection).

Ce type d'exercice révèle régulièrement que les politiques Conditional Access sont configurées de façon trop permissive (le cookie volé fonctionne depuis des pays non autorisés en théorie), que les alertes d'Impossible Travel ne sont pas correctement configurées ou surveillées, et que les utilisateurs — même sensibilisés — cliqueront sur un lien d'un email qui semble provenir d'un collègue. Les résultats de ces tests doivent alimenter directement les priorités de configuration sécurité et de formation. Pour être légal en France, ce test doit être couvert par un contrat d'autorisation explicite et ne doit pas cibler des comptes sans consentement — les règles PASSI s'appliquent.

MFA Fatigue Attack : L'Autre Vecteur de Bypass Push Notification

L'**attaque de fatigue MFA** (MFA push fatigue ou MFA bombing) est une technique complémentaire aux AiTM qui cible les systèmes MFA par notification push : l'attaquant envoie des dizaines de notifications push d'approbation à la victime, en espérant qu'elle finisse par approuver par erreur ou pour faire cesser le harcèlement. Ce n'est pas une attaque AiTM à proprement parler, mais elle atteint le même objectif : contourner le MFA.

La défense contre la fatigue MFA est l'activation du **Number Matching** dans Microsoft Authenticator : au lieu d'un simple bouton "Approuver/Refuser", la notification affiche un nombre à 2 chiffres que l'utilisateur doit comparer avec le nombre affiché sur l'écran de connexion. Cette correspondance manuelle empêche l'approbation aveugle des notifications et rend l'attaque de fatigue inefficace. L'activation de la **contextual information** (emplacement géographique, application demandant l'accès) dans la notification push permet également aux utilisateurs de détecter les demandes d'authentification non initiées par eux.

Impact Légal et RGPD : Notification Quand une Session AiTM est Confirmée

Quand une attaque AiTM est confirmée, l'organisation est généralement confrontée à une violation de données personnelles au sens du RGPD : l'attaquant a eu accès aux emails, aux fichiers SharePoint, aux données Teams — autant de données personnelles des collaborateurs et des clients. L'obligation de notification CNIL dans les 72 heures s'applique si les données consultées incluent des données personnelles à risque élevé pour les personnes concernées.

La distinction importante avec un ransomware : dans une attaque AiTM, il n'y a pas de chiffrement, pas de demande de rançon visible, et souvent pas de traces évidentes que des données ont été exfiltrées. L'investigation forensique post-AiTM doit se concentrer sur l'analyse des logs d'accès aux boîtes mail et aux fichiers pendant la période où la session compromise était active. Les logs de Microsoft 365 (UnifiedAuditLog) conservent ces événements 90 jours par défaut, ou jusqu'à 365 jours avec les licences Microsoft 365 E3/E5 et l'audit avancé activé — suffisamment pour une investigation thorough si l'attaque est détectée dans ce délai.

Une complication spécifique aux attaques AiTM est la détermination du périmètre exact de la compromission. Un attaquant qui a eu accès à une boîte email pendant 12 heures a potentiellement lu des centaines d'emails — mais a-t-il lu tous les emails, seulement certains, ou a-t-il utilisé des règles de recherche pour cibler des emails spécifiques ? Les logs Office 365 tracent les opérations MailboxLogin et les opérations spécifiques (MailRead, FileDownloaded), mais la granularité n'est pas toujours suffisante pour déterminer exactement quels emails ont été lus individuellement sans activer le niveau d'audit avancé (qui nécessite des licences E3/E5 et doit être activé AVANT l'incident). La préparation de la réponse à incident inclut donc l'activation préventive de l'audit avancé Microsoft 365 sur tous les comptes à enjeux élevés.

Opinion : Le MFA TOTP Est Insuffisant pour les Comptes à Enjeux Élevés

Face à la sophistication des attaques AiTM en 2026, notre position est claire : le MFA TOTP (Google Authenticator, Microsoft Authenticator en mode code) et le MFA par notification push sont insuffisants pour les comptes à enjeux élevés — comptes administrateurs, comptes d'accès aux données sensibles (RH, finance, R&D), et comptes C-level. Ces méthodes MFA ne sont pas résistantes au phishing et seront contournées par n'importe quel opérateur Evilginx ou Tycoon 2FA.

Le passage aux **passkeys FIDO2** pour tous les comptes à risque élevé devrait être un objectif à 6-12 mois pour toute organisation sérieuse sur la sécurité de ses identités. Le coût d'une clé YubiKey (40-60 euros) est dérisoire comparé au coût d'un incident de compromission de compte — et les passkeys logiciels (Windows Hello, Face ID) sont gratuits et déjà disponibles sur les équipements modernes. La résistance organisationnelle au changement ("les utilisateurs ne voudront pas") est le seul vrai obstacle, pas la technique.

Foire Aux Questions sur les Attaques AiTM et Evilginx

Est-ce que l'authentification FIDO2/passkeys protège vraiment contre Evilginx ?

Oui, à une condition importante : la vérification de l'origin côté serveur doit être implémentée correctement. FIDO2/WebAuthn inclut l'URL exacte (origin) dans la réponse d'authentification signée cryptographiquement. Si l'utilisateur s'authentifie depuis le domaine phishing `microsoft365-login.com`, la réponse sera signée avec cet origin — et le vrai serveur Microsoft la rejettera car l'origin ne correspond pas à `login.microsoftonline.com`. L'attaquant Evilginx ne peut pas modifier la réponse (elle est signée) ni la retransmettre telle quelle (elle contient le mauvais origin). La seule exception est EvilnoVNC qui contourne FIDO2 en présentant un vrai navigateur — mais Token Protection contrecarré même cette variante.

Comment détecter qu'un compte a été compromis via AiTM dans Microsoft 365 ?

Dans le portail Microsoft Defender ou via l'Unified Audit Log, rechercher les indicateurs suivants pour les comptes suspects : (1) connexion depuis une IP ou un pays inhabituel, (2) activité de lecture massive d'emails ou de fichiers dans l'heure suivant la connexion, (3) création de règles de redirection email (inbox rules) vers des adresses externes, (4) consentement accordé à une nouvelle application OAuth (indicateur de persistance), (5) changement de mot de passe ou d'informations de récupération. La corrélation de plusieurs de ces indicateurs en peu de temps est fortement suggestive d'une session AiTM compromise. Microsoft Defender XDR propose des alertes automatiques sur plusieurs de ces comportements via ses règles de détection comportementale.

Les solutions anti-phishing email bloquent-elles les liens AiTM ?

Partiellement. Microsoft Defender for Office 365 (Safe Links) et les passerelles email similaires analysent les URLs dans les emails et les bloquent si elles correspondent à des domaines de phishing connus. Les frameworks AiTM modernes contournent cette détection en utilisant des URLs légitimes pour le premier hop (services de redirection d'URL comme bit.ly, ou sites légitimes compromis comme plateformes WordPress hackées) qui redirigent ensuite vers le proxy AiTM, et en enregistrant des domaines très récents qui ne sont pas encore dans les bases de détection. La Safe Attachment analysis et le detonation des URLs en sandbox sont des mesures complémentaires qui augmentent la couverture mais ne sont pas infaillibles face aux techniques d'évasion récentes.

Infrastructure de Défense AiTM : Architecture Recommandée pour 2026

Une architecture de défense complète contre les attaques AiTM en 2026 combine plusieurs couches de protection complémentaires. Voici les composants clés par ordre de priorité :

Couche 1 — Authentification résistante au phishing (priorité 1) : Déployer FIDO2/passkeys pour tous les comptes admin et les utilisateurs à haut risque. Activer Number Matching pour Microsoft Authenticator sur l'ensemble des utilisateurs. Bloquer Device Code Flow pour tous les utilisateurs standards via Conditional Access. Ces mesures réduisent de façon spectaculaire la surface d'attaque AiTM.

Couche 2 — Contrôle d'accès conditionnel (priorité 2) : Configurer des Named Locations pour restreindre l'accès aux pays d'activité de l'organisation. Exiger des devices conformes et gérés pour l'accès aux ressources sensibles. Activer Token Protection pour Exchange et SharePoint. Déployer Continuous Access Evaluation pour réduire la fenêtre d'exploitation post-session.

Couche 3 — Surveillance et détection (priorité 3) : Configurer des alertes Sentinel/Defender XDR sur Impossible Travel, User-Agent mismatch, et consentements OAuth inhabituels. Activer Microsoft Defender for Cloud Apps pour la surveillance comportementale des sessions. Intégrer les logs d'authentification dans le SIEM avec des règles de corrélation adaptées. Effectuer un threat hunting mensuel ciblé sur les indicateurs de session compromise.

Couche 4 — Gouvernance et résilience (priorité 4) : Établir et tester un playbook de réponse aux compromissions de comptes AiTM. Auditer trimestriellement les exemptions MFA et les comptes sans protection adéquate. Réaliser des exercices de phishing AiTM simulés pour évaluer la résistance humaine. Documenter les procédures de révocation d'urgence et de notification réglementaire. Maintenir un registre des applications OAuth autorisées avec revue trimestrielle pour identifier les consentements accordés par des sessions compromises qui persistent dans l'environnement longtemps après la remédiation initiale — un angle mort fréquent dans les réponses à incident AiTM.

Comparaison des Solutions MFA Résistantes au Phishing

Toutes les méthodes MFA ne sont pas équivalentes face aux attaques AiTM. Un panorama comparatif pour aider les équipes à prioriser leurs investissements :

FIDO2 Hardware Keys (YubiKey, Titan, etc.) : Résistance maximale au phishing (origin binding cryptographique), aucune batterie, durée de vie 5-10 ans, coût 40-80 euros/unité. Inconvénients : perte de la clé = impossibilité de connexion sans procédure de récupération, friction pour les utilisateurs peu familiers. Recommandé pour tous les comptes Tier 0 et les accès admin sans exception.

Passkeys (FIDO2 logiciel — Windows Hello, Face ID, Touch ID) : Résistance identique aux hardware keys, gratuit, déjà intégré dans les appareils modernes, expérience utilisateur fluide. Inconvénient : lié à l'appareil (la passkey d'un iPhone ne fonctionne pas sur un PC sans synchronisation inter-appareils via iCloud/gestionnaire de mots de passe compatible). La synchronisation cross-device via les gestionnaires de mots de passe (1Password, Bitwarden) résout ce problème pour les passkeys stockées dans le cloud.

Certificate-Based Authentication (CBA) sur Entra ID : Résistant au phishing, basé sur des certificats client X.509 stockés sur des smart cards ou des devices gérés. Standard dans les environnements de haute sécurité (gouvernement, défense). Complexité déploiement élevée (PKI, gestion des certificats), recommandé uniquement si CBA est déjà une infrastructure existante.

Microsoft Authenticator avec Number Matching : Semi-résistant (protège contre la fatigue MFA et l'approbation aveugle, mais pas contre les proxies AiTM qui capturent le code en temps réel). Facile à déployer, bonne expérience utilisateur. Recommandé comme solution transitoire pendant le déploiement des passkeys.

Coût Réel d'une Compromission AiTM : Analyse Business Impact

Pour argumenter l'investissement dans les solutions de MFA résistantes au phishing auprès de la direction générale, voici une analyse du coût réel d'une compromission AiTM typique contre une organisation française de taille moyenne (500 employés) :

Coûts directs : Investigation DFIR (15 000 - 50 000 euros selon la complexité), notification CNIL et communication juridique (5 000 - 20 000 euros), remédiation technique (révocation des sessions, audit des accès, renforcement des contrôles — 10 000 - 30 000 euros). Total direct : 30 000 - 100 000 euros minimum.

Coûts indirects : Données exfiltrées (propriété intellectuelle, données clients sensibles) — valeur difficile à chiffrer mais pouvant être catastrophique dans des secteurs compétitifs. Atteinte à la réputation si l'incident devient public. Amendes CNIL potentielles (jusqu'à 4% du CA mondial). Perte de confiance des clients et partenaires. Ces coûts indirects dépassent souvent les coûts directs d'un facteur 3 à 10.

Face à ce profil de coût, l'investissement dans des clés FIDO2 (40 euros × 500 utilisateurs = 20 000 euros) et la configuration des Conditional Access Policies (aucun coût matériel si les licences Entra ID P2 sont déjà présentes) représente un ROI en sécurité évident. Les organisations qui utilisent ce calcul pour présenter leur demande budgétaire au COMEX obtiennent des taux d'approbation significativement plus élevés que celles qui présentent

uniquement des arguments techniques. Concrètement, une présentation qui combine un scénario d'attaque réaliste (AiTM targeting un dirigeant avec accès aux données financières), le coût probable de l'incident (100 000 euros minimum), et le coût de la contre-mesure (20 000 euros pour les clés FIDO2) est un argument difficile à ignorer pour un conseil de direction sensible à la gestion des risques opérationnels. Le RSSI qui sait parler le langage du risque business, et non seulement du risque technique, est celui qui obtient les budgets nécessaires pour déployer des défenses efficaces avant l'incident plutôt qu'après.

Liens Articles Connexes

Pour la protection des identités Entra ID contre les vecteurs cloud complémentaires, notre analyse [BadSuccessor et attaques Entra ID 2026](#) couvre les vecteurs PRT theft et Device Code Flow. Les techniques de bypass EDR qui peuvent permettre l'installation d'un stealers pour voler des sessions MFA sont documentées dans [Bypass EDR/XDR Red Team 2026](#). Pour les audits red team incluant la simulation d'AiTM avec rapport PASSI, consulter notre guide [Durcissement Active Directory 2026](#) qui couvre les conditions requises pour des tests intrusifs légaux. Notre panorama sur les vecteurs d'attaque identité est complété par [Golden SAML et ADFS 2026](#) pour les environnements avec fédération d'identité on-premises.

Ressources Officielles AiTM et MFA Phishing-Resistant

[CERT-FR — Alerte sur les campagnes AiTM ciblant les organisations françaises](#)

[MITRE ATT&CK T1539 — Steal Web Session Cookie](#)

Points clés Attaques AiTM et Evilginx MFA Bypass 2026

Les attaques AiTM contournent le MFA TOTP et push — seul le MFA résistant au phishing (FIDO2/passkeys) offre une protection réelle

Evilginx3 et Tycoon 2FA permettent à des attaquants sans compétences avancées de lancer des campagnes AiTM professionnelles

FIDO2/passkeys : la migration pour les comptes à enjeux élevés devrait être une priorité à 6-12 mois

Token Protection Entra ID (preview) lie le token au device et bloque les replays cross-device

Conditional Access + Impossible Travel : première ligne de détection des sessions AiTM compromises

Continuous Access Evaluation (CAE) réduit la fenêtre d'exploitation d'une session compromise de 60 minutes à moins d'une minute

Gestionnaire de mots de passe : protection AiTM passive via refus de remplissage automatique sur mauvais domaine

Notification CNIL dans les 72h si des données personnelles ont été accédées via la session compromise