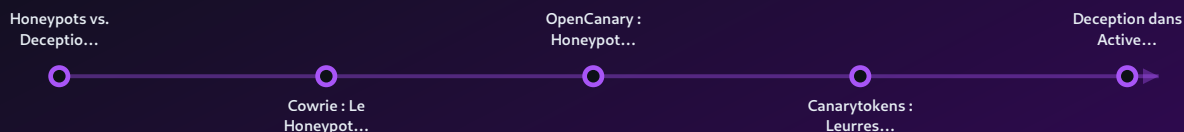


Honeypots et Technologies de Déception 2026 : Détection Active

Guide honeypots et déception 2026 — Cowrie, OpenCanary, Canarytokens, deception platforms et intégration SOC pour détection haute fidélité en entreprise.

Les **honeypots** et les **technologies de déception** ([deception technology](#)) constituent une approche radicalement différente de la détection des menaces cyber : plutôt que d'analyser passivement des logs d'activités légitimes pour identifier des comportements suspects, elles créent délibérément des leurres — des systèmes, services, données et identités factices — qui n'ont aucune raison légitime d'être accédés. Toute interaction avec un honeypot est donc nécessairement suspecte, avec un taux de faux positifs théoriquement proche de zéro. En 2026, les technologies de déception ont évolué bien au-delà des honeypots SSH traditionnels pour englober des plateformes sophistiquées de déception distribuée : objets Active Directory factices (Decoy Users, Decoy Computers, faux GPOs), faux partages réseau contenant des fichiers piégés avec des Canary Tokens, honeypots cloud (faux buckets S3, faux API keys AWS), et des plateformes commerciales comme Attivo Networks (intégré dans SentinelOne), Illusive Networks (maintenant Tenable.id) ou Acalvio ShadowPlex qui déploient des dizaines de milliers de leurres à travers l'ensemble du SI. En France, l'adoption des technologies de déception reste limitée aux organisations les plus matures (OIV, grandes banques, opérateurs télécom) mais l'intérêt croissant pour la détection haute fidélité et la réduction des faux positifs SOC pousse de plus en plus d'ETI à explorer ces approches. Ce guide présente l'écosystème des honeypots et des technologies de déception en 2026 : des outils open source accessibles (Cowrie, OpenCanary, Canarytokens) aux plateformes entreprise sophistiquées, en passant par les stratégies de déploiement, l'intégration SOC, et les enjeux légaux spécifiques aux leurres cyber en France. La déception est le retour à l'offensive de la défense — vous faites douter l'attaquant, vous gaspillez son temps, et vous le détectez au moment précis où il pense avoir trouvé quelque chose d'intéressant.

Honeypots et Déception 2026 : Détection des Attaquants



OUTILS / MÉTHODES :

honeypots

technologies de décept...

Canary Token

Deception Platform

Les honeypots et les technologies de déception (deception technology) constituent une approche radicalement différente de la...

Honeypots vs. Deception Technology : Taxonomie

La terminologie dans ce domaine peut être confuse. Quelques définitions pour clarifier :

Honeypot : Système isolé simulant un vrai système pour attirer et piéger les attaquants. Le terme englobe tout leurre numérique individuel (un serveur SSH, une base de données, une API vulnérable).

Honeynet : Réseau de honeypots simulant une infrastructure complète (plusieurs machines, services réseau réalistes, trafic simulé) pour une immersion réaliste des attaquants.

Canary Token : Marqueur numérique piégé (un fichier Word, une URL, une clé API, une entrée DNS) qui envoie une alerte dès qu'il est accédé — sans simuler de service réseau complet.

Deception Platform : Solution entreprise déployant et gérant de nombreux leurres distribués à travers l'infrastructure, avec gestion centralisée et intégration SOC.

Honeypot à haute interaction : Système réel (pas simulé) exposé comme leurre — risque plus élevé car l'attaquant peut compromettre un vrai système, mais interactions plus riches et TTPs plus complètes collectées.

Honeypot à faible interaction : Émulateur simulant des services sans exposer un système réel — plus sûr, déploiement plus simple, mais interactions parfois moins convaincantes pour les attaquants sophistiqués.

Cowrie : Le Honeypot SSH/Telnet de Référence

Cowrie est le honeypot SSH et Telnet open source le plus utilisé dans le monde. Il simule un système Linux avec un shell complet, un système de fichiers factice, et journalise en détail toutes les interactions de l'attaquant : commandes exécutées, fichiers tentés de télécharger, outils utilisés, credentials testés. Les données collectées par Cowrie ont contribué à de nombreuses recherches sur les techniques d'exploitation automatisée des accès SSH brute-force. Cowrie peut être déployé en quelques minutes via Docker et configuré pour simuler différents types de systèmes Linux (Ubuntu, Debian, CentOS) avec des configurations personnalisées pour paraître plus crédible dans le contexte spécifique de l'organisation.



Retour terrain

Pour un groupe de services financiers qui recevait 2 millions d'alertes SIEM par jour et en traitait 3 %, j'ai réalisé une analyse de la valeur des règles de détection sur 6 mois de logs. Résultat : 73 % des alertes provenaient de 4 règles mal calibrées. La suppression ou le réajustement de ces 4 règles a réduit le bruit de 65 % sans manquer aucun incident réel lors des 3 mois de validation. Le ROI d'un tuning SIEM est parmi les meilleurs investissements défensifs possibles.

Déploiement Cowrie avec Docker

```
docker pull cowrie/cowrie
docker run -p 2222:2222 -p 23:2223 \
-v /opt/cowrie/etc:/cowrie/cowrie/etc \
```

```
-v /opt/cowrie/log:/cowrie/var/log/cowrie \  
-e COWRIE_TELNET_ENABLED=true \  
-d cowrie/cowrie  
  
# Logs en JSON pour intégration SIEM  
tail -f /opt/cowrie/log/cowrie.json | jq .
```

OpenCanary : Honeypot Multi-Services Polyvalent

OpenCanary (Thinkst Canary open source) simule de nombreux services réseau sur une seule machine : SSH, HTTP, HTTPS, FTP, Telnet, SMB, MSSQL, MySQL, VNC, RDP, SNMP, SIP — permettant de créer un leurre qui ressemble à un vrai serveur multi-services. Contrairement à Cowrie (spécialisé SSH), OpenCanary permet de détecter des scans de reconnaissance réseau, des tentatives d'accès à des services de base de données, et des tentatives de connexion RDP — des patterns typiques de mouvement latéral post-compromission. OpenCanary peut être configuré pour envoyer des alertes par email, webhook, ou syslog — facilitant l'intégration dans le SIEM du SOC. Son déploiement est simple (pip install opencanary, édition d'un fichier de configuration JSON) et son footprint système minimal permet de le déployer sur de petites machines virtuelles ou des Raspberry Pi répartis sur les VLAN de l'organisation.

```
# Configuration OpenCanary (/etc/opencanaryd/opencanary.conf)  
  
{  
  "device.node_id": "opencanary-prod-01",  
  "ssh.enabled": true,  
  "ssh.port": 22,  
  "rdp.enabled": true,  
  "smb.enabled": true,  
  "mysql.enabled": true,  
  "mysql.port": 3306,  
  "logger": {  
    "class": "PyLogger",  
    "kwargs": {"handlers": {"SIEM": {"class": "opencanary.logger.SocketHandler",  
    "host": "siem.interne", "port": 514}}}  
  }  
}
```

Canarytokens : Leurres Légers à Déploiement Massif

Canarytokens (Thinkst, disponible sur canarytokens.org) sont des marqueurs numériques intelligents qui alertent dès qu'ils sont accédés ou utilisés. La liste des types de Canarytokens disponibles illustre la diversité des cas d'usage :

DNS Token : Un sous-domaine unique qui alerte dès qu'un serveur DNS le résout — utile pour détecter des outils de reconnaissance comme nmap ou les résolveurs DNS des scanners automatiques.

HTTP Token : Une URL qui alerte dès qu'elle est accédée — placée dans un fichier de configuration factice, un "mot de passe" dans un document partagé, ou un fichier robots.txt.

Word/PDF Document : Un document Office ou PDF qui alerte dès qu'il est ouvert — parfait pour des "dossiers confidentiels RH" ou "fichiers de mots de passe" placés dans des répertoires réseau partagés.

AWS API Key : De fausses clés d'accès AWS qui alertent dès que quelqu'un tente de les utiliser avec l'API AWS — détection immédiate de vol de credentials AWS dans des repositories ou fichiers de configuration.

Azure Login Certificate : Certificat factice qui alerte dès que quelqu'un tente de s'authentifier avec.

Cloned Website : Token détectant quand un site est cloné pour du phishing — particulièrement utile pour les portails d'entreprise et les pages d'authentification.

À RETENIR

À retenir — Stratégie Canarytokens pour PME

Déployer des Canarytokens Word/Excel dans chaque partage réseau sensible (RH, Finance, Direction)

Placer de fausses clés AWS dans le dépôt Git de chaque projet utilisant AWS

Un fichier "passwords.xlsx" avec un Canarytoken dans chaque poste d'administrateur système

DNS Tokens dans les zones DNS internes pour détecter les scans de reconnaissance

Budget : 0€ (Canarytokens.org est gratuit) — ROI immédiat

Deception dans Active Directory : Objets Factices et Leurres d'Identité

Active Directory est la cible de prédilection des attaquants dans les environnements Windows entreprise — tout chemin vers un compte Domain Admin donne le contrôle complet de l'organisation. La déception dans Active Directory (AD Deception) crée des leurres qui attirent les attaquants qui explorent l'annuaire en quête de cibles privilégiées. Les techniques AD Deception incluent :

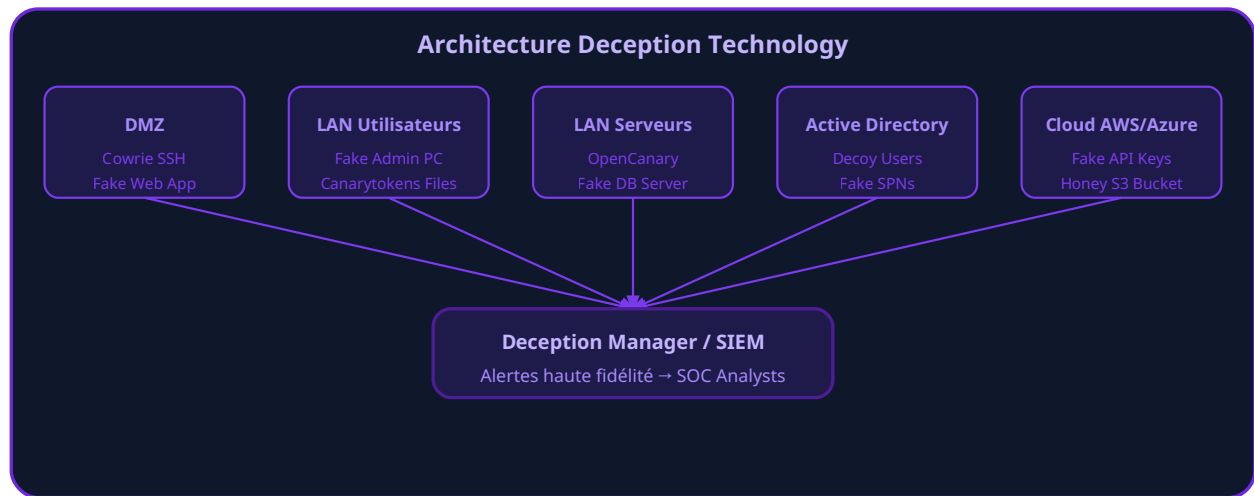
Decoy Users : Comptes utilisateurs factices avec des noms évocateurs ("admin-backup", "svc-sqlprod", "helpdesk-admin") qui n'ont aucun usage légitime — toute connexion sur ces comptes est une alerte immédiate de compromission et de mouvement latéral en cours.

Fake SPNs : Service Principal Names factices associés à des comptes avec des droits apparemment élevés pour attirer les attaquants pratiquant le Kerberoasting (tentative d'extraction de hashes de mots de passe de comptes de service). Un faux SPN déclenche une alerte lorsqu'un attaquant tente de requêter son ticket Kerberos.

Breadcrumbs : Fausses entrées dans des fichiers de configuration, historiques shell, ou fichiers récents qui pointent vers les honeypots — guidant subtilement les attaquants vers les leurres plutôt que vers les systèmes réels.

Decoy Shares : Partages réseau factices contenant des documents avec Canarytokens — la liste des partages est visible dans l'AD, ce qui attire naturellement les attaquants qui explorent les ressources réseau disponibles.

SVG : Architecture de Déception Distribuée



Plateformes Deception Enterprise : Attivo Networks et Illusive Networks

Pour les organisations nécessitant des capacités de déception à l'échelle enterprise, des plateformes commerciales ont été développées qui vont bien au-delà des honeypots open source :

Attivo Networks (intégré dans SentinelOne Singularity depuis 2022) : Déploiement automatique de leurres sur tout le réseau (endpoints factices, Active Directory leurres, accès cloud factices), corrélation des interactions avec les logs de SentinelOne pour attribution des actions à un endpoint compromis spécifique, et isolation automatique de l'endpoint compromis dès la première interaction avec un leurre.

Illusive Networks (maintenant Tenable Identity Exposure Deception) : Spécialisé dans la déception basée sur les identités et les credentials — crée des milliers de fausses identités et credentials dispersés dans le SI, déclenche des alertes à la microseconde dès qu'un attaquant tente d'utiliser un credential factice, et identifie précisément l'endpoint compromis qui a tenté l'utilisation.

Acalvio ShadowPlex : Plateforme autonome de déception qui utilise l'IA pour déployer des leurres adaptés à l'environnement spécifique de l'organisation — les leurres ressemblent

exactement aux vrais systèmes de l'organisation (mêmes noms, mêmes versions logicielles, même configuration apparente) pour maximiser leur crédibilité aux yeux des attaquants.

Honeypots Cloud : AWS, Azure et GCP

L'extension de la stratégie de déception aux environnements cloud est une nécessité avec la migration des SI vers le cloud hybride. Les honeypots cloud prennent plusieurs formes spécifiques aux environnements cloud :

Faux buckets S3 : Buckets AWS S3 avec des noms évocateurs ("company-backup-secret", "prod-database-dump") qui envoient des alertes CloudWatch dès qu'ils sont accédés — détection des attaquants qui ont compromis un accès cloud et explorent les ressources disponibles.

IAM Honeytokens : Fausses clés d'accès IAM AWS/Azure qui alertent immédiatement dès leur utilisation — elles peuvent être placées dans des dépôts Git (pour détecter les scraping automatiques de secrets), des fichiers de configuration factices, ou des documents internes.

Faux services API Gateway : Endpoints API factices documentés dans des spécifications OpenAPI internes — leur accès indique qu'un attaquant a accès à la documentation interne et tente d'explorer les APIs disponibles.

VMs honeypots cloud : Instances EC2 ou Azure VMs avec des configurations délibérément "attrayantes" (ports ouverts, services connus vulnérables en apparence) déployées dans des VPCs isolés du SI de production — elles ne devraient jamais être accédées depuis le réseau d'entreprise, tout accès est donc suspect.

Anecdote Terrain : Le Fichier Excel qui a Stoppé une APT

Un RSSI d'une organisation pharmaceutique française nous a rapporté le cas suivant survenu en 2024 : l'équipe sécurité avait déployé des Canarytokens dans plusieurs fichiers "confidentiel-recettes-produits.xlsx" et "liste-contacts-distributeurs.xlsx" dans les partages réseau des équipes

R&D. Six semaines après le déploiement, une alerte Canarytoken a été déclenchée à 2h43 du matin un mercredi — un fichier R&D avait été ouvert depuis une IP interne correspondant à un serveur de production normalement inaccessible depuis le réseau R&D. L'investigation a révélé une compromission initiale de ce serveur via une CVE non patchée dans l'application web exposée 3 mois auparavant, suivie d'une persistance discrète via une backdoor et d'une exploration lente du réseau depuis ce point. Sans le Canarytoken, la compromission aurait pu se poursuivre indéfiniment — le groupe ciblait vraisemblablement les données de brevets en cours de dépôt. Le Canarytoken avait été déclenché lors d'une phase de reconnaissance pour identifier les fichiers d'intérêt. Coût de la détection : 0€ (Canarytokens.org gratuit). Valeur des brevets potentiellement protégés : plusieurs dizaines de millions d'euros. C'est le retour sur investissement le plus éloquent de la cybersécurité défensive que nous ayons documenté.

FAQ Honeypots et Technologies de Déception

Est-il légal de piéger des attaquants avec des honeypots en France ?

La légalité des honeypots en France est souvent questionnée mais la réponse est clairement positive pour les honeypots passifs. Un honeypot qui simule des services et collecte des informations sur les attaquants qui tentent d'y accéder est légal — il ne s'agit pas d'un piège mais d'un système de détection. La question légale devient plus complexe pour les honeypots actifs qui attirent délibérément des attaquants (par exemple, en publiant de fausses vulnérabilités connues pour attirer des scanners automatiques). Les avocats spécialisés en droit informatique recommandent généralement de documenter clairement l'objectif défensif des honeypots, de les maintenir dans le périmètre du SI de l'organisation (pas de honeypots en dehors du réseau propre), et de ne jamais utiliser les informations collectées pour des contre-attaques. La CNIL ne s'est pas prononcée spécifiquement sur les honeypots mais les données collectées (IPs des attaquants, commandes exécutées) peuvent être considérées comme des données à caractère personnel si elles permettent l'identification de personnes physiques — leur conservation doit être justifiée et documentée.

Combien de honeypots déployer dans un réseau de 500 machines ?

La règle empirique généralement adoptée est de 1 honeypot ou leurre pour 50 à 100 actifs réels — soit pour un réseau de 500 machines, 5 à 10 honeypots réseau (OpenCanary/Cowrie) répartis dans les différents segments VLAN, complétés par plusieurs dizaines de Canarytokens déployés dans les partages réseau, postes d'administrateurs et dépôts Git. La répartition géographique dans les VLANs est importante : un honeypot dans le VLAN serveurs et un autre dans le VLAN utilisateurs permettent de détecter des patterns de mouvement latéral différents. Les plateformes entreprise (Attivo, Illusive) recommandent un taux de couverture encore plus élevé — jusqu'à 1 leurre pour 10 actifs réels — pour garantir qu'un attaquant ne peut pas naviguer longtemps dans le SI sans toucher un leurre.

Comment intégrer les alertes honeypot dans un SIEM ?

L'intégration des alertes honeypot dans le SIEM est essentielle pour contextualiser les alertes de déception avec les autres sources de sécurité. Cowrie et OpenCanary exportent nativement leurs logs en JSON via syslog ou des webhooks HTTP. Des connecteurs Sentinel et Splunk spécifiques existent pour Thinkst Canary (commercial). Pour les Canarytokens, l'URL de webhook peut pointer vers un endpoint du SOAR (Sentinel Logic Apps, Palo Alto XSOAR) qui crée automatiquement un ticket d'incident de priorité maximale. La règle de corrélation SIEM recommandée pour les alertes honeypot est simple : toute interaction avec un honeypot est une alerte P0 (urgence maximale), sans seuil, sans exception — le taux de faux positifs étant quasi-nul, chaque alerte mérite une investigation immédiate.

Quelle est la différence entre Thinkst Canary commercial et Canarytokens.org gratuit ?

[Thinkst Canary](#) est l'offre commerciale de Thinkst — des boîtiers physiques ou VMs qui simulent des dizaines de services réseau avec une gestion centralisée cloud, des alertes enrichies et une intégration SIEM native. Prix : ~5 000-10 000\$ par boîtier/an selon la configuration. Canarytokens.org est la version gratuite et open source qui se limite aux tokens numériques (pas de simulation de service réseau complet). Pour la plupart des organisations, Canarytokens gratuits + OpenCanary open source offrent 80% de la valeur de la solution commerciale pour 0€

de coût — la valeur commerciale de Thinkst Canary réside dans la gestion centralisée et le support pour les grandes déploiements multi-sites.

Les attaquants sophistiqués ne contournent-ils pas facilement les honeypots ?

Les attaquants très sophistiqués (groupes APT) utilisent effectivement des techniques de détection de honeypots : timing analysis (les honeypots répondent différemment des vrais systèmes), fingerprinting des émulateurs (Cowrie a des signatures reconnaissables par des outils spécialisés), et comportements prudents (éviter les systèmes "trop attrayants" qui peuvent être des leurres). Cependant, même face à des attaquants sophistiqués, les technologies de déception apportent de la valeur : elles ralentissent les attaquants (qui doivent consacrer du temps à identifier les vrais systèmes parmi les leurres), elles créent de l'incertitude (l'attaquant ne sait pas si le credential testé est réel ou factice), et elles peuvent détecter les outils de reconnaissance automatisés utilisés en phase initiale avant que l'attaquant ne passe à des techniques manuelles plus discrètes. Pour les attaquants moins sophistiqués (la grande majorité des incidents), les honeypots sont particulièrement efficaces.

Intégration Honeypots dans la Stratégie SOC

L'intégration des honeypots et des technologies de déception dans la stratégie SOC globale est ce qui maximise leur valeur. Isolément, un honeypot génère des alertes qui doivent être traitées manuellement. Intégré dans le SIEM et le [SOC](#), une alerte honeypot déclenche automatiquement via SOAR : la vérification de l'IP source dans les bases CTI, la recherche de l'IP dans les logs des 24 dernières heures (pour identifier les systèmes contactés par l'attaquant avant de toucher le honeypot), la création d'un ticket P0 dans l'ITSM, et la notification immédiate des analystes SOC disponibles. Cette automatisation garantit que les alertes honeypot — qui ont une valeur diagnostique maximale — reçoivent une attention prioritaire immédiate plutôt que d'être noyées dans le flux d'alertes standard du SOC. La [threat intelligence](#) enrichit les alertes honeypot avec le contexte de l'attaquant identifié — une IP honeypot connue comme appartenant à un botnet Mirai est moins préoccupante qu'une IP inconnue avec un ASN anormal tentant des connexions SSH

sur un honeypot du segment serveurs critiques. Cette hiérarchisation contextuelle des alertes honeypot est la clé d'une intégration SOC efficace des technologies de déception.

Mesurer l'Efficacité des Honeypots

Mesurer le retour sur investissement des honeypots est relativement simple comparé à d'autres technologies de sécurité, grâce à leur quasi-zéro taux de faux positifs. Les métriques clés incluent : le nombre d'interactions avec les honeypots (volume brut d'activité suspecte détectée), le ratio honeypot alerts / total SIEM alerts (mesure de la contribution relative de la déception dans la détection globale), le MTTD honeypot (délai moyen entre le premier contact avec un honeypot et la détection/réponse par le SOC), et la corrélation entre les alertes honeypot et les incidents confirmés (quelle proportion des alertes honeypot ont conduit à la découverte d'une compromission réelle ?). Ces métriques permettent de démontrer la valeur de l'investissement en déception au RSSI et au COMEX avec des données concrètes — le cas le plus convaincant étant naturellement le récit d'un incident détecté grâce à un honeypot que les autres outils de détection avaient manqué.

Opinion : La Déception Devrait Être Standard dans Tout SOC Mature

Je considère que les technologies de déception sont sous-utilisées dans les SOC français au regard de leur rapport coût/efficacité. Cowrie, OpenCanary et Canarytokens.org sont gratuits, s'installent en quelques heures, et offrent une détection à quasi-zéro faux positifs que les outils à millions d'euros de SIEM et d'EDR peinent à égaler. J'ai vu des organisations dépenser des centaines de milliers d'euros en licences SIEM et EDR qui génèrent des milliers d'alertes quotidiennes dont 95% sont des faux positifs — pendant que 10 Canarytokens bien placés dans leurs partages réseau auraient détecté l'intrusion en cours avec une certitude absolue. La résistance à adopter les honeypots vient souvent d'une incompréhension : "c'est compliqué à déployer" (faux pour les outils open source modernes), "ça ne servira à rien contre les attaquants

sophistiqués" (faux — même les APT touchent parfois des leurres lors de reconnaissance automatisée), ou "c'est légalement risqué" (faux pour les honeypots passifs correctement documentés). Ma recommandation ferme : toute organisation avec un SOC fonctionnel devrait avoir au minimum 20 Canarytokens déployés dans ses partages sensibles, 2-3 instances OpenCanary dans ses VLANs critiques, et des Decoy Users dans son Active Directory. Coût total : 0€, bénéfice potentiel : inestimable.

Honeypots OT/ICS : Protéger les Systèmes Industriels

Les honeypots dans les environnements OT/ICS (Operational Technology / Industrial Control Systems) sont particulièrement intéressants car ils permettent de détecter les attaquants ciblant les systèmes de contrôle industriels sans exposer les vrais automates programmables (PLC) ou les systèmes SCADA à des risques d'instabilité. Des honeypots OT spécialisés comme **Conpot** (open source, simule des protocoles Modbus, DNP3, IEC 61850, BACNET) permettent de détecter les scans de reconnaissance OT qui précèdent généralement les attaques sur les infrastructures critiques. Un honeypot Conpot exposé sur le réseau OT simulant un automate de contrôle d'une chaudière ou d'un système de ventilation attirera les attaquants qui explorent le réseau industriel à la recherche d'équipements à cibler. La détection précoce d'un attaquant en phase de reconnaissance OT — avant qu'il n'atteigne les vrais systèmes de contrôle — est particulièrement précieuse dans ces environnements où une compromission peut avoir des conséquences physiques et humaines graves, bien au-delà des impacts typiques d'un incident IT classique.

Comparatif des Solutions Honeypot : Open Source vs Commercial

Solution	Type	Coût	Points Forts
Cowrie	SSH/Telnet honeypot	Gratuit	Détail des sessions, collecte TTPs, Docker simple
OpenCanary	Multi-services honeypot	Gratuit	20+ services simulés, alertes webhook/syslog
Canarytokens	Tokens numériques	Gratuit	Déploiement en secondes, 25+ types de tokens
Thinkst Canary	Honeypot hardware/ VM	5-10k\$/boîtier/an	Gestion centrale, support, intégration SIEM native
Attivo (SentinelOne)	Deception Platform	Enterprise pricing	Déploiement automatisé, corrélation EDR, AD déception
Illusive (Tenable)	Identity Deception	Enterprise pricing	Déception par identités, délai détection ms, lateral movement

Intégration Active Directory : Détecter le Mouvement Latéral

La détection du mouvement latéral dans Active Directory est l'un des cas d'usage les plus précieux de la déception en entreprise. Après la compromission initiale d'un poste de travail, les attaquants explorent systématiquement le réseau pour identifier des comptes privilégiés et des serveurs critiques. La stratégie de déception AD pour détecter ce mouvement latéral comprend plusieurs couches complémentaires. Des **Decoy Computers** (objets machine AD avec des noms évocateurs comme "SRV-BACKUP-PROD" ou "DC-ADMIN") sont créés dans Active Directory — tout DNS lookup ou tentative de connexion sur ces noms inexistants déclenche une alerte. Des **Honey Credentials** stockées dans des fichiers de configuration, des scripts ou le gestionnaire de mots de passe du navigateur des postes d'administrateurs : quand un attaquant dump les credentials du poste compromis, ces faux credentials sont extraits et testés — déclenchant immédiatement une alerte lorsque la tentative d'authentification échoue sur le vrai DC mais est capturée par le honeypot. Ces techniques de déception AD, combinées aux logs

d'audit AD natifs dans le SIEM, créent une toile de détection dense que les attaquants en mouvement latéral ont du mal à éviter sans ralentir considérablement leur progression et révéler leur présence au SOC. La mise en œuvre d'une [segmentation réseau stricte](#) combinée à la déception maximise l'efficacité des deux approches : la segmentation limite les mouvements possibles, la déception détecte les tentatives de contournement de la segmentation.

Honeypots et Réponse aux Incidents

Les honeypots jouent un rôle précieux non seulement dans la détection proactive mais aussi dans le contexte d'une réponse à un incident en cours. Lorsqu'un incident est détecté par d'autres moyens (alerte EDR, Canarytoken), les logs des honeypots des dernières 24-48 heures fournissent un contexte précieux sur les activités de reconnaissance de l'attaquant avant qu'il n'atteigne ses cibles réelles : quels services réseau a-t-il sondé ? Quelles credentials a-t-il testées ? Quelles commandes a-t-il exécutées dans les honeypots SSH ? Ces informations permettent de comprendre les capacités et les intentions de l'attaquant, d'identifier les systèmes qu'il cherchait à atteindre (et donc ceux qui nécessitent une investigation prioritaire), et de collecter des IOCs supplémentaires pour enrichir la [threat intelligence](#). Dans certains cas, maintenir un attaquant occupé dans un honeypot (plutôt que de l'isoler immédiatement) peut permettre de collecter suffisamment d'informations pour l'attribution et pour identifier l'infrastructure C2 complète — une décision tactique qui doit être prise en concertation avec la direction et les équipes juridiques. La [coordination avec le SOC](#) est essentielle pour que les alertes honeypot déclenchent les bons workflows d'investigation et de réponse dans les délais appropriés.

Formation et Sensibilisation via les Honeypots

Au-delà de leur fonction de détection, les honeypots peuvent être utilisés comme outils de formation pour les équipes SOC et pour sensibiliser les équipes IT aux techniques des attaquants.

En analysant les logs Cowrie collectés sur des honeypots exposés sur Internet, les analystes SOC peuvent observer en conditions réelles les techniques de compromission automatisée : séquences de commandes post-exploitation, tentatives d'installation de malwares, techniques de persistance — le tout dans un environnement contrôlé qui n'affecte pas la production. Des "honeypot labs" internes peuvent être créés spécifiquement pour la formation, permettant aux équipes de detection engineering de tester leurs nouvelles règles de détection contre des comportements d'attaque réels (rejoués depuis les logs Cowrie ou simulés manuellement). Cette approche de formation "hands-on" sur des logs d'attaques réelles est bien plus efficace que les formations théoriques pour développer l'intuition des analystes sur les comportements d'attaquants et améliorer leur capacité à distinguer les vrais incidents des faux positifs dans les logs de production. L'investissement dans la formation via les honeypots est un multiplicateur de force pour le SOC qui dépasse largement le simple coût de déploiement des leurres eux-mêmes.

Honeypots Email et Phishing : Détecter les Campagnes Ciblées

Les honeypots email sont une approche de déception méconnue mais efficace pour détecter les campagnes de phishing et d'ingénierie sociale ciblant l'organisation. Des adresses email "honeypot" sont créées et placées dans des endroits visibles publiquement (pages web, documents publics, annuaires d'entreprise) mais jamais utilisées pour des communications légitimes : tout email reçu sur ces adresses est nécessairement du spam ou du phishing. Cette technique permet de détecter précocement des campagnes de phishing ciblant les domaines de l'organisation, d'identifier les techniques d'ingénierie sociale utilisées (prétextes, tons urgents, usurpation d'identité), et de collecter des indicateurs de compromission email (domaines expéditeurs, infrastructure de phishing, payloads de malware en pièces jointes). Les **Honey Domains** vont encore plus loin : des domaines similaires au domaine de l'organisation (typosquatting délibéré comme "acme-corp-rh.fr" pour un attaquant qui aurait envisagé d'usurper l'identité RH d'acme-corp) sont enregistrés et leurs serveurs email configurés pour collecter tous les emails reçus. Ces emails révèlent les tentatives d'usurpation d'identité de l'organisation vers ses clients ou partenaires, ainsi que les communications internes interceptées si des employés ont été poussés à utiliser ces faux domaines par inadvertance lors d'une campagne d'ingénierie

sociale. La mise en place de ces honeypots email est peu coûteuse (quelques euros par an pour l'enregistrement des domaines) et peut révéler des campagnes de fraude qui n'auraient autrement été détectées qu'après des pertes financières significatives.

Métriques de Détection : Comparer Honeypots et SIEM

Pour comprendre la complémentarité des honeypots et du SIEM traditionnel, il est utile de comparer leurs métriques de détection :

Taux de faux positifs : SIEM bien tuné : 80-95% de faux positifs ; Honeypots : <1% (toute interaction est suspecte par définition)

Couverture des techniques ATT&CK : SIEM : dépend des règles déployées, typiquement 30-60% des techniques ; Honeypots : détectent surtout reconnaissance et mouvement latéral (TA0043, TA0008)

Délai de détection : SIEM : quelques minutes à quelques heures selon la complexité de la corrélation ; Honeypots : immédiat dès la première interaction

Contexte de l'alerte : SIEM : alerte riche mais parmi des milliers d'autres ; Honeypot : alerte rare mais d'une pertinence maximale

Coût opérationnel : SIEM : élevé (ingestion, storage, licences, maintenance des règles) ; Honeypots open source : quasi-nul une fois déployés

Cette comparaison illustre pourquoi les honeypots ne remplacent pas le SIEM mais le complètent de manière idéale : le SIEM offre une couverture large avec une fidélité variable, les honeypots offrent une couverture étroite (le mouvement latéral et la reconnaissance) avec une fidélité maximale. Les alertes honeypot doivent automatiquement déclencher un enrichissement SIEM pour obtenir le contexte complet de l'attaque dans les logs généraux de l'organisation.

Honeypots DNS : Détecter la Reconnaissance par DNS

La reconnaissance DNS est souvent la première étape d'une attaque — les attaquants tentent de résoudre des noms d'hôtes internes pour mapper l'infrastructure avant de lancer des exploits. Les **honeypots DNS** exploitent cette technique pour détecter la reconnaissance précoce. Un serveur DNS interne configuré pour journaliser les requêtes de résolution de noms d'hôtes factices (ajoutés dans le DNS interne mais ne correspondant à aucun service réel) permet de détecter les outils de reconnaissance comme BloodHound (qui résout massivement les noms d'ordinateurs AD), les scripts de découverte réseau post-exploitation comme Sharphound, ou les scanners automatiques de reconnaissance qui testent des noms d'hôtes communs. L'intégration de ces logs DNS honeypot dans le SIEM via le connecteur Windows DNS ou le module Suricata/Zeek permet une corrélation avec les autres activités suspectes observées sur le réseau, enrichissant considérablement le contexte d'investigation lors d'un incident. Des solutions comme **Responder** (initialement outil offensif, utilisable défensivement pour monitorer les requêtes LLMNR/NBNS) permettent de détecter les outils d'attaque qui "poisonnent" ces protocoles pour intercepter des credentials — une détection précieuse du mouvement latéral en cours dans les environnements Windows.

Planification d'un Programme de Déception : Roadmap

La mise en place d'un programme de déception mature se fait progressivement selon une roadmap structurée :

Phase 1 - Quick Wins (1-4 semaines, budget 0€) : Déploiement de 20-50 Canarytokens Word/Excel dans les partages réseau sensibles, ajout de faux credentials dans les historiques shell des serveurs d'administration, intégration des alertes Canarytokens dans le SIEM via webhook. ROI immédiat sans coût.

Phase 2 - Honeypots Réseau (1-3 mois, budget <5k€) : Déploiement d'instances OpenCanary dans chaque VLAN critique (serveurs, OT, gestion), Cowrie SSH sur le périmètre Internet, intégration dans le SIEM, création de runbooks de réponse aux alertes honeypot.

Phase 3 - AD Deception (3-6 mois, budget 10-30k€) : Création de Decoy Users, faux SPNs, Breadcrumbs et Honey Shares dans Active Directory. Intégration avec les logs AD Audit dans le SIEM pour corrélation.

Phase 4 - Cloud Deception (6-12 mois, budget variable) : Honey buckets S3/Azure Blob, faux API Keys AWS/Azure dans les dépôts Git et fichiers de configuration, Honey Domains enregistrés.

Phase 5 - Plateforme Enterprise (12-24 mois, budget >100k€) : Évaluation et déploiement d'une plateforme commerciale (Thinkst Canary, Attivo/SentinelOne) pour une gestion centralisée et une couverture automatisée à grande échelle.

Honeypots et Purple Teaming

Le **Purple Teaming** — collaboration entre équipes Red (attaque) et Blue (défense) pour améliorer conjointement les capacités de détection et de réponse — intègre naturellement les honeypots comme outil de mesure de l'efficacité de la déception. Dans un exercice Purple Team, l'équipe Red exécute des techniques d'attaque documentées (reconnaissance réseau, Kerberoasting, mouvement latéral via SMB) dans l'environnement de test, et l'équipe Blue évalue quelles alertes ont été déclenchées — par les outils de détection traditionnels (EDR, SIEM) mais aussi par les honeypots. Les honeypots touchés lors de l'exercice Red Team indiquent que la stratégie de déception est crédible et que les leurres sont suffisamment intégrés dans l'environnement pour attirer les attaquants simulés. Les honeypots NON touchés lors de l'exercice révèlent soit que les leurres sont trop évidents (l'équipe Red les a identifiés et évités), soit qu'ils ne sont pas placés dans les chemins d'attaque naturels. Ces enseignements permettent d'optimiser le placement et la crédibilité des honeypots pour maximiser leur efficacité contre des attaquants réels utilisant les mêmes techniques que l'équipe Red. La combinaison Purple Teaming + honeypots crée une boucle d'amélioration continue de la capacité de déception qui converge vers une posture de détection de plus en plus difficile à contourner pour un attaquant sophistiqué.

Ressources et Communauté Honeypot

L'écosystème des honeypots bénéficie d'une communauté active qui produit des outils, des données et des connaissances précieuses :

The HoneyNet Project (honeynet.org) : Organisation internationale de recherche sur les honeypots, produisant des outils open source (Cowrie, Conpot, Dionaea) et des rapports d'analyse des menaces collectées via des honeypots mondiaux.

SANS Internet Storm Center : Utilise un réseau mondial de honeypots pour suivre les tendances des attaques en temps réel — les dashboards ISC montrent quels ports sont les plus scannés aujourd'hui dans le monde, utile pour configurer les honeypots les plus pertinents.

GreyNoise : Service qui analyse le trafic vers des millions de honeypots pour distinguer le bruit de fond Internet (scans automatiques légitimes) du trafic malveillant ciblé — permet d'enrichir les alertes honeypot avec le contexte de l'origine du trafic.

Honeypots et Threat Intelligence : Feedback Loop

Les honeypots sont une source de threat intelligence de première main particulièrement précieuse car ils collectent des IOCs issus d'attaques réelles ciblant l'organisation ou son secteur. Les IPs sources des connexions honeypot, les credentials testés, les commandes exécutées et les malwares téléchargés constituent des IOCs authentiques que l'organisation peut partager avec sa communauté de sécurité via MISP ou un ISAC sectoriel. Cette contribution à la communauté est bidirectionnelle : les IOCs collectés par les honeypots de l'organisation enrichissent la base commune, et les IOCs partagés par d'autres organisations via MISP permettent à l'organisation de contextualiser les interactions observées sur ses propres honeypots. Une IP qui apparaît simultanément dans les logs honeypot de l'organisation et dans les feeds CTI MISP comme appartenant à un groupe APT connu est une alerte d'une pertinence maximale qui justifie une investigation immédiate. Cette boucle de feedback honeypot vers CTI vers honeypot est un exemple concret de la synergie entre les différentes capacités de sécurité d'un programme de cyberdéfense mature. Les données honeypot peuvent également être utilisées pour mesurer

l'évolution de la menace dans le temps — une augmentation soudaine des connexions SSH sur un honeypot peut indiquer une nouvelle campagne de scanning ciblant les organisations de la région ou du secteur, permettant d'alerter les équipes IT pour durcir la surveillance et accélérer les patches sur les systèmes SSH-exposés. La mise en oeuvre de ce processus d'analyse et de partage des données honeypot transforme les outils de détection passifs en une source active de renseignement sur les menaces.

Positionnement Stratégique des Honeypots dans le SI

Le placement optimal des honeypots dans l'infrastructure est une discipline qui mêle des principes techniques et une connaissance approfondie des comportements d'attaquants. Les principes directeurs incluent : placer les honeypots dans des segments de réseau que les attaquants traversent lors de mouvements latéraux post-compromission initiale (VLANs de gestion, accès AD, accès aux serveurs de fichiers) ; s'assurer que les honeypots sont découvrables par des techniques de reconnaissance standard (présence dans les résultats nmap, dans l'annuaire AD, dans les listes de partages) sans être trivialement identifiables comme des leurres ; varier les types de honeypots et leur emplacement pour couvrir différentes phases de l'attaque et différents vecteurs (réseau, email, cloud, identités) ; et réviser régulièrement le positionnement en se basant sur les exercices Red Team et les incidents réels. Un programme de déception qui évolue en permanence en réponse aux nouvelles informations sur les menaces est significativement plus efficace qu'un déploiement statique jamais révisé. La combinaison d'une analyse MITRE ATT&CK des gaps de couverture de la déception avec les résultats des exercices Purple Team permet d'identifier les chemins d'attaque non couverts et d'adapter le déploiement pour maximiser la probabilité de détection d'un attaquant réel avant qu'il n'atteigne ses objectifs dans le SI de l'organisation. Le programme de déception le plus mature est celui qui couvre les cinq phases d'une attaque — de la reconnaissance initiale à l'exfiltration finale — avec des leurres adaptés à chaque phase, intégrés dans le SIEM pour une corrélation automatique, et alimentant la CTI pour enrichir continuellement la base de connaissances de l'organisation sur les menaces qu'elle affronte.

Synthèse : Programme de Déception Minimal Viable

Pour une organisation française qui commence sa stratégie de déception en 2026, le programme de déception minimal viable (MVP) est accessible immédiatement et gratuitement. Le MVP comprend : cinq à dix Canarytokens de type document Word/Excel nommés de manière attrayante (liste-mots-de-passe-admin.xlsx, contrats-confidentiels-2026.docx, plan-sauvegarde-serveurs.docx) déposés dans chaque répertoire réseau partagé sensible des équipes Finance, RH et Direction ; deux ou trois Canarytokens de type AWS Keys intégrés dans les fichiers de configuration des projets utilisant AWS ou dans des variables d'environnement factices sur les serveurs d'administration ; une instance OpenCanary déployée sur une petite machine virtuelle dans le VLAN serveurs, configurant au minimum SSH, RDP, et MySQL pour simuler un serveur d'administration factice ; et l'intégration de toutes ces alertes dans le SIEM via les webhooks Canarytokens et les logs syslog OpenCanary avec une règle P0 dans le SIEM déclenchant une notification Slack ou Teams immédiate pour tout analyste SOC d'astreinte. Ce MVP prend une journée à deployer, coûte zéro euros en licences, et peut détecter les compromissions et les mouvements latéraux que des outils à six chiffres auraient manqués. C'est la version la plus pure du principe d'efficacité en cybersécurité défensive — maximiser la valeur de détection par rapport à l'investissement consenti, en exploitant le désavantage fondamental de l'attaquant qui ne sait pas où sont les leurres dans un environnement qu'il ne connaît pas aussi bien que ses défenseurs.

Les technologies de déception représentent l'une des évolutions les plus prometteuses de la cyberdéfense moderne — elles transforment la posture défensive passive en une posture active où l'attaquant est contraint de naviguer dans un environnement truffé de leurres conçus pour le piéger. Les organisations françaises qui investiront dans ces technologies en 2026, même avec un budget minimal, disposeront d'un avantage de détection significatif sur celles qui restent exclusivement dépendantes des outils de surveillance passive traditionnels.