

Préparer vos réseaux à la « tempête de l'IA »

Durcir votre réseau contre les
attaques accélérées par l'IA

© 2026 Cloud Security Alliance – Tous droits réservés. Vous pouvez télécharger, enregistrer, afficher sur votre ordinateur, consulter, imprimer et créer un lien vers la Cloud Security Alliance à l'adresse <https://cloudsecurityalliance.org>, sous réserve des conditions suivantes : a) le document préliminaire ne peut être utilisé qu'à des fins personnelles, informatives et non commerciales; b) il ne peut être modifié ni altéré de quelque manière que ce soit; c) il ne peut être redistribué; et d) toute mention de marque de commerce, de droit d'auteur ou autre ne peut être supprimée. Vous pouvez citer des passages du document préliminaire dans la mesure permise par les dispositions relatives à l'usage équitable de la loi des États-Unis sur le droit d'auteur, à condition d'en attribuer les passages à la Cloud Security Alliance.

Remerciements

Auteur

Rich Mogull

Conception graphique

Stephen Smith

À propos du commanditaire

Cisco (NASDAQ : CSCO) est le chef de file mondial des technologies qui transforme la façon dont les organisations se connectent et se protègent à l'ère de l'IA. Depuis plus de 40 ans, Cisco connecte le monde en toute sécurité. Grâce à ses solutions et à ses services alimentés par l'IA à la fine pointe de l'industrie, Cisco permet à ses clients, à ses partenaires et à ses communautés d'exploiter les innovations, d'accroître leur productivité et de renforcer leur résilience numérique.

<https://www.cisco.com/site/ca/fr/solutions/artificial-intelligence/security-in-ai-era.html>



Table des matières

Remerciements.....	3
Introduction.....	5
Moderniser les réseaux pour une défense réactive.....	6
Les fondements et le processus	7
Phase 1 : remplacement et durcissement prioritaires et ciblés	10
Phase 2 : ajouter des frontières prioritaires et ciblées	12
Phase 3 : élargir et affiner.....	14
Une transition difficile, un avenir prometteur	15

Introduction

Les avancées des modèles d'IA transforment la relation entre les attaquants et les défenseurs, et remettent en cause les hypothèses historiques de risque sur lesquelles nous avons bâti nos programmes de sécurité. Comme l'ont montré la sortie de modèles aux capacités avancées, comme Mythos et ChatGPT 5.5 (et comme l'explique [le document « The AI Vulnerability Storm »](#)), les modèles évoluent rapidement et améliorent leur capacité à découvrir de nouvelles vulnérabilités, à créer de nouveaux exploits et à mener des attaques automatisées complexes en plusieurs étapes. Les dernières années nous ont montré que les capacités des modèles de frontières actuels seront rapidement égalées par des modèles fondamentaux largement accessibles, puis éventuellement par des modèles locaux à poids ouverts.

Les modèles de risque de base sur lesquels nous avons bâti nos programmes de sécurité ne tiennent plus. Découvrir de nouvelles vulnérabilités et créer de nouveaux exploits prenait autrefois des mois, voire des années, mais cela se fait maintenant en quelques heures. Ces découvertes, et leur exploitation dans des attaques, exigeaient autrefois des compétences techniques poussées; ces barrières disparaissent rapidement. L'économie de la cybersécurité change : à mesure que le temps et les compétences nécessaires à l'offensive diminuent, nos délais défensifs historiques ne sont plus efficaces et nous devons réorienter nos programmes.

L'IA soutient à la fois les attaquants et les défenseurs, mais pas de manière équivalente. Comme l'explique [Core Collapse](#), il existe une asymétrie mathématique fondamentale : les attaquants peuvent concentrer tous leurs efforts sur un seul problème (trouver une vulnérabilité *ici*, développer un exploit pour *celle-ci* ou attaquer *cette* cible), tandis que les défenseurs doivent composer avec la complexité combinatoire de devoir protéger toutes les ressources contre tous les attaquants, contre tous les exploits possibles, pour toujours.

Cette asymétrie se constate aisément lorsqu'on examine l'incidence d'une seule vulnérabilité nouvellement divulguée. Les attaquants peuvent utiliser l'IA pour développer rapidement un exploit et lancer des attaques automatisées. Les défenseurs, même avec un correctif en main, doivent identifier chaque ressource à corriger, la tester, puis la déployer sans perturber les applications en cours d'exécution. C'est une course qui ne se joue pas selon les mêmes règles, peu importe notre niveau de compétence.

Il existe des stratégies qui permettent déjà de rééquilibrer la situation. Avec le temps, à mesure que tous les nouveaux logiciels sont évalués et durcis avant leur mise en production à l'aide de grands modèles de langage, nous pouvons nous attendre à une baisse globale du nombre de vulnérabilités dans les logiciels. Et, en tant que défenseurs aujourd'hui, nous pouvons aussi ajouter des frontières de sécurité *pour accroître la complexité combinatoire des attaques réussies*.

Le renforcement de la résilience face aux attaquants propulsés par l'IA touche tous les aspects de la technologie et dépasse largement ce que les équipes de sécurité peuvent gérer à elles seules. Nous devons accroître les barrières de sécurité dans tous les domaines, des identités et des charges de travail aux conteneurs, aux API et aux points de terminaison.

Les réseaux sont l'un de nos outils les plus importants pour faire respecter ces frontières, mais les modèles de sécurité, de conception et d'exploitation des réseaux réels en production sont souvent insuffisants. Malgré des années d'avancement et d'investissement, nous avons encore trop de réseaux plats, nous comptons souvent sur du matériel désuet, et nos processus d'exploitation et de sécurité réseau ne sont pas prêts à réagir rapidement ni à

s'automatiser. Une mesure aussi simple que la modification d'une règle de pare-feu peut prendre des semaines, voire des mois, dans les organisations aux réseaux complexes.

Pour accroître la complexité, et les coûts, imposés aux attaquants, nos réseaux devraient :

- accélérer rapidement la cadence des correctifs et viser un correctif continu;
- réduire la surface d'attaque en améliorant la segmentation du réseau et en limitant les emplacements où les données à forte valeur sont exposées;
- ajouter plusieurs couches défensives *en série, plutôt qu'en parallèle*, afin d'obliger les attaquants à franchir plusieurs frontières; par exemple, plusieurs pare-feu à différentes couches à l'intérieur d'une même pile applicative;
- être soutenus par des processus opérationnels réactifs capables de s'adapter aussi rapidement que l'environnement de menace évolue.

L'utilisation de capacités d'IA avancées à des fins défensives permettra éventuellement aux organisations à l'aise avec l'IA d'atteindre un rythme suffisamment compétitif pour résister à la vitesse à laquelle les attaquants développent des exploits. D'ici là, les frontières de sécurité seront l'outil qui nous permettra de ralentir les attaques et de gagner le temps nécessaire pour mettre à jour les déploiements.

Le présent document fournit des indications aux administrateurs réseau et aux administrateurs de la sécurité réseau. Nous ne pouvons pas répondre à nos nouvelles exigences de résilience uniquement au moyen des défenses réseau, mais le renforcement des frontières de sécurité réseau est l'un des points de départ les plus efficaces.

Moderniser les réseaux pour une défense réactive

À l'échelle de l'industrie, nous modernisons nos réseaux depuis plus d'une décennie, en allant du renforcement des défenses périmétriques à l'accroissement de la segmentation réseau, et cette tendance s'est accélérée ces dernières années. La migration vers le nuage nous a poussés encore plus loin, puisque les réseaux infonuagiques sont, par défaut, définis par logiciel, et que les réseaux définis par logiciel facilitent grandement la création et la modification des frontières et de la segmentation.

Cependant, nous avons dû faire face à de réels défis pour mettre à jour ces réseaux à grande échelle. Beaucoup de réseaux physiques, sinon la plupart, demeurent relativement plats à l'intérieur. Même les réseaux infonuagiques, pourtant très souples, faciles à segmenter et toujours définis par logiciel, ne sont pas toujours meilleurs, principalement en raison des migrations de type lift-and-shift qui ont reproduit les anciennes conceptions de centres de données. Les mises en œuvre du modèle Zero Trust se sont trop souvent concentrées sur la protection du personnel et du bord du campus, et moins sur le trafic entre les applications et les centres de données (est-ouest).

Nous ne possédons pas et ne gérons pas nécessairement nos propres réseaux, et nous faisons face à des limites en matière d'observabilité de tiers.

De nombreuses organisations ont commencé à mettre en œuvre la microsegmentation, mais généralement seulement pour de petits sous-ensembles de leur réseau, en raison de la complexité de la mise en œuvre. Une grande partie de notre matériel est plus ancien, et difficile et lent à mettre à jour. Et nos processus opérationnels sont délibérément méthodiques : à l'extérieur du périmètre, nous n'avons jamais vraiment eu besoin de corriger rapidement le réseau interne, donc nos processus étaient axés sur le bord, tandis que les correctifs internes progressaient beaucoup plus lentement.

Rien de tout cela n'était faux; c'était une réponse raisonnable aux risques auxquels nous étions confrontés. La plupart de nos recommandations dans ce document sont bien connues et déjà en cours de mise en œuvre. Le problème, c'est que l'IA change les risques et les échéances. Ce n'est pas seulement le périmètre, mais aussi les barrières internes de notre réseau qui deviennent une expansion critique des frontières de sécurité, essentielle pour freiner l'avantage que l'IA confère aux attaquants. Les réseaux plats signifient une surface d'attaque plus grande lorsqu'un système est compromis. Les changements lents et méthodiques ne peuvent pas suivre le rythme du développement des exploits, qui se fait maintenant en quelques heures. Et un matériel que vous ne pouvez pas corriger rapidement, même s'il se trouve derrière les pare-feu, devient une responsabilité beaucoup plus importante.

La bonne nouvelle, c'est que nous savons déjà quoi faire; le nouveau défi, c'est l'exécution. Nous devons avancer rapidement, car les hypothèses de risque sous-jacentes que nous avons utilisées pour justifier nos priorités et nos échéances historiques ne s'appliquent plus. Nos fondements en sécurité réseau (segmentation, possibilité de correction et opérations réactives) peuvent être très efficaces pour rééquilibrer la situation, mais nous devons maintenant exécuter avec plus de granularité, selon un ordre priorisé et dans des délais serrés.

L'ingénierie et les opérations réseau, d'une part, et la sécurité réseau, d'autre part, doivent aussi travailler main dans la main pour régler ces problèmes. Les solutions touchent plusieurs silos et exigent une combinaison d'architectures fondamentales, de correctifs et de gestion des vulnérabilités, ainsi que de défenses spécialisées comme les NGFW et les pare-feu d'applications Web. Le reste du présent document présente les éléments de la modernisation, puis la façon d'y parvenir par étapes.

Les fondements et le processus

La modernisation du réseau comporte trois dimensions pour répondre aux risques adverses liés à l'IA :

- **Modèle d'exploitation et processus** : puisque nous avons commencé par tirer des câbles vers des boîtiers et à défendre des adresses IP peu changeantes, les processus réseau et de sécurité réseau ont tendance à être méthodiques, surtout pour les mises à jour matérielles et les changements structurels. Ces processus ont été conçus pour des environnements à dépendances physiques. Les correctifs étaient surtout appliqués au bord du réseau, le matériel interne étant jugé secondaire puisqu'il présentait moins de risques. La plupart des réseaux mettaient en place une ou deux couches de défense, axées sur le trafic nord-sud afin de se protéger des menaces provenant d'Internet. Le Cloud est plus dynamique, mais de façon inégale, selon la maturité, la dépendance aux migrations de type lift-and-shift et les dépendances hybrides. Le virage consiste à passer de *projets périodiques* à *des opérations continues*. Pas seulement pour les correctifs et

les remplacements, mais aussi pour la segmentation elle-même : il faut ajouter et maintenir des frontières au fur et à mesure que nous bâtissons, plutôt que de les concevoir comme une architecture en une seule étape. Voyez cela comme deux boucles qui fonctionnent en parallèle :

- **Garder le réseau corrigé et à jour** : l'objectif est un correctif continu, combiné à une cadence de remplacement priorisée, afin de ne jamais avoir à défendre quelque chose que l'on ne peut pas corriger. Historiquement, cela a été extrêmement difficile pour les équipements réseau, ce qui explique pourquoi nous recommandons de traiter cet enjeu comme un nouveau processus opérationnel et de commencer par les ressources exposées à l'extérieur. Cela peut et doit être intégré à votre programme *VulnOps*. Décrit plus en détail dans le rapport *Vulnerability Storm*, *VulnOps* traite la gestion des vulnérabilités comme un programme de type DevOps avec automatisation continue.
- **Mettre à jour la segmentation et l'isolation** : aller vers la microsegmentation et ajouter en continu plusieurs frontières, nord, sud, est et ouest, à mesure que les applications changent, en combinant topologie et défenses de sécurité.
- **Infrastructure** : le réseau lui-même est une cible, et nous nous attendons à voir augmenter les volumes de vulnérabilités et d'exploits visant le matériel et les logiciels réseau utilisés dans des attaques automatisées. C'est déjà le cas aujourd'hui; au cours des semaines où nous avons rédigé ce document, plusieurs campagnes ont [ciblé l'infrastructure réseau](#). Les réseaux devraient être :
 - **À jour** : du matériel actualisé qui n'est pas en fin de vie et qui est encore pris en charge.
 - **Corrigeable** : non seulement du matériel actualisé, mais aussi intégré à un programme de correctifs capable de déployer des mises à jour en quelques heures, et non en quelques semaines.
 - **Défini par logiciel** : comme le démontrent si bien les réseaux infonuagiques, les réseaux définis par logiciel (SRST) sont plus agiles et souvent à accès refusé par défaut (selon la plateforme, par exemple un [périmètre défini par logiciel, ou SDP](#)), la segmentation devenant l'état normal plutôt qu'un ajout.
- **Architecture** : l'architecture définit les frontières qui augmentent la complexité et les coûts pour les attaquants. Nous ne pouvons pas nécessairement passer instantanément à un modèle où chaque actif serait protégé par un pare-feu à chaque niveau, mais nous nous appuyons sur ces fondations :
 - **Sécurité périmétrique** : c'est ce que nous maîtrisons généralement le mieux aujourd'hui, et ce besoin ne disparaît pas. Mais, comme vous le verrez, nous devons l'aligner sur notre modèle d'exploitation plus continu, et dans certains cas, nous pouvons vouloir une plus grande diversité de plateformes ainsi qu'une segmentation périmétrique accrue.
 - **Macrosegmentation** : meilleure segmentation des unités d'affaires, des piles applicatives et des environnements d'exploitation. L'objectif est de limiter la surface d'attaque à une seule « pile » au sein du réseau.
 - **Microsegmentation** : ajout de frontières à l'intérieur des couches d'une pile applicative, tout en augmentant aussi les frontières entre les applications et les services.
 - **Zero Trust** : combine ces frontières avec des politiques fondées sur l'identité et le contexte, vérifiées en continu, de sorte que l'accès est accordé à la demande plutôt qu'en fonction de l'emplacement réseau. La confiance cesse d'être une fonction d'une adresse IP ou d'un VLAN et devient dynamique, fondée sur le moindre privilège, et liée à l'identité et au risque plutôt qu'à la topologie. Cela comprend des technologies comme le [périmètre défini par logiciel \(SDP\)](#).

Au cœur de tout cela se trouvent la visibilité, que possédons-nous, où et comment est-ce configuré, l'attribution, qui possède quoi et qui est responsable des décisions d'affaires et techniques, et, idéalement, la télémétrie en temps réel. *Nous ne sous-estimons pas la complexité liée à la mise en œuvre de ces recommandations*, mais les technologies modernes rendent cela plus réalisable, surtout à l'échelle d'un projet.

Nous pouvons aussi utiliser nous-mêmes les outils d'IA pour effectuer des tests et des améliorations en continu. Appliquez vos propres LLM à votre réseau, fournissez-leur des outils d'évaluation sûrs et utilisez les résultats pour voir vos forces et vos faiblesses du point de vue d'un attaquant.

Tout cela prend du temps, parfois des années, et les attaques adverses propulsées par l'IA ont déjà commencé. Nous recommandons une approche par phases, priorisée et fondée sur le risque, qui peut plus réalistement produire rapidement des résultats tangibles tout en jetant de bonnes bases pour atteindre les objectifs à long terme.

Point important : les phases suivantes peuvent se dérouler en parallèle, dans le cadre d'un processus opérationnel continu. Il n'est pas nécessaire de terminer complètement une phase avant de passer à la suivante.

Phase 1 : remplacement et durcissement prioritaires et ciblés

Les attaquants commenceront par ce qu'ils peuvent atteindre, et, en tant que défenseurs, nous nous soucions surtout de ce qui nous fera le plus mal s'il est compromis. La première étape consiste donc à cerner votre surface d'attaque externe et à l'aligner sur les priorités de votre organisation.

Pour cette phase, pensez « de l'extérieur vers l'intérieur, en commençant par les actifs clés ». Commencez par déterminer votre surface d'attaque exposée à Internet, tout en identifiant vos actifs numériques les plus importants, les « joyaux de la couronne ». Lorsque nous parlons de surface d'attaque, nous parlons de *tout*. Chaque appareil et chaque service, *y compris tous vos appareils et services de sécurité et réseau*. Tout ce qui touche à Internet et qui peut potentiellement être joint de l'extérieur. N'oubliez pas des actifs comme les serveurs VPN, les bureaux distants, les routeurs de bordure et les appliances de sécurité virtuelles infonuagiques.

Les actifs numériques critiques, ou « surface à protéger », sont les applications et services dont la compromission dans une violation de données ou une attaque destructrice nuirait le plus à l'entreprise. Pour ces services, il vous faut une carte complète de la connectivité et des flux de données. Ce sont souvent des questions d'affaires auxquelles les équipes réseau et sécurité ne peuvent pas répondre seules.

Ensuite, faites la corrélation croisée. Pour les actifs critiques visés, quels sont, de l'extérieur vers l'intérieur, les dispositifs réseau qui se trouvent sur le chemin? Pare-feu, pare-feu d'applications Web, routeurs, commutateurs, VPN, même ceux qui ne sont pas évidents, tout y passe. Utilisez ce tableau pour déterminer vos prochaines actions :

État	Action recommandée
Fin de vie d'appareil (« EOL », pour « End of Life »)	Renouveler le matériel : retirer et remplacer
Appareil en soutien prolongé	Remplacer dès que possible
En retard par rapport au niveau de correctif courant	Corriger immédiatement
Appareil non déployé avec haute disponibilité (physique ou virtuelle)	Ajouter une paire à haute disponibilité pour permettre les correctifs sans temps d'arrêt
Appareil ne prenant pas en charge les correctifs le jour même en raison de limites matérielles ou logicielles du processeur	L'ajouter à un programme VulnOps afin de prendre en charge les correctifs le jour même, ou renouveler le matériel pour qu'il les prenne en charge

Il existe aussi des mesures tactiques à prendre sur tous ces appareils. Et, encore une fois, nous ne supposons pas que ces mesures soient faciles, mais ce sont des défenses très utiles qui aident à prévenir plusieurs techniques d'attaque en usage :

- Exigez une authentification multifacteur pour la gestion, au minimum, et l'administration devrait idéalement se servir de modules d'authentification enfichables.
- Désactivez les interfaces de gestion exposées à Internet ou à la zone démilitarisée. Les interfaces de gestion comprennent l'accès de gestion aux appareils, aux contrôleurs SRST et aux autres stations de gestion. Elles devraient rejeter tout trafic provenant de réseaux non approuvés.
- Activez le filtrage du trafic sortant. Cela aide à réduire les retours de commande et contrôle des attaquants. N'oubliez pas que cela doit inclure le filtrage fondé sur le DNS, et pas seulement sur les adresses IP.
- Si vous utilisez des scripts ou du code d'automatisation qui touchent à la gestion réseau ou de sécurité, évaluez-les pour y repérer les secrets stockés, et utilisez un LLM pour en évaluer les défauts de sécurité.
- Placez un pare-feu d'applications Web devant toutes les applications exposées à Internet (appareil, appareil virtuel ou service infonuagique).

Cette première phase vise à durcir ce que nous avons et à nous assurer de pouvoir le maintenir à jour à l'avenir. Nous retirons les outils réseau et de sécurité à haut risque qui sont aussi les plus susceptibles d'être compromis eux-mêmes. Une frontière de sécurité ne vaut rien si elle devient le point d'entrée de l'attaque.

Double protection : protéger les technologies opérationnelles et les réseaux hautement critiques

Les technologies opérationnelles et les systèmes de contrôle industriels peuvent être difficiles, voire impossibles, à corriger selon une cadence régulière, si tant est qu'ils puissent l'être. Bien que les technologies opérationnelles aient longtemps évolué dans des environnements isolés, nous les avons, depuis des décennies, de plus en plus reliées aux réseaux TI, ce qui a ouvert des voies à partir d'Internet. Elles sont des cibles de premier plan, constituent des infrastructures critiques et se prêtent mal à la défense contre le rythme rapide des vulnérabilités, des exploits et des attaques découverts par l'IA.

Ces réseaux devraient être déconnectés de tout chemin vers Internet. Si ce n'est pas possible, envisagez des technologies à sens unique pour recueillir la télémétrie, comme les diodes de données. Si ce n'est pas possible, utilisez plusieurs pare-feu de nouvelle génération en ligne (deux ou plus) provenant de différents fournisseurs et maintenus à jour, de sorte qu'un attaquant doive exploiter plusieurs zero-day pour atteindre le réseau OT.

Phase 2 : ajouter des frontières prioritaires et ciblées

La phase 1 porte sur le durcissement des actifs connus et sur le retrait du matériel et des logiciels désuets qui ne peuvent pas être corrigés ou qui ne peuvent pas suivre le rythme des correctifs propulsés par l'IA. La phase 2 commence à ajouter des frontières de sécurité pour accroître la complexité combinatoire imposée aux attaquants et réduire la surface d'attaque des attaques réussies.

La première étape consiste essentiellement en une macrosegmentation : isoler la pile de sorte qu'une compromission de ressources moins précieuses et moins défendues sur le même réseau ne devienne pas un chemin facile vers la pile prioritaire. Il existe plusieurs techniques réseau que vous pouvez appliquer ici, notamment l'ajout de pare-feu, l'isolement des sous-réseaux et des routes, le passage à un réseau virtuel isolé, le cloisonnement au moyen de contrôles SRST, ou encore le passage au Zero Trust. L'isolement contre les attaques est-ouest est la première priorité, à supposer que vous ayez déjà durci vos défenses contre les attaques venant du nord (Internet) à la phase 1.

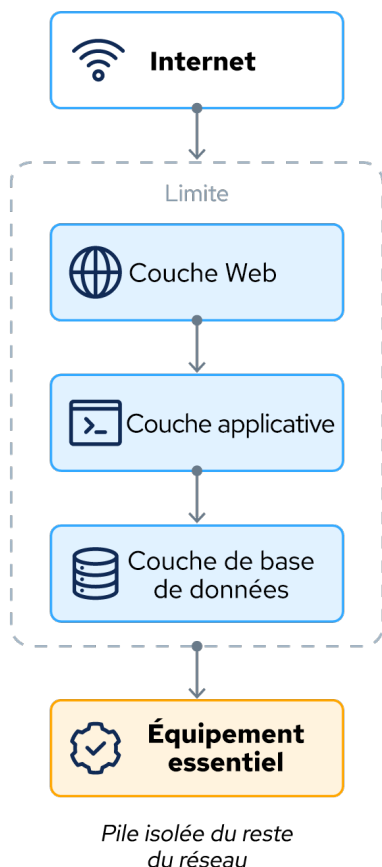
Si vous choisissez d'isoler au moyen d'outils à l'intérieur de l'hôte, ceux-ci doivent fonctionner à un niveau bas, idéalement en dehors du système d'exploitation (hyperviseur), ou sous forme d'agent durci géré de l'extérieur, afin qu'un attaquant qui parvient à entrer sur l'hôte ne puisse pas simplement désactiver la sécurité du réseau.

Deux étapes pour ajouter des frontières de sécurité

Commencez par établir un périmètre autour de la pile.
Ajoutez ensuite des frontières entre chaque couche interne.

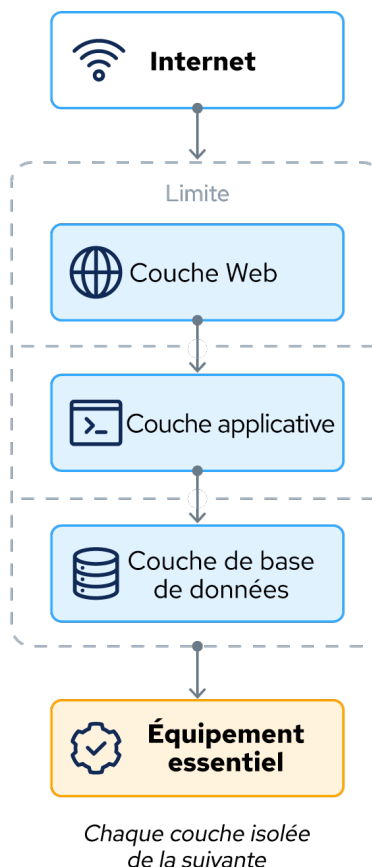
Étape 1

Frontière autour
de la pile



Étape 2

Frontière entre
chaque couche



La deuxième étape, qui prendra plus de temps et nécessitera vraisemblablement de nouvelles technologies, consiste à ajouter des *frontières internes* supplémentaires. Nos étapes précédentes visent à réduire le chemin le plus facile, mais celle-ci ajoute de véritables coûts supplémentaires pour l'attaquant. Chaque couche d'une application devrait communiquer uniquement avec la couche supérieure et la couche inférieure, et seulement sur les ports et protocoles attendus. Le SRST est la meilleure option, et ce modèle est souvent beaucoup plus facile à mettre en œuvre dans les environnements infonuagiques et entièrement virtualisés.

Phase 3 : élargir et affiner

Les phases 1 et 2 visent de façon très ciblée à protéger rapidement nos actifs les plus importants. Nous remplaçons le matériel et les logiciels désuets, nous mettons tout à jour, nous isolons ce qui est moins défendu, puis nous commençons à ajouter des frontières de sécurité pour accroître le coût et la complexité pour l'attaquant. N'oubliez pas qu'il ne s'agit que du volet réseau et sécurité réseau de l'équation; selon nos orientations AIVulnerability Storm Ready, d'autres initiatives parallèles sont tout aussi importantes.

La phase 3 consiste à étendre et à raffiner ces défenses à l'ensemble de nos réseaux. L'idéal est de passer des réseaux plats à la macrosegmentation, puis à la microsegmentation et enfin au Zero Trust. Bien que le passage direct au Zero Trust soit l'idéal, il n'est pas forcément réalisable à court terme. Pour la macrosegmentation, nous commençons par cloisonner les unités d'affaires et les zones de confiance.

La microsegmentation pousse les frontières jusqu'aux charges de travail et aux services individuels, de sorte que chacun ne communique qu'avec ce dont il a réellement besoin. À grande échelle, cela peut être lent et coûteux, et c'est généralement mis en œuvre projet par projet. Cela peut prendre des années, exige une cartographie réelle des flux et des dépendances avant toute modification, et peut perturber la production si l'on va trop vite.

À la CSA, nous définissons la microsegmentation comme faisant partie du Zero Trust, mais pour certains administrateurs réseau et sécurité, il s'agit de concepts distincts qui se recoupent en partie. La microsegmentation peut être mise en œuvre comme une approche réseau seulement, et isole sur une base par actif et par application, tandis que le Zero Trust ajoute l'identité et le contexte aux décisions de connectivité.

Zero Trust est l'endroit où tout cela converge. Une fois vos frontières bien granulaires et vos politiques alignées sur les charges de travail et les identités plutôt que sur l'emplacement réseau, vous autorisez l'accès à chaque connexion. C'est une barrière importante pour les attaquants, comme cela a déjà été démontré dans le nuage. Surveillez vos points de friction au fur et à mesure : les fournisseurs d'identité et les autres points de contrôle consolidés deviennent les cibles les plus précieuses précisément parce que tout le monde leur fait confiance; ils méritent donc davantage de contrôle et plus de frontières, pas moins.

À cette étape, il est absolument essentiel d'examiner de près vos processus et interfaces de gestion technique. Les réseaux exigent des mises à jour constantes, au-delà des seuls correctifs, et ces interfaces de gestion sont une cible de choix.

Rien de tout cela n'est vraiment « terminé », et c'est bien le but. Les réseaux qui résisteront aux attaques propulsées par l'IA ne seront pas ceux qui auront terminé une phase, mais ceux qui auront rendu la segmentation, les correctifs et la visibilité continus, de sorte que chaque nouvelle application arrive déjà avec ses frontières, et que chaque nouvel exploit se heurte à un réseau conçu pour coûter plus cher à l'attaquant qu'à vous. Des incidents surviendront quand même, et l'amélioration de la télémétrie réseau donne aussi plus de moyens aux équipes de réponse aux incidents.

Une transition difficile, un avenir prometteur

Aucune des techniques décrites dans ce document n'est nouvelle, mais l'échelle, la portée et les échéanciers de mise en œuvre sont différents de ceux auxquels nous étions habitués. Tout ce qui est abordé ici repose sur des principes de sécurité fondamentaux, dont plusieurs ont fait leurs preuves dans divers environnements. Mais nous, professionnels du réseau et de la sécurité, devons toujours trouver le juste équilibre entre la friction et le risque. L'ajout de contrôles de sécurité accroît les coûts et la complexité, et peut ralentir certaines parties de l'entreprise. Nos recommandations ne sont pas toujours faciles, et elles ne sont certainement pas gratuites. Nous ne voyons simplement pas d'autre voie.

L'IA change les hypothèses fondamentales de risque sur lesquelles reposent nos pratiques actuelles. Elle permet à un plus large éventail d'attaquants moins expérimentés d'opérer à un niveau plus proche de celui des chasseurs de bogues et des testeurs d'intrusion d'élite. Nous ne pouvons pas compter sur des semaines ou des mois pour corriger le bord du réseau, ni sur des années pour corriger notre infrastructure interne. Nous ne pouvons pas supposer que les attaquants seront incapables de naviguer dans nos architectures internes complexes. Nous ne pouvons pas nous attendre à ce qu'un serveur VPN non corrigé, utilisé seulement par un bureau de succursale, soit sans danger parce qu'aucun attaquant ne s'y intéresserait.

Les mathématiques peuvent aujourd'hui favoriser la sécurité offensive, mais la voie pour rééquilibrer la situation est claire. À terme, tous les logiciels seront intrinsèquement plus sécuritaires avant leur mise en marché, à mesure que nous intégrerons pleinement les mêmes capacités d'IA aux activités de développement logiciel et d'évaluation de la sécurité. À mesure que nous modernisons nos réseaux et avançons vers le Zero Trust, nous augmentons le nombre de frontières de sécurité que les attaquants doivent franchir, ce qui accroît leur coût et leur complexité combinatoires. Lorsque nous pourrions corriger et mettre à jour en continu, nous réduirons la fenêtre d'action dont ils disposent.

C'est cette combinaison qui fait pencher les mathématiques en notre faveur, et même de petits changements peuvent produire de grands effets. Faites évoluer votre modèle d'exploitation pour qu'il soit plus réactif et intégrez VulnOps. Concentrez-vous d'abord sur la protection de vos actifs les plus critiques. Assurez-vous d'utiliser du matériel et des logiciels actualisés, pouvant être corrigés le jour même de la publication des correctifs. Passez autant que possible aux SRST, qui incluent la plupart des réseaux virtuels et infonuagiques, puisqu'ils rendent la segmentation beaucoup plus simple à mettre en œuvre et à maintenir. Commencez par une segmentation est-ouest pour isoler les piles individuelles, mais visez ensuite des frontières internes nord-sud de plus en plus nombreuses, jusqu'au Zero Trust et au SDP.

Zero Trust ou zero-day, à vous de choisir.

Les organisations qui réussiront à l'ère de l'IA ne seront pas celles qui élimineront toute vulnérabilité. Ce seront celles qui réduiront continuellement les expositions exploitables plus vite que les attaquants ne pourront en tirer parti. Les barrières donnent du temps et augmentent les coûts imposés aux attaquants. C'est cette combinaison de réduction de la surface d'attaque, de rapidité de réaction et de frontières de sécurité qui crée une véritable résilience.