

Nachweis & Bewertung

Nachweise, Tests und Bewertungsleitfaden – dokumentierte Ergebnisse, wie man einen aussagekräftigen Proof of Concept durchführt und was man jeden Präventionsanbieter fragen sollte.

- [Besteht Cyber Crucible Ransomware-Simulationstests?](#)
- [Hat Cyber Crucible Fehlalarme? Wie hoch ist die Fehlalarmquote?](#)
- [Testet Cyber Crucible erschöpfend jede mögliche Ransomware?](#)
- [Welche Belege gibt es dafür, dass Cyber Crucible tatsächlich Angriffe stoppt?](#)
- [Wie sollten wir einen Proof of Concept für ein Präventionstool durchführen?](#)
- [Welche Fragen sollten wir jedem Anbieter von Endpunktsicherheit stellen?](#)
- [Was ist aus dem „dumm sammelnden“ Modell der Sicherheitsbranche geworden?](#)
- [Warum erzeugt intentbasierte Prävention weniger Alarm-Rauschen als signaturbasierte Erkennung?](#)

Besteht Cyber Crucible Ransomware-Simulationstests?

Die beste Antwort ist ... es kommt auf die Qualität und Genauigkeit der Ransomware-Simulationen an, aber wir haben bisher nur wenige hochwertige Tests gesehen, die tatsächlichen Angriffswerkzeugen und -verhaltensweisen entsprechen.

Cyber Crucible untersucht zugrunde liegende Verhaltensmuster beim Dateizugriff, im Speicherverhalten, im Prozessverhalten und im kryptografischen Verhalten, um Datenerpressungsaktivitäten im Zusammenhang mit einem Angriff zu erkennen und zu unterbinden.

Einige Simulationen von Datenerpressungssoftware haben sich im Laufe der Zeit weiterentwickelt und stellen mittlerweile eine genauere Darstellung realer Angriffe dar.

Wir haben Simulationen erlebt, die sich offenbar darauf konzentrierten, die von einigen Anbietern offengelegten Gegenmaßnahmen gegen Datenerpressung zu überprüfen, anstatt die Vorgehensweise der Erpressungstools und der Angreifer selbst zu testen. Man könnte dies auch so ausdrücken: „Es wird die Gegenmaßnahme getestet, nicht der Angriff.“

Als Reaktion darauf scheuen wir keine Tests gegen Angreifer-Emulation, Penetrationstests oder Erpressungstools. Wir greifen unsere eigene Software routinemäßig an und spielen dabei Angreifer-Vorgehensweisen nach, die wir online diskutiert gesehen oder in Kundenumgebungen entdeckt (und gestoppt) haben.

Hat Cyber Crucible Fehlalarme?

Wie hoch ist die Fehlalarmquote?

Cyber Crucible strebt eine Produktumgebung mit 0 Fehlalarmen an. Dennoch kommt es gelegentlich zu Fehlalarmen. Lassen Sie uns besprechen, woher diese kommen.

Derzeit liegt die Fehlalarmquote bei etwa 1 Reaktion pro Monat pro 1800 Bereitstellungen/Agenten, mit vollständiger Nachvollziehbarkeit, warum diese Reaktionen auftreten.

Beschädigtes Programm mit schnellen Dateioperationen

Ob absichtlich oder durch fehlerhafte Programmierung – manchmal beschädigen sich Programme selbst oder werden während der Ausführung durch andere Programme beschädigt. Wir sehen dies häufig bei neuen Funktionen in großen Programmpaketen wie Microsoft Office oder proprietärer, individuell entwickelter Software.

Dabei wird der Speicher des Programms beschädigt, was den Umfang der Überprüfung des Programms erhöht. Folgt darauf ein erpressungsähnliches Verhalten beim Dateizugriff, ist Cyber Crucible gezwungen, das Programm zu beenden. Wir werden immer besser darin, bekannte Fehler zu identifizieren, sodass Sie Ihr Programm zumindest bis zum Absturz ausführen können (manche starten automatisch neu, manche nicht).

In einer Welt, in der Angreifer nicht das Dateisystem nutzen, sondern stattdessen den Speicher laufender Programme kapern, müssen wir vorsichtshalber im Zweifel eingreifen.

Cyber Crucible verfügt nun über die Möglichkeit, eine Berechtigung zu erstellen: Wenn Programm A mit Argumenten Argumente A eine Beschädigung in Programm B mit Argumenten Argumente B verursacht, können wir eine temporäre Ausnahme erstellen für:

(Programm A + Argumente A) öffnet (Programm B + Argumente B)

Auf diese Weise wird kein zusätzliches Risiko angenommen, dass ein Hacker beteiligt ist, es sei denn, dieser verwendet zufällig genau dieselben Argumente.

Ein hervorragendes Beispiel dafür war, als Microsoft erstmals das Öffnen von Office-Dokumenten aus ihrer Desktop-Chat-Software einführte. Dies beschädigte das Office-Programm... jedes... einzelne... Mal.... Anschließend begann das Office-Programm in etwa 25 % der Fälle, alle Dokumente auf dem Rechner zu scannen. Das... erforderte eine Ausnahme.

Wenn Sie einen Absturz erleben, nutzen Sie die Support-Schaltfläche, und wir helfen Ihnen dabei, eine Ausnahme zu erstellen, bis der Anbieter das Problem behebt.

Sicherheitsprogramme

Cyber Crucible funktioniert sehr gut im Zusammenspiel mit Sicherheitsprogrammen. Entgegen der Annahme der meisten neuen Kunden nutzen fortschrittlichere Sicherheitstools tendenziell fortschrittlichere Systeminteraktionen als Antivirenprogramme, die ältere Technologien verwenden. Tatsächlich führen einige kostenlose oder kostengünstige Sicherheitstools eine Vielzahl unsicherer PowerShell-Skripte als zentralen Bestandteil ihres Schutzes aus, anstatt tatsächliche Programme.

Als ein hochgradig widerstandsfähiges, kernelbasiertes Sicherheitstool wird Cyber Crucible am besten als „tiefer“ oder „näher“ an der Hardware angesiedelt betrachtet als die meisten anderen Tools. Cyber Crucible identifiziert automatisch andere Sicherheitstools und fügt sie den Verhaltensmodellen hinzu. Solange die Sicherheitstools keine Anzeichen einer Kompromittierung zeigen, dürfen sie normal weiterarbeiten.

Gelegentlich verursacht ein Sicherheitstool Systeminstabilität, wenn dessen Versuche, die Cyber Crucible-Software abzuschalten oder zu beeinträchtigen, fehlschlagen. In solchen Fällen ist es am besten, Cyber Crucible im jeweils anderen Sicherheitstool auf die Whitelist zu setzen, um zu verhindern, dass dieses versucht (und scheitert), Cyber Crucible zu deaktivieren oder zu beeinträchtigen.

Zögern Sie nicht, über das Webportal ein Support-Ticket bei Cyber Crucible zu eröffnen, um die Angelegenheit zu besprechen.

Administrations- oder Backup-Tools

Unter bestimmten Umständen muss ein maßgeschneidertes Verhalten für ein Administrations- oder Backup-Tool erstellt werden, während das Cyber Crucible-Team ein verbessertes Verhaltensmodell entwickelt. In diesem Fall unterstützt das Cyber Crucible-Team bei der Erstellung einer Ausnahme im Verhaltensmodell, die das Programm zusammen mit den Argumenten für ein äußerst kleines

Zeitfenster eines möglichen Angriffs herausfiltert.

Testet Cyber Crucible erschöpfend jede mögliche Ransomware?

Cyber Crucible arbeitet auf Basis kernelseitiger Verhaltensmodellierung, um Datendiebstahl, Zugangsdatendiebstahl und Ransomware-Verschlüsselungsverhalten sehr schnell zu erkennen. Die Verhaltensentscheidungen basieren auf Informationen aus Prozessverhalten, speicherbasierten Verhaltensweisen und bestimmten dateibasierten Verhaltensweisen, um eine Entscheidung genau in dem Moment zu treffen, in dem die Erpressung unmittelbar bevorsteht.

Trotz der sehr großen Anzahl an Ransomware- und Erpressungssoftware-Beispielen lassen sich diese verhaltensbasiert in eine relativ kleine Anzahl von Kategorien einteilen.

Oberflächlich betrachtet müssen jedoch die „oberflächlichen“ Abwehrmechanismen, die von einer Vielzahl von Sicherheitstools verwendet werden, einer sehr großen Anzahl von Erpressungswerkzeugen entgegenwirken. Es ist wichtig zu verstehen, dass unsere Verhaltensanalytik viel tiefer in die Werkzeuge und Vorgehensweisen der Angreifer eindringt und dadurch die Notwendigkeit einer (unmöglichen) erschöpfenden Testung sämtlicher potenzieller Schadsoftware zu jedem Zeitpunkt drastisch reduziert.

Die Entwickler von Cyber Crucible kategorisieren die kernelseitigen Speicher-, Prozess- und Dateiverhaltensweisen eines Erpressungswerkzeugs und stellen anschließend sicher, dass es in eine der bekannten Abwehrfähigkeiten passt. Gelegentlich kann die Formel, ähnlich wie bei einem Impfstoff, leicht angepasst werden, um eine präzisere Reaktion zu erzielen.

Nur sehr selten tritt etwas völlig Neues auf.

Das Cyber-Crucible-Team arbeitet unermüdlich daran, neuartige Techniken vorbeugend aufzudecken, die in unserem Kundenstamm oder in Threat-Intelligence-Daten noch nicht beobachtet wurden.

Der Effekt?

1. Die Abwehr von Erpressungswerkzeugen durch Cyber Crucible ist gegenüber allen bekannten Ransomware-Varianten vollständig.
2. Neue Ransomware-Varianten werden blockiert, noch bevor sie die ersten Kunden erreichen. Bei den meisten Programmen, gegen die wir schützen, wissen wir etwa 120 Tage lang nicht einmal, wie sie heißen sollen, bis die Forscher „aufgeholt“ haben.

3. Entwickler von Ransomware- und Erpressungswerkzeugen rufen uns manchmal an, um herauszufinden, womit wir uns beschäftigen, in der Hoffnung, einen Ansatzpunkt zu finden. Leider scheint die Frustration über unsere Präsenz manchmal ihre Weiterentwicklung anzuspornen, wodurch andere Werkzeuge noch weniger wirksam werden.

4. Wir begrüßen die Prüfung unserer Software. Wenn Penetrationstester, Malware-Analysten oder Fachleute für Adversary Emulation in den Verkaufsprozess einsteigen, freuen wir uns sehr!

Welche Belege gibt es dafür, dass Cyber Crucible tatsächlich Angriffe stoppt?

Kurze Antwort: Dokumentierte Kundenimplementierungen, unabhängige Tests durch eine große Wirtschaftsprüfungs- und Beratungsgesellschaft sowie wiederholte Fälle, in denen Zero-Day-Angriffe Monate vor anderen Anbietern gestoppt wurden.

Dokumentierte Ergebnisse

- **Finanzdienstleistungen:** fast 10.000 bösartige Prozesse wurden innerhalb von 60 Tagen autonom abgefangen, 98 % davon auf einer hochprivilegierten Microsoft-SQL-Serverfarm — ohne jegliche Warnmeldungen von drei eingesetzten EDRs, einem MDR oder einem ausgelagerten Top-50-SOC.
- **Unabhängige Tests:** eine der größten Wirtschaftsprüfungs- und Beratungsgesellschaften Nordamerikas unterzog Cyber Crucible einer eigenen Evaluierung gegen Ransomware-, Datendiebstahl- und Identitätsbetrugsszenarien.
- **Der Branche voraus:** Cyber Crucible hat mehrere Zero-Day-Angriffe auf Kundennetzwerke gestoppt, und zwar bis zu 90 Tage bevor ein anderer Anbieter diese meldete oder darauf reagierte.
- **Browserbasierte Angriffe:** ab dem vierten Quartal 2023 stiegen automatisierte Reaktionen auf Chrome-, Edge- und Chromium-Aktivitäten von null auf mehrere Tausend, wobei entsprechende CVEs später von Google und Microsoft veröffentlicht wurden.
- **Lösegeldzahlungen:** rund 90 % der Ransomware-Opfer zahlten im Jahr 2023. Cyber-Crucible-Kunden zahlten 0 \$.

Widerstandsfähigkeit unter direktem Angriff

Bei einem maritimen Einsatz eskalierten die Angreifer, nachdem herkömmliche Versuche fehlgeschlagen waren, zur Übernahme von Windows-Anmeldedaten und sperrten Systeme auf BIOS-Ebene. Bei zwei weiteren Kunden versuchten Angreifer, die die Kernel-Entscheidungs-Engine nicht überwinden konnten, stattdessen die Kundenbenachrichtigung zu blockieren, indem sie den Netzwerk-Stack des Betriebssystems angriffen — und scheiterten, da die Netzwerkkommunikation von Cyber Crucible unabhängig vom Betriebssystem erfolgt.

Wie sollten wir einen Proof of Concept für ein Präventionstool durchführen?

Kurze Antwort: Setzen Sie es parallel zu Ihrem bestehenden Stack in einer echten Produktionsumgebung ein und vergleichen Sie, was jedes Tool erkennt. Das aussagekräftigste Ergebnis liefert kein Labortest — es ist die Lücke zwischen dem, worauf Ihre aktuellen Tools alarmieren, und dem, was tatsächlich abgefangen wird.

Warum ein paralleler Einsatz am besten funktioniert

Das Finanzdienstleistungsunternehmen in unserem am besten dokumentierten Fall ging genau so vor. Sie haben nichts entfernt. Sie haben Cyber Crucible gezielt hinzugefügt, um herauszufinden, was ihre bestehende Investition übersah, und hatten innerhalb von 60 Tagen eine eindeutige Antwort.

Was gemessen werden sollte

- **Interceptions, auf die Ihr aktueller Stack nie alarmiert hat.** Dies ist das zentrale Signal.
- **Time to Decision.** Prävention muss abgeschlossen sein, bevor Schaden entsteht — nicht erst, bevor ein Ticket geschlossen wird.
- **Betriebsstörungen.** Zählen Sie Ausfallzeiten, Neustarts und durch Fehlalarme verursachte Unterbrechungen.
- **Aufgewendete Analystenstunden.** Autonome Prävention sollte die Arbeitslast reduzieren, nicht eine zusätzliche Warteschlange schaffen.

Ein Hinweis zu Labortests

Tests mit Samples aus Malware-Datenbanken liefern oft nur sehr wenig Aussagekraft, da Angreifer Malware so gestalten, dass sie inaktiv bleibt, sofern sie kein Validierungssignal von einem noch aktiven, live betriebenen Command-and-Control-Server erhält. Ein Sample, das in Ihrem Labor

nichts unternimmt, beweist nichts über einen echten Angriff.

Welche Fragen sollten wir jedem Anbieter von Endpunktsicherheit stellen?

Kurze Antwort: Fragen Sie, wo Entscheidungen getroffen werden, was ohne Konnektivität passiert, wovon das Tool zur Funktion abhängig ist, und ob es verhindert oder lediglich meldet. Die Antworten trennen Architektur schnell von Marketing.

Fragen, die es sich zu stellen lohnt

1. **Wo wird die schützende Entscheidung getroffen — auf dem Endpunkt oder in Ihrer Cloud?** Ein Netzwerk-Roundtrip im kritischen Pfad kann einen Angriff im Millisekundenbereich nicht schlagen.
2. **Was passiert, wenn das Gerät offline oder air-gapped ist?** Wenn der Schutz nachlässt, war er nicht lokal.
3. **Ist Ihr Agent zur Funktion auf Windows-Systembibliotheken angewiesen?** Falls ja, kann ein Angreifer, der diese Bibliotheken im Arbeitsspeicher kompromittiert, ihn blenden oder zum Schweigen bringen.
4. **Benötigen Sie Signaturen oder Threat-Intelligence-Feeds?** Falls ja, hinkt der Schutz dem Angreifer per Definition hinterher.
5. **Erfordert eine Reaktion einen Menschen?** Falls ja, ist die Reaktionszeit durch die menschliche Biologie begrenzt.
6. **Speichern Sie unsere Verschlüsselungsschlüssel oder laden Sie sensible Daten zur Analyse hoch?** Zentralisierte Speicherung schafft ein Ziel und ein Risiko der erzwungenen Offenlegung.
7. **Wie hoch ist Ihre Falsch-Positiv-Rate, und was bewirkt eine Reaktion tatsächlich?** Eine vollständige Systemsperre als einzige Eindämmungsoption ist in einem Krankenhaus oder einer Anlage kostspielig.

Warum diese Fragen wichtig sind

Die meisten Endpunkt-Tools wirken in einem Datenblatt ähnlich. Diese Fragen bringen die architektonischen Entscheidungen zutage, die darüber bestimmen, ob ein Tool rechtzeitig handeln kann — oder nur im Nachhinein erklären kann, was geschehen ist.

Was ist aus dem „dumm sammelnden“ Modell der Sicherheitsbranche geworden?

Kurze Antwort: Jahrelang rieten Anbieter, dass lokale Endpunktverarbeitung veraltet sei, und setzten stattdessen auf leichtgewichtige Agenten, die Rohmetriekdaten an Cloud-Analysen weiterleiten. Moderne speicherresidente Angriffe auf zentrale Betriebssystembibliotheken haben dieses Modell entlarvt — die Agenten laufen weiter, werden aber faktisch blind und stumm.

Warum das Modell übernommen wurde

Es war günstiger zu entwickeln. Kernel-Level-Engineering ist schwierig und teuer; das Versenden von Metriekdaten in die Cloud ermöglichte es Anbietern, kostengünstige Endpunktsoftware auszuliefern, große Cloud-Speicherpools zu monetarisieren und Cloud-KI-Fähigkeiten zu vermarkten. Das Modell war auf Entwicklungsökonomie optimiert, nicht auf Verteidigung.

Warum es jetzt scheitert

Diese Agenten sind vollständig auf native Betriebssystemkomponenten angewiesen, um zu funktionieren und Daten zu sammeln. Wenn Angreifer diese Komponenten im Speicher kompromittieren, stürzt der Agent nicht ab — er läuft weiter und meldet nichts, sodass ein Dashboard anzeigt, alles sei in Ordnung, während Daten exfiltriert werden.

Es gibt zudem dokumentierte, gezielte Sabotage: Angreifer und, in einigen verifizierten Fällen, Wettbewerber, die eigene Schwächen verschleiern wollen, nutzen aktiv diese Bibliotheksabhängigkeiten aus, um Endpunktschutzmaßnahmen zwangsweise zu beenden oder zu sabotieren.

Die architektonische Sackgasse

Um dies zu beheben, müssten etablierte Anbieter cloud-first-Telemetrie-Pipelines aufgeben und für lokale Verarbeitung auf Kernel-Ebene am Edge neu aufbauen — ein mehrjähriges Vorhaben. Aktuelle Übernahmen und Venture-Investitionen zeigen, dass sich die Branche stattdessen weiter in Richtung Cloud-Analysen und zentralisierter Data Lakes bewegt.

Warum erzeugt intentbasierte Prävention weniger Alarm-Rauschen als signaturbasierte Erkennung?

Kurze Antwort: Weil dabei eine engere Fragestellung gestellt wird. Signatur- und heuristische Tools schlagen Alarm, sobald etwas einem *bekannten schädlichen Muster ähnelt*, und überlassen es Menschen, die Spreu vom Weizen zu trennen. Intentbasierte Prävention fragt, ob ein bestimmtes Programm gerade jetzt an einem bekannten Einstiegspunkt für Datendiebstahl etwas Böses tut – eine Frage mit weitaus weniger mehrdeutigen Antworten.

“ Die gemessene Falsch-Positiv-Rate von Cyber Crucible finden Sie im bestehenden Artikel „**Haben Sie Fehlalarme? Wie hoch ist Ihre Fehlalarmrate?**“. Diese Seite erläutert den architektonischen Grund, warum sich das Rauschprofil unterscheidet.

Warum sich das Rauschprofil unterscheidet

Signatur- und heuristische Tools erzeugen Alarme, wenn etwas *einem bekannten schädlichen Muster ähnelt*, und verlassen sich anschließend auf manuelle Triage, um Echtes von Irrelevantem zu unterscheiden. Genau das führt in vielen Umgebungen zu über 100 Alarmen pro Endpunkt und Tag.

Die intentbasierte Bewertung stellt eine engere Frage: Tut dieses Programm gerade jetzt etwas Böses an einem bekannten Einstiegspunkt für Identitäts- oder Datendiebstahl? Diese Frage hat weitaus weniger mehrdeutige Antworten.

Was das operativ bedeutet

- Analysten müssen keinen Alarmen mit geringer Zuverlässigkeit nachjagen.

- Es ist keine YARA-Regel-Feinabstimmung oder manuelle Bedrohungssuche erforderlich, um das Rauschen beherrschbar zu halten.
- Da die Reaktion nur den betreffenden Prozess unterbricht, führt selbst eine fehlerhafte Intervention nicht zum Ausfall eines Systems.

Für aktuelle gemessene Werte in einer Umgebung wie der Ihren wenden Sie sich an Ihren Cyber-Crucible-Ansprechpartner – und lesen Sie den bestehenden Artikel „Haben Sie Fehlalarme?“ für Details auf Produktebene.