

Prova e Avaliação

Evidências, testes e orientações de avaliação - resultados documentados, como realizar uma prova de conceito significativa e o que perguntar a qualquer fornecedor de prevenção.

- [A Cyber Crucible passa nos testes de simulação de ransomware?](#)
- [A Cyber Crucible apresenta falsos positivos? Qual é a taxa de falsos positivos?](#)
- [A Cyber Crucible testa exaustivamente todos os ransomwares possíveis?](#)
- [Que evidências existem de que a Cyber Crucible realmente impede ataques?](#)
- [Como devemos conduzir uma prova de conceito para uma ferramenta de prevenção?](#)
- [Que perguntas devemos fazer a qualquer fornecedor de segurança de endpoint?](#)
- [O que aconteceu com o modelo de "coletor burro" do setor de segurança?](#)
- [Por que a prevenção baseada em intenção gera menos ruído de alertas do que a detecção baseada em assinaturas?](#)

A Cyber Crucible passa nos testes de simulação de ransomware?

A melhor resposta é... depende da qualidade e da precisão das simulações de ransomware, mas não vimos muitos testes de alta qualidade que correspondam às verdadeiras ferramentas e comportamentos de ataque.

A Cyber Crucible examina padrões comportamentais subjacentes no acesso a arquivos, nos comportamentos de memória, nos comportamentos de processos e nos comportamentos criptográficos para identificar e suspender atividades de extorsão de dados em conjunto com um ataque.

Algumas simulações de software de extorsão de dados evoluíram ao longo do tempo, tornando-se uma representação mais precisa de ataques reais.

Já passamos por simulações que pareciam focar em verificar as contramedidas de extorsão de dados divulgadas por alguns fornecedores, e não o modo de operação das próprias ferramentas de extorsão e dos atacantes. Uma forma de expressar isso seria: “Testar a contramedida, não o ataque”.

Em resposta a isso, nunca evitamos testes contra emulação de atacantes, testes de penetração ou ferramentas de extorsão. Atacamos rotineiramente nosso próprio software e reproduzimos o modo de operação de atacantes que observamos ser discutido online ou que encontramos (e interrompemos) em ambientes de clientes.

A Cyber Crucible apresenta falsos positivos? Qual é a taxa de falsos positivos?

A Cyber Crucible busca proporcionar um ambiente de produto com 0 falsos positivos. Dito isso, às vezes ocorrem falsos positivos. Vamos discutir de onde eles vêm.

Atualmente, a taxa de falsos positivos é de aproximadamente 1 resposta por mês, a cada 1800 implantações/agentes, com total responsabilização sobre o motivo pelo qual essas respostas ocorrem.

Programa Corrompido Realizando Operações Rápidas em Arquivos

Seja de propósito ou por má programação, às vezes os programas se corrompem sozinhos, ou são corrompidos por outros programas durante a execução. Vemos isso com frequência em novos recursos de grandes suítes de programas, como o Microsoft Office, ou softwares proprietários personalizados.

O que acontece é que a memória do programa é corrompida, o que aumenta o nível de inspeção sobre o programa. Se isso for seguido por um comportamento de acesso a arquivos semelhante à extorsão, a Cyber Crucible é forçada a suspender o programa. Estamos melhorando na identificação de bugs conhecidos, e você consegue pelo menos continuar executando seu programa até que ele trave (alguns reiniciam automaticamente, outros não).

Em um mundo onde os invasores não usam o sistema de arquivos e, em vez disso, sequestram a memória de programas em execução, precisamos ser cautelosos.

A Cyber Crucible agora tem a capacidade de criar uma permissão, na qual, se o Programa A, com os Argumentos A, causar uma corrupção no Programa B, com os Argumentos B - podemos criar uma exclusão temporária para:

(Programa A + Argumentos A) abre (Programa B + Argumentos B)

Dessa forma, não se assume um risco adicional de que um hacker esteja envolvido, a menos que ele, por acaso, use exatamente os mesmos argumentos.

Um excelente exemplo disso ocorreu quando a Microsoft introduziu pela primeira vez a abertura de documentos do Office diretamente a partir de seu software de chat para desktop. Isso corrompia o programa Office... toda... única... vez... e então o programa Office começava a escanear todos os Documentos na máquina cerca de 25% das vezes. Isso... precisou de uma exclusão.

Se você tiver uma falha, use o botão de suporte, e ajudaremos você a criar uma exclusão, até que o fornecedor corrija o problema.

Programas de Segurança

A Cyber Crucible coexiste muito bem com programas de segurança. Ao contrário do que a maioria dos novos clientes pensa, ferramentas de segurança mais avançadas tendem a usar interações de sistema mais avançadas do que os programas antivírus que utilizam tecnologias mais antigas. De fato, algumas ferramentas de segurança gratuitas ou de baixo custo executam uma variedade de scripts inseguros do Powershell como parte fundamental de sua proteção, em vez de programas reais.

Como uma ferramenta de segurança altamente resiliente baseada em kernel, a Cyber Crucible é melhor compreendida como estando “mais abaixo”, ou “mais próxima” do hardware do que a maioria das outras ferramentas. A Cyber Crucible identifica e adiciona automaticamente outras ferramentas de segurança aos modelos comportamentais. Desde que as ferramentas de segurança não apresentem sinais de comprometimento, elas podem operar normalmente.

Ocasionalmente, uma ferramenta de segurança causa instabilidade no sistema quando suas tentativas de desativar ou interferir no software da Cyber Crucible falham. Nesses casos, o ideal é colocar a Cyber Crucible na lista de permissões da outra ferramenta de segurança, para impedir que ela tente (e falhe) desativar ou interferir na Cyber Crucible.

Sinta-se à vontade para abrir um chamado de suporte junto à Cyber Crucible, por meio do portal web, para discutir o assunto.

Ferramentas Administrativas ou de Backup

Em algumas circunstâncias, é necessário criar um comportamento personalizado para uma ferramenta administrativa ou de backup, enquanto a equipe da Cyber Crucible desenvolve um

modelo comportamental aprimorado. Nesse caso, a equipe da Cyber Crucible ajudará na criação de uma exceção no modelo comportamental, que isolará o programa mais os argumentos para uma janela de oportunidade extremamente pequena para um possível ataque.

A Cyber Crucible testa exaustivamente todos os ransomwares possíveis?

A Cyber Crucible opera com base em modelagem comportamental em nível de kernel, para descobrir rapidamente comportamentos de roubo de dados, roubo de credenciais e criptografia por ransomware. As decisões comportamentais são tomadas utilizando informações do comportamento de processos, comportamentos derivados de memória e certos tipos de comportamentos originados de arquivos, para tomar uma decisão exatamente no momento em que a extorsão está prestes a ocorrer.

Apesar do número muito grande de amostras de ransomware e de software com temática de extorsão, elas podem ser categorizadas comportamentalmente em um número relativamente pequeno de conjuntos.

Na superfície, porém, as defesas “superficiais” utilizadas por diversas ferramentas de segurança têm um número muito grande de ferramentas de extorsão a tentar combater. É importante entender que nossas análises comportamentais vão muito mais a fundo nas ferramentas e nas táticas dos atacantes, reduzindo drasticamente a necessidade de algum tipo de teste (impossível) exaustivo de todos os malwares possíveis o tempo todo.

Os desenvolvedores da Cyber Crucible categorizam os comportamentos de memória, processos e arquivos em nível de kernel de uma ferramenta de extorsão e, em seguida, garantem que ela se enquadre em uma das capacidades defensivas conhecidas. Ocasionalmente, de forma semelhante a uma vacina, a fórmula pode ser ajustada ligeiramente para produzir uma resposta mais precisa.

Muito raramente surge algo completamente novo.

A equipe da Cyber Crucible trabalha exaustivamente para descobrir preventivamente técnicas inéditas ainda não observadas em nossa base de clientes ou em inteligência de ameaças.

O efeito?

1. A defesa da Cyber Crucible contra ferramentas de extorsão é completa contra todas as variantes conhecidas de ransomware.
2. Novas variantes de ransomware são bloqueadas antes mesmo de atingirem os primeiros clientes. Nem sequer sabemos como chamar a maior parte do software contra o qual nos defendemos por cerca de 120 dias, até que os pesquisadores “se atualizem”.

3. Desenvolvedores de ransomware e de ferramentas de extorsão às vezes nos ligam para ver o que estamos fazendo, na tentativa de encontrar uma brecha. Infelizmente, a frustração causada pela nossa presença parece, às vezes, estimular a evolução por parte deles, o que torna outras ferramentas ainda menos eficazes.

4. Recebemos com satisfação os testes de nosso software. Quando testadores de intrusão, analistas de malware ou profissionais de emulação de adversários entram no processo de venda — ficamos entusiasmados!

Que evidências existem de que a Cyber Crucible realmente impede ataques?

Resposta curta: Implementações documentadas de clientes, testes independentes realizados por uma grande empresa de contabilidade e consultoria, e instâncias repetidas de interrupção de ataques de dia zero meses antes de outros fornecedores os terem reportado.

Resultados documentados

- **Serviços financeiros:** quase 10.000 processos maliciosos interceptados de forma autônoma em 60 dias, 98% num conjunto de servidores Microsoft SQL altamente privilegiados — sem qualquer alerta por parte de três EDRs implementados, um MDR, ou um SOC subcontratado do Top-50.
- **Testes independentes:** uma das maiores empresas de contabilidade e consultoria da América do Norte submeteu a Cyber Crucible à sua própria avaliação contra cenários de ransomware, roubo de dados e fraude de identidade.
- **À frente do setor:** a Cyber Crucible impediu vários ataques de dia zero em redes de clientes até 90 dias antes de qualquer outro fornecedor os ter reportado ou respondido a eles.
- **Ataques baseados no navegador:** a partir do quarto trimestre de 2023, as respostas automatizadas contra atividade do Chrome, Edge e Chromium aumentaram de zero para milhares, com os CVEs relacionados publicados posteriormente pela Google e pela Microsoft.
- **Pagamentos de resgate:** aproximadamente 90% das vítimas de ransomware pagaram em 2023. Os clientes da Cyber Crucible pagaram \$0.

Resiliência sob ataque direto

Numa implementação marítima, os adversários escalaram para o sequestro de credenciais do Windows e o bloqueio de sistemas ao nível do BIOS quando as tentativas convencionais falharam. Com outros dois clientes, os atacantes que não conseguiram derrotar o motor de decisão do kernel tentaram, em vez disso, bloquear a notificação ao cliente atacando a stack de rede do sistema operativo — e falharam, porque as comunicações de rede da Cyber Crucible são independentes do sistema operativo.

Como devemos conduzir uma prova de conceito para uma ferramenta de prevenção?

Resposta curta: Implemente-a em conjunto com sua pilha de ferramentas existente em um ambiente de produção real e compare o que cada ferramenta detecta. O resultado mais informativo não é um teste em laboratório — é a lacuna entre o que suas ferramentas atuais alertam e o que realmente é interceptado.

Por que implementar em conjunto funciona melhor

A empresa de serviços financeiros do nosso caso mais bem documentado fez exatamente isso. Eles não removeram nada. Adicionaram o Cyber Crucible especificamente para descobrir o que seu investimento existente estava deixando passar, e obtiveram uma resposta definitiva em até 60 dias.

O que medir

- **Interceptações sobre as quais sua pilha atual nunca alertou.** Este é o sinal principal.
- **Tempo até a decisão.** A prevenção precisa ser concluída antes do dano ocorrer, não antes do fechamento de um chamado.
- **Interrupção nos negócios.** Conte o tempo de inatividade, reinicializações e interrupções causadas por falsos positivos.
- **Horas de analistas consumidas.** A prevenção autônoma deve reduzir a carga de trabalho, não gerar uma fila adicional.

Um alerta sobre testes em laboratório

Testar amostras retiradas de bancos de dados de malware muitas vezes mede muito pouco, pois os invasores projetam o malware para permanecer inativo, a menos que receba um sinal de validação de um servidor de comando e controle ativo que eles ainda operam. Uma amostra que

não faz nada em seu laboratório não prova nada sobre um ataque real.

Que perguntas devemos fazer a qualquer fornecedor de segurança de endpoint?

Resposta curta: Pergunte onde as decisões são tomadas, o que acontece sem conectividade, do que a ferramenta depende para funcionar e se ela previne ou apenas reporta. As respostas separam rapidamente a arquitetura do marketing.

Perguntas que vale a pena fazer

1. **Onde é tomada a decisão de proteção — no endpoint ou na sua nuvem?** Uma ida e volta pela rede no caminho crítico não pode superar um ataque em escala de milissegundos.
2. **O que acontece quando a máquina está offline ou isolada da rede (air-gapped)?** Se a proteção se degrada, ela não era local.
3. **O seu agente depende de bibliotecas do sistema Windows para funcionar?** Se sim, um atacante que comprometa essas bibliotecas em memória pode cegá-lo ou silenciá-lo.
4. **Vocês exigem assinaturas ou feeds de inteligência de ameaças?** Se sim, a proteção fica atrás do atacante por definição.
5. **Uma resposta requer intervenção humana?** Se sim, o tempo de resposta é limitado pela biologia humana.
6. **Vocês armazenam nossas chaves de criptografia ou enviam dados sensíveis para análise?** O armazenamento centralizado cria um alvo e um risco de divulgação compulsória.
7. **Qual é a sua taxa de falsos positivos, e o que uma resposta realmente faz?** O bloqueio total do sistema como única opção de contenção é dispendioso em um hospital ou em uma fábrica.

Por que essas perguntas importam

A maioria das ferramentas de endpoint parece semelhante em uma ficha técnica. Essas perguntas revelam as decisões arquitetônicas que determinam se uma ferramenta consegue agir a tempo — ou se só consegue explicar, depois, o que aconteceu.

O que aconteceu com o modelo de "coletor burro" do setor de segurança?

Resposta curta: Durante anos, os fornecedores recomendaram que o processamento local no endpoint estava obsoleto e promoveram agentes leves que encaminham telemetria bruta para análises em nuvem. Ataques modernos em memória contra bibliotecas centrais do sistema operacional expuseram esse modelo — os agentes continuam em execução, mas ficam efetivamente cegos e mudos.

Por que o modelo foi adotado

Era mais barato de construir. A engenharia em nível de kernel é difícil e cara; enviar telemetria para a nuvem permitiu que os fornecedores lançassem software de endpoint de baixo custo, monetizassem grandes pools de armazenamento em nuvem e comercializassem capacidades de IA em nuvem. O modelo foi otimizado para a economia de desenvolvimento, não para a defesa.

Por que está falhando agora

Esses agentes dependem inteiramente de componentes nativos do sistema operacional para funcionar e coletar dados. Quando os invasores comprometem esses componentes em memória, o agente não trava — ele continua em execução e não relata nada, deixando um painel que indica que está tudo bem enquanto os dados são exfiltrados.

Há também interdição deliberada documentada: invasores, e em alguns casos verificados concorrentes que buscam mascarar suas próprias falhas, exploram ativamente essas dependências de bibliotecas para forçar o encerramento ou sabotar as defesas de endpoint.

O beco sem saída arquitetônico

Corrigir isso exigiria que os fornecedores tradicionais abandonassem os pipelines de telemetria voltados para a nuvem e reconstruíssem para computação de borda local, em nível de kernel — um esforço de vários anos. Aquisições recentes e investimentos de capital de risco mostram o setor avançando ainda mais em direção a análises em nuvem e data lakes centralizados, em vez disso.

Por que a prevenção baseada em intenção gera menos ruído de alertas do que a detecção baseada em assinaturas?

Resposta curta: Porque ela faz uma pergunta mais restrita. As ferramentas de assinatura e heurística emitem alertas sempre que algo *se assemelha* a um padrão conhecido como malicioso e deixam para os humanos a tarefa de resolver a questão. A prevenção baseada em intenção pergunta se um programa específico está, neste exato momento, fazendo algo malicioso em um ponto de entrada conhecido de roubo de dados — uma pergunta com muito menos respostas ambíguas.

“ Para a taxa de falsos positivos medida da Cyber Crucible, consulte o artigo existente **"Vocês têm falsos positivos? Qual é a sua taxa de falsos positivos?"** Esta página explica a razão arquitetural pela qual o perfil de ruído é diferente.

Por que o perfil de ruído é diferente

As ferramentas de assinatura e heurística geram alertas quando algo *se assemelha* a um padrão conhecido como malicioso e, em seguida, dependem da triagem humana para separar o que é real do que é irrelevante. É isso que produz mais de 100 alertas por endpoint por dia em muitos ambientes.

A avaliação baseada em intenção faz uma pergunta mais restrita: este programa está, neste exato momento, fazendo algo malicioso em um ponto de entrada conhecido de roubo de identidade ou de dados? Essa pergunta tem muito menos respostas ambíguas.

O que isso significa na prática operacional

- Os analistas não precisam perseguir alertas de baixa confiança.
- Não é necessário ajustar regras YARA nem realizar caça manual a ameaças para manter o ruído sob controle.
- Como a resposta suspende apenas o processo infrator, mesmo uma interceptação incorreta não derruba o sistema.

Para as taxas atuais medidas em um ambiente semelhante ao seu, pergunte ao seu representante da Cyber Crucible — e consulte o artigo existente "Vocês têm falsos positivos?" para detalhes em nível de produto.