



Inżynieria detekcji cyberzagrożeń w praktyce

Planowanie, tworzenie i walidacja
mechanizmów wykrywania zagrożeń



MEGAN RODDIE
JASON DEYALSINGH
GARY J. KATZ

Tytuł oryginału: Practical Threat Detection Engineering: A hands-on guide to planning, developing, and validating detection capabilities

Tłumaczenie: Radosław Meryk

ISBN: 978-83-289-0902-1

Copyright © Packt Publishing 2023. First published in the English language under the title 'Practical Threat Detection Engineering' – (9781801076715).

Polish edition copyright © 2024 by Helion S.A.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from the Publisher.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiegokolwiek postaci jest zabronione. Wykonywanie kopii metodą kserograficzną, fotograficzną, a także kopiowanie książki na nośniku filmowym, magnetycznym lub innym powoduje naruszenie praw autorskich niniejszej publikacji.

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autor oraz wydawca dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autor oraz wydawca nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Drogi Czytelniku!

Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres

<https://helion.pl/user/opinie/indecy>

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Helion S.A.

ul. Kościuszki 1c, 44-100 Gliwice

tel. 32 230 98 63

e-mail: helion@helion.pl

WWW: <https://helion.pl> (księgarnia internetowa, katalog książek)

Printed in Poland.

- Kup książkę
- Poleć książkę
- Oceń książkę

- Księgarnia internetowa
- Lubię to! » Nasza społeczność

Spis treści |

O autorach	11
O recenzentach	12
Przedmowa	13

CZĘŚĆ 1. Wprowadzenie do inżynierii detekcji

ROZDZIAŁ 1

Podstawy inżynierii detekcji	19
Podstawowe pojęcia	19
Unified Kill Chain	20
Framework MITRE ATT&CK	22
Piramida bólu	24
Rodzaje cyberataków	25
Motywacja dla inżynierii detekcji	28
Definicja inżynierii detekcji	30
Ważne cechy wyróżniające	32
Wartość programu inżynierii detekcji	32
Potrzeba zapewnienia lepszej wykrywalności	33
Cechy dobrego wykrywania zagrożeń	33
Korzyści z programu inżynierii detekcji	35
Przewodnik korzystania z tej książki	37
Struktura książki	37
Ćwiczenia praktyczne	38
Podsumowanie	39

ROZDZIAŁ 2

Cykl życia inżynierii detekcji	40
Faza 1. Odkrywanie wymagań	41
Charakterystyka kompletnego wymagania mechanizmu detekcji	42
Źródła wymagań dla mechanizmów detekcji	43

Ćwiczenie. Źródła wymagań dotyczących mechanizmów detekcji w Twojej organizacji	48
Faza 2. Selekcja	49
Dotkliwość zagrożenia	50
Dopasowanie mechanizmu detekcji zagrożenia do organizacji	50
Pokrycie zagrożeń mechanizmami detekcji	50
Aktywne exploity	51
Faza 3. Analiza	52
Określenie źródła danych	52
Ustalenie typów wskaźników wykrycia	53
Kontekst badawczy	53
Ustalenie kryteriów walidacji	55
Faza 4. Programowanie	55
Faza 5. Testowanie	56
Rodzaje danych testowych	57
Faza 6. Wdrażanie	58
Podsumowanie	59

ROZDZIAŁ 3

Budowa laboratorium testowego inżynierii detekcji	60
Wymagania techniczne	61
Elastic Stack	61
Wdrażanie systemu Elastic Stack za pomocą Dockera	63
Konfiguracja Elastic Stack	68
Konfiguracja narzędzia Fleet Server	72
Instalacja i konfiguracja systemu Fleet Server	73
Dodatkowe konfiguracje dla komponentu Fleet Server	75
Dodawanie hosta do laboratorium	77
Zasady komponentu Elastic Agent	83
Tworzenie pierwszego mechanizmu detekcji	85
Dodatkowe zasoby	87
Podsumowanie	88

CZĘŚĆ 2. Tworzenie mechanizmów detekcji

ROZDZIAŁ 4

Źródła danych inżynierii detekcji	91
Wymagania techniczne	92
Źródła danych i telemetrii	92
Nieprzetworzona telemetria	92
Narzędzia zabezpieczeń	100

Źródła danych MITRE ATT&CK	101
Identyfikacja źródeł danych	102
Analiza problemów i wyzwań związanych ze źródłami danych	104
Kompletność	104
Jakość	105
Terminowość	105
Pokrycie	106
Ćwiczenie. Więcej informacji o źródłach danych	106
Dodawanie źródeł danych	107
Ćwiczenie. Dodawanie źródła danych serwera WWW	107
Podsumowanie	116
Lektura uzupełniająca	117

ROZDZIAŁ 5

Analiza wymagań dla mechanizmów detekcji	118
Przegląd faz wymagań dla mechanizmów detekcji	118
Odkrywanie wymagań dla mechanizmów detekcji	119
Narzędzia i procesy	120
Ćwiczenie. Odkrywanie wymagań w organizacji	122
Selekcja wymagań dla mechanizmów detekcji	124
Dotkliwość zagrożenia	124
Dopasowanie zagrożenia do organizacji	125
Pokrycie wymagań dla mechanizmów detekcji	126
Aktywne eksploity	126
Obliczanie priorytetu	127
Analiza wymagań dla mechanizmów detekcji	130
Podsumowanie	132

ROZDZIAŁ 6

Tworzenie mechanizmów detekcji przy użyciu wskaźników naruszeń zabezpieczeń	133
Wymagania techniczne	134
Wykorzystanie wskaźników naruszenia zabezpieczeń	134
Przykładowy scenariusz. Identyfikacja kampanii IcedID przy użyciu wskaźników	137
Ćwiczenie	146
Instalacja i konfigurowanie systemu Sysmon jako źródła danych	146
Wykrywanie skrótów	148
Mechanizmy detekcji wskaźników sieciowych	151
Podsumowanie ćwiczenia	155

Podsumowanie	155
Lektura uzupełniająca	155

ROZDZIAŁ 7

Opracowywanie mechanizmów detekcji opartych na wskaźnikach behawioralnych	156
Wymagania techniczne	156
Wykrywanie narzędzi przeciwnika	156
Przykładowy scenariusz. Użycie narzędzia PsExec	157
Wykrywanie taktyk, technik i procedur (TTP)	173
Przykładowy scenariusz. Technika omijania kontroli znacznika sieci	174
Podsumowanie	179

ROZDZIAŁ 8

Tworzenie dokumentacji i potoki mechanizmów detekcji	181
Dokumentowanie mechanizmu detekcji	181
Ćwiczenie. Dokumentowanie mechanizmu detekcji	184
Analiza repozytorium mechanizmów detekcji	187
Mechanizm detekcji jako kod	190
Wyzwania związane z tworzeniem potoku mechanizmu detekcji	199
Ćwiczenie. Publikowanie reguły przy użyciu projektu mechanizmów detekcji Elastic	200
Podsumowanie	209

CZĘŚĆ 3. Walidacja mechanizmów detekcji

ROZDZIAŁ 9

Walidacja mechanizmów detekcji	213
Wymagania techniczne	214
Czym jest proces walidacji?	214
Na czym polegają ćwiczenia zespołu purple team?	216
Symulowanie aktywności przeciwnika	217
Atomic Red Team	218
CALDERA	219
Ćwiczenie. Walidacja mechanizmów detekcji dla pojedynczej techniki z wykorzystaniem Atomic Red Team	220
Ćwiczenie. Walidacja mechanizmów detekcji dla wielu technik z wykorzystaniem systemu CALDERA	226

Korzystanie z wyników walidacji	232
Pomiar pokrycia zagrożeń mechanizmami detekcji	234
Podsumowanie	241
Lektura uzupełniająca	241

ROZDZIAŁ 10

Wykorzystanie wiedzy o zagrożeniach	242
Wymagania techniczne	242
Przegląd zagadnień związanych z wiedzą o zagrożeniach	243
Wiedza o zagrożeniach typu open source	243
Wewnętrzne źródła wiedzy o zagrożeniach	245
Zbieranie wiedzy o zagrożeniach	245
Wiedza o zagrożeniach w cyklu życia inżynierii detekcji	246
Odkrywanie wymagań	246
Selekcja	246
Analiza	248
Wiedza o zagrożeniach na potrzeby inżynierii detekcji w praktyce	248
Przykład. Wykorzystywanie na potrzeby inżynierii detekcji wpisów na blogach z informacjami o zagrożeniach	249
Przykład. Wykorzystanie systemu VirusTotal na potrzeby inżynierii detekcji	252
Ocena zagrożeń	255
Przykład. Wykorzystanie oceny zagrożeń na potrzeby inżynierii detekcji	256
Zasoby i dalsza lektura	262
Źródła i pojęcia związane z wiedzą o zagrożeniach	262
Skanery online i piaskownice	263
MITRE ATT&CK	263
Podsumowanie	263

CZĘŚĆ 4. Metryki i zarządzanie

ROZDZIAŁ 11

Zarządzanie wydajnością	267
Wprowadzenie do zarządzania wydajnością	267
Ocena dojrzałości mechanizmu detekcji	268
Pomiar wydajności programu inżynierii detekcji	270
Pomiar skuteczności programu inżynierii detekcji	272
Priorytetyzacja prac związanych z detekcją	274
Trafność, hałaśliwość i czułość	276

Obliczanie skuteczności mechanizmu detekcji	280
Metryki pokrycia o niskiej wierności	280
Automatyczna walidacja	282
Metryki pokrycia o wysokiej wierności	282
Podsumowanie	296
Lektura uzupełniająca	296

CZĘŚĆ 5. Kariera w inżynierii detekcji

ROZDZIAŁ 12

Wskazówki dotyczące kariery w inżynierii detekcji	301
Zdobycie pracy w branży inżynierii detekcji	301
Oferty pracy	302
Rozwijanie umiejętności	303
Inżynier detekcji jako zawód	307
Role i obowiązki inżyniera detekcji	309
Przyszłość inżynierii detekcji	310
Powierzchnie ataku	310
Widoczność	311
Możliwości urzędzeń zabezpieczeń	311
Uczenie maszynowe	312
Współdzielenie metodologii ataków	313
Przeciwnik	313
Człowiek	313
Podsumowanie	314

Podstawy inżynierii detekcji

Rozdział 1

Najważniejszą troską kadry kierowniczej i członków zarządu organizacji w niemal każdej branży jest bezpieczeństwo jej zasobów cyfrowych. Wziąwszy pod uwagę fakt, że współczesne firmy jak nigdy wcześniej zależą od komunikacji i technologii, jest to zrozumiała obawa. Zasoby cyfrowe i ich infrastruktura pomocnicza stanowią coraz większą część zasobów typowej organizacji. Ponadto coraz więcej procesów zależy od niezawodnych technologii komunikacyjnych. W większości przypadków technologie informatyczne umożliwiają firmom bardziej efektywne działanie. Zarządzanie i obrona cyfrowego krajobrazu może jednak być wyzwaniem dla organizacji każdej wielkości.

Ponadto, o ile kiedyś wyrafinowane ataki były ograniczone do działań agentów z obcych państw, o tyle zwiększone wzajemne powiązania technologii w połączeniu z pojawieniem się kryptowalut tworzą niemal idealne środowisko do działania cyberprzestępców. Dodanie wyrafinowanych aktorów zagrożeń motywowanych zyskiem finansowym, a nie tylko wykradaniem informacji państwowych, znacznie poszerzyło liczbę organizacji, które muszą identyfikować i reagować na takie zagrożenia. Powstrzymanie tych ataków wymaga od organizacji zwiększonej elastyczności w walce z przeciwnikiem. Program inżynierii detekcji, który poprawia tempo działania organizacji w dziedzinie operacjonalizacji informacji o nowych zagrożeniach, zapewnia tę elastyczność. Głównym celem inżynierii detekcji jest opracowanie reguł lub modeli algorytmicznych w celu automatycznego identyfikowania obecności podmiotów stanowiących zagrożenie lub ogólnej złośliwej aktywności, tak aby odpowiednie zespoły mogły podjąć działania łagodzące.

W tym rozdziale znajdziesz kilka tematów, które dostarczą Ci wiedzy niezbędnej to studiowania dalszej części tej książki:

- Podstawowe pojęcia: frameworki ataków, ich typowe rodzaje i definicja inżynierii detekcji.
- Wartość programu inżynierii detekcji.
- Przegląd informacji zaprezentowanych w tej książce.

Podstawowe pojęcia

Podstawowa umiejętność śledzenia i kategoryzowania działań przeciwnika pozwala ustalić priorytety i zrozumieć zakres lub obszar oddziaływania mechanizmów wykrywania zagrożeń. W poniższym punkcie opisano popularne frameworki i modele, które

będą wykorzystywane w dalszej części tej książki. Zapewniają one model wyjściowy do rozpoznawania cyberataków i ich szczegółowych składników oraz sposobów obrony przed nimi.

Unified Kill Chain

Cyberataki zwykle są przeprowadzane zgodnie z przewidywalnym wzorcem, który powinien być zrozumiały dla obrońców. Wzorec ten został początkowo udokumentowany jako słynny już Cyber Kill Chain (cybernetyczny łańcuch zabójstw) opracowany przez firmę Lockheed Martin. Z czasem model ten został zaadaptowany i zmodernizowany przez wiele firm. Znaczącą modernizację tego modelu stanowi **Unified Kill Chain** (zunifikowany łańcuch zabójstw). Model ten definiuje 18 ogólnych taktyk na przestrzeni trzech ogólnych celów. Dzięki niemu obrońcy uzyskują rozsądny framework projektowania działań obronnych zgodnie z celami atakujących. Przyjrzyjmy się tym celom:

- **Wejście** (ang. *in*). Celem atakującego na tym etapie jest zbadanie potencjalnej ofiary, odkrycie możliwych wektorów ataku oraz uzyskanie i utrzymanie niezawodnego dostępu do środowiska docelowego.
- **Wewnątrz** (ang. *through*). Po uzyskaniu dostępu do środowiska docelowego napastnik musi się zorientować w środowisku i zebrać dodatkowe zasoby wymagane do dalszej części ataku, np. zdobyć poświadczenia umożliwiające uprzywilejowany dostęp.
- **Wyjście** (ang. *out*). Taktyki z tej grupy koncentrują się na osiągnięciu celów cyberataku. W przypadku oprogramowania ransomware z podwójnym wymuszeniem (ang. *double extortion ransomware*) obejmowałyby to przygotowanie plików do wymuszenia, skopiowanie tych plików do infrastruktury atakującego i wreszcie wdrożenie oprogramowania ransomware na dużą skalę.

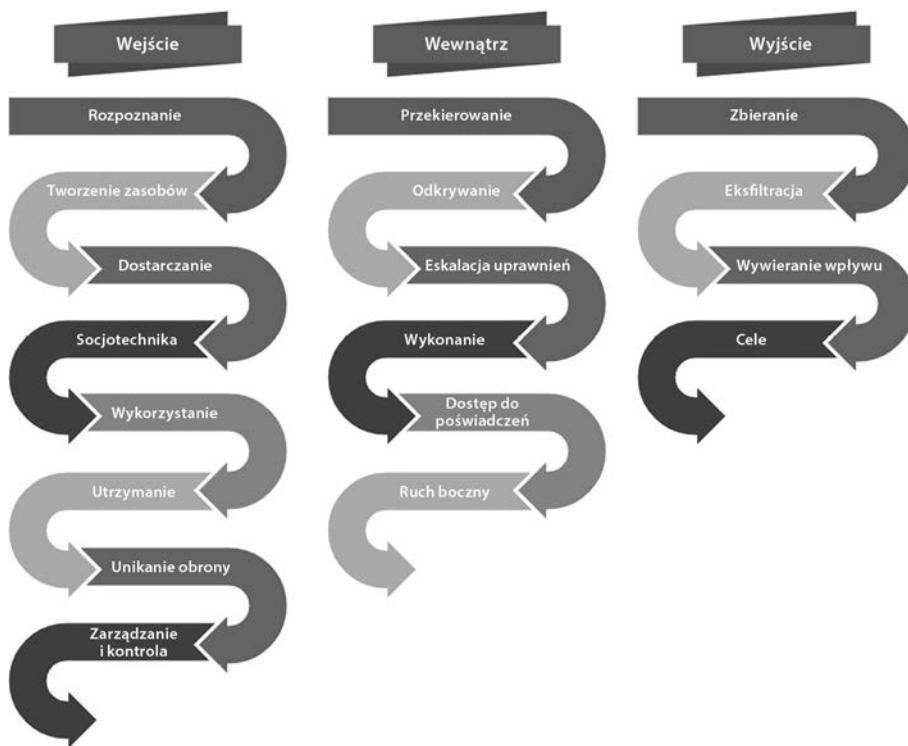
Poszczególne taktyki w każdej fazie „łańcucha zabójstw” pokazano na rysunku 1.1, opartym na „białej księdze” modelu *Unified Kill Chain* autorstwa Paula Polsa.

Aby lepiej zrozumieć, w jaki sposób model Unified Kill Chain ma zastosowanie do cyberataków, przyjrzyjmy się, jak przekłada się on na dobrze znany atak. W szczególności przyjrzymy się kampanii ataków Emotet. Emotet to złośliwy ładunek często rozpowszechniany za pośrednictwem poczty elektronicznej i wykorzystywany do dostarczania dodatkowych ładunków, które mają realizować ostateczne cele napastników. Konkretna kampania, którą przeanalizujemy, została opisana w raporcie *The DFIR Report* w listopadzie 2022 roku: <https://thedfirreport.com/2022/11/28/emotet-strikes-again-lnk-file-leads-to-domain-wide-ransomware/>.

Listę etapów ataku podaną w artykule oraz ich odwzorowanie na fazy modelu Unified Kill Chain zestawiono w tabeli 1.1.

Jak widać w tabeli 1.1, nie wszystkie fazy są wykorzystywane w każdym ataku, a poszczególne fazy nie muszą występować po kolei.

Pełny tekst „białej księgi” modelu *Unified Kill Chain* można znaleźć pod adresem: <https://www.unifiedkillchain.com/assets/The-Unified-Kill-Chain.pdf>.



Rysunek 1.1. Model Unified Kill Chain

Tabela 1.1. Odzworowanie etapów modelu Unified Kill Chain dla łańcucha ataków Emotet

Zdarzenie ataku	Grupa faz modelu Unified Kill Chain	Faza modelu Unified Kill Chain
Emotet uruchamiany za pośrednictwem załącznika złośliwego spamu LNK	Wejście	Dostarczanie
Emotet wysyła przez SMTP wychodzące wiadomości spamowe	Wewnątrz	Przekierowanie
Wyliczenie domen za pomocą narzędzia Cobalt Strike	Wewnątrz	Odkrywanie
Ruch boczny do stacji roboczej użytkownika	Wewnątrz	Przekierowanie
Wyliczenie udziałów SMB	Wewnątrz	Odkrywanie
Próba wykorzystania eksploatu zerologon	Wejście	Wykorzystanie
Zainstalowanie agenta zdalnego zarządzania	Wejście	Dowodzenie i kontrola/utrzymanie
Eksfiltracja za pomocą narzędzia Rclone do chmury Mega	Wyjście	Eksfiltracja
Uruchamianie oprogramowania ransomware	Wyjście	Oddziaływanie

Chociaż przedstawiony scenariusz jest zgodny z przebiegiem typowego cyberataku, to jak zaprezentowano w artykule i jak pokazują późniejsze przykłady, nierzadko zdarza się, że atakujący wykonuje pewne taktyki poza oczekiwaną kolejnością. Podczas gdy Unified Kill Chain określa model przeprowadzania ataków przez aktorów zagrożeń, nie zagłębia się w szczegółowe techniki, które mogą być wykorzystane do osiągnięcia celów każdej fazy „łańcucha zabójstw”. Dokładniejszy opis taktyk, technik i procedur wykorzystywanych przez aktorów zagrożeń zapewnia framework MITRE ATT&CK.

Framework MITRE ATT&CK

Framework MITRE ATT&CK jest bazą wiedzy opracowaną przez MITRE Corporation. Framework klasyfikuje cele aktorów zagrożeń i kataloguje szczegółowe narzędzia i działania związane z osiągnięciem tych celów.

ATT&CK to skrót od **Adversarial Tactics, Techniques i Common Knowledge** (taktyki przeciwnika, techniki i powszechna wiedza). Framework MITRE ATT&CK grupuje techniki przeciwnika w wysokopoziomowe kategorie zwane taktykami. Każda taktyka reprezentuje bardziej szczegółowy bezpośredni cel w ramach ogólnego cyberataku. Ten framework, który określa skuteczny model projektowania i walidacji mechanizmów wykrywania zagrożeń, będzie często przywoływany w dalszej części tej książki. Wysokopoziomowe taktyki zawarte w ramach frameworka Enterprise ATT&CK szczegółowo opisano w poniższych punktach:

- **Rozpoznanie** (ang. *Reconnaissance*). Ta taktyka wchodzi w zakres fazy „wejście” określanej także fazą „uchwycenia przyczółku” (ang. *initial foothold*) modelu Unified Kill Chain. Na tym etapie napastnik zbiera informacje o wybranym celu. Może używać narzędzi do pasywnego zbierania szczegółowych informacji technicznych o celu, takich jak infrastruktura dostępna publicznie, e-maile, partnerzy podatni na ataki itp. W idealnych przypadkach aktor zagrożenia może zidentyfikować publicznie dostępne i podatne na ataki interfejsy, ale rozpoznanie może również obejmować zbieranie informacji o pracownikach organizacji w celu zidentyfikowania potencjalnych celów socjotechnicznych i zrozumienia sposobu działania różnych wewnętrznych procesów biznesowych.
- **Rozwijanie zasobów** (ang. *resource development*). Taktyka wchodzi w zakres fazy „wejście” modelu Unified Kill Chain. Po zidentyfikowaniu wiarygodnego wektora ataku aktorzy zagrożeń projektują odpowiedni atak i programują techniczne zasoby w celu ułatwienia jego przeprowadzenia. Faza ta obejmuje tworzenie, zakup lub kradzież poświadczeń, infrastruktury lub zdolności przeznaczonych do wsparcia operacji przeciwko celowi.
- **Początkowy dostęp** (ang. *initial access*). Taktyka wchodzi w zakres fazy „wejście” modelu Unified Kill Chain. Aktor zagrożenia próbuje uzyskać dostęp do zasobu w środowisku kontrolowanym przez ofiarę. W tej fazie mogą być wykorzystywane różne narzędzia, począwszy od inteligentnie zaprojektowanych kampanii phishingowych, a skończywszy na wdrażaniu kodu, który wykorzystuje nieujawnione luki w odsłoniętych interfejsach oprogramowania (są to tzw. ataki dnia zerowego — ang. *zero-day attacks*).

- **Wykonanie** (ang. *execution*). Taktyka wchodzi w zakres faz „wejście” oraz „wewnątrz” (nazywanej także fazą propagacji sieciowej — ang. *network propagation*) modelu Unified Kill Chain. Celem atakującego jest uruchomienie swojego kodu na docelowym zasobie. Kod wykorzystywany w tej fazie zazwyczaj próbuje zebrać dodatkowe szczegóły na temat docelowej sieci, zrozumieć kontekst bezpieczeństwa, w którym działa kod, lub zebrać dane i zwrócić je do infrastruktury kontrolowanej przez aktora zagrożenia.
- **Utrzymanie** (ang. *persistence*). Taktyka wchodzi w zakres fazy „wejście” modelu Unified Kill Chain. Początkowy dostęp do obcego środowiska może być niestabilny. Aktorzy zagrożeń preferują solidny i możliwy do utrzymania dostęp do docelowych systemów. Techniki utrzymania koncentrują się na zachowaniu dostępu pomimo restartów systemu lub modyfikacji tożsamości bądź infrastruktury.
- **Eskalacja uprawnień** (ang. *privilege escalation*). Taktyka wchodzi w zakres fazy „wewnątrz” modelu Unified Kill Chain. Po uzyskaniu dostępu do środowiska kontrolowanego przez ofiarę aktor zagrożeń zazwyczaj próbuje uzyskać najwyższy możliwy poziom uprawnień. Uprzywilejowany dostęp pozwala na skorzystanie z niemal każdej opcji dostępnej dla administratorów systemu ofiary. W ten sposób aktor zagrożeń usuwa wiele przeszkód, które mogą zablokować aktorom zagrożeń możliwość podjęcia działań zmierzających do osiągnięcia ich celów. Posiadanie uprzywilejowanego dostępu może również utrudnić wykrycie działań aktorów zagrożeń.
- **Unikanie obrony** (ang. *defense evasion*). Taktyka ta wchodzi w zakres fazy „wejście” modelu Unified Kill Chain. Aktorzy zagrożeń muszą poznać systemy obronne ofiar, aby opracować odpowiednie metody ich unikania. Skuteczne unikanie obrony zwiększa prawdopodobieństwo powodzenia ataku. Taktyki te koncentrują się w szczególności na znalezieniu sposobów na pokonanie mechanizmów obrony lub unikanie ich w inny sposób.
- **Dostęp do poświadczeń** (ang. *credential access*). Taktyka ta wchodzi w zakres faz „wejście” i „wyjście” (nazywanej też fazą „oddziaływania na cele” — ang. *action on objectives*) modelu Unified Kill Chain. Dostęp do systemów jest kontrolowany przez tożsamości. Zbieranie poświadczeń lub materiałów uwierzytelniających jest w istocie niezbędne do całkowitego zdominowania środowiska ofiary. Dostęp do wielu systemów i danych uwierzytelniających ułatwia poruszanie się w środowisku, a w przypadku modyfikacji danych uwierzytelniających pozwala napastnikom na przeprowadzenie przekierowania.
- **Odkrywanie** (ang. *discovery*). Taktyka wchodzi w zakres fazy „wewnątrz” modelu Unified Kill Chain. Techniki wchodzące w skład tej taktyki koncentrują się na zapoznaniu się z wewnętrznym środowiskiem ofiary. Do zaplanowania pozostałych faz ataku potrzebne jest rozpoznanie wewnętrznego układu sieci, konfiguracji infrastruktury, informacji o tożsamości i mechanizmów obronnych.
- **Ruch boczny** (ang. *lateral movement*). Ta taktyka wchodzi w zakres fazy „wyjście” modelu Unified Kill Chain. Systemy, do których aktorzy zagrożeń uzyskują dostęp po raz pierwszy, często nie zawierają informacji lub zasobów (narzędzi, materiałów uwierzytelniających, bezpośredniej łączności lub

widoczności) wymaganych do realizacji celów. Po odkryciu powiązanych systemów i zdobyciu odpowiednich danych uwierzytelniających napastnik może i często musi przenieść się z bieżącego systemu do innych, połączonych systemów. Wszystkie te techniki koncentrują się na poruszaniu się po środowisku ofiary.

- **Zbieranie** (ang. *collection*). Ta taktyka wchodzi w zakres fazy „wyjście” modelu Unified Kill Chain. Techniki wchodzące w jej skład koncentrują się na przeprowadzaniu wewnętrznego rozpoznania. Dostęp do nowych środowisk zapewnia nową widoczność, a do planowania kolejnych faz ataku niezbędne jest zrozumienie środowiska technicznego.
- **Zarządzanie i kontrola** (ang. *command and control*). Ta taktyka wchodzi w zakres fazy „wejście” modelu Unified Kill Chain. Pozwala na wdrożenie systemów, dzięki którym można zdalnie kontrolować środowisko ofiary.
- **Eksfiltracja**. Ta taktyka wchodzi w zakres fazy „wyjście” modelu Unified Kill Chain. Nie wszystkie ataki obejmują działania wymuszające, ale taktyki należące do tej kategorii stały się bardziej popularne wraz ze wzrostem liczby ataków ransomware z podwójnym wymuszeniem. Więcej informacji na temat ataków ransomware z podwójnym wymuszeniem można znaleźć na stronie <https://www.zscaler.com/resources/security-terms-glossary/what-is-double-extortion-ransomware>. Techniki wchodzące w skład tej taktyki mają na celu skopiowanie danych ze środowiska ofiary do infrastruktury kontrolowanej przez atakującego.
- **Oddziaływanie** (ang. *impact*). Ta taktyka wchodzi w zakres fazy „wyjście” modelu Unified Kill Chain. W tym momencie aktor zagrożenia może podjąć kroki w celu finalizacji ataku. Na przykład w przypadku ataku ransomware do tej fazy należy szyfrowanie danych na dużą skalę.

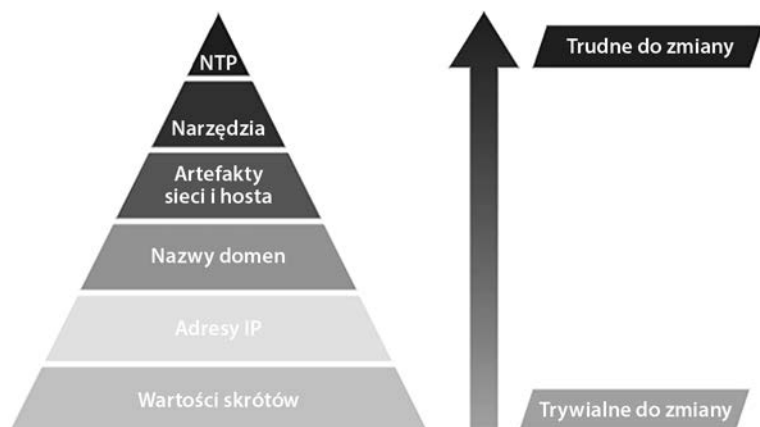
Zachęcamy do dokładniejszego zapoznania się z frameworkiem MITRE ATT&CK na stronie <https://attack.mitre.org/>. W tej książce skupimy się na frameworku Enterprise ATT&CK, ale MITRE zapewnia również frameworki dla ataków ICS i ataków mobilnych. Do szybkiego wyszukiwania i kwalifikowania taktyk przez obrońców niezwykle przydatny jest również Nawigator ATT&CK, dostępny na stronie <https://mitre-attack.github.io/attack-navigator/>.

Większość publikacji dokumentujących obserwacje reakcji na incydenty zazwyczaj odnosi się do taktyk modeli Unified Kill Chain i MITRE ATT&CK. Rozumienie tych taktyk pomaga obrońcom w lepszym zaprojektowaniu mechanizmów wykrywania oraz innych mechanizmów prewencyjnych.

Piramida bólu

Innym schematem pomocnym dla obrońców jest model **Pyramid of Pain** (piramida bólu). Jest to model opracowany przez Davida Bianco, który wizualizuje związek między kategoriami wskaźników i trudnością obrony każdego z nich. Ta trudność jest wyrażana

jako wysiłek aktora zagrożeń potrzebny do zmodyfikowania ataku po wdrożeniu skutecznej obrony dla wskaźnika danej kategorii. Pojęcie modelu piramidy bólu przedstawiono na rysunku 1.2.



Rysunek 1.2. Model piramidy bólu Davida Bianco

Jak widać, uniknięcie mechanizmów zaprojektowanych do działania na statycznych wskaźnikach, takich jak nazwy domen, adresy IP i wartości skrótów, jest dla napastników trywialne. Na przykład modyfikacja binarnego skrótu sprowadza się po prostu do zmiany pojedynczego bitu. Przeciwnikowi znacznie trudniej jest zmodyfikować **narzędzia, taktyki i procedury (NTP)**, które w istocie stanowią podstawę jego podręcznika ataków. Mechanizmy obronne ukierunkowane na warstwę NTP są „złotym standardem”. Zwykle są one jednak trudniejsze do zaimplementowania i wymagają wiarygodnych danych z chronionych zasobów, a także dogłębnego zrozumienia stosowanej taktyki i możliwości przeciwnika. Mechanizmy obronne zaprojektowane w odniesieniu do statycznych wskaźników są skuteczne w krótkoterminowej, taktycznej obronie. Z pełnym wpisem na blogu Davida Bianco można zapoznać się pod adresem: <https://detect-respond.blogspot.com/2013/03/the-pyramid-of-pain.html>.

W dalszej części tej książki będziemy często odwoływać się do omówionych powyżej pojęć. W późniejszych rozdziałach zilustrujemy sposób, w jaki można wykorzystać te modele do zrozumienia cyberataków, przełożenia wysokopoziomowych celów biznesowych w zakresie obrony na wykrywanie i mierzenie pokrycia dla znanych ataków.

Po zapoznaniu się z modelem opisywania cyberataków przyjrzymy się najczęstszym ich rodzajom.

Rodzaje cyberataków

Aby inżynierowie detekcji mogli wykrywać cyberataki, muszą posiadać podstawową wiedzę na temat ataków, z którymi chcą walczyć. Poniżej zestawiono niektóre z najbardziej rozpowszechnionych ataków w momencie pisania tego tekstu. Lista ta może posłużyć za wstępną klasyfikację ataków, przed którymi próbujemy się bronić.

Przechwytywanie biznesowej poczty e-mail

W 2021 roku FBI zgłosiło otrzymanie łącznie 19 954 skarg związanych z incydentami ataków przechwytywania **biznesowej poczty e-mail** — ang. *business email compromise* — **BEC**. Skumulowane straty poniesione w wyniku tych ataków wyniosły 2,4 miliarda dolarów (USD). Pełny raport jest dostępny pod adresem https://www.ic3.gov/Media/PDF/AnnualReport/2021_IC3Report.pdf.

Ataki BEC są wymierzone w użytkowników najpopularniejszego i najbardziej dostępnego narzędzia do współpracy — poczty elektronicznej. W wielu organizacjach normalną częścią operacji biznesowych jest elektroniczny transfer funduszy. Aktorzy zagrożeń badają organizacje i identyfikują personel, który może być zaangażowany w korespondencję związaną z wymianą funduszy. Po zidentyfikowaniu celu napastnik wykorzystuje jedną z kilku technik, aby uzyskać dostęp do skrzynki pocztowej celu (lub kogoś bliskiego z perspektywy procesu biznesowego). Zdobycie takiego dostępu pozwala przenieść cel aktora zagrożeń. Zaczyna on obserwować wymianę wiadomości e-mail w celu zrozumienia wewnętrznych procesów. W tym czasie aktor zagrożeń dąży do zapoznania się z kanałami komunikacji i najważniejszymi jej podmiotami. W idealnych przypadkach stara się zidentyfikować zewnętrznych kontrahentów, z którymi organizacja prowadzi rutynowe interesy. Zwykle są to osoby, które wysyłają korespondencję w sprawie płatności, oraz osoby, które zatwierdzają te płatności w imieniu organizacji. Gdy nadarzy się odpowiednia okazja, aktor zagrożeń próbuje przechwycić i zmodyfikować konwersacje e-mail dotyczące płatności (np. poprzez zmianę numerów docelowych rachunków). Jeśli to działanie pozostanie niezauważone, może dojść do przelania środków na rachunek napastnika zamiast na konto zamierzonego odbiorcy.

Odmowa usługi (DoS)

Celem ataków typu **Denial of service (DoS** — odmowa usługi) jest zablokowanie usługi dla jej prawowitych użytkowników poprzez przeciążenie usługi lub osłabienie w inny sposób infrastruktury, od której ta usługa zależy. Istnieją trzy główne rodzaje ataków DoS: ataki wolumetryczne, ataki na protokoły i ataki na aplikacje.

Ataki wolumetryczne polegają na wysyłaniu nadmiernej ilości ruchu do systemu docelowego. Jeśli atak trwa wystarczająco długo, może doprowadzić do pogorszenia jakości usługi bądź jej całkowitego zablokowania. Ataki na protokoły koncentrują się na warstwie sieciowej i transportowej i próbują wyczerpać te zasoby urządzeń sieciowych, które zapewniają dostęp do docelowej usługi. Ataki na aplikacje wysyłają duże ilości żądań do usługi docelowej. Usługa próbuje przetworzyć każde żądanie, co zużywa moc obliczeniową jej podstawowych systemów. W końcu dochodzi do wyczerpania dostępnych zasobów, a czas reakcji usługi wydłuża się do takiego stopnia, że usługa staje się niedostępna. Te rodzaje ataków można dalej kategoryzować według stopnia ich automatyzacji oraz stosowanych technik.

Zwiększenie liczby systemów wykonujących atak może znacznie zwiększyć jego siłę oddziaływania. Za pomocą zainfekowanych systemów można przeprowadzać zsynchronizowane ataki DoS na pojedynczy cel, znane jako **rozproszone ataki DoS** (ang. *Distributed Denial of Service* — **DDoS**).

Epidemia złośliwego oprogramowania

Skutki uniknięcia mechanizmów obrony przez złośliwe oprogramowanie (określane również terminem **malware** — od ang. *malicious software*) mogą być w zależności od konkretnej rodziny złośliwego oprogramowania bardzo poważne. W przypadkach o niskiej sile oddziaływania użytkownik docelowy może być bombardowany niechcianymi wyskakującymi reklamami, a w bardziej ekstremalnych scenariuszach złośliwe oprogramowanie może dać zdalnemu aktorowi zagrożenia pełną kontrolę nad systemem. Obecność złośliwego oprogramowania w środowisku korporacyjnym zwykle wskazuje na możliwy brak mechanizmów zabezpieczeń. Pozornie mało szkodliwe infekcje złośliwym oprogramowaniem mogą prowadzić do poważniejszych incydentów, w tym ataków ransomware.

Zagrożenia wewnętrzne

Złośliwe działania przeciwko organizacji wykonywane przez jej pracowników są określane jako zagrożenia wewnętrzne. Mogą one występować na każdym poziomie organizacji i wynikać z różnych motywacji. Obrona przed tego rodzaju złośliwymi intruzami jest trudna, ponieważ organizacja obdarzyła ich pewnym stopniem zaufania.

Phishing

Ataki phishingowe należą do kategorii działań socjotechnicznych, w których aktorzy zagrożeń atakują narzędzia komunikacji i współpracy, takie jak poczta elektroniczna, komunikatory internetowe, systemy przesyłania wiadomości tekstowych (SMS), a nawet zwykłe rozmowy telefoniczne. Podstawowym celem we wszystkich przypadkach jest nakłonienie użytkowników do ujawnienia poufnych informacji, takich jak dane uwierzytelniające lub informacje bankowe. Ataki BEC zazwyczaj wykorzystują techniki phishingu.

Ransomware

O ile krajobraz współczesnych zagrożeń obejmuje różnych aktorów, z różnymi celami — od szpiegostwa cybernetycznego, po oszustwa związane z pomocą techniczną — najbardziej rozpowszechnionymi i najmocniej oddziałującymi na ofiary są nowoczesne ataki ransomware.

Celem ataku ransomware jest przerwanie kluczowych operacji biznesowych poprzez wyłączenie najważniejszych systemów i zażądanie od organizacji płatności lub okupu. W zamian za udaną płatność aktorzy zagrożeń obiecują przywrócenie systemów do normalnego działania.

W ostatnim czasie niektórzy operatorzy oprogramowania ransomware dodali do swojego podręcznika osobny komponent wymuszeń. Podczas ataku ransomware aktorzy zagrożeń przenoszą wrażliwe dane ze środowiska organizacji do kontrolowanych przez siebie systemów. Następnie operatorzy oprogramowania ransomware grożą upublicznieniem tych danych, jeśli nie zostanie zapłacony okup. Tego rodzaju ataki są powszechnie określane jako ataki ransomware z podwójnym wymuszeniem.

Udane operacje ransomware stawiają zaatakowane firmy w przerażającej sytuacji. Nie tylko muszą one podjąć decyzję, czy zapłacić okup, ale także odzyskać dane, co może zająć wiele miesięcy, a czasem nawet lat.

Te złośliwe operacje z czasem stały się coraz bardziej wyrafinowane i skuteczne. Według CrowdStrike pierwszy przypadek nowoczesnego ataku ransomware został zarejestrowany w 2005 roku. Od tego czasu częstotliwość, skala i wyrafinowanie ataków ransomware tylko wzrosły. Podsumowanie ewolucji oprogramowania ransomware można znaleźć w artykule *History of Ransomware* w serwisie CrowdStrike. Pełny artykuł można przeczytać pod adresem: <https://www.crowdstrike.com/cybersecurity-101/ransomware/history-of-ransomware/>.

Motywacja dla inżynierii detekcji

Skutki udanych naruszeń zabezpieczeń mogą być kosztowne, a ich usunięcie często wymaga tysięcy roboczogodzin. Raport IBM „2022 Cost of a Data Breach” wykazał, że średni całkowity koszt udanego ataku na dane wyniósł 4,35 mln dolarów. Zazwyczaj im wcześniej zostanie wykryte zagrożenie, tym koszt jego usunięcia jest niższy. W każdej kolejnej fazie łańcucha zabójstw koszt naprawy wzrasta. O ile polowanie na zagrożenia (ang. *threat hunt*) pozwala organizacji szukać przeciwników już w jej środowisku, o tyle identyfikacja następuje wtedy, gdy zostanie przeprowadzone wyszukiwanie. Takie wykrywanie pozwala jednak organizacji zidentyfikować złośliwe zachowanie w momencie jego wykonywania, co przyczynia się do skrócenia średniego czasu do wykrycia zagrożenia. Biorąc pod uwagę, że we wspomnianym wcześniej raporcie IBM „Cost of a Data Breach” ustalono, że średni czas identyfikacji i powstrzymania naruszenia wynosił 277 dni, jest wiele do zrobienia, aby skrócić ten czas.

Aby przekonać się, że czas wykrycia ataku w znacznym stopniu determinuje siłę jego oddziaływania na firmę, rozważmy scenariusz, w którym aktor zagrożenia uzyskał dostęp do stacji roboczej podłączonej do internetu za pośrednictwem udanej kampanii phishingowej. Ten nieautoryzowany dostęp został natychmiast wykryty przez zespół ds. bezpieczeństwa organizacji. Członkowie zespołu szybko odizolowali zaatakowaną stację roboczą i przeprowadzili pełne odtworzenie obrazu jej zawartości do znanego, dobrego stanu. Wykonali również pełny reset poświadczeń użytkownika oraz wszystkich użytkowników, którzy wchodzili w interakcję z tą stacją roboczą. Administratorzy zidentyfikowali wiadomość phishingową w korporacyjnym programie do zarządzania pocztą e-mail, a wszyscy odbiorcy ponownie zainstalowali swoje stacje robocze i zresetowali swoje dane uwierzytelniające.

W tym scenariuszu kroki podjęte przez zespół ds. bezpieczeństwa były stosunkowo proste do wykonania i prawdopodobnie wystarczyłyby do usunięcia zagrożenia ze środowiska. W przeciwieństwie do tego scenariusza, gdyby aktorzy zagrożeń zdołali uzyskać uprzywilejowany dostęp, wydobyć dane, a następnie zainstalować oprogramowanie ransomware we wszystkich systemach, zadanie stałoby się znacznie bardziej uciążliwe. Zespół ds. bezpieczeństwa stanąłby przed podwójnym zadaniem obejmującym zrozumienie, co się stało, a jednocześnie proponowanie najlepszego sposobu przywrócenia zdolności firmy do bezpiecznego działania. Ilość zagrożonych zasobów, trudności dochodzenia i typowe wysiłki naprawcze w zależności od celów modelu Unified Kill Chain zestawiono w tabeli 1.2.

**Tabela 1.2. Uogólnione oddziaływanie na zasoby i wysiłek naprawczy
a cele łańcucha zabójstw**

	Uchwycenie przyczółku	Propagacja w sieci	Oddziaływanie na cele
Wpływ na aktywa	Niski. Zazwyczaj dotyczy urządzeń brzegowych, serwerów publicznych lub stacji roboczych użytkowników. Ze względu na ich pozycję w nowoczesnych architekturach urządzenia te zazwyczaj są domyślnie niezaufane.	Średni. Niektóre systemy wewnętrzne. Zazwyczaj na tym etapie aktor zagrożeń ma dostęp do niektórych serwerów członkowskich w środowisku i ma ustanowiony niezawodny kanał C2.	Wysoki. Kluczowe serwery, takie jak kontrolery domeny usługi Active Directory, serwery kopii zapasowych lub serwery plików.
Poziom kontroli aktora zagrożeń	Niski. Aktor zagrożeń ma niestabilny dostęp do systemu lub próbuje uzyskać dostęp do systemu, zazwyczaj za pomocą phishingu lub ataków na usługi dostępne publicznie. Zazwyczaj w tej fazie obrońcy mają największe szanse na usunięcie zagrożenia.	Średni. Aktor zagrożeń ma wystarczającą kontrolę, aby poruszać się po sieci, ale nie ma wystarczającej kontroli, aby realizować swoje cele. W tym momencie aktorzy zagrożeń zazwyczaj dysponują pewnymi danymi uwierzytelniającymi i mają dostęp do niezawodnego kanału C2.	Wysoki. Aktor zagrożeń czuje się w pełni komfortowo w danym środowisku. Znalazł wszystkie zasoby potrzebne do realizacji swoich celów. W tym momencie prawdopodobnie dysponuje najwyższym dostępnym w środowisku poziomem uprawnień.
Dane wymagane do przeprowadzenia dochodzenia	Stosunkowo mało. Zazwyczaj oddziaływanie na tym etapie jest ograniczone do niewielkiej liczby aktywów. Dane zidentyfikowane na tym etapie wymagane do pełnego określenia zakresu zdarzenia są ograniczone do jednego hosta.	Sporo. Zdolność do poruszania się po wewnętrznej sieci zazwyczaj wskazuje na obecność niezawodnego kanału C2. Do identyfikacji zagrożonych zasobów wymagana jest większa ilość danych historycznych i danych odbieranych w czasie rzeczywistym.	Dużo. Śledczy potrzebują dostępu do danych historycznych i danych odbieranych w czasie rzeczywistym ze wszystkich połączonych zasobów.

**Tabela 1.2. Uogólnione oddziaływanie na zasoby i wysiłek naprawczy
a cele łańcucha zabójstw — ciąg dalszy**

	Uchwycenie przyczółku	Propagacja w sieci	Oddziaływanie na cele
		W tym momencie osoby reagujące na incydenty, aby w pełni śledzić ruch boczny, muszą mieć wgląd we wszystkie połączone zasoby.	Dodatkowo w przypadkach, w których celem jest eksfiltracja danych, wymagana jest również telemetria dostępu i przepływu danych. Te dane są trudne do zebrania i zazwyczaj nie są śledzone.
Wysiłek wymagany do naprawy	Niski. Działania na tym etapie zazwyczaj dotyczą urządzeń brzegowych lub dostępnych publicznie zasobów. Zwykle takie zasoby są traktowane jako niezaufane, więc zazwyczaj dostępne są mechanizmy pozwalające na ich szybkie odizolowanie.	Średni. Poruszanie się po sieci wymaga więcej pracy dochodzeniowej w celu zidentyfikowania poszczególnych zasobów, do których uzyskano dostęp, stopnia ich wykorzystania oraz wymagań dotyczących środków zaradczych.	Wysoki. W prawie każdym przypadku potrzebna jest odbudowa kluczowej infrastruktury. Często trzeba to zrobić pod dodatkową presją przywrócenia firmy do minimalnego stanu operacyjnego pozwalającego zminimalizować poniesione straty.

Wyraźnie widać, jak ważne jest uświadomienie sobie cyberataków przeprowadzonych w środowisku, a tym bardziej, jak ważne jest ich jak najwcześniejsze wykrycie. Odpowiednie osoby muszą otrzymywać istotne informacje o cyberatakach w odpowiednim czasie. Jest to główny cel inżynierii detekcji.

Definicja inżynierii detekcji

Najwyższy priorytet dla zespołów ds. bezpieczeństwa mają szybka identyfikacja incydentów, ich kwalifikacja i łagodzenie potencjalnych skutków. Szybka identyfikacja potencjalnych incydentów zabezpieczeń jest dość skomplikowana. Ogólnie rzecz biorąc, członkowie zespołów ds. zabezpieczeń muszą mieć możliwość wykonywania następujących czynności:

1. Odbieranie zdarzeń dotyczących zasobów wymagających ochrony, a także zasobów, które mogą wpływać na nie pośrednio.

2. Identyfikacja zdarzeń mogących wskazywać na incydent zabezpieczeń, najlepiej natychmiast po jego wystąpieniu.
3. Zrozumienie sposobu oddziaływania potencjalnego incydentu.
4. Przekazanie istotnych szczegółów zdarzenia właściwym zespołom w celu zbadania go i złagodzenia skutków.
5. Otrzymywanie informacji zwrotnych od zespołów dochodzeniowych w celu ustalenia sposobu usprawnienia całego procesu.

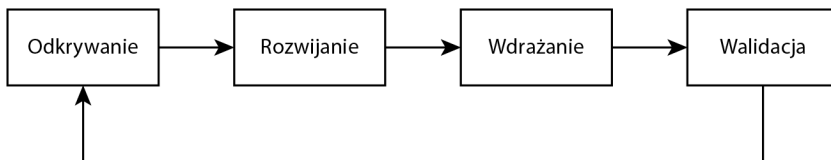
W małych środowiskach każda z tych czynności może być trudna do wykonania. Przy każdym wzroście rozmiaru zarządzanego środowiska złożoność radykalnie wzrasta.

Definicja inżynierii detekcji

Inżynierię detekcji można zdefiniować jako zbiór procesów umożliwiających wykrywanie w środowisku potencjalnych zagrożeń. Procesy te obejmują kompleksowy cykl życia incydentów — od zbierania wymagań dotyczących wykrywania zagrożeń, agregowania danych telemetrycznych systemu, wdrażania i utrzymywania logiki wykrywania, po walidację skuteczności programu.

Aby osiągnąć te cele, w dobrym programie inżynierii detekcji zazwyczaj trzeba wdrożyć cztery podstawowe procesy:

- **Odkrywanie.** Obejmuje zbieranie informacji dotyczących wykrywania zagrożeń. W ramach tego procesu należy określić, czy istniejące mechanizmy wykrywania spełniają wymagania. Należy również ocenić istotność wykrycia zagrożeń, a także odbiorców i ramy czasowe alertów.
- **Projektowanie, programowanie i testowanie.** Wymagania systemu wykrywania zagrożeń zostały zinterpretowane i na tej podstawie jest formułowany plan wdrożenia mechanizmu wykrywania. Zaprojektowany mechanizm jest najpierw wdrażany w środowisku testowym, gdzie jest sprawdzany pod kątem oczekiwanych rezultatów.
- **Wdrażanie i monitoring powdrożeniowy.** Mechanizm detekcji jest wdrażany w środowisku produkcyjnym. W tym przypadku monitorowana jest wydajność mechanizmu wykrywania i systemów detekcji.
- **Walidacja.** Wykonywane są rutynowe testy mające na celu określenie skuteczności programu inżynierii detekcji jako całości. Cykl życia procesów inżynierii detekcji przedstawiono na rysunku 1.3.



Rysunek 1.3. Procesy wchodzące w skład inżynierii detekcji

Dokładniejszy opis każdego z tych procesów zawiera rozdział 2. „Cykl życia inżynierii detekcji”.

Ważne cechy wyróżniające

Inżynieria detekcji może być źle rozumiana. Częściowo wynika to z faktu pokrywania się niektórych procesów z innymi funkcjami wchodzącymi w skład organizacji bezpieczeństwa. Inżynierię detekcji można scharakteryzować za pomocą następujących cech wyróżniających:

- **Poszukiwanie zagrożeń** (ang. *threat hunting* — polowanie na zagrożenia). Proces poszukiwania zagrożeń polega na proaktywnym rozwijaniu analiz śledczych w oparciu o hipotezy zakładające udane, niewykryte naruszenie zabezpieczeń. Proces poszukiwania zagrożeń pozwala zidentyfikować w środowisku aktywne zagrożenia, którym udało się ominąć istniejące mechanizmy zabezpieczeń. Proces ten zapewnia dane wejściowe do programu inżynierii detekcji, ponieważ pozwala zidentyfikować niedociągnięcia w mechanizmach wykrywania. Dane, które są dostępne dla inżynierii detekcji, są zazwyczaj tymi samymi danymi, z których korzystają łowcy zagrożeń. W związku z tym w procesie poszukiwania zagrożeń można również zidentyfikować niedociągnięcia w istniejącej infrastrukturze gromadzenia danych, które będą musiały zostać rozwiązane i zintegrowane z infrastrukturą wykrywania zagrożeń.
- **Działania centrum zabezpieczeń** (ang. *security operations center* — **SOC**). Zespoły SOC zazwyczaj koncentrują się na monitorowaniu środowiska zabezpieczeń, podczas gdy inżynieria detekcji dostarcza tym zespołom danych wejściowych. Podczas gdy zespoły SOC korzystają z produktów funkcji inżynierii detekcji, zazwyczaj ściśle z nimi współpracują w celu uzyskania wniosków dotyczących usprawnień w mechanizmach wykrywania zagrożeń lub gromadzenia danych.
- **Inżynieria danych**. Inżynierowie danych projektują, implementują i utrzymują systemy do gromadzenia, przekształcania i dystrybucji danych zazwyczaj w celu spełnienia wymagań analizy danych i analizy biznesowej. Jest to zgodne z kilkoma celami inżynierii detekcji, jednak program inżynierii detekcji jest silnie skoncentrowany na bezpieczeństwie i opiera się na inżynierii danych w celu uzyskania informacji potrzebnych do tworzenia mechanizmów wykrywania zagrożeń.

W tym podrozdziale przeanalizowaliśmy kilka podstawowych pojęć z zakresu cyberbezpieczeństwa, które będą przydatne w dalszej części tej książki podczas omawiania procesów inżynierii detekcji. Ponadto ustaliliśmy definicję inżynierii detekcji. Pamiętając o tej definicji, w kolejnym podrozdziale przeanalizujemy wartość, jaką wnosi do organizacji program inżynierii detekcji.

Wartość programu inżynierii detekcji

Aby można było uzyskać finansowanie programu inżynierii detekcji, przed wprowadzeniem go należy uzasadnić sensowność jego wdrożenia interesariuszom w organizacji. W niniejszym podrozdziale zostanie omówione znaczenie inżynierii detekcji. W szczegól-

ności przyjrzymy się rosnącemu zapotrzebowaniu na dobre mechanizmy wykrywania zagrożeń, sposoby definiowania ich jakości oraz poziom spełnienia tego zapotrzebowania przez program inżynierii detekcji.

Potrzeba zapewnienia lepszej wykrywalności

Postępy w rozwoju oprogramowania, takie jak ruch open source, przetwarzanie w chmurze, **infrastruktura jako kod (IaC)** oraz **potoki ciągłej integracji (ciągłego wdrażania) (CI/CD)** przyniosły organizacjom wymierne korzyści. Wspomniane postępy pozwalają organizacjom na łatwe korzystanie z technologii innych podmiotów, częste wdrażanie nowych wersji oprogramowania, szybkie uruchamianie i wyłączanie infrastruktury oraz dostosowywanie się do zmian w otoczeniu.

Niestety te same postępy pomogły również aktorom zagrożeń. Repozytoria open source dostarczają mnóstwo zaawansowanych narzędzi. Przetwarzanie w chmurze i IaC umożliwiają przeciwnikom szybkie wdrażanie i niszczenie infrastruktury C2, podczas gdy postępy w procesach oprogramowania i automatyzacji zwiększyły tempo operacji aktualizacji i tworzenia nowych funkcjonalności. Wspomniane zmiany jeszcze bardziej pogorszyły wartość statycznych wskaźników i spowodowały konieczność opracowania lepszych, bardziej wyrafinowanych metod wykrywania. W związku z tym dziedzina inżynierii detekcji zaczyna ewoluować w kierunku wspierania wysiłków na rzecz opracowywania bardziej wyrafinowanych mechanizmów wykrywania zagrożeń. Dzięki skutecznemu programowi inżynierii detekcji organizacje mogą wyjść poza wykrywanie statycznych wskaźników i zamiast tego wykrywać złośliwą aktywność na poziomie pojedynczych technik.

Cechy dobrego wykrywania zagrożeń

Nie ma jednej, dobrej definicji wykrywania zagrożeń. Poszczególne organizacje zajmujące się cyberbezpieczeństwem stosują różne progi wskaźników fałszywie dodatnich — czyli wskaźników wykryć zagrożeń, które inicjują się, choć nie powinny. Co więcej, napastnicy, z którymi muszą się mierzyć, różnią się pod względem wyrafinowania, a także widoczności i wykorzystywanych narzędzi. Inżynierowie ds. wykrywania zagrożeń muszą zidentyfikować metryki i kryteria oceny zgodne z potrzebami konkretnej organizacji. Przegląd procesów i podejść, które pomagają kierować tymi decyzjami, zaprezentowano w rozdziale 9. Kryteria oceny można podzielić na trzy obszary:

- zdolność wykrycia przeciwnika,
- koszty tej zdolności dla organizacji zajmującej się cyberbezpieczeństwem,
- koszt uniknięcia mechanizmów wykrywania przez przeciwnika.

Zdolność do wykrycia przeciwnika można podzielić ze względu na **pokrycie wykrywania** (ang. *coverage*) lub zakres aktywności identyfikowanych przez mechanizmy detekcji. Aby to lepiej wytłumaczyć, można posłużyć się kategoriami modelu MITRE ATT&CK. Jak wspomniano wcześniej, ten framework zawiera definicje na różnych poziomach szczegółowości — począwszy od taktyk będących najbardziej ogólną grupą, w których skład wchodzi techniki podzielone na najbardziej szczegółowe komponenty, czyli procedury.

Większość behawioralnych mechanizmów detekcji koncentruje się na wykrywaniu jednej lub większej liczby procedur zastosowanych przez napastnika w celu wdrożenia techniki.

Zwiększenie zasięgu wykrywania poprzez wykrywanie wielu procedur powiązanych z techniką lub stworzenie mechanizmu wykrywania skutecznego w odniesieniu do wielu technik zwiększa złożoność wykrywania, ale może również poprawić jego **trwałość**.

O ile zasięg wykrywania można traktować jako obszar powierzchni w warstwie NTP modelu MITRE ATT&CK, o tyle trwałość mechanizmu wykrywania określa czas, przez jaki ten mechanizm pozostaje skuteczny. W przewidywaniu trwałości mechanizmu wykrywania może pomóc zrozumienie zmienności infrastruktury, narzędzi i procedur przeciwnika oraz względnego kosztu ich zmiany.

Te dwa kryteria oceny określają procent ataków możliwych do wykrycia i oczekiwany czas, przez jaki zastosowane mechanizmy detekcji zachowają skuteczność. Niestety, kwantyfikacja tych kryteriów oceny do postaci metryk wymaga pełnej wiedzy na temat zdolności przeciwnika i jego tempa operacyjnego w zmianie tych zdolności. Mimo to w procesie dążenia do poprawy zdolności wykrywania przeciwnika można wykorzystać te kryteria do oceny skuteczności i jakości mechanizmów wykrywania zagrożeń.

Można jednak obliczyć historyczną skuteczność organizacji poprzez obliczenie średniego czasu do wykrycia zagrożenia jako czasu od rozpoczęcia ataku na organizację do chwili wykrycia napastnika.

Zdolność do wykrywania przeciwników wymaga **kosztów ponoszonych przez organizacje zajmujące się cyberbezpieczeństwem**. Koszty te są związane z tworzeniem, uruchamianiem i utrzymywaniem mechanizmów wykrywania zagrożeń, zasobami wymaganymi do przeglądania powiązanych z nimi alertów oraz działaniami podejmowanymi w odpowiedzi na te alerty. Przegląd informacji o przepływie pracy inżynierii detekcji zamieścimy w dalszej części tego rozdziału. **Czas utworzenia** tego przepływu pracy definiuje koszty utworzenia mechanizmu wykrywania zagrożeń. Na przykład, aby poprawić zasięg i trwałość mechanizmu wykrywania, konieczne jest zbadanie podejść do określonej techniki, co zwiększa koszty tworzenia tego mechanizmu detekcji. Zrozumienie **złożoności** wykrywania zagrożeń przez inżynierów detekcji wpływa na zdolność zrozumienia i utrzymania mechanizmu detekcji przez przyszłych analityków. Wpływa również na wydajność wykrywania zagrożenia (zarówno pozytywnie, jak i negatywnie). Utrzymanie mechanizmu detekcji w organizacji jest procesem ciągłym. Do określenia skuteczności lub wartości mechanizmu detekcji można wykorzystać stopień jego **przestarzałości** (ang. *staleness*). Czy określona technika lub narzędzie są nadal aktywnie wykorzystywane? Czy mechanizm wykrywania służy do wykrywania zagrożenia, które już jest w pełni załatwane, lub do ochrony infrastruktury (oprogramowania), którego już nie ma w sieci?

Każdy alert, który analityk musi przejrzeć, wiąże się z kosztami. **Poziom zaufania** (ang. *confidence*) do mechanizmu detekcji mierzy prawdopodobieństwo słuszności alertu — to znaczy wyzwolenia go w oczekiwanych warunkach. Dostrojenie mechanizmu wykrywania w celu zmniejszenia wskaźnika fałszywych alarmów może jednak zmniejszyć zasięg wykrywania i doprowadzić do niewykrycia ataku. Z kolei **hałaśliwość** (ang. *noisiness*) mechanizmu wykrywania określa częstość generowania alertu, który nie skutkuje wyeliminowaniem zagrożenia. Hałaśliwość mechanizmu wykrywania może wynikać z niskiego

poziomu zaufania — tzn. wysokiego wskaźnika fałszywych alarmów — ale może być również związana ze **skutkami** (ang. *impact*) wykrycia zagrożenia. Zrozumienie potencjalnych skutków wykrycia zagrożenia pozwala zmierzyć znaczenie lub dotkliwość wykrytego zagrożenia.

Na przykład mechanizm wykrywania zagrożeń może zidentyfikować rozpoznawcze skanowanie sieci. Brak możliwości podjęcia działania w odpowiedzi na to zagrożenie pomimo pewności jego wykrycia może skutkować zignorowaniem go. W procesie dostrajania mechanizmów wykrywania zagrożeń każda organizacja musi określić swoją tolerancję na wyniki fałszywie dodatnie. Jednak poziom zaufania do wykrycia zagrożenia i związane z tym potencjalne skutki można wykorzystać do ustalenia priorytetów alertów organizacji. W rozdziale 5. wskażemy korzyści z wykonywania detekcji o niskiej precyzji bez znaczącego wpływu na wydajność pracy analityków.

Zdolność do działania (ang. *actionability*) uzyskana dzięki mechanizmom detekcji określa łatwość, z jaką analityk SOC może wykorzystać mechanizmy detekcji do dalszej analizy zagrożenia lub naprawy jego skutków. Nie oznacza to, że wykrycie każdego zagrożenia musi wiązać się z natychmiastowym działaniem lub reakcją. Wykrycie zagrożenia może mieć na tyle niską wiarygodność, że natychmiastowe zbadanie go lub zareagowanie na nie jest nieefektywne. Zamiast tego działanie związane z alertem ma na celu zwiększenie zaufania do mechanizmu detekcji dzięki innym powiązanym i zidentyfikowanym działaniom lub ułatwienie analizy potencjalnej przyczyny źródłowej. Wiedza, na którą nie da się odpowiednio zareagować, ma jednak ograniczoną wartość.

Specyfika wykrywania zagrożeń zwiększa możliwości działania dzięki wyjaśnieniu tego, co zostało wykryte. Na przykład, o ile model uczenia maszynowego może zapewnić zwiększony zasięg wykrywania zagrożeń, z wysokim poziomem ufności, to może nie być w stanie podać konkretnych powodów utworzenia alertu. Taki brak szczegółowych informacji, np. rodziny wykrytego złośliwego oprogramowania, a tym samym brak identyfikacji jego możliwości, mechanizmów utrwalania lub innych szczegółów wymaganych do prawidłowej segregacji lub naprawy zagrożenia może ograniczyć możliwości działania.

Na koniec podczas oceny mechanizmu detekcji trzeba uwzględnić **koszty ponoszone przez przeciwnika**. Chociaż w większości przypadków nie ma możliwości wglądu w szczegółowe koszty związane z przeprowadzeniem ataku, przy określaniu kosztów przeciwnika można przyjąć się pośrednim dowodom. Wskazówek dotyczących poziomu kosztów ponoszonych przez przeciwnika może dostarczyć wiedza na temat łatwości, z jaką przeciwnik może uniknąć mechanizmów wykrycia, np. odniesienie do modelu piramidy bólu. Na przykład koszt zmiany skrótu złośliwego oprogramowania jest znacznie niższy niż koszt zmiany protokołu C2 tego oprogramowania. Nietrwałość infrastruktury, narzędzi i procedur napastnika jest miarą częstotliwości, z jaką atakujący zmienia swój atak po to, aby zmniejszyć możliwości wykrycia. Identyfikacja części ataku o niższej ulotności pozwala obrońcy zwiększyć trwałość mechanizmów wykrywania.

Korzyści z programu inżynierii detekcji

Podczas prezentowania pojęcia programu inżynierii detekcji kadry kierowniczej liczy się tylko jedno uzasadnienie: taki program znacząco zmniejsza ryzyko przeniknięcia zaawansowanego przeciwnika do sieci i dokonywania zniszczeń w firmie. Chociaż powinno to

dotyczyć każdego aspektu organizacji zajmującej się cyberbezpieczeństwem, każda firma osiąga ten cel w inny sposób. Program inżynierii detekcji różni się od innych komponentów programu cyberbezpieczeństwa, ponieważ umożliwia organizacjom szybkie reagowanie na nowe ataki. W celu dostosowania mechanizmu detekcji pozwala on wykorzystać wewnętrzne informacje o przeciwnikach atakujących branżę i specyfikę sieci firmy.

Podczas gdy mechanizmy detekcji zagrożeń od dowolnego dostawcy są zwykle dostarczane w pakiecie z informacjami dotyczącymi zagrożeń, jakie mogą być wykryte przez te mechanizmy, to są one tworzone w oparciu o podejście do wykrywania zagrożeń niezależne od klienta. Są one napisane w taki sposób, że można je masowo dystrybuować na urządzenia klienckie bez wpływu na działalność biznesową. W związku z tym gotowe mechanizmy detekcji dostępne na rynku koncentrują się na regułach i sygnaturach, które można zastosować w każdym środowisku. Nie pokrywają one jednak przypadków szczególnych — nimi zajmują się w organizacji działy inżynierii detekcji. Opracowanie programu inżynierii detekcji pozwala kontrolować ukierunkowanie mechanizmów detekcji na konkretne zagrożenia i planować ich wykorzystanie w taki sposób, aby obejmowały przypadki użycia specyficzne dla danego środowiska. Na przykład sprzedawcy gotowych rozwiązań detekcji zagrożeń nie mogą blokować wszystkich logowań z innych krajów, ponieważ miałyby to negatywny wpływ na bazę ich klientów. Z kolei wewnętrzny zespół inżynierów ds. wykrywania zagrożeń może zdecydować o zablokowaniu wszystkich prób logowania z innych krajów i odpowiednio zmodyfikować mechanizmy detekcji. Tematykę projektowania mechanizmów wykrywania dostosowanych do specyfiki danego środowiska opiszemy szczegółowo w rozdziałach 2. i 5.

Oprócz tej podstawowej korzyści istnieją dodatkowe korzyści, których organizacje zajmujące się cyberbezpieczeństwem mogą oczekiwać od dobrze zaprojektowanego programu inżynierii detekcji. W szczególności omówimy następujące kluczowe korzyści:

- Ustandaryzowany kod mechanizmów detekcji znajdujący się pod kontrolą systemu kontroli wersji.
- Automatyczne testowanie.
- Oszczędność kosztów i czasu.

Tym zaletom przyjrzymy się w poniższych podpunktach.

Ustandaryzowany kod mechanizmów detekcji znajdujący się pod kontrolą systemu kontroli wersji

W ramach budowania programu inżynierii detekcji możesz ustalić standardy kodu mechanizmów wykrywania zagrożeń. Dzięki temu ich kod jest łatwy do zrozumienia i kompatybilny z rozwiązaniami mechanizmów detekcji niezależnie od autora. Bez takiej standaryzacji autorzy poszczególnych mechanizmów detekcji mogliby programować je według własnego uznania, co mogłoby doprowadzić do dezorientacji innych programistów próbujących je zinterpretować.

Ponadto dzięki wykorzystaniu repozytoriów kodu mechanizmów detekcji cały kod przed wdrożeniem go w środowisku produkcyjnym może być kontrolowany przez system kontroli wersji, poddawany przeglądowi i testowany. Utrzymywanie scentralizowanego repozytorium kodu mechanizmów detekcji zmniejsza ryzyko wprowadzenia

nieprzetestowanych zmian lub reguł do środowisk produkcyjnych i ułatwia śledzenie kodu sprawiającego problemy. Utrzymywanie repozytorium kodu mechanizmów detekcji omówimy w rozdziale 5.

Automatyczne testowanie

Dzięki automatyzacji testowania mechanizmów detekcji zmniejszamy ryzyko wprowadzenia do środowisk produkcyjnych błędów spowodowanych wprowadzeniem nowego lub zmodyfikowanego kodu mechanizmów detekcji. Co więcej, im więcej mechanizmów automatyzacji, tym mniej czasu inżynierowie ds. wykrywania zagrożeń muszą poświęcić na ręczne testowanie kodu. Proces walidacji mechanizmów wykrywania zagrożeń zostanie dokładnie omówiony w części III.

Oszczędność kosztów i czasu

Oszczędność kosztów i czasu związana z inżynierią wykrywania zagrożeń to najważniejsze argumenty przemawiające do interesariuszy. Dla każdego finansowania programu interesariusze i kierownictwo poszukują uzyskania jak najszybszego **zwrotu z inwestycji** (ang. *return of investment* — **ROI**). Wspomniany zwrot z inwestycji ma postać oszczędności kosztów i czasu wynikających z wielu czynników. Na przykład automatyczne testowanie poprawia jakość wykrywania zagrożeń. Skraca ono czas potrzebny na testowanie mechanizmów detekcji, a także czas, jaki musieliby poświęcić analitycy na reakcję na fałszywe alarmy.

Największe oszczędności kosztów i czasu wynikają ze zmniejszenia prawdopodobieństwa włamania do sieci. Zmniejszenie ryzyka włamania poprzez wdrożenie dobrze zaprogramowanych mechanizmów wykrywania zmniejsza ryzyko poniesienia kosztów związanych z tymi włamaniami.

W tym podrozdziale pokazaliśmy wartość programu inżynierii detekcji i korzyści wynikające z jego zastosowania dla organizacji, która wdraża taki program. Następny podrozdział będzie podsumowaniem tego rozdziału. Zostanie w nim przedstawiony materiał, który będzie omówiony w dalszej części książki.

Przewodnik korzystania z tej książki

Poprzednie podrozdziały tego rozdziału dostarczyły podstawowej wiedzy, niezbędnej do dokładnego zrozumienia zawartości tej książki. W ostatnim podrozdziale przedstawimy krótki przegląd pozostałej części tej książki i tematów poruszanych w każdym rozdziale.

Struktura książki

Celem tej książki jest dokładne zaprezentowanie etapów procesu tworzenia programu inżynierii detekcji. Oprócz dogłębnej wiedzy na temat różnych aspektów cyklu życia inżynierii detekcji książka zawiera ćwiczenia pozwalające uczyć się narzędzi i stosować w praktyce omawiane zagadnienia. Książkę podzielono na cztery części, z których każda zapewnia wgląd w inny aspekt inżynierii detekcji.

Część I zawiera podstawową wiedzę niezbędną do studiowania pozostałej części tej książki. W poprzednich podrozdziałach tego rozdziału przedstawiono kluczowe pojęcia i terminologię, które będą wykorzystywane w pozostałej części książki. Omówiono również uzasadnienie zastosowania programu inżynierii wykrywania zagrożeń i korzyści, jakie przynosi on organizacji. W rozdziale 2. zagłębimy się we wszystkie fazy cyklu życia inżynierii detekcji i przedstawimy ogólny przegląd działań podejmowanych w każdej z nich. Na koniec w *rozdziale 3.* zaprezentujemy tematykę budowy laboratorium inżynierii detekcji. Wspomniane laboratorium będzie wykorzystywane w dalszej części tej książki do wykonywania praktycznych ćwiczeń.

Część II koncentruje się na aspekcie tworzenia w cyklu życia inżynierii detekcji. Ta część zaczyna się od rozdziału 4., który koncentruje się na identyfikacji i ocenie źródeł danych dostępnych dla inżynierów wykrywania zagrożeń. Rozdział zawiera opis laboratorium z rozdziału 3., wzbogaconego o dodatkowe źródła detekcji. Rozdział 5. pomaga zrozumieć wymagania dotyczące wykrywania zagrożeń oraz definiuje procedury i metody przechowywania kodu mechanizmów detekcji. Część II kończy się rozdziałem 6., w którym znajdziesz praktyczny przewodnik po przekształcaniu ustalonych wcześniej wymagań dotyczących wykrywania zagrożeń w kod mechanizmów detekcji, który można przetestować w utworzonym laboratorium.

W części III przechodzimy do pojęcia testowania i walidacji mechanizmów wykrywania zagrożeń. Najpierw, w rozdziale 7., zaprezentowano praktyczne wskazówki dotyczące walidacji wykrywania zagrożeń przy użyciu istniejących danych oraz generowania danych symulacyjnych. Dodatkowo w rozdziale zamieszczono wprowadzenie do udowadniania za pomocą wyników walidacji pokrycia NTP. W rozdziale 8. wprowadzono koncepcję wykorzystania w programie inżynierii detekcji wiedzy o zagrożeniach jako źródła mechanizmu wykrywania, wymogu wykrywania i metody zrozumienia pokrycia. Rozdział 9. zamyka część III omówieniem zagadnień związanych z zarządzaniem wydajnością. Obejmuje ono metody pomiaru skuteczności mechanizmów wykrywania zagrożeń, a także programu inżynierii detekcji jako całości. Ponadto dowiesz się, jak wdrożyć do programu inżynierii detekcji mechanizmy ciągłego doskonalenia.

Część IV kończy tę książkę rozdziałem 10. Ten rozdział jest przeznaczony dla osób, które chcą dowiedzieć się więcej o inżynierii detekcji w kategoriach możliwej kariery. Omówione zostaną umiejętności potrzebne do rozpoczęcia kariery w inżynierii detekcji oraz codzienne obowiązki inżyniera detekcji. Z tego rozdziału dowiesz się, dokąd zmierza przyszłość inżynierii detekcji oraz jak można zaangażować się w działania społeczności inżynierii detekcji.

Ćwiczenia praktyczne

Celem autorów tej książki jest dostarczenie wiedzy nie tylko w postaci tekstu, ale także praktycznych ćwiczeń, które pozwolą Ci zdobyć doświadczenie w procesie inżynierii detekcji. Te laboratoria rozpoczynają się w rozdziale 3., w którym zbudujemy środowisko testowe zawierające potrzebną infrastrukturę i narzędzia wymagane do realizacji pozostałych laboratoriów w tej książce.

Dzięki temu środowisku testowemu większość rozdziałów będzie zawierać ćwiczenia pozwalające napisać od podstaw mechanizmy wykrywania zagrożeń i ocenić ich jakość.

Wspomniane laboratoria obejmują zarówno ćwiczenia odnoszące się do konkretnych technologii wykrywania zagrożeń, jak i te, które analizują pokrycie środowiska jako całości.

Kod związany ze wspomnianymi ćwiczeniami jest dostępny publicznie w serwisie GitHub pod adresem <https://github.com/PacktPublishing/Practical-Threat-Detection-Engineering>.

Mamy nadzieję, że praktyczna wiedza dostarczona w tej książce umożliwi inżynierom zajmującym się wykrywaniem zagrożeń wyciągnięcie praktycznych wniosków i wdrożenie poznanych strategii i technik w swoich środowiskach.

Podsumowanie

W tym rozdziale przedstawiliśmy ogólne wprowadzenie do niektórych podstawowych pojęć cyberbezpieczeństwa, które będą wykorzystywane w całej książce. Po zaprezentowaniu tej wiedzy uzasadniliśmy sens tworzenia programu inżynierii detekcji, a konkretnie wskazaliśmy korzyści i wartość dla organizacji wynikające z wdrożenia takiego programu. Rozdział zamknęliśmy przedstawieniem przeglądu treści pozostałej części tej książki.

W następnym rozdziale zaczniemy omawiać sposoby identyfikacji i planowania specyficznych wymagań organizacji w zakresie wykrywania zagrożeń. Zagłębimy się również w etapy cyklu życia inżynierii wykrywania zagrożeń, które zostały wprowadzone w tym rozdziale.

PROGRAM PARTNERSKI

— GRUPY HELION —



1. ZAREJESTRUJ SIĘ
2. PREZENTUJ KSIĄŻKI
3. ZBIERAJ PROWIZJĘ

Zmień swoją stronę WWW w działający bankomat!

Dowiedz się więcej i dołącz już dzisiaj!

<http://program-partnerski.helion.pl>

GRUPA
Helion

Nie oczekuj, że wróg się nie zjawi. Przygotuj się, aby go odpowiednio przyjąć!

Efektywny potok detekcji zagrożeń jest niezbędnym elementem programu cyberbezpieczeństwa. W procesach inżynierii detekcji szczególną uwagę należy poświęcić technikom tworzenia i walidacji mechanizmów detekcji. To oczywiste — od ich jakości zależy skuteczność zabezpieczeń w organizacji. Trzeba więc zrozumieć, czym jest inżynieria detekcji i jakie ma znaczenie dla cyberbezpieczeństwa.

Ta książka jest przewodnikiem dla profesjonalistów w dziedzinie cyberbezpieczeństwa. Przedstawia podstawowe zasady reagowania na incydenty bezpieczeństwa i szczegółowo, na przykładach, omawia proces tworzenia zdolności szybkiej i skutecznej reakcji na takie zdarzenia. Zaprezentowano tu techniki informatyki śledczej, od pozyskiwania dowodów i badania pamięci ulotnej po badanie dysku twardego i dowodów pochodzących z sieci. Szczególną uwagę poświęcono zagrożeniom atakami ransomware. Nie zabrakło omówienia roli analizy zagrożeń w procesie reagowania na incydenty, a także zasad sporządzania raportów dokumentujących reakcję na incydent i wyniki analizy. Pokazano również, w jaki sposób prowadzi się polowania na zagrożenia.

W książce:

- przebieg procesu inżynierii detekcji
- budowa laboratorium testowego
- utrzymywanie mechanizmów detekcji w formie ustandaryzowanego kodu
- tworzenie mechanizmów detekcji
- wczesne wykrywanie cyberataków i złośliwej aktywności
- ścieżki kariery w inżynierii detekcji

Megan Roddie od lat pracuje w branży bezpieczeństwa informacji. Jest też autorką kursów i instrukcją w SANS Institute, gdzie regularnie publikuje analizy śledcze i badania dotyczące reagowania na incydenty w chmurze.

Jason Deyalsingh jest doświadczonym konsultantem w dziedzinie cyberbezpieczeństwa. Specjalizuje się w kryminalistyce cyfrowej i reagowaniu na incydenty.

Gary J. Katz zajmuje się problematyką cyberbezpieczeństwa, a swoimi przemyśleniami dzieli się w artykułach i książkach.

 Helion	KOD KORZYŚCI Sięgnij po więcej! ▶ 
 helion.pl  HELION SA ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl	ISBN 978-83-289-0902-1  9 788328 909021
Cena: 89,00 zł	

<packt>