
Wojny w cyberprzestrzeni

Koncepcje, strategie i taktyki,
dzięki którym przetrwasz
i ocalisz swoją organizację

Chase Cunningham



Tytuł oryginału: Cyber Warfare - Truth, Tactics, and Strategies: Strategic concepts and truths to help you and your organization survive on the battleground of cyber warfare

Tłumaczenie: Grzegorz Kowalczyk

ISBN: 978-83-283-7248-1

Copyright © Packt Publishing 2020. First published in the English language under the title 'Cyber Warfare – (9781839216992)'

Polish edition copyright © 2021 by Helion SA
All rights reserved.

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 Helion SA 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 Helion SA nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Helion SA
ul. Kościuszki 1c, 44-100 Gliwice
tel. 32 231 22 19, 32 230 98 63
e-mail: helion@helion.pl
WWW: <http://helion.pl> (księgarnia internetowa, katalog książek)

Drogi Czytelniku!
Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres
<http://helion.pl/user/opinie/wojwcy>
Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

- [Kup książkę](#)
- [Poleć książkę](#)
- [Oceń książkę](#)

- [Księgarnia internetowa](#)
- [Lubię to! » Nasza społeczność](#)

Spis treści

Przedmowa	9
O autorze	13
O recenzencie	14
Wprowadzenie	15
Rozdział 1. Krótka historia zagrożeń cybernetycznych i pojawienie się APT	19
Hakerzy nie są tacy, jak pokazuje nam Hollywood	19
Bitwa na wiązki	22
Hakowanie z użyciem modemów	22
Rozwój oprogramowania antywirusowego	23
Początki zagrożeń APT	24
Wczesne ataki APT	28
Zamieszanie w cyberbronie	30
Amerykańskie i sojusznicze agencje do spraw bezpieczeństwa cybernetycznego	30
Atak cybernetyczny słyszany na całym świecie	31
Cyfrowa wojna oko za oko	34
Puszka Pandory zostaje otwarta	35
Podsumowanie	37
Źródła	38
Rozdział 2. Perymetr sieciowy odchodzi do lamusa...	39
Scenariusz, który ilustruje luki w modelu opartym na perymetrze sieciowym	40
Upadek perymetrów globalnych	41
Perymetr sieciowy nie sprawdza się nawet w organizacjach spełniających wymogi bezpieczeństwa	44
Perymetry systemów rządowych też zawodzą	46

Użytkownicy, BYOD i całkowity zanik perymetru sieciowego	48
Aplikacje zwiększają poziom zagrożenia	51
Niepowodzenia metod uwierzytelniania	52
Urządzenia IoT robią dziury w każdym perymetrze	56
Nie możesz naprawić głupoty ani złośliwości	57
Podsumowanie	61
Źródła	62
Rozdział 3. Nowe taktyki i trendy — co nas czeka w niedalekiej przyszłości?	65
Ataki zmieniają kierunek	66
Pojazdy autonomiczne... Niepoprawne dane, pechowy dzień	68
Drony... Śmierć nadchodzi z nieba	72
Napastnicy integrują działania i taktykę w celu optymalizacji skuteczności ataków	78
Ransomware wchodzi na platformy mobilne	82
Ataki DDoS osiągają poziom broni	85
Podsumowanie	87
Źródła	88
Rozdział 4. Ataki z wywieraniem wpływu — wykorzystywanie mediów społecznościowych do złośliwych celów	89
Nowy atak cybernetyczny	90
Cyfrowe pole bitwy ulega zmianom	91
Tylko #hashtag czy już amunicja?	92
Wywieranie wpływu na osoby wpływowe	99
Podsumowanie	104
Źródła	104
Rozdział 5. DeepFakes, uczenie maszynowe i sztuczna inteligencja w cyberprzestrzeni	105
Z dużego ekranu na smartfony — narodziny DeepFakes	106
Czym są DeepFakes?	106
Sieci GAN napędzają DeepFakes	107
DeepFakes w praktyce, czyli DeepMastersPrints	112
Podrabianie głosu za pomocą uczenia maszynowego — technologia DeepVoice	116
ReadFakes	119
Wiadomości z ostatniej chwili mogą być źródłem problemów	122
Kiedy „badania” danych i sztucznej inteligencji idą na marne	123
Podsumowanie	126
Źródła	127
Rozdział 6. Zaawansowane działania w wojnie cybernetycznej	129
Kampanie wojny cybernetycznej	130
Kampania ataków na indyjskie elektrownie jądrowe	131
Chińska kampania „transferu technologii”	133
Kampania ingerowania w przebieg wyborów w USA i Libii	135

Operacje pod fałszywą flagą mylą tropy w cyberprzestrzeni	137
Klasyfikacja etapów ataków cybernetycznych	138
Grupy hakerów celowo unikają rozgłosu	144
Kampanie cyberataków w nadchodzącej dekadzie	147
Fałszerstwa	148
Podsumowanie	152
 Rozdział 7. Planowanie strategiczne dla przyszłych działań w wojnie cybernetycznej	 155
Każdy ma plan, dopóki nie zostanie w zęby	156
Jaka strategia?	157
Kiedy charakter konfliktu wymaga zmiany strategii	158
Infiltracja nie jest równoznaczna z dominacją	159
Liderzy również muszą umieć „pobrudzić sobie ręce”	161
To środowisko, a nie sprzęt, decyduje o tym, co działa, a co nie	161
Gromadzenie informacji i wywiad to nie zawsze jedno i to samo	163
Zbyt wiele to czasem naprawdę zbyt wiele	165
Wielkie mury mogą powodować wielkie problemy	166
Misja nie została zrealizowana...	168
Jak wygląda skuteczna strategia w cyberprzestrzeni?	172
Zmiana koncepcji strategicznych	173
Strategiczna obrona „granic”	174
Małymi krokami do przodu	176
Orkiestracja umożliwia realizację strategii	178
Podsumowanie	179
 Rozdział 8. Innowacje strategiczne i mnożniki siły w cyberprzestrzeni	 181
Narzędzia obronne i czynniki strategiczne	182
Poznaj Małpę	182
Co jeszcze potrafi Infection Monkey?	185
Zaawansowane możliwości pakietu Infection Monkey	187
Perymetr definiowany programowo	189
Białe listy aplikacji	194
Narzędzia ofensywne i czynniki strategiczne	197
Po co nam hasła?	197
Pakiet WhatBreach	198
Pakiet SNAP_R	202
Fałszywe komentarze w celu wywierania wpływu	203
Podsumowanie	205
Źródła	205
 Rozdział 9. Przygotowanie do uderzenia	 207
Wyłączenie odpowiedzialności	208
Mikrosegmentacja jest kluczem do przetrwania	208
Czym jest mikrosegmentacja?	208
Narzędzia i technologie mikrosegmentacji	211

Pragmatyczne zastosowanie sieci SDN	212
Możliwe pułapki w mikrosegmentacji	213
Odzyskanie „wzgórza”	215
Hasła z wozu, koniom lżej	218
Zbieranie informacji	224
Podsumowanie	228
Źródła	229
 Rozdział 10. Zdolność przetrwania w wojnie cybernetycznej	
i potencjalne skutki ataków	231
Po co komu prawa w czasie wojny?	232
Reguła 1. — „Domyślny” oznacza „martwy”	233
Reguła 2. — Myśl strategicznie, działaj taktycznie	237
Reguła 3. — Szczegóły, szczegóły	239
Reguła 4. — Pozbądź się haseł	241
Reguła 5. — Ogranicz promień rażenia	244
Wpływ porażki	248
Ataki na system opieki zdrowotnej	248
Ataki na systemy ICS (przemysłowe systemy sterowania)	249
Zagrożenie dla losu całych narodów	249
Podsumowanie	252
 Dodatek. Poważne incydenty cybernetyczne wykryte lub ujawnione w 2019 roku	253

Perymetr sieciowy odchodzi do lamusa...

Przez ostatnie ponad 30 lat zdecydowana większość planów zabezpieczenia sieci i infrastruktury cyfrowej bazowała na koncepcji bezpieczeństwa opartego na perymtrze sieciowym. Większość organizacji na całym świecie uznało tę koncepcję za słuszną i doszło do wniosku, że gdy mury będą wystarczająco wysokie, a zewnętrzne perymetry sieci wystarczająco twarde, napastnik nie będzie w stanie dostać się do środka. Cała globalna architektura sieciowa została zaprojektowana i wdrożona zgodnie z tą koncepcją, a miliardy dolarów zostały wydane na zaimplementowanie wielopoziomowych mechanizmów „głębokiej obrony” i wdrożenie zabezpieczeń przypominających zamek otoczony fosą. Okazało się jednak, że to wszystko na nic.

Model bezpieczeństwa bazujący na silnie bronionym perymtrze sieciowym zdecydowanie nie nadąża za ewolucją internetu, eksplodującą liczbą urządzeń sieciowych, łatwością dostępu, gwałtownie rosnącą popularnością usług w chmurach obliczeniowych i coraz większą liczbą mobilnych pracowników, wykorzystujących koncepcję BYOD (ang. *Bring Your Own Device* — przynieś swoje własne urządzenie). Obecnie nie ma już w zasadzie żadnego perymetru sieciowego. W momencie, kiedy użytkownik może zabrać do domu służbowego laptopa, załogować się z domowego komputera lub skorzystać z urządzenia mobilnego czy aplikacji i uzyskać dostęp do firmowej sieci, silnie kiedyś umocniony perymetr sieciowy zostaje praktycznie pocięty na kawałki.

W tym rozdziale wyjaśnię dlaczego współcześnie używane systemy są zbudowane w sposób, który w pewnym sensie ułatwia przeprowadzenie skutecznego ataku i potencjalne wycieki danych:

- Pokażę, dlaczego model bezpieczeństwa oparty na perymtrze sieciowym jest z definicji zawodny.
- Omówię ograniczenia, jakie obecna technologia nakłada na infrastrukturę sieciową.

- Przeanalizuję zjawisko lawinowego wzrostu liczby incydentów bezpieczeństwa i awarii wynikających z wzajemnego połączenia sieci.
- Przedstawię sposób, w jaki wrogie państwa i inni agresorzy wykorzystują wady i luki w zabezpieczeniach takiej architektury.

Na początek omówię ciekawy scenariusz, który w bardzo trafny sposób pokazuje, że model bezpieczeństwa oparty na perymetrze sieciowym należy już do przeszłości.

Scenariusz, który ilustruje luki w modelu opartym na perymetrze sieciowym

Przyjrzyjmy się następującemu scenariuszowi. Załóżmy, że mamy użytkownika, który pracuje w domu, ma uprawnienia administratora na swoim komputerze (jak to zwykle bywa w sytuacji, gdy jest to jego własne urządzenie) i od czasu do czasu pozwala z niego korzystać swojemu dziecku, które używa go do odrabiania zadań domowych. Rezolutny dzieciak otrzymuje tym samym dostęp do wydajnego, najeżonego aplikacjami i zwykle pozbawionego nadzoru komputera swojego rodzica, i zamiast połączyć się z bezpieczną i sprawdzoną witryną sieciową swojej szkoły, ulega pokusie i łączy się z pozornie niewinnie wyglądającą stroną internetową, o której opowiadali w szkole jego koledzy.

Wiedziony ciekawością młody człowiek chce zobaczyć, co się tam znajduje, ale okazuje się, że na dzień dobry strona prosi o pobranie i zainstalowanie w przeglądarce sieciowej odpowiedniej wtyczki, która rzekomo jest niezbędna do wyświetlenia pełnej zawartości strony. Pamiętaj, że skoro dziecko może wykonać taką operację, to z pewnością to zrobi (przecież działa zwykle na koncie rodzica, który ma uprawnienia administratora na tym komputerze).

Po pobraniu i zainstalowaniu wtyczki strona wyświetla się prawidłowo, wszystko działa dobrze, program antywirusowy nie wyświetla żadnych alarmów (dzieje się tak, ponieważ rzekoma „wtyczka”, która została zainstalowana, to nowa wersja złośliwego oprogramowania i program antywirusowy nie posiada jeszcze wykrywających jego kod sygnatur). Zadowolony młody użytkownik kończy przeglądanie strony, kończy odrabiać zadanie domowe, wyłączy komputer rodzica i wszystko kończy się szczęśliwie. Albo tak im się przynajmniej wydawało.

Zainstalowane i gotowe do działania złośliwe oprogramowanie cierpliwie czeka na ponowne uruchomienie komputera, a następnie spokojnie działa sobie w tle, ukradkiem pobiera dodatkowy moduł rejestrujący naciśnięcia klawiszy oraz małą, złośliwą aplikację, która wyszukuje nazwy kont oraz dane uwierzytelniające dla połączeń VPN, a także inne hasła dostępu i ich skróty.

Następnego dnia rano, nasz przykładowy rodzic zamienia się w użytkownika biznesowego i uruchamia swój komputer w celu połączenia się z siecią firmową i rozpoczęcia pracy. Jednak w tej sytuacji pomyślne nawiązanie połączenia VPN otwiera dla zainstalowanego dzień wcześniej złośliwego oprogramowania szeroką furtkę prowadzącą do infrastruktury sieciowej firmy, gdzie takie oprogramowanie może działać z uprawnieniami tego użytkownika. Po nawiązaniu połączenia, złośliwe oprogramowanie (a w zasadzie kryjący się za nim napastnik) może dokonać

próby ustanowienia bardziej stabilnego przyczółka w sieci i często może to zrobić, jeżeli uwierzytelniony użytkownik ma odpowiednie uprawnienia (np. administratora). Złośliwy program, który dostał się do firmowej w sieci, ma teraz takie same uprawnienia.

Wzajemne połączenia pomiędzy wirtualnymi sieciami LAN i innymi segmentami sieci, a także zwykle słabe mechanizmy uwierzytelniania użytkowników, powszechnie spotykane w systemach wewnętrznych, znacząco ułatwiają rozprzestrzenianie się złośliwego oprogramowania. Dzieje się tak zwykle dlatego, że wewnętrzne sieci firmy są zazwyczaj traktowane jako sieci zaufane, co w takiej sytuacji pozwala na prawie nieograniczone działanie złośliwego kodu.

Po dostaniu się do sieci wewnętrznej, złośliwe oprogramowanie zazwyczaj kontynuuje działanie w celu znalezienia najbardziej wartościowych celów, poufnych dokumentów, wrażliwych danych i innych informacji przydatnych dla napastnika. Następnie, korzystając z wybranego protokołu sieciowego, nawiązuje ukryte połączenie o niskiej przepustowości ze swoim serwerem macierzystym, do którego przesyła zdobyte informacje i z którego otrzymuje kolejne polecenia. Wykradzione dane mogą zostać sprzedane, użyte do szantażu lub przygotowania kolejnego, bardziej wyrafinowanego ataku, którego celem będzie dokonanie próby zablokowania systemów firmy lub zaszyfrowania krytycznych dla jej działania danych i próba wymuszenia okupu.

Jeżeli na zaatakowanej maszynie lub w jej środowisku sieciowym napastnik nie znajdzie niczego wartościowego, metoda dostępu oraz informacje o strukturze skompromitowanej sieci mogą zostać odsprzedane na czarnym rynku w sieci Dark Net, co może umożliwić innym hakerom wykorzystanie tego systemu do przeprowadzania swoich ciemnych operacji w przyszłości. W takim scenariuszu skompromitowana sieć może stać się punktem pośrednim pozwalającym na atakowanie innych systemów i inne działania cyberprzestępców. Jednak bez względu na to, jakie będą skutki takiego włamania, można śmiało zaryzykować twierdzenie, że koncepcja wydzielania zaufanych stref w sieciach wewnętrznych oraz poleganie na przestarzałym modelu bezpieczeństwa, bazującym na utwardzonym zewnętrznym perymetrze sieciowym, nie tylko przyczyniła się do przeprowadzenia takiego ataku, ale wręcz je umożliwiła.

Opisany wyżej scenariusz pokazuje, że nawet bardzo rozbudowane perymetry sieciowe, powszechnie stosowane przez firmy i organizacje na całym świecie, nie zapewniają bezpieczeństwa i mogą zostać szeroko otwarte dla napastnika tylko dlatego, że użytkownik zabrał swój komputer do domu. W prawdziwym świecie takie sytuacje są powszechne i wiele firm, organizacji i agend rządowych odczuło ich skutki na własnej skórze. W następnym rozdziale omówimy prawdziwą sytuację, w której pewna firma padła ofiarą tego, co początkowo wydawało się być infekcją typu ransomware o dosyć ograniczonym zasięgu.

Upadek perymetrów globalnych

Kolejnym przykładem tego, jak model bezpieczeństwa oparty na perymetrze sieciowym znacząco ułatwia przeprowadzanie ataków i jest rażąco nieskuteczny w zwalczaniu współczesnych zagrożeń, jest analiza tego, co przytrafiło się okrętowemu gigantowi Maersk.

W 2017 roku ukraińska spółka tworząca oprogramowanie dla księgowości — grupa Linkos — działała normalnie. Serwery firmy, które były połączone z setkami komputerów klientów i odpowiedzialne za aktualizację ich oprogramowania księgowego, bez jej wiedzy stały się punktem startowym dla ataku z użyciem złośliwego oprogramowania typu ransomware o nazwie NotPetya.

Grupa Linkos, która nie była tutaj niczemu winna, poza tym, że jej siedziba znajdowała się w kraju aktywnie atakowanym przez wojskowe skrzydło sił cybernetycznych rządu rosyjskiego, padła ofiarą prowadzonej w ukryciu od wielu miesięcy tajnej operacji, mającej na celu uzyskania przewagi militarnej w regionie.

Rosyjska grupa zajmująca się wojną cybernetyczną stworzyła pierwsze w swoim rodzaju oprogramowanie typu ransomware, będące połączeniem narzędzia NSA o nazwie EternalBlue, które wyciekło w 2017 r., oraz znanego narzędzia o nazwie Mimikatz, służącego do pozyskiwania i wykorzystywania haseł, które jest używane od 2011 roku.

Wspomniana rosyjska grupa specjalistów połączyła te narzędzia w bardzo efektywną broń cybernetyczną, pozwalającą na całkowite zablokowanie zaatakowanych komputerów i szybkie rozprzestrzenianie infekcji po całej sieci środowiska celu. Nadmiernie rozbudowane przywileje użytkowników, niewielki stopień złożoności stosowanych haseł połączony z powszechnym używaniem tych samych haseł w różnych systemach i wieloma współdzielonymi zasobami sieciowymi, stały się perfekcyjnym środowiskiem znakomicie ułatwiającym sprawne funkcjonowanie tej cyberbroni.

Po otrzymaniu odpowiedniego polecenia, złośliwe oprogramowanie (w tym przypadku typu ransomware) zostało zdalnie uruchomione. W ciągu kilku godzin, połączenia z serwerów grupy Linkos do wszystkich obsługiwanych przez nie klientów biznesowych były zagrożone, a dzięki powiązaniom między tymi klientami a innymi podmiotami i sieciami, atak zaczął rozprzestrzeniać się w najlepsze.

Microsoft wkrótce wypuścił odpowiednią poprawkę bezpieczeństwa zamykającą lukę wykorzystywaną przez exploit EternalBlue, ale po raz kolejny okazało się, że połączony charakter sieci w globalnym internecie, w połączeniu z nieefektywnymi procesami zarządzania aktualizacjami i brakiem obowiązkowych procedur instalowania poprawek bezpieczeństwa, ułatwił rozprzestrzenianie się tego złośliwego oprogramowania i skierował ransomware NotPetya między innymi do sieci firmy Maersk. Innymi słowy, tysiące połączonych ze sobą sieci, opierających się na zbliżonych założeniach technicznych, a do tego setki ludzkich pomyłek i biznesowych błędów, wszystko to razem tworzyło idealny grunt dla tej infekcji.

Dowód na użycie ransomware NotPetya jako broni, a nie narzędzia wymuszenia, pojawił się, gdy ofiary zdały sobie sprawę, że powiadomienie o możliwości odzyskania zaszyfrowanych danych po wpłaceniu okupu było kłamstwem. Złośliwe oprogramowanie ingerowało bezpośrednio w najgłębiej położone struktury zaatakowanej maszyny, jej główny rekord rozruchowy (ang. *Master Boot Record* — MBR), w którym komputer przechowuje szereg informacji niezbędnych do uruchomienia systemu operacyjnego. Opłacenie okupu przez ofiarę ataku nie przynosiło żadnego rezultatu i nie pozwalało na odszyfrowanie danych, co powodowało, że zaatakowany komputer stawał się po prostu mocno przeplaconym i nieco zbyt dużym przyciskiem do papieru.

Złośliwe oprogramowanie użyte do ataku nie zawierało nawet żadnego prawdziwego klucza deszyfrującego, który mógłby zostać użyty do odzyskania zaszyfrowanych informacji; była to po prostu broń zbudowana wyłącznie w celu zniszczenia użyteczności zainfekowanego systemu, a firma Maersk miała niebawem dołączyć do szybko rosnącej grupy ofiar.

Pierwotna infekcja w grupie Maersk była w pewnym sensie skutkiem powszechnie stosowanych praktyk biznesowych, a nie tylko elementów technicznych. W jednym z odległych biur firmy Maersk, znajdującym się w Odessie na Ukrainie, administrator działu IT otrzymał zadanie zainstalowania oprogramowania księgowego M.E.Doc na jednym z komputerów w dziale finansowym. Oprogramowanie to zostało zakupione w firmie Linkos i było przez nią zarządzane, co spowodowało, że złośliwe oprogramowanie miało do dyspozycji wszystko, co było niezbędne do jego aktywacji.

Gdy ransomware NotPetya uzyskał dostęp do sieci Maersk, łatwość i szybkość, z jaką infekcja zaczęła się rozprzestrzeniać, były ogromnym zaskoczeniem dla wszystkich. W ciągu kilku godzin cała sieć warta miliard dolarów, z setkami milionów wydanych na urządzenia i technologie zabezpieczeń, runęła jak domek z kart pod naporem cyfrowego huraganu spowodowanego przez ransomware NotPetya.

Jakby nieszczęścia było mało, po wykryciu infekcji pracownicy działu IT firmy Maersk zaczęli zdawać sobie sprawę z tego, że ich praktyki w zakresie konfiguracji i zarządzania rozległej infrastruktury sieciowej umożliwiły szybkie przeniesienie się złośliwego oprogramowania na kontrolery domeny. Stało się tak, ponieważ postępując zgodnie z najlepszymi praktykami branżowymi, pracownicy działu IT firmy Maersk skonfigurowali kontrolery domeny tak, aby korzystały ze współdzielonego modelu konfiguracji, ponieważ są one mózgiem całego systemu uwierzytelniania w sieciach opartych na systemie Windows. Oznaczało to jednak, że dzięki mechanizmowi replikacji infekcja mogła się rozprzestrzenić niemal równocześnie na wszystkie kontrolery domeny w całej sieci, co pomogło radykalnie zwiększyć zasięg ataku i systematycznie blokowało kolejne krytyczne elementy wewnętrznej infrastruktury zarządzania i kontroli firmy Maersk.

Resztki infrastruktury sieciowej udało się uratować jedynie dzięki zupełnemu przypadkowi. Tuż przed rozpoczęciem ataku w jednym z odległych biur Maersk w Ghanie doszło do przerwy w dostawie prądu i dzięki temu sieć komputerowa tego biura nie została zainfekowana, a jej kontroler domeny mógł później posłużyć do odzyskiwania danych. Gdyby nie to niewiarygodne zrządzenie losu, sprawy przybrałyby zupełnie inny obrót i prawdopodobieństwo, że firma mogłaby podźwignąć się po takim ataku, byłoby prawie zerowe.

Niemal wszystkie terminale portowe tego giganta transportu morskiego zostały zainfekowane i stały się bezużyteczne, co wpłynęło na logistykę i transport towarów na całej planecie. Aby uchronić firmę przed całkowitym paraliżem i umożliwić jej jakiekolwiek funkcjonowanie, operatorzy byli zmuszeni polegać na papierowych dokumentach, kontaktach w serwisach takich jak Gmail i prywatnych telefonach komórkowych. Tysiące komputerów i serwerów działających w sieci korporacyjnej zostało całkowicie zablokowane, a działanie światowych sieci logistycznych, dostawców, kierowców ciężarówek i użytkowników w firmie Maersk było ogromnie utrudnione przez wiele kolejnych tygodni.

Straty poniesione przez samą tylko firmę Maersk w wyniku tego ataku zostały oszacowane na co najmniej ćwierć miliarda dolarów, nie licząc kosztów usuwania szkód, odzyskiwania danych i przywracania funkcjonalności poszczególnych systemu. Całkowite koszty ataku poniesione przez firmę Maersk szacowane są na blisko miliard dolarów (Greenberg, 2018), a wszystko to z powodu jednego pakietu zainstalowanego w systemie oprogramowania księgowego, które było bezpośrednio podłączone do zewnętrznych serwerów aktualizacji i klientów.

W skali całego świata koszty były liczone w dziesiątkach miliardów dolarów. Zainfekowane zostały sieci i komputery tysięcy przedsiębiorstw, szpitali i organizacji cywilnych. Pacjenci i karetki pogotowia były odsyłane do innych placówek służby zdrowia, ponieważ działanie ich szpitali macierzystych zostało zablokowane przez złośliwe oprogramowanie. Ofiarą ataku padły nawet rządowe sieci komputerowe Departamentu Obrony USA. Jeżeli ktokolwiek miał jeszcze jakieś wątpliwości, że powszechnie przyjęte na całym świecie praktyki i mechanizmy bezpieczeństwa całkowicie zawiodły, to NotPetya jest idealnym studium przypadku pokazującym, że powtórzenie się takich katastrof w przyszłości jest najzupełniej możliwe i że możliwości broni cybernetycznych w ostatnich latach radykalnie wzrosły.

Opisany przypadek dobitnie pokazuje, że globalny gigant transportu morskiego i setki innych organizacji na całym świecie poniosło poważne straty z powodu nieskuteczności i fiaska starych praktyk bezpieczeństwa. W następnym podrozdziale zobaczymy, jak pozornie szczelny i doskonale zabezpieczony perymetr sieciowy organizacji spełniającej wszelkie wymogi bezpieczeństwa poniósł klęskę z powodu nieadekwatności starych praktyk.

Perymetr sieciowy nie sprawdza się nawet w organizacjach spełniających wymogi bezpieczeństwa

Włamanie do sieci firmy Equifax stanowi kolejne studium przypadku wskazujące na przestarzałość rozwiązań i nieskuteczność obecnych praktyk w zakresie bezpieczeństwa sieciowego firm i organizacji. Nawet wielkie przedsiębiorstwa i korporacje, które wydały miliony dolarów na systemy bezpieczeństwa i są w pełni świadome wymagań, jakie stawiają wdrożone plany bezpieczeństwa danych, mogą paść ofiarą epickiej katastrofy, gdy w ich modelu bezpieczeństwa opartym na perymetrach sieciowych zostaną wykryte jakiekolwiek słabe punkty.

Weźmy pod uwagę techniczne i menedżerskie aspekty włamania do sieci firmy Equifax. Firma dysponowała ogromnym budżetem dla swojego zespołu bezpieczeństwa sieciowego, wszystkie wymagane i zgodne z przepisami rozwiązania były wdrożone, a także zaimplementowano na szeroką skalę systemy monitoringu i analizy bezpieczeństwa. Mimo to, w ciągu trwającego blisko rok włamanie i penetracji sieci firmy, napastnicy przechwycili niemal całe repozytoria danych firmy oraz dane ponad 140 milionów Amerykanów i ponad 800 000 obywateli Wielkiej Brytanii.

Początkowym wektorem ataku była luka w zabezpieczeniach publicznie dostępnego w sieci serwera, który był odpowiedzialny za przyjmowanie wniosków i odwołań w sprawach kredytowych.

Na tym serwerze zainstalowana była nieco przestarzała wersja frameworka Apache Struts. Co ciekawe, w tym samym tygodniu, w którym doszło do pierwszego ataku, amerykański zespół CERT udostępnił odpowiednią poprawkę bezpieczeństwa dla tego pakietu.

Napastnicy w tym przypadku po prostu użyli publicznie dostępnego exploita, a po uzyskaniu dostępu do sieci firmy dokonali udanej próby eskalacji swoich uprawnień i rozpoczęli dalszą eksplorację sieci. Jest to najbardziej klasyczny sposób przeprowadzania ataku, którego metodologia jest doskonale znana wszystkim warty swojej nazwy zespołom ds. bezpieczeństwa sieciowego, a mimo to w przypadku firmy Equifax taki atak zakończył się sukcesem.

Po uzyskaniu dostępu do sieci atakujący wykorzystali zdobyte poświadczenia i prawa na poziomie administratora systemu do utworzenia ukrytego punktu dostępowego, który pozostał w systemie niewykryty przez wiele miesięcy. Środowisko firmy Equifax było wyposażone w zapory sieciowe i systemy wykrywania i monitorowania włamań oraz analizy sieci. Niestety okazało się, że dzięki wygasłemu certyfikatowi system nie działał optymalnie, a przesłanki, które powinny skłonić zespół IT do podjęcia działań naprawczych, najzwyczajniej w świecie nie zostały zauważone.

Certyfikat tego krytycznego elementu monitorowania był nieważny od ponad 10 miesięcy i został zaktualizowany dopiero po upływie długiego czasu od wykrycia włamania. Atak, który początkowo był ograniczony tylko do kilku serwerów, po krótkim czasie rozprzestrzenił się na ponad 50 ogromnych baz danych przechowujących wrażliwe dane osobowe setek milionów ludzi (Ng, 2018).

Miary nieszczęścia dopełnił fakt, że oprócz niedociągnięć w zakresie monitorowania i segmentacji sieci, firma stosowała niewystarczające, zupełnie podstawowe praktyki zarządzania danymi. Administratorzy systemów nie korzystali z mechanizmów uwierzytelniania wieloskładnikowego, a przeprowadzona później analiza powłamaniowa wykazała, że administratorzy używali bazy danych przechowującej niezaszyfrowane nazwy kont użytkowników i ich hasła dostępu (Schwartz, 2018). Odkrycie tej bazy danych przez napastników znakomicie ułatwiło im szybkie i bardziej skuteczne ataki na inne bazy danych firmy.

Co więcej, napastnicy po uzyskaniu dostępu wykonali ponad 9000 zapytań do baz danych (Government Accountability Office (GAO), 2018). Tak ogromna ilość dodatkowych zapytań sama w sobie powinna być więcej niż wystarczająca, aby wywołać alarm i uruchomić analizę aktywności, ale znów, dzięki wspomnianemu już wcześniej problemowi z certyfikatem i nadmiernie rozbudowanej infrastrukturze sieci wykorzystującej wiele wzajemnych, zdublowanych połączeń, aktywność napastników pozostała zupełnie niezauważona.

Kiedy przyjrzymy się całej historii od strony menedżerskiej, okazuje się, że kierownictwo firmy próbowało obarczyć całkowitą odpowiedzialnością za tę spektakularną katastrofę tylko jednego pracownika (Brandom, 2017). Choć oczywiście z całą pewnością ktoś był odpowiedzialny za zarządzanie mechanizmem instalowania poprawek bezpieczeństwa i aktualizacji oprogramowania, w rzeczywistości głównymi przyczynami, które doprowadziły do nieszczęścia, były błędy techniczne przy projektowaniu i wdrażaniu systemów bezpieczeństwa połączone z brakiem świadomości faktu, że pozornie bezpieczna architektura całego środowiska może znakomicie ułatwiać potencjalnemu napastnikowi przeprowadzenie skutecznego ataku i przez długi czas sprzyjać skrytemu podejmowaniu złośliwych działań.

Podobnie jak w poprzednich przypadkach, nadmierne uprawnienia użytkowników, niepoprawna segmentacja sieci, zbyt szerokie prawa dostępu i niewłaściwy nadzór nad bezpieczeństwem danych w połączeniu z modelem, który po przełamaniu zewnętrznego perymetru pozwalał na niemal nieograniczony dostęp i szybkie przemieszczanie się napastników w sieci wewnętrznej, były tym, co ostatecznie przesądziło o skuteczności włamania do systemów firmy Equifax.

Choć koszty finansowe tej ogromnej porażki nadal nie zostały jeszcze do końca oszacowane, można śmiało powiedzieć, że niemal wszyscy obywatele Stanów Zjednoczonych zostali narażeni na ujawnienie informacji kredytowych, co mogło mieć ogromny wpływ na ich zdolność do ubiegania się o nowe kredyty. Szacuje się również, że oprócz tego ucierpiało ponad 15 milionów obywateli Wielkiej Brytanii i dziesiątki tysięcy obywateli Kanady. Sama firma Equifax, która jest odpowiedzialna za przetwarzanie informacji dotyczących wiarygodności kredytowej prawie połowy populacji amerykańskiej, szacuje obecnie swoje straty na 1,3 miliarda dolarów. Warto zauważyć, że przedstawione wyliczenia nie obejmują całkowitych kosztów, jakie firma musiała ponieść na modernizację swojej sieci korporacyjnej i wdrażanie w niej zmian.

Jak zauważyliśmy już w poprzednim podrozdziale, bardzo często ogromne systemy korporacyjne są zbudowane tak, że stanowi to swego rodzaju zaproszenie do spektakularnej katastrofy. Opisany paradygmat ma zastosowanie nawet do — wydawałoby się — znakomicie zabezpieczonych sieci różnych agend rządowych. W następnym podrozdziale pokażę, jak powszechnie taki nieskuteczny model bezpieczeństwa jest stosowany w wielu organizacjach rządowych; omówię ponadto przykładowe skutki, jakie przyniosły włamania do takich systemów.

Perymetry systemów rządowych też zawodzą

Nawet agencje rządowe mogą paść ofiarą nieudanego podejścia do zagadnień bezpieczeństwa. Amerykańskie biuro ds. zarządzania personelem (US Office of Personnel Management — OPM) jest jedną z najważniejszych agencji w systemie federalnym Stanów Zjednoczonych. Agencja ta jest odpowiedzialna za przechowywanie i przetwarzanie informacji kadrowych dotyczących wszystkich pracowników zatrudnionych przez rząd federalny Stanów Zjednoczonych. Obejmuje to dane osobowe milionów obecnych i byłych pracowników federalnych, personelu wojskowego, a także pełne informacje o wszystkich dochodzeniach przeprowadzanych w celu wydania poświadczeń bezpieczeństwa wydawanych przez Departament Obrony do weryfikacji dostępu pracowników do najbardziej tajnych agencji, programów i operacji rządowych. W zasadzie należałoby założyć, że przy przetwarzaniu tego rodzaju wrażliwych danych o bezcennej wprost wartości, system komputerowy tej agencji powinien być jednym z najlepiej zabezpieczonych systemów w całym Departamencie Obrony. Nic bardziej mylnego.

Podobnie jak w przypadku firm Equifax i Maersk, włamanie do sieci agencji US Office of Personnel Management było rezultatem wielu błędów, sukcesywnie popełnianych przez długie lata podczas projektowania i wdrażania, efektem czego był wszechstronny system zabezpieczeń, który jednak mógł runąć jak domek z kart, gdyby tylko potencjalnemu napastnikowi udało

się w jakiś sposób przedostać przez fosę i potężne mury obronne zamku. W przypadku agencji OPM taka „dziura w murze” miała postać pozornie niewinnej wiadomości e-mail, która zawierała złośliwego trojana PlugX, pozwalającego na zdalny dostęp do zainfekowanego komputera. Złośliwe oprogramowanie, nieświadomie uruchomione przez jednego z odbiorców wiadomości, zostało w ten sposób wprowadzone „za mury”, bezpośrednio do wewnętrznej sieci OPM.

Po otwarciu załącznika w wiadomości pocztowej, feralny użytkownik nie miał żadnego pojęcia o uruchomieniu jego ukrytej, złośliwej aktywności. Trojan, który został specjalnie zmodyfikowany tak, aby uniknąć wykrycia przez oprogramowanie antywirusowe, zaczął się mnożyć i pobierać, a następnie instalować złośliwe pliki DLL i inne pliki binarne pozwalające na przejście przez napastnika kontroli nad zainfekowanym komputerem.

Podobnie jak w przypadku wielu innych, podobnych scenariuszy ataków wykorzystywanych w przeszłości, złośliwe oprogramowanie zrobiło dokładnie to, co złośliwe oprogramowanie zazwyczaj robi w podobnych przypadkach — wykorzystywało dostęp na prawach danego użytkownika oraz słabą wewnętrzną segmentację do przeszukiwania sieci w celu znalezienia bardziej wartościowych celów. W przypadku ataku na agencję OPM, takim celem były komputery wyposażone w narzędzia typu PAM (ang. *Privileged Access Management*) lub inaczej narzędzia do zarządzania uprzywilejowanym dostępem dla administratorów (Koerner, 2016). W uproszczeniu mówiąc, są to komputery, które gromadzą i przechowują w swoich repozytoriach dane uwierzytelniające uprzywilejowanych kont użytkowników, którzy mają uprawnienia do zarządzania poszczególnymi systemami w całym środowisku sieciowym.

Szczegółowa analiza powłamaniowa pozwoliła na zidentyfikowania feralnego użytkownika, który jako pierwszy uruchomił złośliwy załącznik z otrzymanego e-maila, i wykazała, że zainfekowana wiadomość nadeszła z konta pocztowego zewnętrznego kontraktora, który pracował jako administrator systemu w sieci OPM. W trakcie dochodzenia okazało się, że sieć komputerowa firmy, w której pracował kontraktor, została skutecznie zaatakowana i była w ukryty sposób kontrolowana przez napastników co najmniej od roku przed rozpoczęciem ataku na sieć agencji OPM.

Napastnicy działali w dyskretny sposób i bardzo skutecznie zacierali za sobą ślady. Udało im się usunąć pliki dzienników aktywności systemu, a przed przesyłaniem wykradzionych danych dzielili użyteczne dla nich pliki na małe fragmenty, co skutecznie pozwoliło na uniknięcie wykrycia przez systemy OPM monitorujące działanie sieci, które powinny zapobiegać właśnie takim sytuacjom. Cierpliwość i spryt, którymi wykazali się napastnicy podczas włamania do OPM, pozwoliły im na wykonanie kopii wielu krytycznych i najbardziej utajnionych danych, z których korzystał rząd federalny we wszystkich swoich agencjach.

W żadnym z omawianych przez nas przypadków napastnicy nie używali cudownych supernarzędzi ani wyrafinowanych, innowacyjnych technologii rodem z filmów science-fiction, pozwalających im w niemal magiczny sposób pokonywać zabezpieczenia zaatakowanych systemów. Niemal we wszystkich przypadkach skutecznych ataków i włamań, które miały miejsce w ciągu ostatnich czterech dekad, porażka systemów zabezpieczających była rezultatem ich uzależnienia od modelu bezpieczeństwa zbudowanego w oparciu o umocnione, zewnętrzne perymetry sieciowe, w połączeniu z nieodpowiednimi lub w najlepszym przypadku mało skutecznymi praktykami zarządzania.

To właśnie ze względu na możliwość łatwego przeskakiwania napastników z jednego systemu na inne, nadmiernych uprawnień użytkowników i braku odpowiedniego monitorowania tego, co działo się w ciemnych zakątkach infrastruktury sieciowej, to, co powinno być tylko niewielką, chwilową uciążliwością, stało się prawdziwą katastrofą o epickich rozmiarach. Perymetryczny model bezpieczeństwa cybernetycznego poniósł zdecydowaną porażkę w swojej najbardziej podstawowej funkcji: obronie zewnętrznego pogranicza infrastruktury sieciowej.

Istnieje jednak większe i jeszcze bardziej zagmatwane zagadnienie, które w przyszłości będzie spędzać sen z powiek wielu przedsiębiorstwom, małym i dużym firmom, a być może kiedyś nawet całym państwom. Zdobywająca coraz większą popularność koncepcja BYOD („przynies swoje własne urządzenie”) stawia zupełnie nowe wyzwania, które szeroko otwierają napastnikom drzwi do nowych rodzajów ataków. W następnym podrozdziale będziemy omawiać konsekwencje tego ciekawego zjawiska.

Użytkownicy, BYOD i całkowity zanik perymetru sieciowego

W ciągu ostatniej dekady moc obliczeniowa dostępna dla użytkowników, urządzeń i aplikacji rosła niemal wykładniczo, a wraz nią pojawiała się coraz większa, wieloaspektowa mozaika potencjalnych rodzajów przyszłych awarii infrastruktury sieciowej. Dodajmy do tego rosnącą złożoność i popularność usług oferowanych przez chmury obliczeniowe oraz problemy z utrzymaniem kontroli i zarządzania tymi wszystkimi komponentami, które domyślnie funkcjonują poza granicami naszego perymetru sieciowego, a okaże się, że sytuacja zaczyna się pogarszać z prędkością światła.

Jeszcze nie tak dawno temu użytkownicy musieli być fizycznie obecni w miejscu pracy, aby mieć możliwość połączenia się z siecią swojej firmy czy skorzystania z jej zasobów komputerowych. Jednak na przestrzeni ostatnich dwóch dekad zmniejszenie kosztów komputerów osobistych i ogromny wzrost mocy, jaką dysponują te urządzenia, przyniosło znaczące korzyści dla całej populacji użytkowników, ale jednocześnie poważnie zaburzyło bezpieczeństwo infrastruktury sieciowej. Konieczność przyjęcia przez korporacje i agencje rządowe nowej koncepcji funkcjonowania w coraz bardziej mobilnym i nastawionym na pracę zdalną ekosystemie sieciowym, stwarza poważne problemy dla tych, których zadaniem jest opracowanie mechanizmów kontroli bezpieczeństwa takiego rozproszonego środowiska i zarządzania nim.

W większości środowisk korporacyjnych preferowaną metodą zabezpieczenia zdalnego dostępu dla tych użytkowników, którzy pracują zdalnie na swoich urządzeniach (BYOD) bądź znajdują się w odległych geograficznie lokalizacjach, są sieci VPN (ang. *Virtual Private Network* — wirtualna sieć prywatna). Takie rozwiązania są dostępne od początku lat 90. i choć pozwalają na znaczne zminimalizowanie ryzyka związanego z potencjalnymi błędami w zabezpieczeniach, to jednak nie czyni ich to odpornymi na różnego rodzaju ataki.

Sieci VPN używane przez korporacje i innych użytkowników komercyjnych to nic innego jak proste aplikacje, które wykorzystują protokoły tunelowania do nawiązania połączenia. Odbywa

się to na różne sposoby. Większość sieci VPN, zarówno korporacyjnych, jak i komercyjnych, używa określonego protokołu do przesyłania i szyfrowania danych. Protokół docelowy wybrany dla połączenia VPN jest wynikiem porozumienia uzgodnionego między dwoma punktami końcowymi na podstawie zdefiniowanego wcześniej zestawu reguł transmisji danych i algorytmów szyfrowania. Wielu komercyjnych dostawców rozwiązań VPN zapewnia możliwość wyboru spośród kilku różnych protokołów VPN w zależności od potrzeb użytkowników w zakresie bezpieczeństwa. Co ciekawe, z reguły takiej możliwości nie posiada większość rozwiązań VPN przeznaczonych dla sieci korporacyjnych i rządowych. Najpopularniejsze protokoły używane w sieciach VPN to między innymi:

- PPTP (ang. *Point-to-Point Tunneling Protocol*)
- L2TP (ang. *Layer Two Tunneling Protocol*)
- IPSec (ang. *Internet Protocol Security*)
- SSL/TLS OpenVPN

Podstawowym zadaniem sieci VPN jest zabezpieczenie przesyłanych danych przed odczytaniem przez osoby trzecie poprzez zastosowanie zaawansowanych narzędzi szyfrujących i odpowiednich protokołów sieciowych. Dzieje się tak, gdy strumień danych przekazywane otwartym tekstem są szyfrowane i zamieniane w nieczytelny strumień zaszyfrowanych danych. Każde rozwiązanie sieci VPN wykorzystuje odpowiednio dobrany algorytm do szyfrowania i odszyfrowywania transmitowanych strumieni danych. Oczywiście każdy z protokołów używanych w sieciach VPN ma swoje mocne i słabe strony. Siła protokołu opiera się na algorytmie kryptograficznym użytym do szyfrowania danych.

Ataki na połączenia VPN wykorzystują jedną z dwóch taktyk. Haker może przełamać szyfrowanie wykorzystując słabości algorytmu bądź luki w zabezpieczeniach, może też wejść w posiadanie klucza szyfrowania w inny, zwykle mało etyczny sposób. Hakerzy i pracownicy komórek kryptoanalitycznych, którzy zawodowo zajmują się łamaniem szyfrów, mogą również wykorzystywać różnego rodzaju ataki kryptograficzne do odszyfrowania zaszyfrowanych danych bez posiadania klucza.

Łamanie szyfrów jest jednak zadaniem bardzo czasochłonnym i wymagającym dużej mocy obliczeniowej. Przełamanie szyfrowania może wymagać wielu lat pracy silnego komputera (choć ten czas można znacznie skrócić dzięki zastosowaniu chmury obliczeniowej lub technologii komputerów kwantowych). Z tego względu zdecydowana większość ataków na sieci VPN polega na kradzieży kluczy. Biorąc pod uwagę, że obliczenia niezbędne do przeprowadzenia próby złamania klucza są bardzo złożone (a zastosowanie komputerów kwantowych czy chmury obliczeniowej nie zawsze jest możliwe), kradzież klucza wydaje się być znacznie łatwiejszym zadaniem. Pomyślne przeprowadzenie ataku na sieci VPN jest zazwyczaj rezultatem połączenia skutecznego podstępu, wyrafinowanej inżynierii społecznej, odpowiednio przygotowanego kłamstwa i dostępu do dużej mocy obliczeniowej.

Wszystko, czego potrzebuje napastnik do przygotowania ataku, to proste skanowanie portów środowiska celu. Większość rozwiązań VPN używanych przez przedsiębiorstwa i konsumentów można rozpoznać według numerów portów, których używają do realizacji połączeń. Skanowanie portów w sieci docelowej może wykryć m.in. porty, które pokazano w poniższym zestawieniu:

- OpenVPN wykorzystuje następujące porty:
 - UDP: 1194, 1197, 1198, 8080, 9201
 - TCP: 502, 501, 443
 - L2TP: 1701
 - UDP 500, 1701 i 4500
- IKEv2 wykorzystuje następujące porty:
 - UDP 500
- PPTP wykorzystuje następujące porty:
 - TCP 1723 dla protokołu 47 (GRE)

Analiza wyników skanowania portów może od razu stanowić dla napastnika wskazówkę, że środowisko celu używa takiego czy innego rozwiązania VPN, i skłonić go do przygotowania ataku mającego na celu pozyskanie w taki czy inny sposób odpowiednich kluczy szyfrowania. Jeszcze łatwiejszą metodą ataku na połączenie VPN, która wbrew pozorom jest często wykorzystywana w praktyce, jest po prostu obserwacja wybranego użytkownika korporacyjnego w miejscach publicznych, takich jak kawiarnia, biblioteka czy parking, oczekiwanie, aż użytkownik zaloguje się do swojej sieci VPN, a następnie odwrócenie jego uwagi i kradzież komputera, gdy takie połączenie VPN jest nadal aktywne.

W większości przypadków, jeżeli użytkownik się nie wyloguje lub komputer nie zostanie zablokowany, połączenie pozostanie aktywne przez długi czas i napastnik będzie mógł je wykorzystać w odpowiednim dla niego momencie.

Dostawcy usług VPN również mogą być celem ataków, czego przykładem mogą być ataki na firmy Avast i NordVPN w 2019 roku. Podczas tych ataków napastnicy byli w stanie wykorzystać tymczasowe dane uwierzytelniające, co stało się możliwe dzięki luce w zabezpieczeniach narzędzia do zdalnego zarządzania u dostawcy tymczasowego centrum obliczeniowego. Wykorzystanie tej luki zapewniło napastnikom niemal nieograniczony dostęp do zlokalizowanych w tym centrum obliczeniowym serwerów, które zarządzały szyfrowaną komunikacją sieci VPN oferowanych przez te firmy.

W trakcie włamania napastnikom udało się wykraść klucz szyfrowania TLS (ang. *Transport Layer Security*), co teoretycznie pozwalało im na przechwytywanie danych przesyłanych przez dowolnego z 12 milionów komercyjnych użytkowników firmy za pomocą kryptograficznego ataku typu „man-in-the-middle” (Kan, 2019). Od razu nasuwa się tutaj pytanie o to, jak wielu z tych klientów używało tych maszyn do łączenia się z innymi kontaktami biznesowymi i ilu z nich używało tych samych haseł dostępu do logowania do zasobów firmy?

Badania przeprowadzane przez wielu niezależnych ekspertów, którzy przeanalizowali setki dostawców VPN, wykazały, że:

- Większość rozwiązań VPN korzystających z protokołu SSL (ang. *Secure Sockets Layer*) nadal używa jego przestarzałej wersji SSLv3, która liczy sobie już ponad dwie dekady i nie jest już obsługiwana przez wiele nowoczesnych rozwiązań.

- Wiele sieci VPN używających protokołu SSL korzysta z niezaufanych i niezweryfikowanych certyfikatów SSL, które pozwalają na przeprowadzenie ataku typu „man-in-the-middle”.
- Algorytm SHA-1, którego bezpieczeństwo jest już od dawna kwestionowane, jest nadal w powszechnym użyciu w wielu rozwiązaniach VPN.
- Prawie 50% VPN-ów SSL do swoich certyfikatów RSA używa kluczy o rozmiarze 1024 bitów. W obecnych czasach klucze RSA o rozmiarach poniżej 2048 bitów nie są już uważane za bezpieczne ze względu na ich słabości kryptograficzne.
- Jedna na dziesięć sieci SSL VPN nadal opiera się na implementacji OpenSSL, a większość z nich jest nadal podatna na zagrożenia związane z exploitem Heartbleed, który liczy sobie już ponad pół dekady.
- Tylko około 5% VPN-ów SSL jest zgodnych z wymogami bezpieczeństwa PCI (ang. *Payment Card Industry*).
- Nie znaleziono ani jednego dostawcy VPN, który spełniałby wytyczne NIST (amerykańskiej organizacji rządowej odpowiedzialnej za standardy i regulacje dla przedsiębiorstw; odpowiednik naszego Głównego Urzędu Miar).

Na podstawie tych wyników badań można z bardzo dużym prawdopodobieństwem założyć, że większość z tych narzędzi, wykorzystywanych przez setki tysięcy użytkowników, przedsiębiorstw czy nawet rządów do pracy zdalnej (np. BYOD), posiada poważne luki w zakresie bezpieczeństwa.

Aplikacje zwiększają poziom zagrożenia

Kiedy uświadomimy sobie zagrożenia, jakie technologia sieci VPN wprowadza do modelu bezpieczeństwa opartego na zewnętrznym perymtrze sieciowym, z pewnością dojdziemy do wniosku, że stanowi to bardzo poważny problem. Dodatkowym elementem, ściśle powiązanym z pracą zdalną i koncepcją BYOD, jest kwestia bezpieczeństwa aplikacji. Aplikacje są tym, czego każdy użytkownik, wszędzie, na każdym urządzeniu, używa do wykonywania zarówno swojej pracy, i jak i wielu innych zadań w codziennym życiu. Niestety w wielu przypadkach aplikacje są tworzone pod presją czasu (plany dystrybucji, terminy określone w umowie itp.), a kwestie ich bezpieczeństwa schodzą na dalszy plan. Fakt ten w praktyce oznacza, że wiele wspaniałych aplikacji, z których korzystamy na co dzień, nie zostało napisanych z myślą o wymogach bezpieczeństwa.

Według badań przeprowadzonych wspólnie przez Instytut Ponemon i firmę IBM, ponad 50% przedsiębiorstw nie posiada żadnego budżetu na bezpieczeństwo aplikacji (Instytut Ponemon, 2016). Ponad 40% przedsiębiorstw nie sprawdza aplikacji pod kątem bezpieczeństwa przed wprowadzeniem ich do użytku produkcyjnego, a mniej więcej jedna trzecia aplikacji używanych w firmach i korporacjach nigdy nie była testowana pod kątem znanych luk w zabezpieczeniach. Zgodnie z raportem Hewlett Packard Enterprises (HPE) z 2016 roku, mniej więcej jedna na dziesięć aplikacji posiada hasła dostępu zakodowane na sztywno w swoim kodzie lub

w plikach konfiguracyjnych. Wreszcie, prawie połowa wszystkich aplikacji jest używana w przedsiębiorstwach, które przyznały, że nie posiadają żadnego lub prawie żadnego programu zarządzania instalowaniem aktualizacji i poprawek bezpieczeństwa.

Innymi słowy, organizacje te otwarcie przyznały badaczom, że nie mają wdrożonych żadnych metod wykrywania i identyfikowania słabych punktów nowych aplikacji, a większość z nich nie posiada również żadnych konkretnych planów radzenia sobie z lukami w zabezpieczeniach aplikacji, które już zostały wdrożone.

Można stąd wyciągnąć smutny wniosek, że bardzo często aplikacje, które są wykorzystywane przez użytkowników w ich codziennej pracy w przedsiębiorstwach, agencjach rządowych i w prywatnych zastosowaniach konsumenckich, nie są do końca bezpieczne. Oznacza to, że prędzej czy później każdy użytkownik wejdzie w interakcję z aplikacją posiadającą takie czy inne luki w zabezpieczeniach, co może doprowadzić do powstawania różnego rodzaju incydentów bezpieczeństwa. Incydenty takie mogą przybierać różne formy, począwszy od ataków typu „man-in-the-middle”, poprzez komplikacje związane z **bezpieczeństwem warstwy transportowej** (TLS), trudności z obsługą danych binarnych, problemy z bezpieczeństwem hasła, a skończywszy na wielu innych potencjalnych działaniach, które mogą prowadzić do poważnych konsekwencji związanych z modelem bezpieczeństwa opartym na zewnętrznym perymetrze sieciowym.

Choć aplikacje są w zasadzie z definicji budowane z różnymi błędami umieszczonymi na stałe w ich kodzie, istnieje jednak bardziej oczywisty problem, który od lat nęka praktyków bezpieczeństwa: hasła. W następnym podrozdziale omówimy podstawowe błędy powszechnie spotykane w przypadku tego najstarszego modelu uwierzytelniania i bezpieczeństwa dostępu, z którego kiedykolwiek korzystał człowiek.

Niepowodzenia metod uwierzytelniania

Hasła, jeden z najbardziej efektywnych sposobów uwierzytelniania stosowanych przez firmy, organizacje, korporacje i użytkowników w milionach systemów na naszej planecie, są jednocześnie główną przyczyną katastrof i niepowodzeń w cyberprzestrzeni. Prawie wszystko na pewnym etapie używa hasła. W zasadzie możemy powiedzieć, że niemal każda aplikacja, której używamy, a także każda sieć VPN, a nawet każdy komputer na świecie, używa hasła do uwierzytelniania. W podobny sposób korzystają z nich narzędzia administracyjne, bazy danych czy systemy zapór sieciowych. Wszystko wszędzie korzysta z hasła.

Chociaż wydaje się to stosunkowo prostym, sprawdzonym i użytecznym sposobem uwierzytelniania, hasła są bezpieczne tylko wtedy, gdy pozostają nieznane osobom, które nie są ich użytkownikami.

W ciągu połowy ostatniej dekady niemal każde większe repozytorium sieciowe, przechowujące dane użytkowników i ich hasła dostępu, odnotowało mniej lub bardziej poważne incydenty bezpieczeństwa skutkujące wyciekami danych. W 2019 roku niezależny badacz opublikował listę ponad 700 milionów kont e-mail i nazw użytkowników, które można było łatwo skorelować z ponad 20 milionami złamanych hasła.

Wspomniane nazwy kont użytkowników i hasła wyciekły w wyniku ataków hakerów na takie serwisy i firmy jak Yahoo, Equifax, OMB, Target, Home Depot i setki innych. Serwis Have I Been Pwned (HIBP) twierdzi, że jest w posiadaniu ponad 8 miliardów rekordów, które są wynikiem ponad 400 wycieków danych na całym świecie.

Ze względu na ilości ujawnionych danych uwierzytelniających możemy z niemal całkowitą pewnością założyć, że praktycznie każdy użytkownik na naszej planecie ma co najmniej jedno skompromitowane konto. Oczywiście jest fakt, że w internecie nie ma 8 miliardów użytkowników, a już z całą pewnością nie ma 8 miliardów użytkowników w żadnym systemie korporacyjnym, co wykładniczo zwiększa prawdopodobieństwo, że wiele z tych poświadczeń będzie nadawało się do wykorzystania w kolejnych atakach.

Hakerzy z powodzeniem mogą korzystać z taktyki nazywanej „credential stuffing” (z ang. zapychanie poświadczeniami), która polega po prostu na przeprowadzanie ataku typu „brute-force” na system docelowy, w celu uzyskania dostępu za pośrednictwem pozyskanych wcześniej danych uwierzytelniających. Wiele aplikacji nie posiada ograniczeń co do liczby nieudanych prób logowania, a jeżeli nawet takie ograniczenia zostały zaimplementowane, to można użyć prostego skryptu, który będzie oczekiwał na upłynięcie limitu czasu i ponawiał kolejne próby logowania. Takie rozwiązanie pozwala napastnikom na przeprowadzanie ciągłych, zautomatyzowanych ataków na środowisko celu, aż do momentu, kiedy nie zostanie znaleziony prawidłowy, działający zestaw poświadczeń.

Tajemnicą poliszynela jest fakt, że podziemie przestępcze, jak również wiele mniej lub bardziej oficjalnych organizacji sponsorowanych przez państwa narodowe, posiada ogromne bazy danych kont użytkowników i złamanych hasel, które były wielokrotnie wykorzystywane podczas ataków na różne systemy. W większości przypadków znalezienie zestawu prawidłowych danych uwierzytelniających do atakowanego systemu jest tylko kwestią czasu.

W ciągu 17 miesięcy zespół ds. bezpieczeństwa w firmie Akamai, która dysponuje zasobami sieciowymi rozmieszczonymi na całym świecie, w ostatnich latach wykrył ponad 50 miliardów prób ataków typu *credential stuffing* na różne cele (Constantin, 2019). Każda z tych miliardów prób mogła zakończyć się (a w niektórych przypadkach tak się właśnie stało) uzyskaniem dostępu do sieci i baz danych przechowujących wrażliwe dane korporacyjne lub rządowe. Wystarczy jedna para ważnych poświadczeń z wielu miliardów prób i wszystkie wyrafinowane mechanizmy zabezpieczające perymetru sieciowego firmy zaczynają się wykruszać.

Musimy także wziąć pod uwagę smutny fakt, że ogromna rzesza użytkowników korzysta z bardzo prostych i krótkich hasel. Według badań opublikowanych jeszcze w 2019 roku, dwa najbardziej rozpowszechnione hasła używane na świecie to *password* i *123456*. W 2015 roku niezależny instytut badawczy SplashData opublikował zestawienie najgorszych hasel stosowanych przez użytkowników. Co ciekawe, lista tych hasel praktycznie nie zmieniała się przez kolejnych kilka lat.

Miejsce	2018	2017	2016	2015
1	123456	123456	123456	123456
2	password	password	password	password
3	123456789	12345678	12345	12345678
4	12345678	qwerty	12345678	qwerty
5	12345	12345	football	12345
6	111111	123456789	qwerty	123456789
7	1234567	letmein	1234567890	football
8	sunshine	1234567	1234567	1234
9	qwerty	football	princess	1234567
10	iloveyou	iloveyou	1234	baseball
11	princess	admin	login	welcome
12	admin	welcome	welcome	1234567890
13	welcome	monkey	solo	abc123
14	666666	login	abc123	111111
15	acb123	abc123	admin	1qaz2wsx

Jak widać, choć użytkownicy są zwykle mniej lub bardziej świadomi ważności hasel i dostępów, które zapewnia ten mechanizm uwierzytelniania, mimo to nadal używają tych samych, łatwych do odgadnięcia, rażąco ignoranckich hasel w wielu aspektach swojego codziennego życia.

Do niezdolności użytkowników do stosowania odpowiednio złożonych hasel dochodzą inne przykłady nieskutecznych praktyk bezpieczeństwa i modeli zabezpieczeń opartych na perymtrze sieciowym, a zwłaszcza to, że wszystko obraca się wokół wykorzystywania hasel do uwierzytelniania dostępów, a w większości małych i średnich firm i organizacji takie trywialne, niebezpiecznie proste hasła nie są umieszczane na systemowej czarnej liście zapobiegającej możliwości ich użycia. Jak się okazało, nawet w tak dużej firmie jak Equifax hasło do jednego z niezmiernie ważnych zasobów sieciowych brzmiało po prostu... *admin*.

Stwierdzono, że nawet członkowie Kongresu i wiele znanych osobistości świata mediów używają słabych i niepewnych metod uwierzytelniania i hasel. Członka Izby Reprezentantów USA, Lance’a Goodena z Teksasu, który był współwnioskodawcą projektu ustawy *Cybersecurity and Financial System Resilience Act of 2019*, widziano, jak podczas posiedzenia komisji kongresowej odblokował swój telefon za pomocą kodu 7777777. Uważny obserwator mógł również zauważyć, że podczas telewizyjnego spotkania z prezydentem Donaldem Trumpem amerykański raper i producent muzyczny Kanye West do odblokowania telefonu używał kodu 0000000. Można by pomyśleć, że tak wysoko postawione osoby, zwłaszcza zajmujące się opracowywaniem przepisów dotyczących bezpieczeństwa cybernetycznego w sektorze bankowym, będą zdawały sobie sprawę z powagi sytuacji i będą używały silnych, złożonych hasel i bezpiecznych metod uwierzytelniania, ale — niestety — bardzo często tak się nie dzieje.

Prosta logika i zdrowy rozsądek sugerowałyby, że jeżeli istnieje jakiekolwiek hasło, które jest praktycznie niemożliwe do złamania, utworzone przy użyciu skomplikowanych schematów zapobiegających możliwości nieautoryzowanego użycia chronionego systemu, to będzie ono stosowane w programie amerykańskich rakiet balistycznych Minuteman z głowicami atomowymi. Jednak w notatce z 2004 roku, dr Bruce Blair, były oficer wojsk rakietowych, stwierdził, że „amerykańskie Dowództwo Strategicznych Sił Powietrznych co najmniej raz celowo ustawiło kody startowe we wszystkich silosach rakiet atomowych Minuteman na terenie Stanów Zjednoczonych na ciąg składający się z ośmiu zer”.

W 1962 roku prezydent Kennedy nakazał swojemu sekretarzowi obrony, Robertowi McNamarze, zainstalować system zabezpieczający o nazwie PAL we wszystkich rakietach Minuteman z głowicami atomowymi, znajdujących się w amerykańskim arsenale jądrowym. Jednak dzięki opieszałości Sił Powietrznych USA w instalowaniu tych urządzeń oraz ogólnej niechęci dowódców Sił Powietrznych USA do sekretarza McNamary, wdrożenie tych zmian zajęło ponad dwie dekady.

Dr Blair w swojej notatce stwierdził, że według standardowej procedury operacyjnej oficerowie amerykańskich sił strategicznych obsługujący wyrzutnie rakiet Minuteman powinni „zgodnie z listą kontrolną dwukrotnie sprawdzić wprowadzaną sekwencję kodów odblokowujących panel uruchamiający procedurę odpalenia z podziemnego bunkra startowego, aby upewnić się, że nie zostały w niej przypadkowo wprowadzone cyfry inne niż zero”. Innymi słowy, oficerowie kierujący odpalaniem rakiet zostali poinformowani, że sekwencja 00000000 została na stałe umieszczona w kodach umożliwiających odblokowanie paneli sterujących pozwalających na rozpoczęcie procedury wystrzelenia ponad 50 rakiet balistycznych Minuteman z głowicami atomowymi.

Nie oznacza to wprawdzie, że łatwiej było przeprowadzić nieautoryzowane wystrzelenie rakiet balistycznych (istnieje wiele innych, skomplikowanych procedur kontrolnych, które należy przeprowadzić przed odpaleniem), niemniej jeden z krytycznych elementów sekwencji startowej dla strategicznej broni jądrowej USA opierał się na banalnie prostym, ośmiocyfrowym kodzie składającym się z samych zer.

Opisana wyżej historia o programie rakiet Minuteman jest tylko jednym z wielu przykładów, ale chodzi głównie o to, że nawet w tak ściśle zhierarchizowanej i zdyscyplinowanej organizacji jak Siły Powietrzne USA, zarządzanie hasłami może być żałośnie nieudolne. Jeżeli organizacja o tak dużych możliwościach i tak ogromnej odpowiedzialności może ignorować najlepsze praktyki w zarządzaniu hasłami przez ponad 20 lat, to jakie szanse ma przeciętne przedsiębiorstwo lub zwykły użytkownik?

Urządzenia IoT robią dziury w każdym perymetrze

Urządzenia internetu rzeczy (ang. *Internet of Things* — IoT) są obecnie jednymi z najszybciej rozwijających się segmentów rynku urządzeń sieciowych na świecie. Szacuje się, że w 2019 roku do sieci internet podłączonych było ponad 6 miliardów takich urządzeń, a liczba ta lawinowo rośnie. Urządzenia IoT obsługują połączenia i aplikacje sieciowe, wymagają użycia haseł do uwierzytelniania i są często projektowane i wytwarzane w krajach, w których działają sponsorowane przez rząd grupy hakerów. Innymi słowy, takie urządzenia są „genetycznie” obciążone pewnym stopniem niepewności praktycznie od momentu, w którym zejść z linii montażowej producenta, a warto zauważyć, że w sieciach zdecydowanej większości firm i organizacji na świecie, nie wspominając już nawet o sieciach domowych, można znaleźć mniejszą bądź większą liczbę różnych urządzeń IoT.

Niezależnie od tego, czy są to inteligentne telewizory, inteligentne termostaty, drukarki bezprzewodowe, kamery z dostępem do internetu, czy inne urządzenia, możemy mieć pewność, że w niemal każdej infrastrukturze sieciowej działają jakieś urządzenia IoT.

Powszechnie stosowane w urządzeniach IoT niestandardowe, bezprzewodowe protokoły komunikacyjne są jednym z głównych wektorów ataków wykorzystywanych przez hakerów i rządowe agendy cybernetyczne. Istnieje bardzo wiele protokołów komunikacyjnych używanych przez producentów takich urządzeń. Poniżej wymieniono tylko dwa z głównych protokołów i związane z nimi luki w zabezpieczeniach. Pełna lista wszystkich potencjalnych problemów z urządzeniami IoT jest zdecydowanie zbyt długa, aby ją zmieścić w jednej (nawet bardzo grubej) książce:

- **Protokół ZigBee** — nasłuchiwanie i przechwytywanie procesu wymiany kluczy szyfrowania umożliwia ataki typu „man-in-the-middle” i pozwala na wymuszenie przywrócenia ustawień fabrycznych, w wyniku czego urządzenie może automatycznie połączyć się z dowolną dostępną siecią, która może być złośliwą siecią kontrolowaną przez napastnika w celu gromadzenia przesyłanych danych (Zillner, 2015).
- **Protokół NFC** — posiadając odpowiednią wiedzę techniczną, można manipulować również połączeniami NFC, zmuszając zaatakowane urządzenie np. do uruchomienia przeglądarki sieciowej, połączenia się ze złośliwą witryną i pobrania złośliwego oprogramowania, przesyłania danych użytkownika, wykonywania niechcianych połączeń, a nawet wysyłania wiadomości SMS.

Warto tutaj wspomnieć, że w ostatnich czasach źródłem wycieku danych stały się nawet nowoczesne żarówki LED sterowane bezprzewodowo, które bywały wykorzystywane przez hakerów do przeprowadzania ataków na sieci bezprzewodowe poza granicami ich budynków. Sam charakter takich urządzeń, których liczba i popularność rosną wręcz lawinowo, i powody, dla których z nich korzystamy (ułatwianie sobie życia i chęć wykonywania wielu czynności „zdalnie”), są jednocześnie czynnikami znakomicie ułatwiającymi przeprowadzenie skutecznego ataku. Łatwość użytkowania, współdzielenie danych, dostępność aplikacji, mniej lub bardziej poważne błędy w oprogramowaniu, a czasem i celowo zaimplementowane tylnie

furtki w kodzie mogą przyczyniać się do powstawania poważnych luk w zabezpieczeniach każdej sieci, w której takie urządzenia są stosowane. Żaden perymetr sieciowy, w którym działają zainstalowane urządzenia IoT, nie powinien być uważany za bezpieczny.

Niestety, niezależnie od tego, jak dobrze bądź źle mogą być skonfigurowane urządzenia IoT, to jednak użytkownicy, którzy z nich korzystają i działają w sieci, bardzo często są na bakier z zasadami bezpieczeństwa i stają się prawdziwym źródłem kłopotów. W następnej części przeanalizujemy kwestie związane z podstawową edukacją użytkowników, szkoleniami i nieskutecznymi praktykami, które utrudniają zarządzanie bezpieczeństwem i prawie uniemożliwiają jego utrzymanie.

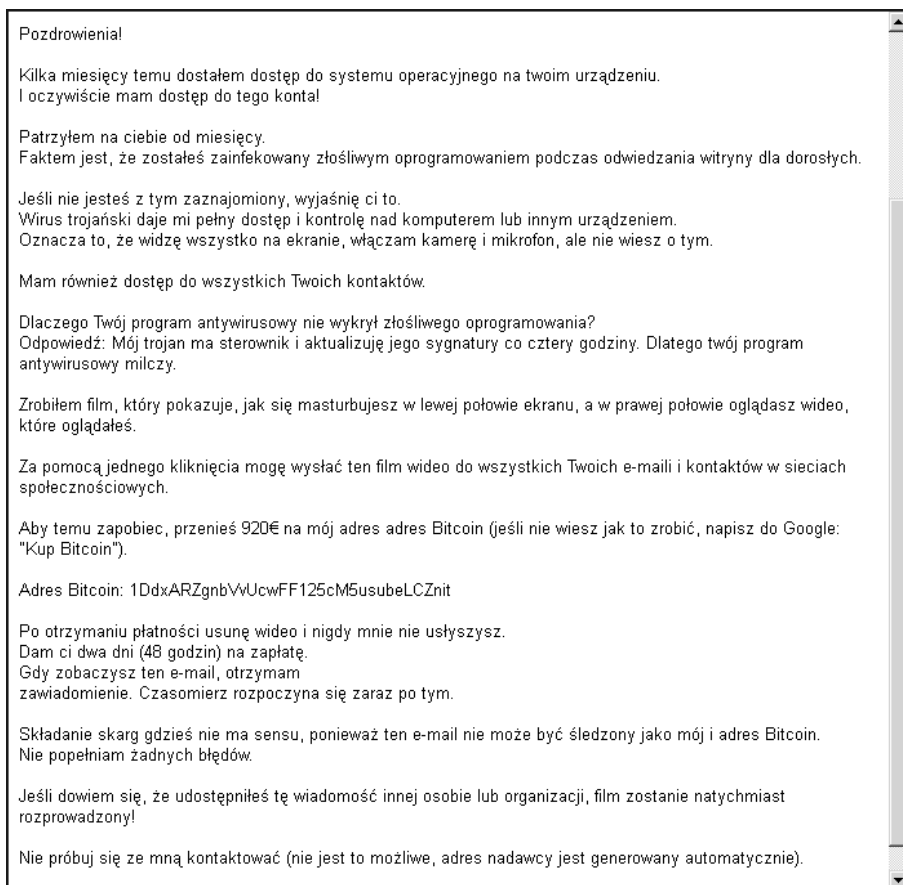
Nie możesz naprawić głupoty ani złośliwości

W idealnym świecie żaden człowiek nigdy nie dotknąłby sieci. Maszyny robiłyby wszystko, a ludzie po prostu korzystaliby z ich usług. Maszyny działają logicznie i z naciskiem na funkcjonalność. Nielatwo je oszukać i zazwyczaj są odporne na sugestie z wykorzystaniem inżynierii społecznej. Ale póki co nie żyjemy jeszcze w świecie przyszłości, w którym maszyny wszystko robią za nas. Wciąż mamy użytkowników, którzy korzystają z sieci, a często ich „radosna twórczość” bądź przypadkowe czy zupełnie bezmyślne działania powodują problemy i sytuacje, które mogą całkowicie sparaliżować coś, co do tej pory wydawało się być względnie bezpieczną siecią. Warto w tym kontekście pamiętać o następujących kwestiach:

- Po pierwsze, najbezpieczniejsza sieć to taka, której nigdy nie dotknął żaden użytkownik. Po drugie, w chwili, kiedy użytkownik położy palce na klawiaturze i połączy się z siecią, zagrożenie związane z czynnikami ludzkimi, atakami z wykorzystaniem inżynierii społecznej, phishingiem i innymi standardowymi metodami „pracy” hakerów, staje się rzeczywistością. Podczas gdy technologia ma stosunkowo binarną naturę, ludzie nie są tacy. Jesteśmy otwarci na wpływy, naciski, strach, głupotę i emocje. Tam, gdzie maszyna po prostu nie otworzyłaby wiadomości e-mail, która z daleka wygląda podejrzanie lub ma podejrzane załączniki, użytkownik może go bezmyślnie kliknąć, ponieważ znajduje się w nim uroczne zdjęcie małego kociaka.
- Obecnie w cyberprzestrzeni nadrzędną metodą zabezpieczenia użytkowników jest szkolenie ich w zakresie rozpoznawania potencjalnie złośliwych działań mających miejsce w ich sieci i systemach. Szkolenia takie odbywają się zazwyczaj poprzez połączenie kontrolowanego phishingu i materiałów edukacyjnych dostępnych online. Jednak pomimo tego, że bardzo często regularne przeprowadzanie takich szkoleń skutkuje wyraźnym, możliwym do sprawdzenia procentowym spadkiem wskaźnika kliknięć, wprowadzenie złośliwego oprogramowania do sieci wymaga tylko jednego podatnego użytkownika i jego jednego kliknięcia. Bez względu na to, jak dobrze przeszkoleni będą użytkownicy i jak aktualne są materiały szkoleniowe, w większości firm i organizacji praktycznie stały współczynnik kliknięć w odpowiednio spreparowane linki przesyłane w kontrolowanych, szkoleniowych

kampaniach phishingowych wynosi zazwyczaj od 3% do 5%. Chociaż na pierwszy rzut oka może się wydawać, że to naprawdę niewiele, to jednak wystarczy wziąć pod uwagę fakt, że w dużej firmie, zatrudniającej np. 500 000 pracowników, wspomniane „niewielkie” 3% przekłada się na naprawdę znaczną liczbę podatnych użytkowników i ich maszyn, będących potencjalnymi wektorami ataku dla napastników.

- Ludzie często stają się również podatni na ataki, gdy w grę wchodzi wstyd i zastraszanie w cyberprzestrzeni. W 2019 roku na światowej scenie pojawiła się taktyka o nazwie „sextortion”, czyli cybernetyczny szantaż seksualny. Jest ona niezwykle prosta, ale jednocześnie bardzo skuteczna. W jednym z takich ataków, napastnicy weszli w posiadanie ponad 400 MB zawierającego adresy e-mail pliku, który wyciekł do sieci podczas jednego z wcześniejszych włamań. Następnie, korzystając z wielu podstawionych, praktycznie niemożliwych do wyśledzenia kont e-mail, wysyłali tysiące sfabrykowanych, szantażujących wiadomości do potencjalnych celów. Przykładowy wygląd takiej wiadomości został przedstawiony poniżej:



Rysunek 1. Przykład wiadomości e-mail użytej w kampanii typu „sextortion”

Jedną z najbardziej aktywnych kampanii była powiązana z botnetem o nazwie Phorpiex, za pomocą którego automatycznie były wysyłane wiadomości e-mail. Badacze z firmy CheckPoint zajmującej się cyberbezpieczeństwem oszacowali, że z zainfekowanych komputerów znajdujących się w tym botniecie wysyłanych było średnio około 30 000 e-maili na godzinę. Botnet Phorpiex korzystał ze złośliwego oprogramowania, które w regularnych odstępach czasu pobierało nową, nieustannie aktualizowaną bazę danych adresów e-mail z serwera C&C, kontrolowanego przez napastników.

Bazy danych używane przez botnet Phorpiex zawierały listy adresów skrzynek e-mail wraz z przypisanymi do nich prawdziwymi hasłami dostępu, które wyciekły podczas innych ataków. Takie rozwiązanie bardzo pomagało napastnikom uwiarygodnić atak i zwiększało prawdopodobieństwo, że ofiara się na niego nabierze. Bardzo często zdarzało się, że nawet osoby, które nie miały żadnego kontaktu z pornografią w sieci, ze strachu płaciły okup, ponieważ naprawdę uwierzyły, że ktoś monitoruje ich zainfekowany komputer lub telefon. Okup był płacony w bitcoinach, co nie zostawiało żadnych śladów finansowych i praktycznie uniemożliwiało wysledzenie sprawców ataku.

W ostatnich miesiącach takie ataki stały się jeszcze bardziej ukierunkowane i złośliwe, ponieważ napastnicy zaczęli odsprzedawać innym hakerom listy osób, które już zapłaciły okup. Po otrzymaniu takiej listy nowa grupa cyberprzestępców ponownie kontaktowała się z tymi samymi osobami, ale zamiast wymuszać bitcoiny, przestępcy „prosilili” o podanie nazw kont użytkowników i haseł dostępu do określonych systemów. Takie metody wykorzystują stres użytkownika i opierają się na większym prawdopodobieństwie, że osoby, które zapłaciły już raz w przeszłości, mają coś do ukrycia i łatwiej będzie z nich „wydusić” poufne informacje dające dostęp do sieci. Jeżeli którakolwiek z tak naciskanych osób będzie dyrektorem wysokiego szczebla lub administratorem systemu z wyższymi uprawnieniami w sieci, to taka sytuacja mogłaby przynieść katastrofalne skutki dla tej firmy czy organizacji.

Zainfekowanie i przejęcie przez napastnika kontroli nawet nad pojedynczym komputerem stanowi oczywiście poważny problem, ale istnieje jednak znacznie poważniejsze niebezpieczeństwo związane z użytkownikami: zagrożenia wewnętrzne. „Krety” to wrogo nastawieni, pozornie „zaufani” użytkownicy — aktualni lub byli pracownicy, kontrahenci lub współpracownicy, posiadający informacje wewnętrzne dotyczące praktyk bezpieczeństwa organizacji oraz danych i systemów komputerowych, mający określoną motywację lub powód, aby w złośliwy czy wręcz wrogi sposób wykorzystywać infrastrukturę firmy od wewnątrz. Motywacje takich użytkowników mogą być bardzo różne, począwszy od wymiaru pieniężnego, poprzez powody polityczne, aż do bardzo emocjonalnych powodów osobistych. Niezależnie jednak od rodzaju motywacji, potencjalne działania takiego złośliwego, wewnętrznego użytkownika mogą mieć skutki wielokrotnie groźniejsze niż nieumyślne kliknięcie linku przez kogoś, kto nie ma złych zamiarów.

Kiedy wrogo nastawiony „kret”, mający dostęp do poufnych informacji, podejmuje decyzję o przeprowadzeniu złośliwego działania przeciwko sieci lub infrastrukturze swojej firmy, z reguły jest już zweryfikowanym, zaufanym użytkownikiem i zazwyczaj jest wyposażony w odpowiednią wiedzę i narzędzia, których potrzebuje, aby naprawdę zaszkodzić. Użytkownicy w niemal każdej firmie posiadają mniejszy bądź większy zakres uprawnień administracyjnych,

mają dostęp do zasobów sieciowych, własności intelektualnej i znają wewnętrzną specyfikę swojej organizacji.

W ciągu ostatniej dekady bardzo często zdarzało się, że złośliwi użytkownicy mający dostęp do poufnych informacji mogli bez przeszkód poruszać się po zasobach sieciowych, ponieważ nie były one dobrze monitorowane. Edward Snowden, Bradley (Chelsea) Manning, Jason Needham, Walter Liew, Robert Hanson i wielu innych byli w stanie zebrać cenne dane z sieci swoich pracodawców, a następnie siać spustoszenie w ich systemach. Nawet NSA, z całą swoją sprawnością techniczną i monitoringiem, nie była w stanie powstrzymać swojego pracownika przed wyniesieniem na zewnątrz wysoce poufnych danych.

Nghia Hoang Pho z Ellicott City w stanie Maryland pracował w jednostce TAO agencji NSA. Pho twierdził podczas procesu, że zabierał pliki do domu, aby „pracować po godzinach i zdobyć awans”. W ciągu pięciu lat był w stanie wykraść (choć, jak twierdził, nieumyślnie) ogromną liczbę ściśle tajnych, chronionych plików, co było możliwe ze względu na posiadany przez niego poziom zaufania i dostępu w sieci NSA. Uważa się, że to jego domowy komputer był prawdopodobnym źródłem opublikowanego przez grupę Shadow Brokers wycieku danych i narzędzi z NSA.

Z kolei Paige Thompson nie była zatrudniona w firmie Capital One, kiedy udało się jej przełamać zabezpieczenia systemów firmy. Thompson pracowała wówczas w małym przedsiębiorstwie dostarczającym usługi sieciowe w ramach infrastruktury chmurowej firmy Amazon dla amerykańskiej bankowej spółki holdingowej Capital One. Pani Thompson została aresztowana w lipcu 2019 roku za włamanie do systemów tej spółki, które dotknęło aż 100 milionów klientów. Dane, które ukradła z Capital One, były przechowywane na podatnym na ataki serwerze firmy Amazon, ponieważ jego zabezpieczenia zostały źle skonfigurowane przez administratorów usług bankowych w chmurze.

Thompson uzyskała dostęp do poświadczeń logowania, wykradzionych z zasobów przechowywanych na serwerach Amazon lub inaczej mówiąc, z tzw. kontenerów S3, a następnie wykorzystała je do pozyskania wrażliwych danych spółki Capital One, a później do przejęcia kontroli nad maszynami w chmurze i wykorzystania ich ogromnej mocy obliczeniowej do wydobywania kryptowaluty.

Thompson bez ogródek wypowiadała się na temat swojej motywacji i charakteru swoich planowanych operacji. Na związanym z usługami Amazon Web Services kanale Slack przyznała, że musiała „pozyskać informacje z ich serwerów”, a w jednym z postów na Twitterze napisała: *W zasadzie sama sobie założyłam wypełnioną dynamitem kamizelkę samobójcy, wykradając dokumenty spółki Capital One i przyznając się do tego. Myślę, że trzeba będzie udostępnić ich zawartość* (Merle 2019).

Thompson była utalentowanym i wysoce wykwalifikowanym inżynierem, posiadającym dogłębną wiedzę na temat wykrywania i wykorzystywania luk w zabezpieczeniach systemów, ale jej faktyczna praca u pracodawcy nigdy na tym nie polegała. Z własnych, wciąż w większości nieznanых powodów, postanowiła jednak w niecny sposób wykorzystać luki w zabezpieczeniach systemów chmurowych Amazon Web Services, które mogły mieć wpływ na wiele różnych organizacji i potencjalnie uderzyć w miliony użytkowników.

Biorąc pod uwagę wszystko, o czym pisaliśmy w tym rozdziale, nasuwa się kilka kluczowych wniosków, które powinniśmy wyciągnąć z bolesnych lekcji, jakich doświadczyły na własnej skórze różne firmy i organizacje będące ofiarami złośliwych napastników wykorzystujących erę Upadku Perymetrów Sieciowych:

- Ludzie są jednym z najsłabszych ogniw łańcucha w dziedzinie cyberbezpieczeństwa. Łatwo nas oszukać, jesteśmy podatni na wpływy i omylni z samej naszej natury.
- W miarę jak infrastruktury teleinformatyczne firm i organizacji stają się coraz większe i coraz bardziej zróżnicowane, z coraz większą ilością urządzeń, coraz lepszą dostępnością i coraz większą prędkością działania usług w chmurze, ludzie będą nadal stanowić najbardziej zawodne i podatne na ataki ogniwa w każdym systemie przetwarzania danych.
- Nawet najbardziej zaawansowane szkolenia i programy edukacyjne dla użytkowników nie przydadzą się na nic, gdy jeden z nich po prostu bezmyślnie kliknie złośliwy link otrzymany w pozornie niewinnie wyglądającej wiadomości e-mail.
- Żadne mechanizmy kontroli i zabezpieczeń nie są w stanie uchronić infrastruktury teleinformatycznej firmy czy organizacji przed zdeterminowanym do działania i posiadającym odpowiednie uprawnienia złośliwym użytkownikiem wewnętrznym.

Jak widać, bez wyspecjalizowanych mechanizmów monitorowania behawioralnego i strategicznie rozmieszczonych punktów kontroli bezpieczeństwa, to użytkownicy nadal będą ponosić główną odpowiedzialność za włamania i wycieki danych w każdym środowisku sieciowym, które będzie ignorowało posiadaną przez nich władzę i szkody, jakie mogą świadomie bądź nawet nieświadomie wyrządzić.

Podsumowanie

Model bezpieczeństwa oparty na perymetrze sieciowym jest już przestarzały i w żaden sposób nie zapewnia bezpieczeństwa firmom i organizacjom. Nie dzieje się tak jednak dlatego, że podstawowa koncepcja perymetru sieciowego jest z góry skazana na porażkę. Zamiast tego to raczej gwałtowny rozwój technologii w połączeniu z coraz bardziej złożonymi i wzajemnie powiązanymi infrastrukturami sieci teleinformatycznych sprawia, że takie podejście do bezpieczeństwa staje się nieskuteczne. Powszechna dostępność coraz bardziej mobilnych systemów i urządzeń sieciowych, która jest prawdziwym dobrodziejstwem dla ludzkości, umożliwiającym prowadzenie działalności gospodarczej i korzystanie z zasobów sieciowych z niemal każdego zakątka naszej planety, jest jednocześnie jej najgorszym wrogiem. Awaria bądź incydent bezpieczeństwa w obrębie jednego perymetru sieciowego może wywołać efekt kuli śnieżnej skutkujący awarią bądź skompromitowaniem wielu kolejnych systemów, które z kolei... i tak dalej, i tak dalej.

Choć oparty na perymetrze sieciowym model bezpieczeństwa okazał się wysoce nieefektywny i stał się przyczyną wielu incydentów i niepowodzeń, już obecnie na horyzoncie pojawiają się poważne problemy wykraczające daleko poza przyziemne sprawy perymetrów; mogą one w znaczący sposób wpłynąć na zmianę krajobrazu bezpieczeństwa cyberprzestrzeni

w nadchodzącym dziesięcioleciu. Nadszedł zatem czas, aby już teraz zrozumieć, z czym są związane te problemy i zbadać, w jaki sposób mogą zostać wykorzystane do złowrogich celów, zanim staną się zjawiskami, które całkowicie wymkną się nam spod jakiegokolwiek kontroli.

W następnym rozdziale przedstawię kilka szczegółowych przykładów incydentów bezpieczeństwa związanych z wykorzystywaniem modelu zabezpieczeń opartego na perymtrze sieciowym. Omówię ponadto wybrane zagadnienia z niedalekiej przyszłości, które będą miały ogromny wpływ na bezpieczeństwo rządów i organizacji. Przedstawię tam również niektóre z nowych i bardziej innowacyjnych typów ataków, które mogą pojawić się już w najbliższej przyszłości.

Źródła

1. R. Brandom, *Equifax CEO blames breach on a single person who failed to deploy patch*, 3 października 2017 r. Pobrań z witryny *theverge.com*: <https://www.theverge.com/2017/10/3/16410806/equifax-ceoblame-breach-patch-congress-testimony>.
2. L. Constantin, *Credential stuffing explained: How to prevent, detect and defend against it*, 30 października 2019 r. Pobrań z witryny *csoonline.com*: https://www.csoonline.com/article/3448558/credentialstuffing-explained-how-to-prevent-detect-and-defendagainst-it.html?utm_source=twitter&utm_medium=social&utm_campaign=organic.
3. Government Accountability Office (GAO), *Actions Taken by Equifax and Federal Agencies in Response to the 2017 Breach*, 1 sierpnia 2018 r. Pobrań z witryny *gao.gov*: <https://www.gao.gov/assets/700/694158.pdf>.
4. A. Greenberg, *The Untold Story of NotPetya, the Most Devastating Cyberattack in History*, 22 sierpnia 2018 r. Pobrań z witryny *wired.com*: <https://www.wired.com/story/notpetya-cyberattack-ukraine-russia-code-crashed-the-world>.
5. HPE, *Cyber Risk Report*, HPE, New York 2016.
6. M. Kan, *NordVPN, TorGuard Hit by Hacks Involving Insecure Servers*, 21 października 2019 r. Pobrań z witryny *pcmag.com*: <https://www.pcmag.com/news/371439/nordvpn-torguard-hit-by-hacks-involving-insecure-servers>.
7. B. I. Koerner, *Inside the Cyberattack that Shocked the US Government*, 23 października 2016 r. Pobrań z witryny *wired.com*: <https://www.wired.com/2016/10/inside-cyberattack-shocked-us-government>.
8. A. Ng, *How the Equifax hack happened, and what still needs to be done*, 7 września 2018 r. Pobrań z witryny CNET: <https://www.cnet.com/news/equifaxs-hack-one-year-later-a-look-back-at-how-it-happened-and-whats-changed>.
9. Ponemon Institute, *The State of Application Insecurity*, Ponemon Institute, New York 2016.

10. M. J. Schwartz, *Postmortem behind the equifax breach multiple failures*, 11 września 2018 r. Pobrańe z witryny *www.bankinfosecurity.com*:
<https://www.bankinfosecurity.com/postmortem-behind-equifaxbreach-multiple-failures-a-11480>.
11. T. Zillner, *ZIGBEE exploited, the good the bad and the ugly*, Blackhat Conference, Las Vegas 2015.

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 

Chcesz pokoju? Szukaj się do cyberwojny!

Rewolucja technologiczna i szalona ekspansja internetu zmieniły wszystko. Coś, co nazywamy rewolucją informatyczną, grozi przesunięciem granic międzypaństwowych, zniszczeniem potężnych i szanowanych instytucji, naruszeniem tkanki społecznej i norm, a także zakwestionowaniem naszej prywatności oraz podważeniem tego, co dobre i co złe. Prywatne porachunki, pospolita przestępczość, wreszcie operacje wojenne przeniosły się do cyberprzestrzeni. Cyfrowa wojna stała się rzeczywistością. Cybernapastnicy posługują się wyrafinowanymi technikami z rosnącą skutecznością.

Dzięki tej książce zrozumiesz realia cyberprzestrzeni i ewolucję cyfrowego pola walki. Pozwoli Ci to na lepsze przygotowanie się do nieuchronnej przyszłości. Omówiono tu szereg zagadnień związanych z narzędziami, taktyką i ze strategią wojny cybernetycznej. Pokazano, w jaki sposób poradzić sobie z permanentnymi zagrożeniami w cyberprzestrzeni wspomaganej technologiami uczenia maszynowego, big data, autonomicznymi pojazdami, dronami i mediami społecznościowymi. Nie zabrakło kwestii politycznych, kulturowych i geograficznych, które mają związek z różnymi metodami ataków. Lektura pomaga zrozumieć motywacje napastników oraz skutki ich ataków, a także przygotować się na obronę własnych systemów i danych.

W książce między innymi:

- prawdziwa historia zagrożeń cybernetycznych
- narzędzia i taktyki walki cybernetycznej w przyszłości
- wykorzystanie mediów społecznościowych w wojnie cybernetycznej
- stosowanie technik AJAX, system płatności, CMS, API RESTful
- minimalizowanie skutków ataku
- innowacyjne narzędzia i technologie poprawiające możliwości obronne organizacji

Chase Cunningham

jest emerytowanym oficerem marynarki wojennej Stanów Zjednoczonych z ponad 20-letnim doświadczeniem w informatyce śledczej. Zajmuje się zwalczaniem zagrożeń i wdrażaniem systemów bezpieczeństwa. Koncentruje się na integracji standardów i narzędzi zabezpieczających z operacyjnym funkcjonowaniem organizacji oraz na wykorzystywaniu zaawansowanych rozwiązań opartych na uczeniu maszynowym.

Helion 	<i>Sprawdź nasze szkolenia!</i>	KOD KORZYŚCI Sięgnij po więcej! ▶	
 helion.pl	SZKOLENIA 	ISBN 978-83-283-7248-1	
 HELION SA ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl	AKADEMIA IT & BUSINESS HELIONSZKOLENIA.PL	 9 788328 372481	
INFORMATYKA W NAJLEPSZYM WYDANIU		Cena: 59,00 zł	

Packt