

Komunikowanie danych i zastosowanie sieci komputerowych w biznesie

WYDANIE XIII

JERRY FITZGERALD

ALAN DENNIS

ALEXANDRA DURCIKOVA



Tytuł oryginału: Business Data Communications and Networking, 13th Edition

Tłumaczenie: Andrzej Grażyński

ISBN: 978-83-283-5767-9

Copyright © 2017, 2015, 2012, 2009, 2007 John Wiley & Sons, Inc.
All rights reserved.

All Rights Reserved. This translation published under license with the original publisher John Wiley & Sons, Inc.

Translation copyright © 2020 by Helion SA.

No part of this publication may be reproduced, stored in a retrieval system, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, scanning, or otherwise without either the prior written permission of 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/przd13>

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

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

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

SPIS TREŚCI

O autorach	13
Przedmowa	15

■ CZĘŚĆ PIERWSZA: WPROWADZENIE

ROZDZIAŁ 1

Wprowadzenie do komunikacji danych	19
1.1. Wstęp	20
1.2. Sieci komunikacji danych	22
1.2.1. Komponenty sieci	25
1.2.2. Typy sieci	26
1.3. Modele sieci	28
1.3.1. Model referencyjny OSI	28
1.3.2. Model internetowy	31
1.3.3. Transmisja komunikatu przez warstwy	32
1.3.4. Zalety i wady modelu warstwowego	35
1.4. Standardy sieciowe	36
1.4.1. Znaczenie standardów	36
1.4.2. Proces standaryzacyjny	36
1.4.3. Powszechne standardy	39
1.5. Trendy przyszłościowe	40
1.5.1. Bezprzewodowe sieci LAN i BYOD	40
1.5.2. Internet rzeczy	41
1.5.3. Masywność online	43
1.6. Implikacje dla cyberbezpieczeństwa	44

■ CZĘŚĆ DRUGA: FUNDAMENTALNE KONCEPCJE

ROZDZIAŁ 2

Warstwa aplikacyjna **55**

2.1. Wstęp	56
2.2. Typy architektur aplikacji	56
2.2.1. Architektura z centralnym hostem	58
2.2.2. Architektura kliencka	59
2.2.3. Odmiany architektury klient-serwer	60
2.2.4. Architektury chmurowe	64
2.2.5. Architektura peer-to-peer	67
2.2.6. Wybór właściwej architektury	68
2.3. Sieć Web	69
2.3.1. Tak działa sieć Web	69
2.3.2. Wewnątrz żądania HTTP	71
2.3.3. Wewnątrz odpowiedzi HTTP	72
2.4. Poczta elektroniczna	74
2.4.1. Tak działa poczta elektroniczna	74
2.4.2. Wewnątrz pakietu SMTP	79
2.4.3. Załączniki i MIME	79
2.5. Inne aplikacje	80
2.5.1. Telnet	80
2.5.2. Komunikatory internetowe	80
2.5.3. Wideokonferencje	82
2.6. Implikacje dla cyberbezpieczeństwa	85

ROZDZIAŁ 3

Warstwa fizyczna **97**

3.1. Wstęp	98
3.2. Obwody	100
3.2.1. Konfigurowanie obwodu	100
3.2.2. Przepływ danych	101
3.2.3. Multipleksowanie	102
3.3. Nośniki transmisyjne	106
3.3.1. Skrętka	106
3.3.2. Kabel koncentryczny	107
3.3.3. Światłowód	107
3.3.4. Radio	109
3.3.5. Mikrofale	110
3.3.6. Satelity	111
3.3.7. Wybór nośnika	112

3.4. Cyfrowa transmisja danych cyfrowych	113
3.4.1. Kodowanie	114
3.4.2. Tryby transmisji	115
3.4.3. Transmisja cyfrowa	116
3.4.4. Transmisja danych przez ethernet	118
3.5. Analogowa transmisja danych cyfrowych	119
3.5.1. Modulacja	120
3.5.2. Przepustowość obwodu	122
3.5.3. Transmisja modemowa	123
3.6. Cyfrowa transmisja danych analogowych	124
3.6.1. Digitalizacja sygnału	124
3.6.2. Rozmowy telefoniczne	126
3.6.3. Komunikatory internetowe	127
3.6.4. VoIP	127
3.7. Implikacje dla cyberbezpieczeństwa	128

ROZDZIAŁ 4

Warstwa łączy danych **141**

4.1. Wstęp	142
4.2. Sterowanie dostępem do nośnika	143
4.2.1. Rywalizacja	143
4.2.2. Nadzorowany dostęp	143
4.2.3. Porównanie metod	144
4.3. Kontrola błędów	145
4.3.1. Źródła błędów	146
4.3.2. Zapobieganie błędom	148
4.3.3. Wykrywanie błędów	149
4.3.4. Retransmisja pakietu	151
4.3.5. Progresywna korekcja błędów	151
4.3.6. Kontrola błędów w praktyce	153
4.4. Protokoły warstwy łączy danych	154
4.4.1. Transmisja asynchroniczna	154
4.4.2. Transmisja synchroniczna	155
4.5. Efektywność transmisji	159
4.6. Implikacje dla cyberbezpieczeństwa	162

ROZDZIAŁ 5

Warstwy sieciowa i transportowa **173**

5.1. Wstęp	174
5.2. Protokoły warstw transportowej i sieciowej	176
5.2.1. TCP — protokół sterowania transmisją	177
5.2.2. IP — protokół internetowy	177

5.3. Funkcje warstwy transportowej	180
5.3.1. Łączność z warstwą aplikacyjną	180
5.3.2. Segmentacja	181
5.3.3. Zarządzanie sesjami	182
5.4. Adresowanie	187
5.4.1. Przydzielanie adresów	188
5.4.2. Rozwiązywanie adresów	195
5.5. Trasowanie	198
5.5.1. Typy trasowania	200
5.5.2. Protokoły trasowania	202
5.5.3. Multicast	206
5.5.4. Anatomia routera	207
5.6. Przykładowa sieć TCP/IP	209
5.6.1. Przypadek nr 1 — znane adresy	211
5.6.2. Przypadek nr 2 — nieznane adresy	213
5.6.3. Połączenia TCP	214
5.6.4. TCP/IP a warstwy modelu sieciowego	215
5.7. Implikacje dla cyberbezpieczeństwa	217

■ CZĘŚĆ TRZECIA: TECHNOLOGIE SIECIOWE

ROZDZIAŁ 6

Projektowanie sieci **241**

6.1. Wstęp	242
6.1.1. Komponenty architektoniczne sieci	242
6.1.2. Tradycyjny styl projektowania	244
6.1.3. Projektowanie modułowe	246
6.2. Analiza wymagań	248
6.2.1. Podział sieci na komponenty architektoniczne	249
6.2.2. Aplikacje w sieci	251
6.2.3. Użytkownicy sieci	251
6.2.4. Kategoryzacja wymagań	252
6.2.5. Dokumenty	252
6.3. Projektowanie technologiczne	254
6.3.1. Projektowanie klientów i serwerów	254
6.3.2. Projektowanie obwodów	254
6.3.3. Narzędzia wspomagające projektowanie sieci	257
6.3.4. Dokumenty	258
6.4. Szacowanie kosztów	258
6.4.1. RFP — prośba o złożenie oferty	258
6.4.2. Prezentacja na forum zarządu	260
6.4.3. Dokumenty	260
6.5. Implikacje dla cyberbezpieczeństwa	260

ROZDZIAŁ 7

Sieci lokalne — przewodowe i bezprzewodowe **267**

7.1. Wstęp	268
7.2. Komponenty sieci LAN	269
7.2.1. Karty sieciowe	270
7.2.2. Obwody sieciowe	270
7.2.3. Koncentratory, przełączniki i punkty dostępowe	272
7.2.4. Sieciowe systemy operacyjne	276
7.3. Ethernet przewodowy	277
7.3.1. Topologia	278
7.3.2. Sterowanie dostępem do nośnika	281
7.3.3. Typy ethernetu	282
7.4. Ethernet bezprzewodowy	283
7.4.1. Topologia	284
7.4.2. Sterowanie dostępem do nośnika	284
7.4.3. Format ramki bezprzewodowej	286
7.4.4. Typy bezprzewodowego ethernetu	286
7.4.5. Bezpieczeństwo	288
7.5. Zalecane praktyki w projektowaniu sieci LAN	290
7.5.1. Przewidywania dotyczące użytkowników przewodowego ethernetu	291
7.5.2. Przewidywania dotyczące użytkowników Wi-Fi	292
7.5.3. Projektowanie centrów danych	295
7.5.4. Projektowanie krawędzi e-handlu	298
7.5.5. Projektowanie środowiska SOHO	299
7.6. Polepszanie wydajności sieci LAN	301
7.6.1. Zwiększanie wydajności serwerów	302
7.6.2. Zwiększanie przepustowości obwodów	304
7.6.3. Redukowanie wymagań	304
7.7. Implikacje dla cyberbezpieczeństwa	305

ROZDZIAŁ 8

Sieci szkieletowe **317**

8.1. Wstęp	318
8.2. Przełączane sieci szkieletowe	319
8.3. Trasowane sieci szkieletowe	322
8.4. Wirtualne sieci LAN	325
8.4.1. Zalety sieci VLAN	327
8.4.2. Jak działają sieci VLAN?	328
8.5. Zalecane praktyki w projektowaniu sieci szkieletowych	332
8.6. Polepszanie wydajności sieci szkieletowych	333
8.6.1. Zwiększanie wydajności urządzeń	334
8.6.2. Zwiększanie przepustowości obwodów	334
8.6.3. Redukowanie wymagań	334
8.7. Implikacje dla cyberbezpieczeństwa	335

ROZDZIAŁ 9

Sieci rozległe**345**

9.1. Wstęp	346
9.2. Sieci obwodów dedykowanych	347
9.2.1. Podstawowa architektura	347
9.2.2. Usługi T-carrier	352
9.2.3. Usługi SONET	353
9.3. Sieci komutacji pakietów	354
9.3.1. Podstawowa architektura	354
9.3.2. Usługi Frame Relay	356
9.3.3. Usługi IP	357
9.3.4. Usługi ethernetowe	357
9.4. Wirtualne sieci prywatne	358
9.4.1. Podstawowa architektura	358
9.4.2. Typy VPN	360
9.4.3. Jak działa VPN?	360
9.5. Zalecane praktyki w projektowaniu sieci WAN	363
9.6. Polepszanie wydajności sieci WAN	365
9.6.1. Zwiększanie wydajności urządzeń	366
9.6.2. Zwiększanie przepustowości obwodów	366
9.6.3. Redukowanie wymagań	367
9.7. Implikacje dla cyberbezpieczeństwa	367

ROZDZIAŁ 10

Internet**381**

10.1. Wstęp	382
10.2. Jak działa internet?	383
10.2.1. Podstawowa architektura	383
10.2.2. Połączenie z dostawcą internetu	385
10.2.3. Dzisiejszy internet	387
10.3. Technologie dostępu do internetu	388
10.3.1. Cyfrowa linia abonencka (DSL)	388
10.3.2. Modem kablowy	390
10.3.3. Światłowód do domu (FTTH)	393
10.3.4. WiMax	393
10.3.5. Coraz szybciej i coraz lepiej, czyli LTE	395
10.3.6. Jeszcze szybciej i jeszcze lepiej, czyli 5G	396
10.4. Przyszłość internetu	399
10.4.1. Zarządzanie internetem	399
10.4.2. Kreowanie przyszłości	401
10.5. Implikacje dla cyberbezpieczeństwa	402

■ CZĘŚĆ CZWARTA: ZARZĄDZANIE SIECIĄ

ROZDZIAŁ 11

Bezpieczeństwo sieci 413

11.1. Wstęp	414
11.1.1. Dlaczego w sieciach konieczne są zabezpieczenia?	416
11.1.2. Typy zagrożeń dla bezpieczeństwa	417
11.1.3. Kontrole sieciowe	419
11.2. Ocena ryzyka	420
11.2.1. Definiowanie kryteriów pomiaru ryzyka	420
11.2.2. Inwentaryzacja zasobów IT	421
11.2.3. Identyfikacja zagrożeń	425
11.2.4. Dokumentowanie istniejących kontroli	429
11.2.5. Poszukiwanie możliwości usprawnień	430
11.3. Zapewnienie ciągłości funkcjonowania	430
11.3.1. Ochrona przed złośliwym oprogramowaniem	430
11.3.2. Ochrona przed atakami DoS	432
11.3.3. Ochrona przed kradzieżą	436
11.3.4. Niwelowanie skutków awarii urządzeń	437
11.3.5. Ochrona przed skutkami katastrof	439
11.4. Zapobieganie włamaniom	443
11.4.1. Polityka bezpieczeństwa	444
11.4.2. Ochrona na granicy sieci i firewalle	444
11.4.3. Ochrona serwerów i klientów	452
11.4.4. Szyfrowanie	457
11.4.5. Uwierzytelnianie użytkowników	466
11.4.6. Obrona przed socjotechniką	471
11.4.7. Systemy zapobiegania włamaniom	474
11.4.8. Odtwarzanie po włamaniu	477
11.5. Zalecenia praktyczne	478
11.6. Implikacje dla Twojego cyberbezpieczeństwa	480

ROZDZIAŁ 12

Zarządzanie siecią 497

12.1. Wstęp	498
12.2. Projektowanie sieci pod kątem wydajności	498
12.2.1. Sieci zarządzane	499
12.2.2. Zarządzanie ruchem sieciowym	503
12.2.3. Redukowanie ruchu sieciowego	505
12.3. Zarządzanie konfiguracją	509
12.3.1. Konfigurowanie sieci i komputerów klienckich	509
12.3.2. Dokumentowanie konfiguracji	510

12.4. Zarządzanie wydajnością i awariami	512
12.4.1. Monitorowanie sieci	514
12.4.2. Kontrolowanie awarii	515
12.4.3. Statystyki wydajności i statystyki awarii	518
12.4.4. Polepszanie wydajności	522
12.5. Wsparcie dla użytkowników	522
12.5.1. Rozwiązywanie problemów	522
12.5.2. Szkolenia dla użytkowników	524
12.6. Zarządzanie kosztami	525
12.6.1. Źródła kosztów	525
12.6.2. Redukowanie kosztów	528
12.7. Implikacje dla cyberbezpieczeństwa	530

Skorowidz**543**

ROZDZIAŁ 11

BEZPIECZEŃSTWO SIECI

W tym rozdziale wyjaśniamy, dlaczego bezpieczeństwo jest nieodłącznym aspektem projektowania i użytkowania sieci komputerowych, przedstawiamy także różne sposoby jego zapewnienia. Pierwszym krokiem w planowaniu zabezpieczeń sieci jest ocena ryzyka: zidentyfikowanie kluczowych zasobów, które wymagają ochrony, i określenie zagrożeń wynikających z braku takiej ochrony w odniesieniu do każdego z nich. Istnieje wiele sposobów unikania, wykrywania i niwelowania problemów z bezpieczeństwem, wynikających z ingerowania w działanie sieci, jej paraliżowania, niszczenia jej zasobów lub nieautoryzowanego dostępu do nich.

CZYTAJĄC TEN ROZDZIAŁ:

- poznasz podstawowe źródła zagrożeń dla bezpieczeństwa sieci,
- zapoznasz się z metodami oceny ryzyka dla bezpieczeństwa,
- zrozumiesz, jak zapewnić ciągłość funkcjonowania,
- nauczysz się, jak zapobiegać włamaniom do sieci.

STRUKTURA ROZDZIAŁU

- 11.1. Wstęp
 - 11.1.1. Dlaczego w sieciach konieczne są zabezpieczenia?
 - 11.1.2. Typy zagrożeń dla bezpieczeństwa
 - 11.1.3. Kontrole sieciowe
- 11.2. Ocena ryzyka
 - 11.2.1. Definiowanie kryteriów pomiaru ryzyka
 - 11.2.2. Inwentaryzacja zasobów IT
 - 11.2.3. Identyfikacja zagrożeń

- 11.2.4. Dokumentowanie istniejących kontroli
- 11.2.5. Poszukiwanie możliwości usprawnień
- 11.3. Zapewnienie ciągłości funkcjonowania
 - 11.3.1. Ochrona przed złośliwym oprogramowaniem
 - 11.3.2. Ochrona przed atakami DoS
 - 11.3.3. Ochrona przed kradzieżą
 - 11.3.4. Niwelowanie skutków awarii urządzeń
 - 11.3.5. Ochrona przed skutkami katastrof
- 11.4. Zapobieganie włamaniom
 - 11.4.1. Polityka bezpieczeństwa
 - 11.4.2. Ochrona na granicy sieci i firewalle
 - 11.4.3. Ochrona serwerów i klientów
 - 11.4.4. Szyfrowanie
 - 11.4.5. Uwierzytelnianie użytkowników
 - 11.4.6. Obrona przed socjotechniką
 - 11.4.7. Systemy zapobiegania włamaniom
 - 11.4.8. Odtwarzanie po włamaniu
- 11.5. Zalecenia praktyczne
- 11.6. Implikacje dla Twojego cyberbezpieczeństwa
- Podsumowanie

11.1. WSTĘP

Przedsiębiorstwa i agencje rządowe zawsze stosowały rozmaite zabezpieczenia, zarówno infrastruktury fizycznej, jak i poufnych informacji, realizowane za pomocą różnych zamków, barier, strażników czy nawet wojska — tak było od zawsze, od kiedy rodzaj ludzki zaczął organizować się w społeczności. Wymyślne sposoby utajniania sekretów datują się od co najmniej 3 500 lat, ostatnie 50 nadało im jednak zupełnie nowego wymiaru, a to za sprawą pojawienia się najpierw komputerów, a potem internetu.

Pojawienie się internetu oznaczało kompletną redefinicję koncepcji bezpieczeństwa informacji. Zręczni włamywacze i kasiarze, którym nie był w stanie oprzeć się żaden sejf ani system alarmowy, ustąpili miejsca specjalistom innej profesji, posługującym się bardziej wyrafinowanymi metodami. Skoro to, co najcenniejsze, zmieniło zasadniczo formę swego istnienia — z papierowych zapisków, opasłych segregatorów i szeleszczących banknotów na bity-bajty, rezydujące na firmowych serwerach, to miejsce precyzyjnych pilników zajęły przyborniki ze złośliwym oprogramowaniem: wirusami, robakami, oprogramowaniem *ransomware* itp. Bardziej wysublimowane stały się także cenne zasoby: dziś rzadziej kradnie się cenne dokumenty i sztaby złota, bo znacznie bardziej opłacalne stały się kradzieże tożsamości i numerów kart kredytowych. I to przede wszystkim one zdolne są zniszczyć reputację firm, którym zaufały tysiące klientów i które zaufanie to zawiodły, nie będąc w stanie należycie chronić swoich zasobów. Prawo nie nadąży — niestety — za rosnącą skalą cyberprzestępczości, mimo iż w większości krajów sam

fakt nielegalnego uzyskania informacji przez przełamywanie zabezpieczeń, nawet bez uszkodzenia tejże informacji, penalizowany jest grzywną czy nawet więzieniem. niespójność tego prawa powoduje, że cyberprzestępstwa przybierają skalę transgraniczną, a ich ściganie staje się utrudnione. Trudno jest uzyskać ekstradycję cyberprzestępcy, a gdy już stanie on przed sądem, traktowany jest łagodniej niż (powiedzmy) rabuś napadający na bank.

Zagadnienie bezpieczeństwa komputerów zyskało na znaczeniu dzięki różnym aktom prawnym wysokiej rangi, między innymi ustawie HIPAA (ang. *Health Insurance Portability and Accountability Act* — „Ustawa dotycząca przenośności i odpowiedzialności za ubezpieczenie zdrowotne”) oraz tzw. Ustawie Sarbanesa-Oxleya, obu uchwalonym przez Kongres USA, w latach (odpowiednio) 1996 i 2002. Mimo to liczba incydentów związanych z włamaniami do sieci wciąż rośnie, w tempie ok. 30% rocznie. W roku 2016 odnotowano ok 50 milionów udanych kradzieży haseł; z anonimowej ankiety, przeprowadzonej w tymże roku na reprezentatywnej próbie 1500 dorosłych obywateli USA, wynika, że 51% z nich padło ofiarami cyberprzestępstwa w jakiejś formie — nie tylko zainfekowania komputera, ale także szpiegostwa przemysłowego, oszustwa, żądania okupu (ang. *ransomware*) i kradzieży tożsamości. W epoce przedinternetowej i we wczesnych latach internetu tworzenie wirusów było raczej odmianą sportu amatorskiego, uprawianego przez młodocianych hakerów z niedosytem poziomu adrenaliny, dziś stało się jednym ze sposobów pozyskiwania (często bardzo dużych) pieniędzy.

Do światowej opinii publicznej co rusz docierają sensacyjne wiadomości o kradzieży tysięcy czy nawet milionów numerów kart kredytowych, których ofiarami padli klienci dużych firm, jak Zappos czy Target, ale tak naprawdę każda firma może stać się celem cyberataku. Zgodnie z raportem firmy Symantec, ponad połowa poszkodowanych w ten sposób firm to firmy średniej wielkości, zatrudniające mniej niż 2500 pracowników. Wytlumaczenie jest proste: nie mają one tak solidnych zabezpieczeń jak duże korporacje.

Akcja wywołuje reakcję — istnieje wiele organizacji, zarówno prywatnych, jak i publicznych — których zadaniem jest pomoc przedsiębiorstwom, organizacjom i użytkownikom indywidualnym w ich (samo)obronie przed cyberzagrożeniami, czyhającymi w internecie. Wśród tych organizacji należy wymienić przede wszystkim CERT (ang. *Computer Emergency Response Team* — zespół reagowania na zagrożenia komputerowe), działającą¹ na Uniwersytecie Carnegie Mellon, APWG (ang. *Anti-Phishing Working Group* — grupa robocza ds. przeciwdziałania podszywaniu się) oraz laboratoria udostępniające narzędzia do walki z zagrożeniami i ochrony przed nimi: Kaspersky Lab., McAfee i Symantec.

Istnieją trzy powody, dla których zwiększanie bezpieczeństwa komputerów nabrało szczególnego znaczenia w ciągu kilku ostatnich lat. Pierwszym jest jakościowe przeobrażenie ataków hakerskich. Włamywanie się do cudzych komputerów, niegdyś stanowiące rodzaj hobby, dziś stało się wyspecjalizowaną gałęzią przemysłu. Profesjonalne organizacje szkolą i zatrudniają specjalistów, których zadaniem jest włamywanie się do wskazanych sieci w celu wykradania cennych informacji. I nie mówimy tu o tzw. etycznym hackingu, czyli testowaniu przez przedsiębiorstwa własnych zabezpieczeń przez wynajętych hakerów, lecz o przestępcach, którzy za opłatą wykradają numery kart kredytowych, dane osobowe lub własność intelektualną. Powodzenie takiego ataku spowodowane jest nie tylko lukami w zabezpieczeniu technicznym, lecz także czynnikiem ludzkim. Stosując **zabiegi socjotechniczne** lub phishingowe e-maile, hakerzy

¹ Zobacz stronę polskiego oddziału <https://www.cert.pl/> — przyp. tłum.

wpływają na zachowanie nieświadomych podstępów członków personelu, podchwytliwie wyłudzać sekretne informacje. Cele takich ataków — zarówno komputery, jak i ludzie — są precyzyjnie określone, więc ataki takie nazywamy **ukierunkowanymi** (ang. *targetted*).

Drugą przyczyną doniosłej roli bezpieczeństwa komputerów jest nowy fenomen socjologiczny, zwany **haktywizmem** (ang. *hacktivism*, połączenie słów *hacking* i *activism*), a polegający na wykorzystywaniu sieci komputerowych do promowania określonych celów społecznych i politycznych². Jest to więc mariaż nielegalnych technik hakerskich i określonych intencji, wymierzony zwykle w duże organizacje i agencje rządowe, zmierzający do skompromitowania ich wityn, w celu zwrócenia uwagi na określoną ideę społeczną lub polityczną. W roku 2011 grupa Anonymous przypuściła atak na witryny serwisów Visa i MasterCard, w proteście przeciwko odmowie obsługi płatności na rzecz WikiLeaks. Co prawda skutki haktywizmu nie bywają tak dotkliwe jak w przypadku wykradania czy niszczenia informacji, jednak z biegiem czasu stają się coraz częstsze i stąd zasługują na szczególną uwagę.

Po trzecie wreszcie, gwałtownie zwiększająca się liczba urządzeń mobilnych połączonych z internetem stanowi niezwykle wdzięczny obszar eksploatacji wszelkich braków w zabezpieczeniach tychże urządzeń. Przecież są one powszechnie wykorzystywane do transakcji bankowych, zakupów internetowych, operacji związanych z prowadzeniem firmy itd.; ich użytkownicy, w większości nieedukowani technicznie i skupiający się raczej na utylitarnych aspektach internetu, nie podejrzewają nawet, że oto właśnie przestali być wyłącznymi właścicielami swych sekretnych informacji. Typowy smartfon jest z reguły zabezpieczony słabiej niż przeciętny komputer — a przy tym bardziej narażony, bo podłączający się do wielu różnych sieci, często o nieznanym reputacji.

Z jednej więc strony nasze osobiste informacje stają się coraz cenniejsze, z drugiej natomiast zwiększa się potencjalne zagrożenie dla ich bezpieczeństwa, czyli prywatności naszej i naszych biznesów. Ważne jest zatem uświadomienie sobie jednego i drugiego oraz zapoznanie się ze sposobami należytej ochrony cennych wartości i unikania czyhających na nie zagrożeń. Temu właśnie poświęcona jest dalsza część tego rozdziału.

11.1.1. Dlaczego w sieciach konieczne są zabezpieczenia?

W ostatnich latach organizacje i firmy stają się coraz bardziej zależne od sieci komunikacji danych, ich rozproszonego przetwarzania, przechowywania w bazach danych i przesyłania ich między połączonymi sieciami LAN. Rozwój internetu, który stworzył niebywałe możliwości łączenia ze sobą urządzeń znajdujących się gdziekolwiek na świecie, stał się jednocześnie przyczyną zwiększonego zagrożenia pod adresem zasobów firmy. A skoro tak, zwiększył się nacisk na bezpieczeństwo przechowywania i przetwarzania tych zasobów, co wyraża się przez bieżące publikowanie informacji o rozpoznanych zagrożeniach oraz ustanawianie oficjalnych wytycznych dotyczących bezpieczeństwa w firmach i agencjach rządowych.

Straty związane z naruszeniem bezpieczeństwa mogą być ogromne. Średni koszt takiego incydentu to 3,5 miliona dolarów — niemało, ale to i tak wierzchołek góry lodowej, bo utrata zaufania klientów, do których trafi informacja o zaistniałym incydencie, może oznaczać dla firmy ostateczną katastrofę.

Firma może ponieść znaczne straty nawet w sytuacji, gdy włamanie do jej sieci nie spowoduje uszkodzenia danych. Jeżeli jednym z elementów świadczonej usługi jest jej permanentna

² Patrz <https://pl.wikipedia.org/wiki/Haktywizm> — przyp. tłum.

dostępność, przerwa w jej świadczeniu może okazać się dla firmy bardzo dotkliwa. Bank of America — jeden z największych banków w USA — ocenia, że 24-godzinna przerwa w funkcjonowaniu jego sieci kosztowałaby 50 milionów dolarów. Wiele innych dużych organizacji poczyniło podobne szacunki.

Ochrona prywatności klientów i ochrona przed kradzieżą ich tożsamości to równie istotne przesłanki na rzecz zwiększonej ochrony sieci. W roku 1998 Unia Europejska uchwaliła szereg aktów prawnych, zgodnie z którymi firmy dopuszczające do ujawnienia chronionych danych swoich klientów, będą obciążane dotkliwymi karami finansowymi. W USA organizacje ochrony zdrowia podporządkowane są ustawie HIPAA w zakresie ochrony danych osobowych, a prawo stanu Kalifornia przewiduje grzywnę do 250 000 dolarów za każdy egzemplarz danych udostępniony w sposób nieautoryzowany — jeżeli więc intruzowi uda się wykraść dane 100 klientów, grzywna może sięgnąć 25 milionów.

Jak można się domyślać, wartość danych przechowywanych w systemach informatycznych firmy oraz wartość samych systemów znacznie przekraczają koszt samej sieci. Z tej racji podstawowym założeniem bezpieczeństwa jest ochrona danych przetwarzanych w sieci komputerowej, a nie ochrona sieci jako takiej.

11.1.2. Typy zagrożeń dla bezpieczeństwa

Bezpieczeństwo komputera zwyczajowo rozumiane jest jako zapobieganie włamaniom, czyli uniemożliwianie hakerowi penetracji danych zapisanych w tym komputerze. Jest to jednak zbytne uproszczenie, bowiem na właściwie pojęte bezpieczeństwo składają się trzy czynniki: **poufność**, **integralność** i **dostępność** (ang. *confidentiality*, *integrity*, *availability* — określane łącznie akronimem CIA).

Poufność oznacza ochronę tajnych danych organizacji i danych na temat ich klientów przed nieautoryzowanym ujawnieniem. **Integralność** to zapewnienie, że wspomniane dane nie zostaną zmodyfikowane albo zniszczone. Pod pojęciem **dostępności** rozumiemy natomiast zapewnienie nieprzerwanego działania sprzętu i oprogramowania, dzięki czemu personel firmy, jej klienci i dostawcy mogą być pewni, że nie wystąpią przerwy w świadczeniu usług przez tę firmę. Na te trzy elementy czyhają rozmaite zagrożenia, którym przeciwdziałać można za pomocą metod dających się zaklasyfikować w dwóch szerokich kategoriach: **zapewnienia ciągłości funkcjonowania** i **zapobiegania nieautoryzowanemu dostępowi**.

Zapewnienie ciągłości funkcjonowania odnosi się przede wszystkim do dostępności i częściowo do integralności danych, z czego wynikają trzy podstawowe kategorie zagrożeń. **Dezorganizacja** usługi lub poważne ograniczenie jej świadczenia, w sposób trwały albo tymczasowy, na przykład awaria przełącznika lub obwodu, oznacza trwałe odcięcie segmentu sieci aż do naprawienia uszkodzenia — z uszczerbkiem dla części użytkowników, bez szkody dla pozostałych. Przerwanie ciągłości funkcjonowania może być także konsekwencją **zniszczenia** danych, na przykład za sprawą wirusa lub fizycznej awarii dysku. Trzecim zagrożeniem dla ciągłości funkcjonowania są **katastrofy** — te powodowane przez naturalne kataklizmy (huragany, tornada, trzęsienia ziemi, pożary, powodzie, lawiny błotne) oraz te będące skutkiem celowych działań człowieka (na przykład ataków terrorystycznych) lub jego działań bezmyślnych³, zdolne niszczyć całe budynki czy nawet kampusy.

³ W 2011 roku pewna 75-letnia pani odcięła od internetu 3 kraje: Gruzję, Azerbejdżan i Armenię; w poszukiwaniu złomu wykopała z ziemi jakiś kabel i przecięła go łopatą — *przyp. tłum.*

Przeciwdziałanie nieautoryzowanemu dostępowi, zwanemu także **intruzją**, ma związek zarówno z poufnością danych, jak i ich integralnością. Intruzje kojarzy się często z atakami z internetu, tymczasem — co niespodziewane — ponad połowa z nich ma związek z pracownikami firmy. Intruz może działać z czystej ciekawości, wykradając dane niestanowiące dla niego większej wartości, lecz równie dobrze taka kradzież może być inspirowana przez szpiegostwo przemysłowe lub działania konkurencji, wykradającej dane na temat opracowywanych technologii i negocjowanych kontraktów, czy też danych osobowych i numerów kart kredytowych. Co gorsza, intruz może zmodyfikować pliki tak, by wykraść firmowe pieniądze na swoje konto lub uszkodzić wrażliwe dane firmy w sposób uniemożliwiający jej funkcjonowanie.

DLA

11.1. Ciągłe to samo...

MENEDŻERA

Każda firma, niezależnie od branży, powinna się liczyć z możliwością potencjalnego ataku hakerskiego. Bolesnie przekonała się o tym firma Target w grudniu 2013 roku. Rosyjscy hakerzy zainstalowali złośliwe oprogramowanie w kasach fiskalnych w punktach sprzedaży, dzięki czemu udało się wykraść dane kart kredytowych ponad 40 milionów klientów.

Hakerzy prawdopodobnie uzyskali dostęp do sieci firmy Target, używając danych uwierzytelniających dostawcy klimatyzatorów. Przeprowadzone śledztwo wykazało, że zainstalowane złośliwe oprogramowanie nie było ani specjalnie wykoncypowane, ani też nowatorskie, co więcej — było wykrywane przez dwa systemy, które firma Target zainstalowała wcześniej w swojej sieci. Dlaczego więc specjaliści od zabezpieczeń zwyczajnie ignorowali ostrzeżenia wyświetlane przez owe systemy?

Cóż, Target — jak wiele innych firm — każdego dnia bombardowana była wręcz tysiącami komunikatów o próbach ataku; prawdopodobieństwo dostrzeżenia tego jednego, nowego, wyjątkowego, który pojawił się pewnego dnia, było więc znikome. Mimo iż ataki sieciowe bywają wyrafinowane, większość należy do tych powszednich, „dobrze znanych” (ang. *well known*), kwitowanych przez obsługę zwyczajowym „ciągłe to samo...”. Można grać różnymi kartami: w brydża, w pokera, w oko — hakerzy grają w prawo wielkich liczb: im bardziej są nieustępliwi i wyrafinowani w swych działaniach, tym większa szansa, że dopną swego, docierając w końcu do sieci i wykradając (na przykład) numery kart kredytowych.

Ta historyjka przypomina nam, że cyberbezpieczeństwo jest problemem globalnym i każdy, kto używa internetu, może być (a być może właśnie jest) celem ataku. Edukacja na temat bezpieczeństwa i należyte inwestowanie w bezpieczeństwo to warunek *sine qua non* przeżycia w zmaganiach się z pułapkami ery internetu.

Na podstawie: Michael Riley, Ben Elgin, Dune Lawrence, Carol Matlack *Missed Alarms and 40 Million Stolen Credit Card Numbers: How Target Blew It*, „Bloomberg Businessweek” (www.businessweek.com) oraz Krebs on Security (krebsonsecurity.com).

11.1.3. Kontrole sieciowe

Bezpieczna sieć to sieć kontrolowana. Elementami kontroli są: sprzęt, oprogramowanie, reguły i procedury, ograniczające lub eliminujące zagrożenia dla bezpieczeństwa. Przedmiotem kontroli jest wykrywanie i ewentualne korygowanie wszelkich zdarzeń (a także próby zapobiegania im), które w jakikolwiek sposób mogą być powiązane z bezpieczeństwem systemów komputerowych.

Kontrole prewencyjne zmierzają do osłabiania lub wręcz powstrzymywania przyczyn potencjalnych zagrożeń. System hasel chroni sieć przed nieuprawnionym dostępem użytkowników do systemu, a system zapasowych obwodów pozwala sieci nadal funkcjonować w obliczu awarii. Kontrola prewencyjna to także odstraszenie, zniechęcanie i powstrzymywanie potencjalnych sprawców przez wzbudzenie w nich obaw lub wątpliwości. Solidny system zamków drzwiowych to właśnie przykład kontroli prewencyjnej, powstrzymującej nieproszonych gości przed wejściem do mieszkania.

Kontrole detekcyjne to zespół środków mających na celu wykrywanie i ujawnianie wszelkich niepożądanych zdarzeń. Przykładem takiego środka jest oprogramowanie ewidencjonujące próby (udane i nieudane) wtargnięcia do sieci. Wszelkie zarejestrowane zdarzenia i sytuacje muszą być dokumentowane, by można było wyciągać konsekwencje w stosunku do osób i firm naruszających reguły lub łamiących prawo, a także podejmować działania korekcyjne, zapobiegające powtórzeniu się podobnych zdarzeń w przyszłości.

Zadaniem **kontroli korekcyjnych** jest naprawianie skutków niepożądanych zdarzeń, wykonywane przez ludzi lub komputery, a polegające na wykrywaniu i korygowaniu błędów. Kontrole korekcyjne to także przywracanie funkcjonowania sieci, unieruchomionej przez zaistniałe błędy lub katastrofy, na przykład oprogramowanie może automatycznie podejmować próby ponownego uruchomienia obwodu, który przestał funkcjonować w wyniku awarii.

Dalszą część rozdziału poświęcamy dyskusji na temat wszystkich trzech rodzajów kontroli. Zaprezentujemy także framework do ogólnej oceny ryzyka, dokonywanej poprzez identyfikowanie zagrożeń i związanych z nimi kontrolami. Framework ten dostarcza menedżerowi sieci ogólnego poglądu bieżących zagrożeń i dostępnych kontroli, zmierzających do ich neutralizowania.

Należy jednak pamiętać, że ustanowienie serii różnych kontroli nie jest bynajmniej wystarczającym działaniem w kierunku zapewnienia bezpieczeństwa, należy bowiem jeszcze wyznaczyć osobę, grupę lub wydział, które będą odpowiedzialne za to bezpieczeństwo — opracowywanie kontroli, monitorowanie ich aktywności oraz decydowanie o ich aktualizacjach i wymianie.

Zestaw funkcjonujących kontroli musi podlegać okresowym przeglądom, weryfikacjom i testom. Przeglądy potwierdzają lub poddają w wątpliwość bieżącą użyteczność poszczególnych kontroli, weryfikacje potwierdzają lub negują ich aktualność, a testowanie bada zgodność ich funkcjonowania z oryginalnymi specyfikacjami.

Niekiedy zdarza się konieczność tymczasowego zastąpienia takiej czy innej kontroli różnymi działaniami administratora, na przykład w sytuacji, gdy któryś komponent sieci — sprzętowy lub programowy — nie funkcjonuje prawidłowo. Tego rodzaju akcje powinny być ściśle kontrolowane, powinny też istnieć formalnie zdefiniowane procedury decydowania o ich podejmowaniu.

11.2. OCENA RYZYKA

Pierwszym krokiem na drodze do wypracowywania zabezpieczeń sieci jest przeprowadzenie **oceny ryzyka**. Istnieje kilka powszechnie używanych frameworków, umożliwiających dokonywanie takiej oceny w odniesieniu do systemów informatycznych i sieci komputerowych, poprzez analizowanie i określanie priorytetów poszczególnych czynników tego ryzyka. Zasady oceniania ryzyka powinny być sformułowane na tyle prosto, by mógł je zrozumieć odbiorca bez wykształcenia technicznego, i na tyle precyzyjnie, by były miarodajne dla personelu zajmującego się zarządzaniem i obsługą sieci. Rezultatem oceny ryzyka powinno być uszeregowanie poszczególnych systemów i komponentów sieci według stopnia narażenia na niebezpieczeństwo, tak by wyraźnie oddzielić komponenty szczególnie zagrożone atakami i próbami nadużycia od tych, dla których zagrożenie to jest znacznie mniejsze. Konsekwencją przeprowadzenia oceny ryzyka powinno być określenie zestawu niezbędnych kontroli, jakie powinny funkcjonować w sieci, i skonfrontowanie go ze stanem obecnym.

Wśród wspomnianych frameworków do oceny ryzyka systemów komputerowych, najpopularniejsze są trzy następujące:

1. OCTAVE (*Operationally Critical Threat, Asset, and Vulnerability Evaluation*) autorstwa CERT (*Computer Emergency Readiness Team*).
2. COBIT (*Control Objectives for Information and Related Technology*) autorstwa *Information Systems Audit and Control Association*.
3. Przewodnik Narodowego Instytutu Standaryzacji i Technologii USA (NIST), zatytułowany *Risk Management Guide for Information Technology Systems*, dostępny do pobrania pod adresem <https://nvlpubs.nist.gov/nistpubs/Legacy/SP/nistspecialpublication800-30.pdf>.

Każdy z tych frameworków oferuje nieco inny proces niż pozostałe, z nieco odmiennym uszeregowaniem priorytetów, w każdym jednak przypadku proces ten sprowadza się do pięciu następujących podstawowych kroków:

1. Zdefiniowanie kryteriów pomiaru ryzyka.
2. Inwentaryzacja zagrożonych zasobów.
3. Zdefiniowanie możliwych zagrożeń.
4. Udokumentowanie bieżącego stanu funkcjonujących kontroli.
5. Zidentyfikowanie koniecznych modyfikacji i możliwych ulepszeń.

11.2.1. Definiowanie kryteriów pomiaru ryzyka

Pomiar ryzyka to ewaluacja konsekwencji, jakie wynikają dla firmy z incydentu naruszenia bezpieczeństwa. Jeśli na przykład hakerowi uda się wykraść z firmowego serwera informacje o kartach kredytowych klientów, bezpośrednie tego konsekwencje będą miały przede wszystkim wymiar *finansowy*, bo wielu klientów zrezygnuje z dokonywania zakupów w tej firmie, przynajmniej w najbliższej przyszłości. Zależnie od lokalizacji firmy, incydent ten może mieć także skutki *prawne* w postaci kar finansowych, jakie ustawodawstwo wielu krajów nakłada na instytucję za niedostateczną ochronę danych osobowych. Gdy informacja o incydencie przedostanie się do publicznej wiadomości, z pewnością odbije się to negatywnie na *reputacji* firmy.

Każda organizacja powinna opracować własny wykaz różnych rodzajów wspomnianych konsekwencji, zapewne jednak znajdzie się w tym wykazie pięć najważniejszych, dotyczących **finansów** (dochodów i kosztów), **produktywności** (efektywności operacji biznesowych), **reputacji** (postrzegania przez klientów), **bezpieczeństwa** (zdrowia klientów i pracowników) i **prawa** (procesów sądowych i kar finansowych). W przypadku konkretnej firmy niektóre z wymienionych rodzajów konsekwencji mogą nie mieć zastosowania. Należy pamiętać, że cały czas mowa o ryzyku związanym z sieciami i systemami informatycznymi, i chociaż życie i zdrowie klientów i pracowników jest sprawą istotną dla większości firm, to zagrożenie w tym obszarze, wynikające z udanego ataku hakerskiego, jest raczej znikome.

Po zidentyfikowaniu rodzajów możliwych konsekwencji należy określić ich priorytety — zależnie od konkretnej firmy, każdy rodzaj konsekwencji może być dla niej znaczenie niewielkie, średnie albo wyjątkowe. Na przykład w szpitalach za najważniejszą konsekwencję włamania do sieci należy uznać możliwe zagrożenie dla życia i zdrowia pacjentów, finansowe konsekwencje, choć — oczywiście — istotne, schodzą na dalszy plan. Dla odmiany, w restauracjach jest raczej nieprawdopodobne, by wskutek kradzieży numerów kart kredytowych z sieci komputerowej zagrożone zostało bezpieczeństwo konsumentów, za to dla reputacji danej restauracji może mieć to skutek katastrofalny. Jako że różnicowanie priorytetów nie jest zadaniem banalnym, może pojawić się pokusa przypisania wszystkim zagrożeniom najwyższego priorytetu — co oczywiście nie ma sensu, bo przeczy samej zasadzie *różnicowania* priorytetów.

Niezależnie od poszczególnych rodzajów konsekwencji incydentu, różny może być stopień *dotkliwości* tych konsekwencji — mały, średni albo duży — i kolejnym krokiem jest określenie ilościowego kryterium rozgraniczenia tego stopnia. Jeżeli na przykład finansową konsekwencją incydentu będzie spadek dochodów ze sprzedaży o 5%, to mamy do czynienia ze spadkiem mało, czy może już średnio dotkliwym? Decyzje w takich sprawach mają naturę biznesową, nie technologiczną, więc powinny być podejmowane przez kierownictwo biznesowe.

Na rysunku 11.1 przedstawiamy rezultat przykładowego procesu pomiaru ryzyka dla księgarni internetowej. Uwzględnione zostały tylko cztery z pięciu rodzajów konsekwencji, niezależnie bowiem od np. niedostępności witryny sprzedażowej czy nawet błędów w naliczaniu należności, klienci w żadnym przypadku nie muszą się obawiać o swoje zdrowie. Notabene, zupełnie inaczej wyglądałyby sprawy w firmie farmaceutycznej: złośliwe oprogramowanie, modyfikujące receptury lub procesy technologiczne, może się okazać cichym zabójcą — dla pacjentów przyjmujących leki lub dla personelu! Z perspektywy zarządu naszej fikcyjnej księgarni, finanse i reputacja mają znaczenie nadrzędne, a produktywność i kwestie prawne są raczej drugorzędne. Na rysunku 11.1 widzimy także ilościową gradację dotkliwości konsekwencji; przykładowo spadek wartości sprzedaży o mniej niż 2% klasyfikowany jest jako mało dolegliwy, natomiast spadek 10-procentowy byłby już poważnym zmartwieniem.

11.2.2. Inwentaryzacja zasobów IT

Zasobem (ang. *asset*) nazywamy cokolwiek, co przedstawia dla firmy wartość zasługującą na ochronę — sprzęt, oprogramowanie, dane, aplikacje itp. Rysunek 11.2 zawiera wyszczególnienie sześciu standardowych typów zasobów.

Obszar konsekwencji	Priorytet	Mała dotkliwość	Średnia dotkliwość	Duża dotkliwość
Finanse	Wysoki	Spadek sprzedaży o mniej niż 2%	Spadek sprzedaży o 2% – 10%	Spadek sprzedaży o więcej niż 10%
Produktywność	Średni	Wzrost rocznych kosztów operacyjnych o mniej niż 3%	Wzrost rocznych kosztów operacyjnych o 3% – 6%	Wzrost rocznych kosztów operacyjnych o więcej niż 6%
Reputacja	Wysoki	Spadek liczby klientów o mniej niż 2%	Spadek liczby klientów o 2% – 15%	Spadek liczby klientów o więcej niż 15%
Prawo	Średni	Narażenie na grzywnę lub karę finansową poniżej 10 000 dolarów	Narażenie na grzywnę lub karę finansową od 10 000 do 60 000 dolarów	Narażenie na grzywnę lub karę finansową powyżej 60 000 dolarów

RYСУNEK 11.1. Wyniki przykładowego procesu określenia kryteriów pomiaru ryzyka

Sprzęt	Serwery: e-mail, WWW, DNS, DHCP Serwery plików w sieciach LAN Komputery klienckie Urządzenia sieciowe (przełączniki i routery)
Obwody	Obwody o zasięgu lokalnym (LAN i sieci szkieletowe) Obwody dzierżawione (WAN) Obwody łączące z dostawcą internetu
Oprogramowanie sieciowe	Systemy operacyjne serwerów Ustawienia systemowe Aplikacje serwerów e-mail i WWW
Oprogramowanie klienckie	Systemy operacyjne Ustawienia systemowe Aplikacje (edytory tekstu, arkusze kalkulacyjne)
Dane firmowe	Bazy danych firmowych
Aplikacje krytyczne	Witryna bankowości internetowej

RYСУNEK 11.2. Typy zasobów

DNS = *Domain Name Service* (usługa nazw domenowych)

DHCP = *Dynamic Host Control Protocol* (protokół dynamicznego konfigurowania hostów)

LAN = *Local Area Network* (sieć lokalna)

WAN = *Wide Area Network* (sieć rozległa)

Na szczególną uwagę zasługuje ostatnia z kategorii wymienionych na rysunku — **aplikacje krytyczne** (ang. *mission-critical application*), czyli systemy i aplikacje mające znaczenie krytyczne z punktu widzenia funkcjonowania i przetrwania organizacji. Takie aplikacje nie mogą ulegać awariom, a gdy jednak któraś z nich jej ulegnie, personel sieciowy porzuca wszystkie inne zadania na rzecz usunięcia awarii. Przykładem takiej aplikacji jest witryna bankowości internetowej: gdy ulegnie awarii, bank nie jest w stanie świadczyć usług swoim klientom — nie mają oni możliwości udania się do innego oddziału, mogą co najwyżej zmienić bank na bardziej przyjazny. Aplikacje tej kategorii są dobrze znane swym użytkownikom i konieczność odpowiedniego ich zabezpieczenia nie jest zagrożona przeoczeniem.

Kolejnym typem wrażliwych zasobów firmy są jej dane organizacyjne. Jeżeli na przykład zniszczeniu uległby komputer mainframe, wart 10 milionów dolarów, można go zastąpić nowym — owszem, to kosztuje, ale po kilku tygodniach fatygi wszystko wraca do stanu poprzedniego. Załóżmy teraz, że na tym komputerze rezydują dane na temat przebiegu edukacji każdego z kilku tysięcy studentów uniwersytetu. Konsekwencje ich utraty są nieporównywalne z opisanym przypadkiem utraty fizycznego komputera, bo suma samych odszkodowań przekroczyłaby znacznie 10 milionów, a odtworzenie danych na podstawie dokumentacji papierowej trwałoby znacznie dłużej niż parę tygodni, nie mówiąc już o ogromie pracy, jaki spadłby na barki personelu.

Gdy cenne zasoby zostaną zinwentaryzowane, należy zróżnicować je pod względem ważności zapotrzebowania na ochronę. W tym celu należy sobie odpowiedzieć na pytanie — w odniesieniu do każdego zasobu — w jaki sposób odbija się na firmie skompromitowanie jego *poufności*, *integralności* lub *dostępności*? Podobnie jak poprzednio, dotkliwość konsekwencji musimy ocenić w kategoriach: mała, średnia i duża, a ponadto dla każdego zasobu należy udokumentować powody, dla których znalazł się on w konkretnej kategorii. Przykładowe zestawienie tego rodzaju, typowe dla większości organizacji, znajduje się na rysunku 11.3.

Zasób	Istotność	Najważniejsze wymagania bezpieczeństwa	Opis	Odpowiedzialność
Baza klientów	Wysoka	Poufność Integralność Dostępność	Baza zawiera dane klientów, włącznie z adresami i numerami kart kredytowych	Wiceprezes ds. Marketingu Dyrektor ds. informatyki
Serwer WWW	Wysoka	Poufność Integralność Dostępność	Wykorzystywany przez klientów do składania zamówień. Musi być koniecznie dostępny 7 dni w tygodniu, 24 godziny na dobę	Dyrektor ds. informatyki
Serwer e-mail	Średnia	Poufność Integralność Dostępność	Wykorzystywany przez pracowników do komunikacji wewnętrznej. Komunikacja ta nie może zostać podsłuchana, bo przesyłane wiadomości mogą zawierać poufne informacje	Dyrektor ds. informatyki
Dane finansowe	Wysoka	Poufność Integralność Dostępność	Wykorzystywane przez dyrektorów zarządzających oraz wiceprezesa ds. operacyjnych. Nikt inny nie może mieć dostępu do tych informacji	Dyrektor ds. finansowych
Komputery pracowników	Niska	Poufność Integralność Dostępność	Każdy pracownik ma w swoim boksie komputer stacjonarny, wykorzystywany do świadczenia usług na rzecz klientów oraz do obsługi firmowej witryny WWW	Dyrektorzy oddziałów

RYСУNEK 11.3. Wynik przykładowej inwentaryzacji zasobów dla księgarni internetowej

DLA

11.1. Podstawowe zasady kontroli w bezpiecznej sieci

INŻYNIERA

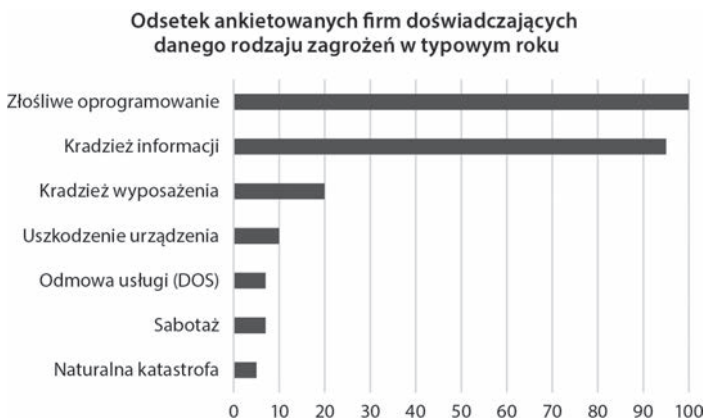
- Im mniej skomplikowane kontrole, tym lepiej.
- Koszt kontroli powinien być adekwatny do zidentyfikowanego ryzyka; ponieważ czasem trudno je oszacować, reguła ta ma po części charakter subiektywny.
- Zapobieganie incydentom jest bardziej pożądane niż niwelowanie ich skutków.
- Adekwatny system wewnętrznych kontroli oznacza wystarczające bezpieczeństwo, z uwzględnieniem równowagi między kosztami a stopniem zagrożenia.
- Kontrole automatyczne (realizowane przez komputer) są bardziej niezawodne niż te manualne, bazujące na omylnej ludzkiej interakcji.
- Kontrolowani powinni być wszyscy, nie tylko wybrane jednostki.
- Gdy dana kontrola wiąże się z nadpisywaniem danych, upewnij się, że procedura nadpisująca posiada swoją własną kontrolę, dla uniknięcia błędnego użycia.
- Ustalaj dostęp do informacji na zasadzie „musi wiedzieć” — jeżeli użytkownik nie wie, do czego potrzebny jest mu dostęp do sieci czy określonych danych, nie powinien tego dostępu uzyskać.
- Dokumentacja kontroli powinna być tajna.
- Nazwy, lokalizacje i sposoby użycia komponentów sieci nie powinny być publicznie dostępne.
- Kontrole powinny umożliwiać audyt sieci, co w praktyce oznacza dokumentowanie historii ich zmian.
- Projektując kontrole zawsze zakładaj, że będą działać w nieprzyjnym środowisku.
- Na kursach i szkoleniach przedstawiaj zawsze wizję wysokiego bezpieczeństwa.
- Dbaj o właściwe oddzielenie obowiązków: kto inny powinien kontrole projektować i instalować, kto inny natomiast używać ich na co dzień do monitorowania sieci.
- W miarę możliwości implementuj kontrole-pułapki na hakerów próbujących włamywać się do sieci.
- Gdy któraś kontrola załamie się, sieć powinna domyślnie przełączyć się do stanu, w którym kontrola ta zakazuje dostępu komukolwiek do czegokolwiek, w przeciwnym razie brak tej kontroli (na przykład kontroli haseł) oznaczałby ekstremalne narażenie sieci na niebezpieczeństwo.
- W przypadku załamania się jednego obszaru sieci, kontrole w pozostałych obszarach powinny nadal funkcjonować. Na przykład załamanie się sieci szkieletowej nie powinno mieć wpływu na funkcjonowanie wewnętrznych kontroli w przyłączonych do niej sieciach LAN.

- Nie zapominaj o sieciach LAN. Tradycyjnie planowanie zabezpieczeń i scenariuszy odzyskiwania sprawności po katastrofach koncentrowało się na komputerach mainframe i sieciach WAN, obecnie także sieci LAN, coraz częściej stanowiące elementy architektoniczne większych sieci, również stają się pełnoprawnymi kandydatami do takich zabezpieczeń — czego, niestety, wiele organizacji wciąż nie bierze pod uwagę.
- Zawsze zakładaj, że Twój przeciwnik jest sprytniejszy od Ciebie.
- Zaplanuj jakiś rodzaj ostatniej deski ratunku, na wypadek, gdyby zawiodły wszystkie kontrole.

11.2.3. Identyfikacja zagrożeń

Zagrożeniem (ang. *threat*) nazywamy potencjalne zdarzenie, zdolne zakłócić pracę sieci, przerwać działanie systemu lub spowodować straty finansowe organizacji.

Na rysunku 11.4 zilustrowano statystykę występowania najbardziej typowych zagrożeń, opracowaną na podstawie ankiet wypełnianych w ostatnich latach. Wartości pokazane na wykresie reprezentują odsetek ankietowanych firm deklarujących doświadczenie danego zagrożenia, choć niekoniecznie połączonego z doznaniem szkody. I tak na przykład wszystkie firmy doświadczyły zaimplementowania obecności jakiegось odmiany **złośliwego oprogramowania** (ang. *malware*) — wirusów lub *ransomware* — które jednak w większości przypadków nie zdążyło wyrządzić szkody, bo zneutralizowane zostało przez oprogramowanie antywirusowe. Drugim najczęściej doświadczanym typem ataku były próby **nieuprawnionego dostępu do poufnych informacji**, czyli po prostu ich **kradzieży**, czy to za pomocą **phishingu** (czyli spreparowanych e-maili, mających skłonić adresata do ujawnienia swych haseł), czy też poprzez **eksploatację luk bezpieczeństwa** bądź zabiegi **socjotechniczne** (ang. *social engineering*). Do zdecydowanie rzadziej występujących należały przypadki kradzieży wyposażenia, uszkodzania urządzeń, ataków odmowy usługi (DOS), sabotażu (czyli zniszczenia lub zmodyfikowania danych przez hakera) i skutki naturalnych katastrof (pożarów lub powodzi).



RYСУNEK 11.4. Prawdopodobieństwo występowania zagrożeń

Prawdopodobieństwo wystąpienia zagrożenia w firmie zależne jest od jej rodzaju i profilu: bank z pewnością jest bardziej prawdopodobnym celem ataku phishingowego niż mała, rodzinna firma cukiernicza. Najczęstszymi celami ataków bywają placówki medyczne, firmy świadczące usługi finansowe i agencje rządowe, bo właśnie one są właścicielami danych na tyle cennych, że perspektywa korzyści płynących z udanego ataku jest wielce obiecująca.

Po zidentyfikowaniu rodzaju możliwych zagrożeń, należy dla każdego z nich przewidzieć **scenariusz**, według którego dane zagrożenie miałyby się ziścić. Ponieważ określony zasób narażony jest na skompromitowanie przez kilka różnych zagrożeń, na przebieg możliwej jego kompromitacji składać się może kilka scenariuszy. Nawiązując do rysunku 11.3: haker, w rezultacie udanego ataku na bazę zawierającą dane klientów, może wykraść z niej informację (czyli skompromitować jej poufność), zafałszować przechowywane w niej dane (kompromitując w ten sposób ich integralność) lub doprowadzić do zniszczenia systemu operacyjnego serwera, na którym baza ta rezyduje (przekreślając w ten sposób jej dostępność). Dla konkretnego scenariusza związanego z konkretnym zasobem, należy ocenić prawdopodobieństwo jego wystąpienia (niskie, średnie albo wysokie) oraz opisać jego spodziewane konsekwencje.

Przykładowy scenariusz skompromitowania bazy zawierającej dane klientów księgarni internetowej, poprzez kradzież poufnych danych, przedstawiamy na rysunku 11.5. Górna część opisuje ryzyko narażenia bazy na taką ewentualność, natomiast w dolnej części (na przyciemnionym tle) wymienione są zaimplementowane kontrole, których zadaniem jest ochrona bazy przed tym zagrożeniem — zajmujemy się tym w następnej sekcji.

Opis scenariusza rozpoczyna się od nazwy chronionego zasobu oraz odnośnego zagrożenia. Zagrożenie klasyfikowane jest według prawdopodobieństwa wystąpienia (niskie, średnie albo wysokie), następnie identyfikowana jest zagrożona właściwość zasobu (poufność, integralność lub dostępność) — zagrożenie może obejmować kilka właściwości.

Kolejna pozycja scenariusza to ocena ilościowa wpływu zagrożenia na poszczególne obszary: reputację, finanse, produktywność, bezpieczeństwo i prawo. Dla każdego obszaru określamy jego priorytet oraz stopień wpływu na ogólną sytuację zagrożenia — oba w kategoriach niski, średni, wysoki, z oceną punktową (odpowiednio) 1, 2 i 3. Iloczyn punktacji priorytetu i stopnia wpływu daje ocenę każdego z obszarów. Sumując poszczególne oceny, otrzymujemy łączną ocenę dla rozpatrywanego zagrożenia. Mnożąc tę ostatnią wartość przez ocenę prawdopodobieństwa wystąpienia zagrożenia (wg punktacji 1, 2 lub 3), otrzymujemy ostateczną wartość ryzyka związanego z zagrożeniem.

W scenariuszu przedstawionym na rysunku 11.5 wartość ryzyka kradzieży informacji z bazy klientów oceniono na 50 punktów. Wartość ta sama z siebie nie niesie jakiegś użytecznej informacji, nabiera natomiast sensu, gdy porównać ją z wartościami ryzyka wynikającego z innych scenariuszy, związanych z innymi zagrożeniami, bo w ten sposób możemy uszeregować scenariusze w kolejności od najczarniejszych do bardziej łagodnych.

Na rysunku 11.6 widoczny jest inny scenariusz — centrum danych naszej księgarni internetowej zostaje zniszczone przez tornado o sile 4 lub 5 w pięciostopniowej skali Fujity (prędkość wiatru 333 – 512 km/h). Co ciekawe, ocena ryzyka związanego z tą katastroficzną wizją jest ponad trzykrotnie mniejsza od oceny ryzyka związanego z kradzieżą poufnych informacji.

Zasób	Baza przechowująca dane klientów		
Istotność zasobu	Wysoka		
Zagrożenie	Kradzież informacji		
Opis zagrożenia	Haker działający z zewnątrz bądź sfrustrowany pracownik (obecny lub były) może uzyskać nieautoryzowany dostęp do informacji o klientach i udostępnić tę informację zainteresowanym podmiotom trzecim		
Prawdopodobieństwo	Średnie (2)		
Zagrożona właściwość zasobu	☒ Poufność ☐ Integralność ☐ Dostępność		
Zagrożone obszary	Priorytet	Wpływ	Ocena
Finanse	Wysoki (3)	Średni (2)	6
Produktywność	Średni (2)	Wysoki (3)	6
Reputacja	Wysoki (3)	Wysoki (3)	9
Prawo	Średni (2)	Średni (2)	4
	Ocena wpływu		25
Ocena ryzyka (prawdopodobieństwo × ocena wpływu)			50
Adekwatność istniejących kontroli	Średnia		
Strategia kontrolowania ryzyka	☐ Akceptacja ☒ Przeciwdziałanie ☐ Dzielenie ☐ Odłożenie		
Kontrole przeciwdziałania ryzyku			
Firewall	Baza jest zaszyfrowana		
Polityka pracownicza	Firewall zainstalowany przed serwerem bazy blokuje próby nieautoryzowanego dostępu		
Szyfrowanie	Każdy pracownik posiada dane uwierzytelniające, ich ważność wygasa w ciągu 24 godzin od odejścia z pracy lub przejścia na inne stanowisko, niewymagające dostępu do bazy		
Szkolenia	Pracownicy zobowiązani są do udziału w corocznych szkoleniach z zakresu poufności informacji, phishingu, a także podstępnych zabiegów socjotechnicznych, mających na celu wyłudzenie haseł		
Automatyczna blokada ekranu	Po pięciominutowym okresie bezczynności użytkownika jego komputer automatycznie zostaje zablokowany, wyjście z blokady wymaga zalogowania się. Niezależnie od tego użytkownik, opuszczając swoje stanowisko pracy nawet na chwilę, zobowiązany jest do jawnego zablokowania swojego komputera. Blokada komputera chroni przed jego użyciem przez intruza, działającego w imieniu zalogowanego aktualnie użytkownika		

RYSUNEK 11.5. Scenariusz kompromitacji bazy przez kradzież poufnych informacji

Zasób	Baza przechowująca dane klientów		
Istotność zasobu	Wysoka		
Zagrożenie	Naturalny kataklizm – tornado		
Opis zagrożenia	W przypadku siły F4 lub F5 zniszczone zostanie centrum danych, w którym znajduje się serwer przechowujący dane klientów		
Prawdopodobieństwo	Niskie (1)		
Zagrożona właściwość zasobu	☐ Poufność ☐ Integralność ☒ Dostępność		
Zagrożone obszary	Priorytet	Wpływ	Ocena
Finanse	Wysoki (3)	Niski (1)	3
Produktywność	Średni (2)	Wysoki (3)	6
Reputacja	Wysoki (3)	Niski (1)	3
Prawo	Średni (2)	Niski (1)	2
	Ocena wpływu		14
Ocena ryzyka (prawdopodobieństwo x ocena wpływu)	14		
Adekwatność istniejących kontroli	Średnia		
Strategia kontrolowania ryzyka	☐ Akceptacja ☒ Przeciwdziałanie ☐ Dzielenie ☐ Odłożenie		
Kontrole przeciwdziałania ryzyku			
Kopia zapasowa	Każdej nocy sporządzana jest kopia bazy danych na serwerze w drugim centrum danych, zlokalizowanym w odległości 1000 km od głównego centrum danych		
Plan odtwarzania po katastrofie	Co dwa lata testowany jest plan odtwarzania w warunkach symulowanego kataklizmu. Celem tych symulacji jest upewnienie się, że zawartość bazy danych może zostać pomyślnie przywrócona z kopii zapasowej, by serwis sprzedaży odzyskał pełną sprawność w ciągu maksymalnie 48 godzin		

RYSUNEK 11.6. Scenariusz kompromitacji bazy w wyniku wystąpienia tornada

W zaprezentowanych scenariuszach zarówno prawdopodobieństwo, jak i priorytety oraz wpływ punktowane były w skali trójstopniowej (1, 2 lub 3 punkty). W bardziej wnikliwej analizie zagrożeń liczba stopni jest nie tylko większa, ale też różna dla różnych wielkości: niektóre organizacje stosują 5-stopniową skalę dla priorytetu, 7-punktową dla wpływu i 100-punktową (procentową) dla prawdopodobieństwa.

11.2.4. Dokumentowanie istniejących kontroli

Gdy wiadomo już wszystko na temat ryzyka zagrożeń związanych z poszczególnymi zasobami, bo ryzyko to zostało nie tylko rozpoznane, ale i ocenione punktowo, pora zastanowić się nad tym, jaką należy przyjąć **strategię jego kontrolowania**. W ogólności, organizacja może ryzyko zagrożeń **zaakceptować, przeciwdziałać, dzielić** lub **odłożyć**.

Akceptując ryzyko zagrożenia, organizacja nie stara się mu w żaden sposób przeciwdziałać i godzi się na (wiadome *a priori*) konsekwencje jego wystąpienia. Tego rodzaju strategia ma rację bytu wyłącznie w odniesieniu do zagrożeń o bardzo małym wpływie na kondycję organizacji.

Przeciwdziałanie ryzyku zagrożenia oznacza zaimplementowanie kontroli, które mają na celu zagrożenie to powstrzymać, a w najgorszym razie minimalizować skutki jego wystąpienia. Do tej kategorii należą wszelkiego rodzaju systemy antywirusowe, inteligentne firewalle, lecz także szkolenia pracowników w celu uczulenia ich na próby manipulacji w postaci podstępów socjotechnicznych.

Dzielenie ryzyka to podstawa funkcjonowania wszelkich instytucji i agencji ubezpieczeniowych. Organizacja decyduje się mianowicie wykupić ubezpieczenie od zagrożenia, które wydaje się mało prawdopodobne, ale jednak nie jest wykluczone; gdy wystąpi, odszkodowanie uzyskane od ubezpieczyciela powinno — w założeniu — pomóc w uporaniu się z jego skutkami. W ten sposób kierowcy samochodów ubezpieczają się od wypadków drogowych, w których mają nadzieję nigdy nie uczestniczyć, na tej samej zasadzie kierownictwo firmy może ubezpieczyć ją od skutków tornada. Ubezpieczenie firmy od skutków danego zagrożenia wcale nie wyklucza podejmowania innych działań łagodzących jego skutki.

Wreszcie, organizacja może **odłożyć** w czasie przeciwdziałanie zagrożeniu, na przykład w celu zebrania dodatkowych informacji i faktów z nim związanych. Postępuje się tak niekiedy z zagrożeniami, które nie mają charakteru nieuchronnych, a ich występowanie nie ma większego wpływu na losy organizacji.

Dla każdego scenariusza wystąpienia zagrożenia należy określić strategię kontroli; jeżeli w odniesieniu do konkretnego scenariusza organizacja decyduje się na zapobieżenie mu lub zminimalizowanie jego skutków, należy wyszczególnić wszystkie kontrole z tym związane. W dwóch następnych sekcjach tego rozdziału opiszemy kilka kontroli tej kategorii.

Gdy udokumentowane zostaną działające obecnie kontrole, konieczna jest ocena ich adekwatności w odniesieniu do powiązanych z nimi czynników ryzyka. Za wysoce adekwatne można uznać te kontrole, które skutecznie potrafią zapobiec odnośnemu scenariuszowi; te, których skuteczność można polepszyć w wyniku pewnych usprawnień, można uznać za średnio adekwatne, natomiast kontrole konieczne wymagające usprawnień, zasługują na miano mało adekwatnych. Oczywiście poszczególne organizacje mogą rozwinąć tę elementarną klasyfikację, na przykład do skali alfabetycznej (A, A–, A+, B — itd. — lub skali procentowej.

Dolne części rysunków 11.5 i 11.6 ukazują wspomnianą adekwatność w formie opisowej. W przypadku scenariusza kradzieży poufnej informacji, wymienione są kontrole zmierzające do przeciwdziałania temu scenariuszowi: szyfrowanie danych, firewall, polityka kontroli dostępu pracowników do sieci, szkolenie pracowników i blokowanie ekranów. Na wypadek tornada firma przygotowana jest przez regularne sporządzanie kopii zapasowych informacji narażonych na zniszczenie. W obu przypadkach adekwatność istniejących kontroli oceniono ogólnie jako średnią.

11.2.5. Poszukiwanie możliwości usprawnień

Ostatnim krokiem oceny ryzyka — i jednocześnie jego ostatecznym celem — jest identyfikacja możliwych i niezbędnych usprawnień w istniejącym systemie kontroli. Większość organizacji staje w obliczu tylu zagrożeń, że niemożliwe staje się przeciwdziałanie im wszystkim na najwyższym poziomie, więc konieczne jest poświęcenie szczególnej uwagi tylko tym, które niosą ze sobą największe ryzyko. Scenariusze o największej ocenie ryzyka są szczególnie analizowane w celu zweryfikowania, czy istnieją dla nich kontrole o przynajmniej średnim stopniu adekwatności, a także — czy należy je chronione są zasoby szczególnie cenne (te o istotności oznaczonej jako „wysoka” na rysunku 11.3). Rozważane są także możliwości implementacji dodatkowych kontroli, mogących poprawić skuteczność minimalizowania ryzyka, oraz ewentualne możliwości dzielenia poszczególnych czynników ryzyka. W podrozdziałach 11.3 i 11.4 opiszemy wiele różnych kontroli, implementowanych w celu zminimalizowania ryzyka utraty ciągłości funkcjonowania oraz ryzyka uzyskania nieautoryzowanego dostępu do sieci.

Kolejnym przedmiotem analizy są te scenariusze, w których adekwatność kontroli mających minimalizować ryzyko oceniona została jako „niska”. W idealnym przypadku kontrole te powinny dotyczyć zagrożeń o niskim ryzyku wystąpienia, co nie zawsze jest prawdą. Sprawdza się także, czy nakłady poniesione na implementację poszczególnych kontroli są adekwatne do stopnia ryzyka, przed którym kontrole te stanowią ochronę.

11.3. ZAPEWNIENIE CIĄGŁOŚCI FUNKCJONOWANIA

Pod pojęciem **ciągłości funkcjonowania** (ang. *business continuity*) rozumiemy funkcjonowanie aplikacji oraz dostępność niezbędnych do tego danych, niezależnie od prób zakłócenia pracy sieci, skompromitowania danych czy też zaistnienia katastrofy. Na plan zachowania ciągłości funkcjonowania składają się dwa elementy: pierwszym jest opracowanie systemu kontroli, którego celem jest zminimalizowanie ryzyka wystąpienia zagrożeń mających poważny wpływ na organizację; drugi dotyczy niwelowania skutków urzeczywistnionego zagrożenia. W niniejszej sekcji skupimy się na zagrożeniach dla ciągłości funkcjonowania: wirusach, kradzieżach danych, odmowach usługi, uszkodzaniu urządzeń i katastrofach, oraz na kontrolach zaprojektowanych w celu ochrony przed tymi zagrożeniami. Firmy często zaniedbują troskę o zapewnienie ciągłości funkcjonowania biznesu, bo włamania do sieci są dla nich jedynie nagłówkami w wiadomościach. Do czasu...

11.3.1. Ochrona przed złośliwym oprogramowaniem

Szczególną uwagę należy zwrócić na ochronę przed złośliwym oprogramowaniem (ang. *malware*): wirusami, robakami, *ransomware* itp. Niektóre egzemplarze *malware* są w gruncie rzeczy nieszkodliwe, a ich działalność ogranicza się do wypisywania różnych (głupawych nieraz) komunikatów; inne przysporzyć mogą niemałych problemów, na przykład szyfrując dane i żądając zapłaty za udostępnienie klucza umożliwiającego odszyfrowanie (tak właśnie działa *ransomware*). W większości przypadków dezorganizacja działania aplikacji lub zniszczenie danych mają rozmiar lokalny — dotyczą niewielkiej liczby komputerów — i względnie łatwo można sobie z nimi

poradzić: wirus zostaje usunięty przez odpowiednie oprogramowanie, a komputer wraca do pełnej sprawności. Ransomware potrafi wyrządzić szkody na znacznie większą skalę.

Większość złośliwego oprogramowania przedostaje się do komputera jako część innych programów oraz na nośnikach pamięci USB, a uaktywnia się w momencie uruchomienia rzeczono-go programu lub w momencie podłączenia urządzenia do gniazda USB; integralnym elementem aktywacji malware jest jego „rozmnażanie się”, czyli infekowanie innych programów znajdujących się na dysku (dyskach) komputera. Jedną z odmian złośliwego oprogramowania są **makrowirusy**, ukrywające się w binarnych dokumentach hipertekstowych, arkuszach kalkulacyjnych, wiadomościach e-mail itp., uaktywniające się w momencie otwarcia pliku czy wiadomości. Niektóre wirusy zmieniają swoją postać przy rozmnażaniu się, co czyni je szczególnie trudnymi do wykrywania.

Robak (ang. *worm*) to specjalny typ wirusa, który rozprzestrzenia się bez interwencji człowieka. Typowy wirus aktywuje się w momencie uruchomienia programu lub otwarcia pliku przez człowieka, natomiast robak samoczynnie kopiuje się („przełazi”) z komputera na komputer, wykorzystując wiadomości e-mail lub luki w zabezpieczeniach (które opisywać będziemy w dalszej części tego rozdziału).

DLA

11.2. Audytorzy to Twoi przyjaciele

MENEDŻERA

Bezpieczeństwo komputerowe to dziś istotna kwestia *biznesowa*, nie tylko techniczna. Ponieważ źle zabezpieczona sieć może nieść ze sobą zgubne dla firmy konsekwencje, cyberbezpieczeństwo stanowi obecnie kluczową część wielu audytów finansowych.

Większość ludzi o nastawieniu technicznym skupia swoją uwagę na technologii, traktując dokumentowanie jako sprawę drugorzędną — *a propos*: ile czasu, w porównaniu z tworzeniem kodu nowej klasy, poświęcasz na tworzenie jego dokumentacji i komentarzy?

Audytorzy natomiast koncentrują się na dokumentacji i powtarzalnych procesach biznesowych; to, co nie jest udokumentowane, nigdy nie miało miejsca.

Mimo iż te dwie tendencje wydają się nie do pogodzenia, to przy właściwym zarządzaniu mogą przynieść niemałe korzyści. Audytorzy, żądając dokumentacji od personelu technicznego, pośrednio zmuszają go do dokładnego weryfikowania, czy wszystkie niezbędne kontrole rzeczywiście zostały zaimplementowane. A gdy pracownik odchodzi z firmy i na jego miejsce przychodzi nowy, dokumentacja okazuje się pomocna we wdrożeniu go w nowe obowiązki.

Audytorzy mogą także przyczyniać się do ulepszenia systemu kontroli. Dostawcy oprogramowania zabezpieczającego dostarczają nie tylko binarne pliki aplikacji, lecz także dokumentację zaimplementowanych kontroli, co ułatwia współpracę audytorów z personelem technicznym. Co ciekawe, projektowanie nowego oprogramowania w firmie często rozpoczyna się od jego specyfikacji w formie przydatnej właśnie dla audytu, dopiero potem rozpoczyna się projektowanie techniczne i implementowanie (z uwzględnieniem funkcjonalnego podziału między komputery klienckie i serwery, co jest częścią szerszego zagadnienia — zarządzania pulpitami, omawianego w dalszej części rozdziału).

Najprostszym sposobem ochrony przed wirusami jest zainstalowanie oprogramowania antywirusowego; większość organizacji automatycznie instaluje takowe na swoich komputerach, lecz wielu użytkowników indywidualnych zaniedbuje tę sprawę w odniesieniu do swych domowych komputerów. Jako że nowe wirusy tworzone są na świecie w tempie zastraszającym (zdaniami specjalistów codziennie powstaje średnio 10 nowych), oprogramowanie antywirusowe może być skuteczne pod warunkiem jego bieżącego aktualizowania; większość pakietów antywirusowych oferuje opcję automatycznej aktualizacji.

Ponieważ złośliwe oprogramowanie często ukrywa się w plikach pobieranych z internetu (plikach muzycznych, filmach, wygaszaczach ekranu itp.), należy unikać pobierania lub kopiowania plików pochodzących z nieznanych źródeł, a przynajmniej sprawdzać (za pomocą oprogramowania antywirusowego) każdy pobrany plik przed jego otwarciem — i to nawet w przypadku plików otrzymanych od znajomych!

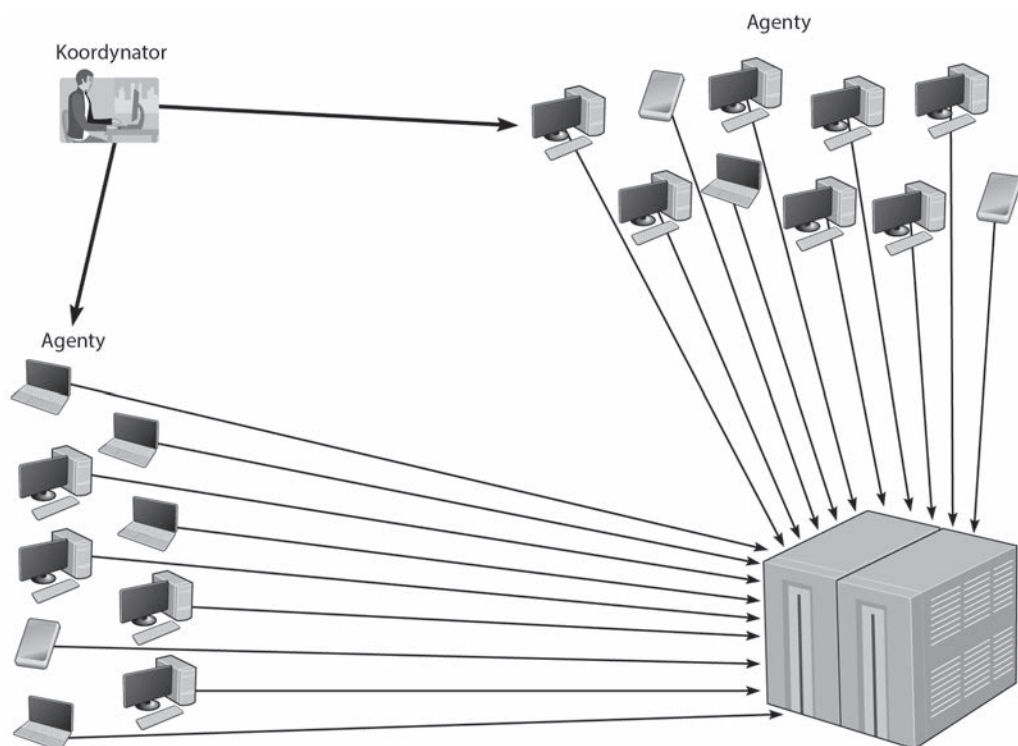
11.3.2. Ochrona przed atakami DoS

Atak **odmowy usługi** (ang. DoS — *Denial of Service*), najczęściej wykonywany w wersji rozproszonej (ang. DDoS — *Distributed Denial of Service*), to zarzucenie serwera ogromną liczbą komunikatów, których nie jest w stanie przetwarzać na bieżąco, w efekcie czego w krótkim czasie przestaje on reagować na żądania legalnych użytkowników, do świadczenia usług na rzecz których został przecież uruchomiony. Atak ten jest szczególnie dotkliwy, gdy jego celem jest serwer o istotnym znaczeniu: serwer WWW, DNS lub e-mail.

Antidotum na taki atak mogłoby się wydawać filtrowanie komunikatów, tak by komunikaty nadchodzące z określonego adresu IP były a priori odrzucane przed osiągnięciem serwera WWW. To może zadziałać; niestety — hakerzy zwykle wykorzystują narzędzia do fałszowania źródłowego adresu IP w wychodzących komunikatach, trudno je więc odróżniać od komunikatów pochodzących z legalnych żądań.

Haker, przygotowując rozproszony atak DoS, przejmuje kontrolę nad tysiącami komputerów i innych inteligentnych urządzeń (na przykład telewizorów), implantując w każdym z nich oprogramowanie zwane **agentem DDoS**, w wyniku czego takie urządzenie przekształca się w potwora („zombie”), ziejącego strumieniem złośliwych komunikatów w stronę upatrzzonego celu. (Takie programy, naśladujące działanie człowieka obsługującego komputer, nazywane są *botami*). Kontrola nad tą armią zainfekowanych urządzeń (zwaną *botnetem*) sprawowana jest przez hakera za pomocą oprogramowania zwanego **koordynatorem DDoS** (ang. *DDoS handler*), na którego komendę urządzenia zaczynają wysyłać zmasowane strumienie komunikatów — do serwera może wówczas docierać nawet miliard komunikatów w ciągu sekundy. Strategia ta przedstawiona jest na rysunku 11.7.

Istnieje kilka sposobów zapobiegania atakom DoS i DDoS. Pierwszy z nich to filtrowanie ruchu (ang. *traffic filtering*), czyli skonfigurowanie głównego routera łączącego sieć z internetem (lub firewallem) tak, aby odrzucane były komunikaty o źródłowym adresie IP niemieszczącym się w ustalonym przedziale lub zbiorze (firewallami zajmiemy się w dalszej części rozdziału). Jeżeli na przykład adres źródłowy IP komunikatu przychodzącego z internetu mieści się w zakresie wewnętrznych adresów sieci, z pewnością nie jest adresem autentycznym. Filtrowanie takie wymaga jednak dodatkowej pracy ze strony routera, co staje się przyczyną opóźnień.

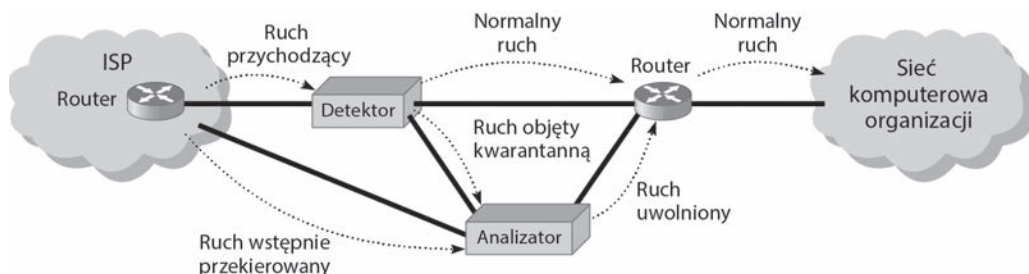


RYSUNEK 11.7. Rozproszony atak odmowy usługi DDoS

Jako że destrukcyjny charakter ataku DoS/DDoS wynika przede wszystkim z ogromnej *liczebności* komunikatów napływających w krótkim czasie do serwera, drugim możliwym podejściem zapobiegawczym jest skonfigurowanie głównego routera łączącego sieć z internetem lub firewallem w taki sposób, by ograniczona została liczba komunikatów wchodzących do sieci w jednostce czasu, i to bez względu na adresy źródłowe tych komunikatów (ang. *traffic limiting*). W zamieszczonej w dalszym ciągu rozdziału ramce „Dla inżyniera, 11.2” znajduje się opis różnych rodzajów ataku Dos/DDoS i pakietów stanowiących jego amunicję. Pakiety te niczym nie różnią się od tych wysyłanych legalnie przez aplikacje; jeżeli więc odrzucać będziemy komunikaty przekraczające określony limit, to ofiarą tego zabiegu padną również legalne żądania użytkowników. Serwery w sieci będą nadal funkcjonować bez zarzutu, ale niektóre żądania użytkowników będą gubione.

Trzecim, bardziej wyrafinowanym podejściem do unieszkodliwiania ataków DoS/DDoS, jest użycie specjalnego urządzenia, zwanego **detektorem anomalii ruchu** (ang. *traffic anomaly detector*), instalowanego przed głównym routerem lub firewallem. Urządzenie to monitoruje normalny ruch, ucząc się jego wzorców. Gdy wykryje nienormalnie wysokie natężenie ruchu adresowanego do konkretnego serwera czy urządzenia, podejrzewając atak DoS, przechwytuje pakiety składające się na ten ruch i przekazuje je do kwarantanny, jednocześnie pozwalając na swobodny przepływ pakietów pochodzących z normalnego ruchu. Pakiety, które znalazły się

w kwarantannie, są następnie przekazywane do **analizatora anomalii ruchu** (ang. *traffic anomaly analyzer*), który próbuje z nich wyłowić normalny ruch, kierując się legalnością źródłowych adresów IP, i uwolnić ów ruch z kwarantanny. Jednocześnie detektor anomalii informuje router dostawcy internetu o wykryciu podejrzanego ruchu i instruuje ów router, by ten podejrany ruch kierowany był od razu do analizatora anomalii, nie do głównego obwodu. Wszystko to odbywa się za cenę niewielkiego dodatkowego obciążenia sieci; oczywiście nie jest to proces doskonały, lecz znacząco lepszy od dwóch poprzednio opisywanych podejść. Jego działanie zilustrowaliśmy na rysunku 11.8.



RYСУNEK 11.8. Analiza ruchu jako przeciwdziałanie atakowi DoS

DLA

11.3. DDoS dla zarobku?

MENEDŻERA

Choć idea ataków DDoS ma już długą historię, ich liczba wzrosła ponad dziesięciokrotnie w porównaniu z rokiem 2010. Przyczyna tego jest prosta: hakowanie przestało być zabawą, a stało się profesją, zorganizowaną i intratną. Na hakerskich forach zorganizowane grupy poszukują hakerów do współpracy, a szukający zatrudnienia hakerzy chwalać się swoimi umiejętnościami. Ciąg dalszy to już tylko negocjowanie wynagrodzenia.

Ataki DDoS stały się także swoistą formą rozmowy kwalifikacyjnej, za pomocą której potencjalny chlebodawca przekonuje się o kwalifikacjach kandydata: jest tu kilka witryn, jedne zabezpieczone słabo, inne solidnie; jeśli uda Ci się sparaliżować je wszystkie, to witamy w zespole.

Na atakach DDoS można nieźle zarobić, gdy ma się wiedzę i doświadczenie, pozwalające na bombardowanie upatrzonego celu z szybkością 300 Gb/s.

Na podstawie: *The New Normal: 200 – 400 Gbps DDoS Attacks*, Krebs on Security (krebsonsecurity.com).

DLA

11.2. Wewnątrz ataku DoS

INŻYNIERA

Typowe ataki DoS opierają się na nadużywaniu standardowych protokołów TCP/IP oraz innych procesów połączeniowych, w wyniku czego cel ataku (serwer lub inne urządzenie) zmuszony jest do reagowania w sposób, który staje się dla niego coraz bardziej uciążliwy. Zależnie od amunicji używanej do bombardowania celu, wyróżnia się pięć najczęściej wykorzystywanych typów ataku.

Atak ICMP polega na zarzucaniu sieci lawiną pakietów *ICMP Echo Request*, w których adresem źródłowym IP jest adres atakowanego celu, a adres docelowy jest używany jako rozgłoszeniowy. Pakiety ICMP trafiają do wszystkich komputerów w sieci, a te — zgodnie z protokołem ICMP — wszystkie naraz zaczynają bombardować cel ataku lawiną pakietów *ICMP Echo Reply*. Ponieważ w domenie rozgłaszania może znajdować się kilkadziesiąt komputerów, w wyniku wysłania jednego pakietu *ICMP Echo Request* do atakowanego celu płynie kilkadziesiąt komunikatów *ICMP Echo Reply*.

Atak UDP oparty jest na tym samym pomysle co atak ICMP, lecz zamiast pakietów ICMP używane są datagramy UDP z poleceniem *Echo Request*.

Floodowanie TCP SYN to zasypywanie celu komunikatami TCP SYN, normalnie wykorzystywanymi w procesie nawiązywania połączenia TCP poprzez uzgadnianie trójstopniowe (patrz sekcja 5.3.3). Gdy cel otrzymuje komunikat SYN, natychmiast odpowiada wysłaniem komunikatów ACK i SYN — najczęściej w powietrze, czyli na spreparowany adres IP (źródłowy w komunikacie SYN), po czym bezskutecznie oczekuje na komunikat ACK. Zanim upłynie limit czasu tego oczekiwania, cel zostaje zasypywany milionami takich „wiszących” komunikatów SYN, tworząc dla każdego blok kontrolny, co w krótkim czasie doprowadza do wyczerpania dostępnej pamięci.

Atak na tablicę procesów uniksowych podobny jest do floodowania SYN, z tą różnicą, że rolę komunikatów SYN pełnią żądania nawiązania połączenia, nigdy niepotwierdzone, bo podany adres zwrotny jest fałszywy. Proces nawiązywania połączenia zawiesza się, a szybko rosnąca liczba wiszących żądań doprowadza wkrótce do wyczerpania pamięci.

Finger of Death („dotyk śmierci”) to atak oparty na wiszących żądaniach uniksowego protokołu *finger* (patrz [https://pl.wikipedia.org/wiki/Finger_\(protok%C3%B3l_C5%82\)](https://pl.wikipedia.org/wiki/Finger_(protok%C3%B3l_C5%82))).

Fałszywa rekurencja DNS polega na rozsyłaniu przez atakującego zapytań do serwerów DNS, zwykle znajdujących się w tej samej sieci co atakowany cel. Adres IP atakowanego celu jest wpisany w żądaniu jako adres źródłowy, w efekcie czego atakowany cel zostaje zarzucony odpowiedziami DNS. Komunikaty DNS są większej objętości niż pakiety ICMP, UDP i SYN, więc efekt rażenia jest silniejszy. Rekurencja (patrz 5.4.2) jest fałszywa, ponieważ adres źródłowy żądania jest adresem atakowanego celu, a nie adresem pytającego klienta.

Na podstawie: *Web Site Security and Denial of Service Protection*, www.nufusion.com.

Ponieważ ataki DDos stały się już zjawiskiem powszechnym, ze strony społeczności internetowej pojawiła się presja na dostawców internetu (ISP), by to właśnie oni dokonywali weryfikacji komunikatów przychodzących od klientów pod kątem poprawności źródłowych adresów IP. Pojedynczy mechanizm filtrujący chroni wtedy wszystkich klientów danego dostawcy przed pakietami o spreparowanych adresach źródłowych, co znacznie utrudnia przeprowadzanie ataków DDoS. A dostawcy z kolei narzucają dodatkowe wymagania na swoich klientów: ponieważ małe i średnie firmy słabo zabezpieczają swoje sieci, w efekcie stając się mimowolnymi współsprawcami ataków DDoS, dostawcy wymagają od nich instalowania dodatkowych zabezpieczeń, między innymi firewalli.

11.3.3. Ochrona przed kradzieżą

Jednym z niedocenianych zagrożeń jest kradzież sprzętu i wyposażenia sieciowego. Według różnych źródeł, straty z tego tytułu szacuje się na miliard dolarów rocznie, a ponieważ komputery i sprzęt sieciowy są rzeczami wartościowymi, często po kradzieży trafiają na aukcje internetowe.

Zabezpieczenie fizyczne sprzętu jest więc kluczowym elementem zabezpieczenia go przed kradzieżą. Jednym ze sposobów realizacji fizycznego bezpieczeństwa jest rygorystyczne kontrolowanie wstępu do pomieszczeń, w których znajduje się cenny sprzęt. Tylko ograniczona grupa osób posiada prawo takiego wstępu, a każde wejście do konkretnego pomieszczenia i wyjście z niego musi być ściśle ewidencjonowane. Ponieważ jednak w niektórych organizacjach — głównie w uniwersytetach — nie sposób zaimplementować takiej polityki, potrzebne są dodatkowe zabezpieczenia, w postaci fizycznego integrowania egzemplarzy sprzętu (głównie laptopów) ze stanowiskami pracy (biurkami) za pomocą linek antykradzieżowych lub innych umocowań. Przykład zabezpieczenia laptopa za pomocą linki antykradzieżowej przedstawia rysunek 11.9.



RYСУNEK 11.9. Przykład zabezpieczenia laptopa za pomocą linki antykradzieżowej

Źródło: dzięki uprzejmości Alexandry Durcikovej

Wśród sprzętu komputerowego najczęstszym obiektem kradzieży są — oczywiście — laptopy. Znikają z domów pracowników, ale przede wszystkim z samochodów, pokoiów hotelowych, na lotniskach, z prostej przyczyny: podróżujący pracownik nie ma możliwości efektywnego zabezpieczenia sprzętu przed kradzieżą, poza — oczywiście — wzmożoną uwagą i niestwarzaniem okazji dla potencjalnych złodziei. Firmy instruują odpowiednio w tym względzie swoich pracowników udających się w podróże służbowe, mimo to właśnie laptopy stają się najczęstszym łupem amatorów cudzej własności.

Kradzież urządzenia mobilnego to dla jego użytkownika sytuacja niewątpliwie kłopotliwa, ale jednak nie beznadziejna. Po pierwsze, dostępne jest mnóstwo oprogramowania, umożliwiającego *zdalne blokowanie* skradzionych (zagubionych) urządzeń. Dzięki takiemu oprogramowaniu użytkownik może swoje urządzenie trwale unieruchomić, poprzez zalogowanie się (za pomocą innego urządzenia) do swojego konta w odpowiedniej usłudze, niekiedy wystarcza nawet wysłanie odpowiednio spreparowanego SMS-a. Rzecz jasna, aby oprogramowanie to mogło w krytycznym momencie spełnić swoją rolę, musi zostać *uprzednio zainstalowane* na rzeczonym urządzeniu — niestety, rzadko który użytkownik zdaje sobie sprawę z tego, że przeciwdziałanie skutkom ewentualnej kradzieży jego urządzenia powinno rozpocząć się zaraz po jego nabyciu.

Po drugie, każde urządzenie mobilne posiada pewne unikalne cechy odróżniające je od innych urządzeń (numer IMEI, numer seryjny płyty głównej czy numer seryjny zainstalowanego systemu operacyjnego), więc każde jego połączenie z internetem stwarza okazję do *namierzenia jego lokalizacji*. W szczególnie korzystnej pod tym względem sytuacji są użytkownicy Windows 10, posiadający konto Microsoft *pod warunkiem jednak, że w skradzionym urządzeniu włączona jest funkcja jego lokalizowania* (zalecane jest więc jej włączenie zawczasu w ustawieniach prywatności). Należy (za pomocą innego urządzenia z Windows 10) wejść na zakładkę *Ustawienia | Aktualizacja i zabezpieczenia*, wybrać z bocznego panelu opcję *Znajdź moje urządzenie*, włączyć funkcję o tejże nazwie (domyślnie jest wyłączona), po czym zalogować się do wspomnianego konta Microsoft (na stronie <https://account.microsoft.com/devices>) i postępować zgodnie z wyświetlanymi wskazówkami, by zobaczyć mapę wskazującą miejsce, w którym nasze urządzenie zostało ostatnio zaobserwowane. Szczegółowa prezentacja opisanego postępowania jest przedmiotem krótkiego filmu instruktażowego dostępnego pod adresem <https://www.youtube.com/watch?v=eS-AckTG4vw>.

11.3.4. Niwelowanie skutków awarii urządzeń

Każde urządzenie lub osprzęt — router, przełącznik, kabel czy nawet dzierżawiony obwód — prędzej czy później może się zepsuć. Oczywiście producenci sprzętu starają się go uczynić jak najbardziej niezawodnym, niemniej jednak każdy menedżer sieci musi być przygotowany na awarię tego czy innego komponentu.

Skoro nie sposób wykluczyć awarii sprzętu, to należy przynajmniej ochronić sieć przed jej skutkami, wbudowując w tę sieć pewien stopień redundancji. Dla każdego komponentu, którego znaczenie dla ciągłości funkcjonowania sieci jest kluczowe, projektant powinien przewidzieć drugi, redundantny komponent, który przejmie funkcję głównego w przypadku jego awarii. Takim newralgicznym komponentem dla wszystkich sieci jest połączenie ich z internetem, dlatego projektanci sieci zapewniają co najmniej dwa takie połączenia, w miarę możliwości pochodzące od różnych ISP — jest mało prawdopodobne, że przerwa w świadczeniu usługi wystąpi

równocześnie u obu dostawców. Oczywiście oznacza to konieczność zainstalowania dwóch routerów dla połączeń z internetem: obsługa obu za pomocą pojedynczego routera oznaczałaby źle pojętą oszczędność, bo w przypadku jego awarii sieć utraciłaby połączenie z internetem mimo niezawodnego funkcjonowania usług obu dostawców.

Podobnie ma się rzecz z wewnętrznymi sieciami organizacji. Jeśli centralna sieć szkieletowa kampusu ma kluczowe znaczenie (a z reguły ma), organizacja musi zainstalować *dwie* takie sieci, każdą bazującą na oddzielnym zestawie urządzeń. Także każda dystrybucyjna sieć szkieletowa, która połączona jest z siecią centralną, powinna być z nią połączona za pomocą dwóch łączy i dwóch routerów.

Schodzimy w dół hierarchii: również każda dostępową sieć LAN powinna być połączona z nadrzędną, dystrybucyjną siecią szkieletową, za pośrednictwem dwóch łączy. Redundancja jest kosztowna, dlatego wiele organizacji stosuje ją wybiórczo: większość ogranicza ją do centralnej sieci kampusowej oraz dostępu do internetu, natomiast w zakresie sieci dystrybucyjnych i dostępowych sieci LAN jest ona stosowana selektywnie, tylko w odniesieniu do tych, które mają naprawdę krytyczne znaczenie dla sieci jako całości, bo na przykład zapewniają połączenia z serwerami spełniającymi krytyczne funkcje.

Redundancja ma zastosowanie również do serwerów, czego wyrazem są farmy serwerowe: gdy ulega awarii jeden z serwerów, pozostałe przejmują jego funkcje i fakt ten jest dla sieci wręcz niezauważalny. Niektóre organizacje wykorzystują serwery **odporne na awarie** (ang. *fault tolerant*) dzięki zdublowaniu w każdym z nich krytycznych komponentów.

Redundantna macierz dyskowa (ang. RAID — *Redundant Array of Independent Disks*) to technologia pamięci masowej zbudowanej (jak sama nazwa wskazuje) na bazie kilku niezależnych napędów dyskowych. Gdy plik zapisywany jest do pamięci RAID, fizycznie zapis ten zostaje rozdzielony między kilka oddzielnych napędów.

Istnieje kilka typów macierzy RAID⁴. RAID 0 wykorzystuje wiele napędów dyskowych, jest więc szybsza od pojedynczego dysku, ponieważ dane są odczytywane i zapisywane równolegle z wielu (na wiele) dysków równolegle. RAID 1 zapisuje duplikaty (przynajmniej jedną kopię) wszystkich danych, na co najmniej dwóch różnych dyskach, dzięki czemu awaria jednego nie powoduje utraty zapisanych na nim danych, bo ich kopie znajdują się na innym dysku (dyskach). Technika taka nazywa się **odbiciem lustrzanym** (ang. *mirroring*), bo taka jest relacja kopii danych do ich oryginału. RAID 2 dodatkowo wyposażona jest w kontrolę błędów w celu zapewnienia, że zapisywanie i odczytywanie danych odbywa się bez błędów; kontrola ta została ulepszona i usprawniona w macierzy RAID 3. RAID 4 oferuje nieco szybszy odczyt w stosunku do RAID 3, bo zapisane dane są bardziej optymalnie rozpartycjonowane na poszczególnych napędach dyskowych. RAID 5 zapewnia szybszy zarówno odczyt, jak i zapis w porównaniu z RAID 4, ze względu na bardziej optymalną strukturę redundantnych danych związanych z kontrolą błędów. Wreszcie, RAID 6 zapewnia zachowanie wszystkich danych nawet w przypadku awarii dwóch napędów dyskowych.

⁴ Czytelnikom zainteresowanym bardziej szczegółowymi informacjami proponuję stronę <https://pl.wikipedia.org/wiki/RAID> — *przyp. tłum.*

Przerwy w dostawie prądu to jedna z najczęstszych awarii sieci komputerowych. **Zasilacze awaryjne** (ang. UPS — *Uninterruptible Power Suppliers*) to urządzenia wykrywające zanik zasilania i zapewniające zasilanie zastępcze dla podłączonych do nich urządzeń, przez pewien czas, zależny od ilości energii zgromadzonej w bateriach. Tanie UPS-y, do użytku domowego, zapewniają zasilanie zastępcze w granicach 15 minut — wystarczająco dużo, by zakończyć bieżące prace i zamknąć lub zahibernować komputer. W większych organizacjach instalowane są mocniejsze zasilacze, zdolne podtrzymać pracę krytycznych serwerów, przełączników i routerów przez około godzinę, do czasu uruchomienia awaryjnego generatora prądotwórczego.

11.3.5. Ochrona przed skutkami katastrof

Katastrofa to zdarzenie, w wyniku którego ulega zniszczeniu znaczna część sieci i infrastruktura obliczeniowa w jednej części organizacji. Katastrofy mogą być powodowane przez siły natury (huragany, powodzie, trzęsienia ziemi, pożary) albo przez działalność człowieka (podpalenia, zamachy bombowe, ataki terrorystyczne).

11.3.5.1. Unikanie katastrofy

Uniknąć katastrofy nie sposób — bo jak powstrzymać trzęsienie ziemi? — lecz realnym założeniem jest minimalizowanie szkód, jakich organizacja doświadczyć może w wyniku katastrofy. Najbardziej fundamentalnym sposobem realizacji tego założenia także w tym przypadku jest redundancja, tym razem w obszarze cennych danych, których kopie zapasowe — jedna lub więcej — powinny być przechowywane w odległych regionach świata. Jeżeli jeden ich egzemplarz zostanie zniszczony w wyniku lokalnego kataklizmu, pozostałe będą bezpieczne.

Podstawowym czynnikiem ułatwiającym minimalizowanie fizycznych szkód katastrof jest przewidywanie i dalekowzroczność przy projektowaniu sieci. Jest zrozumiałe, by w rejonach zagrożonych powodzią nie lokalizować wrażliwych komponentów sieci w pobliżu rzek i jezior, ani też na niskich kondygnacjach (w suterenie czy na parterze). Elementem profilaktyki przeciwpożarowej są systemy zraszania wodą, urządzenia oddymiające i ochronne powłoki termiczne. Prawdopodobieństwo ataku terrorystycznego można zmniejszyć, utrzymując w tajemnicy lokalizację kluczowych komponentów sieci i zapewniając im profesjonalną ochronę.

11.3.5.2. Odtwarzanie po katastrofie

Krytycznym elementem związanym z **usuwaniem problemów** spowodowanych przez katastrofę jest **plan odtwarzania po katastrofie** (ang. DRP — *Disaster Recovery Plan*), obejmujący różne poziomy reagowania na możliwe rodzaje katastrof oraz metody kompletnego lub częściowego odtwarzania danych, programowych komponentów sieci i fizycznych pomieszczeń. Przykład takiego szczegółowego planu wykraczałby poza ramy niniejszej książki, zainteresowanym Czytelnikom możemy polecić *plan ciągłości działania* (*business continuity plan*) Massachusetts Institute of Technology (tak, *business* brzmi przyjemniej niż *disaster*, nie tylko dla Amerykanów), dostępny pod adresem web.mit.edu/security/www/pubplan.htm. Na rysunku 11.10 przedstawiamy natomiast podsumowanie najważniejszych elementów takiego planu.

Dobry plan odtwarzania po katastrofie powinien zawierać:

- Nazwisko decyzyjnego menedżera głównego, władnego wydawać decyzje w związku z operacjami wychodzenia z awarii
- Nazwisko zastępczego drugiego menedżera, na wypadek nieobecności menedżera głównego
- Przydział obowiązków dla poszczególnych pracowników w czasie katastrofy
- Ustaloną a priori listę priorytetów poszczególnych działań naprawczych
- Wskazanie lokalizacji alternatywnych pomieszczeń przeznaczonych dla organizacji lub zewnętrznej firmy ratowniczej i określenie procedur przeniesienia operacji do tychże pomieszczeń, z użyciem kopii zapasowych danych i oprogramowania
- Definicje procedur przywracania obiektów komunikacji danych (sieci szkieletowych, sieci miejskich, sieci rozległych i sieci lokalnych), serwerów i systemów aplikacji, z uwzględnieniem lokalizacji obwodów i urzędzeń, osób dysponujących poszczególnymi kategoriami informacji oraz osób odpowiedzialnych za kontakty z dostawcami
- Listę akcji do podjęcia w przypadku częściowej szkody lub zagrożenia atakami bombowymi, pożarem, zalaniem wodą, porażeniem prądem elektrycznym, sabotażem, niepokojami społecznymi i niewydolnością dostawców
- Listę czynności do ręcznego wykonywania, do czasu przywrócenia operatywności sieci
- Listę procedur zapewniających adekwatne testowanie i aktualizowanie planu odtwarzania po katastrofie

Plan oraz związane z nim dane i oprogramowanie powinny być przechowywane w bezpiecznym miejscu, w którym nie grozi im zniszczenie w wyniku katastrofy. Miejsce to musi być jednak dostępne dla każdego, kto będzie potrzebował skorzystać z planu.

RYSUNEK 11.10. Podsumowanie przykładowego planu odtwarzania po katastrofie

Najważniejsze elementy planu odtwarzania po katastrofie to kontrole **sporządzania kopii zapasowych i odtwarzania** z tych kopii, umożliwiające organizacjom odzyskiwanie danych i restart aplikacji w przypadku awarii sieci lub jej segmentu. Najprostszą metodą sporządzania kopii zapasowych jest nagrywanie ich na taśmy i przesyłanie tychże (przez kurierów) do odległych lokalizacji. W większości organizacji kopie krytycznych danych i aplikacji sporządzane są codziennie, natomiast kopie mniej istotnych (między innymi e-maili) — raz na tydzień. Coraz powszechniej stosowaną alternatywą dla ekspedycji taśm jest przesyłanie kopii danych za pośrednictwem sieci WAN — wychodzi znacznie taniej, nawet gdy zwiększy się częstotliwość sporządzania kopii. W celu uniemożliwienia przejęcia danych przez osoby niepowołane, należy te dane koniecznie zaszyfrować (szyfrowaniem zajmujemy się w dalszej części tego rozdziału).

Ciągła ochrona danych (ang. CDP — *Continuous Data Protection*) to kolejna opcja bezpieczeństwa, którą firmy stosują w uzupełnieniu do regularnych kopii zapasowych. W ramach CDP kopie wszystkich danych i transakcji z wybranych serwerów zapisywane są na dedykowanych serwerach CDP po każdej zatwierdzonej transakcji. CDP jest bardziej elastyczne zarówno od tradycyjnych backupów, stanowiących migawki danych w określonych chwilach, jak i od mirroringu, czyli tworzenia duplikatu zawartości dysku sekunda po sekundzie. CDP umożliwia zarówno przechowywanie kopii danych w dużej odległości od ich oryginałów, jak i odtwarzanie stanu danych ze wskazanego momentu czasowego. Jeżeli na przykład serwer zostanie unieruchomiony przez wirus o godzinie 14:45, administrator sieci po usunięciu wirusa może odtworzyć zawartość serwera z godziny 14:30, po czym ponowić transakcje zatwierdzone po tej godzinie — tak, jakby nic złego się nie wydarzyło.

Kopie zapasowe i CDP zapewniają bezpieczeństwo danych, lecz nie dają gwarancji, że danych tych będzie można użyć. Plan odtwarzania po katastrofie powinien zawierać opis udokumentowanego i przetestowanego podejścia do różnych skutków katastrof. Jeżeli na przykład zniszczony zostanie główny serwer bazy danych, ile czasu potrzebować będzie organizacja do przywrócenia danych i oprogramowania na nowym serwerze? A jeżeli zniszczone zostanie główne centrum danych? Odpowiedzi na te pytania drastycznie różnią się pod względem kosztów: odtworzenie danych i oprogramowania na zapasowym serwerze lub na serwerze dysponującym niewykorzystaną przepustowością to jedno, ale odtworzenie funkcjonalności zniszczonego centrum danych w centrum zapasowym w czasie (powiedzmy) 12 godzin to już całkiem inna propozycja.

Wiele firm dysponuje planem odtwarzania po katastrofie, lecz tylko nieliczne **testują jego adekwatność** w ramach symulowanej katastrofy (ang. *disaster recovery drill*). Takie próbne ćwiczenia, podobne do próbnych alarmów pożarowych, stanowią dla pracowników okazję do sprawdzenia własnych, rzadko wykorzystywanych umiejętności pod kątem tego, czy sprawdzą się one w warunkach rzeczywistego zagrożenia, czy też nie. Bez takich ćwiczeń pierwsza okazja do przetestowania planu pojawia się dopiero wtedy, gdy trzeba go wdrożyć. Gdy pewnego dnia na jednej z wysp Bermudów nastąpił całkowity zanik zasilania, generator prądotwórczy włączył się samoczynnie i firma mogła działać normalnie. Jednak system sterujący otwieraniem drzwi za pomocą kart magnetycznych, który nie był podłączony do generatora, nie mógł otworzyć żadnych drzwi i pracownicy zostali uwięzieni na kilkanaście godzin. Nikt wcześniej nie przetestował planu opracowanego na wypadek braku zasilania.

Organizacje zwykle bardziej troszczą się o backup ważnych danych niż o komputery swoich pracowników. Kiedy ostatnio robiłeś kopię zapasową danych ze swojego komputera? Co zrobiłbyś, gdyby Twój komputer został skradziony lub zniszczony? Dla użytkowników indywidualnych istnieje tania alternatywa dla CDP: usługa **kopii zapasowej online**. Serwisy takie jak *mozy.com* umożliwiają przechowywanie na swoich serwerach kopii zapasowej danych użytkowników, Użytkownik pobiera oprogramowanie klienckie i instaluje je wskazując, które foldery na dysku jego komputera uwzględniane będą w kopii zapasowej. Po sporządzeniu pierwszej kopii zapasowej (co może trochę potrwać) oprogramowanie to co kilka godzin automatycznie wysyła do serwera informacje o zmianach w zawartości folderów. Gdy użytkownik chce odtworzyć część lub całość zarchiwizowanych danych, wchodzi na stronę serwisu i pobiera je.

11.3.5.3. Outsourcing odtwarzania po katastrofie

Większość dużych organizacji planuje odtwarzanie funkcjonowania po ewentualnej katastrofie w sposób dwupoziomowy. Po pierwsze, zapewnia sobie pewną rezerwę sprzętową na wypadek pomniejszej katastrofy, jak awaria głównego serwera lub segmentu sieci (choć można polemizować, czy przymiotnik „pomniejsza” jest w tym przypadku właściwy). Jednak budowa sieci dysponującej rezerwą wystarczającą na odtwarzanie funkcjonowania po poważniejszym incydencie, na przykład po kompletnej utracie centrum danych, przekracza możliwości większości firm. W związku z tym wiele dużych przedsiębiorstw powierza („outsourcuje”) swoje losy firmom wyspecjalizowanym w przywracaniu normalności po większych awariach i katastrofach.

DLA

11.4. Gdy katastrofa staje u drzwi...

MENEDŻERA

„Pali się!” usłyszałem, gdy odebrałem telefon. Zbliżało się południe, gdy zatelefonowała do mnie studentka ze swego biura na najwyższym piętrze szkoły biznesu na Uniwersytecie Georgia. Zaproszenie ognia przez dekarza naprawiającego dach spowodowało pożar, który okazał się największym w okolicy w ciągu ostatnich 20 lat, choć jeszcze wtedy nie zdawaliśmy sobie z tego sprawy. Miałem tyle czasu, by pozbiierać najważniejsze rzeczy z mojego biura na parterze (pamiątki, dyplomy, fotografie z 10 lat pracy na uczelni), gdy włączył się alarm. Nie troszczyłem się o mój komputer, wszystkie pliki były zabezpieczone w zewnętrznej kopii zapasowej.

Dziesięć godzin walki z ogniem, 100 strażaków, 5,5 tysiąca m³ wody; gdy ogień ugaszono, zaczęła się nasza praca. Ogień doszczętnie zniszczył najwyższe piętro budynku, wraz z moją pracownią, w której znajdowało się 20 komputerów. Woda poważnie uszkodziła resztę budynku, wraz z moim biurem, w którym — jak się później dowiedziałem — warstwa wody sięgała ponad pół metra wysokości w kulminacyjnym momencie. Mój komputer, jak zresztą wszystkie inne komputery w budynku, na skutek zalania stał się bezużyteczny.

Moje osobiste pliki nie ucierpiały, odtworzyłem je na nowym komputerze i mogłem dalej pracować, po uprzednim sporządzeniu nowych kopii, o przechowanie których poprosiłem kolegę. Kopia zapasowa serwera WWW, którym administrowałem, wykonana została 2 dni temu (zgodnie z cotygodniowym cyklem) i przechowywana była po przeciwnej stronie kampusu. Odtworzyliśmy wszystkie pliki na serwerze WWW biblioteki uniwersyteckiej, ponawiając wszystkie prace wykonane w ciągu wspomnianych 2 dni, oczywiście modyfikując ustawienia uniwersyteckiego serwera DNS, by żądania kierowane uprzednio do starego serwera, trafiały teraz do nowej lokalizacji. W ciągu niespełna 24 godzin nasza witryna znów była w pełni operatywna.

Nie wszystko jednak układało się tak pomyślnie. Zawartość głównego serwera WWW szkoły biznesu została zarchiwizowana na taśmach w nocy poprzedzającej pożar i chociaż taśmy przechowywane były w bezpiecznym miejscu, to zniszczony został jedyny w kampusie napęd, na którym można było te taśmy odczytać! Sprowadzenie i uruchomienie nowego, odczytanie zawartości taśm i ponowne uruchomienie serwera zajęło nam 5 dni. Przez 30 dni pracowaliśmy w tymczasowych biurach, podłączeni do nowej sieci. W tym czasie udało nam się odtworzyć 90% biurowych komputerów, oczywiście wraz z zawartością.

Takie doświadczenia zmieniają ludzi. Jestem teraz bardziej troskliwy w kwestii kopii zapasowych i odruchowo przyspieszam kroku na dźwięk alarmu pożarowego.

Źródło: Alan Dennis

Firmy takie oferują szereg usług, z których najprostszą jest udostępnianie bezpiecznych magazynów danych na potrzeby kopii zapasowych. Pełny zakres usług obejmuje udostępnianie profesjonalnych centrów danych na użytek klientów, którzy doświadczyli katastrofy. Gdy tylko klient deklaruje katastrofę, firma natychmiast przystępuje do odtwarzania na podstawie posiadanych kopii zapasowych i jest w stanie przywrócić całkowitą operatywność sieci klienta w ciągu kilku godzin, wykorzystując w tym celu własny sprzęt. Usługi takie nie są tanie, lecz wyjątkowo opłacalne wobec perspektywy potencjalnych strat liczonych w milionach dolarów dziennie, wynikających z braku dostępu do krytycznych danych i aplikacji. Co należyście potrafi docenić każdy, komu udało się owe straty zminimalizować dzięki wspomnianym usługom.

11.4. ZAPOBIEGANIE WŁAMANIAM

Włamania (ang. *intrusions*) do sieci to drugi ważny typ problemów z bezpieczeństwem, zasługujący na szczególną uwagę. Nikt przecież nie pragnie obecności intruzów w swojej sieci.

Intruzów usiłujących uzyskać nieautoryzowany dostęp do sieci komputerowych można podzielić na cztery kategorie. Do pierwszej należą intruzy przypadkowi, posiadający jedynie ograniczoną wiedzę na temat bezpieczeństwa komputerowego. Krążą oni po internecie, próbując uzyskać dostęp do różnych komputerów. Ich usiłowania można porównać z próbami dostania się do czyjegoś domu przez szarpanie za klamkę — jeśli właściciel nie zamknie drzwi na klucz, sam sobie jest winien. Niestety, w internecie dostępnych jest coraz więcej narzędzi umożliwiających nawet hakerskim żółtodziobom dokonywanie zaawansowanych niekiedy prób włamań. Polegają oni prawie wyłącznie na owych narzędziach, bez dogłębnej znajomości ich działania, a nie na własnej wiedzy czy doświadczeniu. Z racji tego określani są często mianem *skryptowych dzieciaków* (ang. *script kiddies*).

W kategorii drugiej plasują się intruzy będący ekspertami w dziedzinie bezpieczeństwa, których usiłowania motywowane są przede wszystkim instynktem myśliwego żądnego dreszczu emocji, chęcią zabłyśnięcia przed znajomymi czy też żądzą sprawienia kłopotów temu czy innemu właścicielowi sieci. Kultura informatyczna nazywa ich **hakerami** (ang. *hackers*); bardzo często cechują się oni osobista filozofią, wymierzoną przeciwko własności danych i oprogramowania. Wyrządzone przez nich szkody bywają raczej niewielkie, bo nie są motywowane chęcią zysku; hakerzy stanowią wyjątek od tej zasady, zwani **krakerami** (ang. *crackers*), mogą być jednak bardzo groźni.

Trzecia kategoria to intruzy najbardziej niebezpieczni. To profesjonalni hakerzy, włamujący się do sieci korporacji i agencji rządowych w wyraźnie określonym celu, którym może być szpiegostwo przemysłowe, oszustwa finansowe lub po prostu destrukcja. Według Ministerstwa Obrony USA (DoD — Department of Defense), które rutynowo monitoruje próby ataków na obiekty wojskowe USA, ataki takie przypuszczane są głównie przez indywidualistów lub małe grupy intruzów dwóch poprzednich kategorii. Mimo iż niektóre z nich doprowadziły do przerobienia stron związanych z wojskowością i wywiadem, nie są oceniane jako działania wysokiego ryzyka. W końcu lat 90. minionego stulecia ministerstwo odnotowało jednak rosnącą tendencję ataków testowych, eksplorujących sieci w celu badania efektywności różnych ataków programowych jako broni w wojnie cybernetycznej. Efektem tego było sformułowanie **programu wojny informacyjnej** (ang. *information warfare program*) i powołanie organizacji koordynującej obronę wojskowych sieci komputerowych pod egidą US Space Command.

Nie mniej groźni są intruzy czwartej kategorii: pracownicy firmy lub organizacji, posiadający legalny dostęp do sieci, lecz próbujący uzyskać nieautoryzowany dostęp do danych, których im oficjalnie nie udostępniono. Wykradzione w ten sposób informacje intruz może sprzedać zainteresowanemu podmiotowi trzeciemu, może też zmodyfikować dane w celu uzyskania dodatkowych dochodów.

Podstawową zasadą przeciwdziałania włamaniom jest **ochrona proaktywna**, polegająca na rutynowych testach własnych zabezpieczeń i wykrywaniu ewentualnych słabych miejsc, zanim odkryje je potencjalny intruz. Istnieje wiele różnych metod zapobiegania włamaniom i nieautoryzowanym dostępom, mimo to żadna sieć nie może być uważana za w pełni bezpieczną. Najskuteczniejszym sposobem zapewnienia bezpieczeństwa jest przestrzeganie zasady, którą kierują się projektanci systemów wojskowych: szczególnie wrażliwe dane nie powinny być dostępne online, lecz przechowywane na komputerach odizolowanych od innych sieci.

W kolejnych sekcjach opiszemy najważniejsze kontrole zapobiegające włamaniom i umożliwiające odtwarzanie po włamaniu, gdy się ono wydarzy.

11.4.1. Polityka bezpieczeństwa

Podobnie jak plan odtwarzania po katastrofie ma krytyczne znaczenie dla kontrolowania ryzyka związanego z zaburzeniem pracy sieci, zniszczeniem danych czy katastrofą, tak polityka bezpieczeństwa ma znaczenie krytyczne dla kontrolowania ryzyka związanego z włamaniami. Polityka bezpieczeństwa powinna jasno definiować zbiór zasobów podlegających szczególnej ochronie oraz wykaz najważniejszych kontroli, służących do zapewnienia tej ochrony. Powinna także regulować uprawnienia personelu, definiując co wolno, a czego nie wolno poszczególnym pracownikom, a także określając zasady ich szkolenia pracowników — szczególnie tych z mniejszym doświadczeniem w zakresie pracy z komputerem — na temat kluczowych zasad bezpieczeństwa. Kolejnym elementem polityki bezpieczeństwa powinien być plan rutynowego testowania i ulepszania systemu istniejących kontroli. Dobry zbiór reprezentatywnych przykładów i szablonów polityki bezpieczeństwa dostępny jest pod adresem www.sans.org/resources/policies, przykładowe jej wytyczne przedstawiamy natomiast na rysunku 11.11.

11.4.2. Ochrona na granicy sieci i firewallle

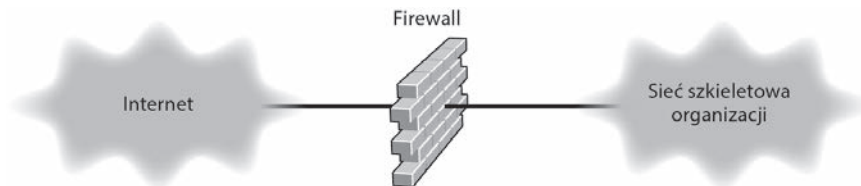
Byłoby wspaniale, gdyby udało się zatrzymać wszystkich intruzów na granicach sieci i uniemożliwić im dostęp do serwerów znajdujących się wewnątrz niej. Istnieją trzy drogi dostępu do typowej sieci: internet, sieci LAN i sieci WAN. Przeprowadzone badania wykazują preferowanie internetu w tej roli — ponad 70% ankietowanych organizacji doświadczyło prób ataku tą właśnie drogą, pozostałe 30% to próby wtargnięcia do sieci za pomocą dostępowych sieci LAN i sieci WAN. Stąd wniosek, że ochronę sieci na jej granicach należy ukierunkować przede wszystkim na jej połączenie z internetem. Choć — oczywiście — fizycznego zagrożenia też nie można lekceważyć.

Wartościowa polityka bezpieczeństwa powinna definiować następujące elementy:

- Nazwisko decyzyjnego menedżera odpowiedzialnego za kwestie bezpieczeństwa
- System raportowania incydentów naruszenia bezpieczeństwa i skład zespołu odpowiedzialnego za szybkie reagowanie na te incydenty
- Ocenę ryzyka w odniesieniu do poszczególnych zasobów, uszeregowaną według priorytetów
- Efektywne kontrole na wszystkich ważniejszych punktach wejścia do sieci, mające na celu powstrzymywanie lub odstraszenie zewnętrznych agentów
- Efektywne kontrole wewnątrz sieci, zapewniające, że żaden użytkownik wewnętrzny nie przekroczy przyznanego mu uprawnień
- Zasady minimalizowania wymaganych kontroli, w celu upraszczania obsługi sieci i minimalizowania uciążliwości dla użytkowników
- Politykę dozwolonego użycia, wyjaśniającą użytkownikom ich zezwolenia i ograniczenia, między innymi w zakresie dostępu do kont innych użytkowników, bezpieczeństwa haseł i reguł posługiwania się pocztą elektroniczną
- Procedurę monitorowania zmian w najważniejszych komponentach sieci – routerach, serwerach DNS itp.
- Plan rutynowych szkoleń personelu w zakresie polityki bezpieczeństwa i świadomości zagrożeń sieci
- Plan rutynowych testów i aktualizacji wszystkich kontroli, włącznie z monitorowaniem doniesień prasowych i raportów dostawców pod kątem informacji o wykrytych lukach bezpieczeństwa
- Plan rocznego audytu i przeglądu praktyk związanych z zapewnieniem bezpieczeństwa

RYSUNEK 11.11. Przykładowe wytyczne polityki bezpieczeństwa

Podstawowym narzędziem ochrony sieci przed nieuprawnionym dostępem z internetu są **zapory sieciowe**, zwane popularnie z angielska **firewallami** (ang. *firewall* — dosł. „ściana przeciwogniowa”). Firewall jest specjalizowaną odmianą routera, analizującą wszystkie pakiety wchodzące do sieci i wychodzące z niej, i selektywnie zezwalającą im na wpływ do sieci albo blokującą je na jej granicach, zgodnie ze zdefiniowanymi regułami. Sieć zostaje zaprojektowana w ten sposób, że każde jej połączenie z internetem chronione jest przez firewall (jak na rysunku 11.12), którego nie sposób ominąć. Niektóre firewalle posiadają zdolność wykrywania i unieszkodliwiania ataków DoS/DDoS oraz prób nieautoryzowanego dostępu do sieci. Zależnie od sposobu działania, firewalle dzielą się na trzy podstawowe grupy: firewalle pakietowe, firewalle aplikacyjne i firewalle NAT.



RYSUNEK 11.12. Ochrona sieci za pomocą firewalla

11.4.2.1. Firewally pakietowe

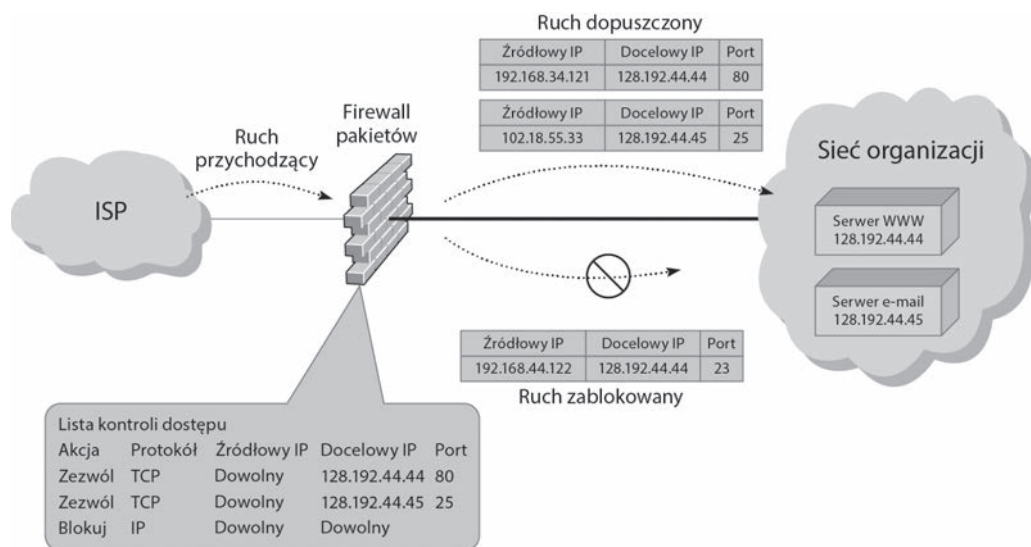
Firewall pakietowy (ang. *packet-level firewall*) analizuje adresy — źródłowy i docelowy — w każdym pakiecie przepływającym do lub z sieci i na podstawie zdefiniowanych reguł pakiet taki jest albo przepuszczany, albo odrzucany. W ogólności analizie podlegają tylko adresy warstwy sieciowej (IP) oraz numery portów warstwy transportowej. Analiza ma charakter bezstanowy — firewall analizując bieżący pakiet nie ma żadnej wiedzy na temat pakietów poprzednio analizowanych, więc decyzja o przepuszczeniu albo odrzuceniu pakietu jest decyzją bezkontekstową, niezależną od historii pakietów dotychczas analizowanych. Firewall pakietowy jest najprostszą odmianą firewalla i jednocześnie najmniej bezpieczną, nie monitoruje bowiem zawartości pakietu ani przyczyny jego wysłania, i zwykle nie prowadzi dziennika, który mógłby być pomocny w analizowaniu pracy sieci.

Reguły postępowania przez firewall z pakietami, zwane **listami kontroli dostępu** (ang. ACL — *Access Control Lists*), ustanawiane są przez menedżera sieci. Jak pamiętamy, pakiet IP zawiera w nagłówku źródłowy i docelowy adres IP, a enkapsulowany w nim segment TCP zawiera numer portu, identyfikujący aplikacje (źródłową i docelową) na poziomie warstwy aplikacyjnej. Większość aplikacji serwerowych używa standardowych numerów portów⁵, na przykład serwer WWW oferuje swe usługi na porcie 80, a serwer e-mail (SMTP) na porcie 25.

Załóżmy, że publiczny serwer WWW organizacji ma adres IP 128.192.44.44, a jej serwer e-mail — adres IP 128.192.44.45 (jak na rysunku 11.13). Menedżer sieci chce zapewnić, że nikt z zewnątrz organizacji nie będzie miał możliwości modyfikowania zawartości serwera WWW (na przykład za pomocą Telnetu lub FTP). Na liście ACL powinna więc znaleźć się reguła, zabraniająca dopływu z internetu do serwera WWW pakietów innych niż HTTP. Dokładniej, reguła ta powinna stanowić, że pakiet o dowolnym adresie źródłowym, adresie docelowym 128.192.44.44 i numerze portu docelowego 80 powinien być do sieci wpuszczony; podobnie powinny być do niej wpuszczane pakiety protokołu SMTP (port docelowy 25) o adresie docelowym 128.192.44.45 (jak na rysunku 11.13). Ostatnia na liście ACL jest (stosowana zwyczajowo) reguła nakazująca odrzucanie wszystkich pakietów, które nie są *explicite* dopuszczone przez inne reguły (niektóre firewally nie wymagają specyfikowania tej reguły, bo stosują ją domyślnie). Jeżeli więc na przykład intruz będzie chciał zmodyfikować zawartość serwera WWW za pomocą Telnetu (port docelowy 23), jego próby spełzną na niczym, bo firewall będzie pakiety Telnetu konsekwentnie odrzucał.

Mimo iż na listach ACL można specyfikować źródłowe adresy IP, to jednak menedżerowie często nie wykorzystują tej możliwości, z prostego powodu: większość hakerów fałszuje źródłowe adresy IP w wysyłanych pakietach (co z angielska nazywane jest „spoofingiem IP”), więc ich wykorzystywanie w regułach zabezpieczeń nie jest warte związanej z tym pracy. Bardzo często zdarza się tak, że sfałszowany pakiet trafiający do sieci ma w polu adresu źródłowego... jej własny adres IP lub adres jednej z jej podsieci, dlatego wielu menedżerów sieci umieszcza w swoich firewallach regułę nakazującą blokowanie takich pakietów.

⁵ Listę numerów portów popularnych protokołów można zobaczyć między innymi pod adresem https://pl.wikipedia.org/wiki/Port_protoko%C5%82u — *przyp. tłum.*



RYSUNEK 11.13. Tak działa firewall pakietowy

11.4.2.2. Firewalles aplikacyjne

Firewalles aplikacyjne (ang. *application-level firewalls*) są bardziej skomplikowane i przez to droższe od firewalli pakietowych, ponieważ analizują zawartość pakietu na poziomie warstwy aplikacyjnej w poszukiwaniu luk w zabezpieczeniu, ułatwiających realizację rozpoznanych wcześniej typów ataku (o lukach bezpieczeństwa piszemy w dalszym ciągu tego rozdziału). Funkcjonowanie firewalli aplikacyjnych sterowane jest regułami odnoszącymi się do poszczególnych aplikacji. Na przykład większość firewalli aplikacyjnych może analizować pakiety sieci Web (HTTP), pakiety e-mail (SMTP) i innych znanych protokołów. Jeżeli organizacja wykorzystuje na własny użytek tworzone przez siebie aplikacje, to konieczne jest wyposażenie firewalli w reguły umożliwiające funkcjonowanie tych aplikacji.

Jak pisaliśmy w sekcji 5.3.3, współdziałanie klienta z serwerem w ramach protokołu TCP odbywa się w trybie połączeniowym — klient ustanawia połączenie z serwerem, zanim zacznie wysyłać żądania. Firewalles aplikacyjne mogą wykonywać funkcję **inspekcji stanowej** (ang. *stateful inspection*), czyli monitorowania i przechowywania informacji o statusie każdego połączenia i wykorzystywania jej w procesie klasyfikowania pakietów jako potencjalnie zagrażających bezpieczeństwu.

Wiele firewalli aplikacyjnych uniemożliwia użytkownikom zewnętrznym upload *plików wykonywalnych*, przez co zarówno intruzy, jak i legalni użytkownicy nie mają możliwości zdalnej modyfikacji oprogramowania, chyba że mają fizyczny dostęp do firewalla. Niektóre firewalles aplikacyjne uniemożliwiają modyfikowanie własnego oprogramowania (chyba że zmiany te dokonywane są przez producenta); monitorują one własne oprogramowanie i natychmiast blokują wszelkie połączenia zewnętrzne, próbujące dokonywać takich modyfikacji.

11.4.2.3. Firewall NAT

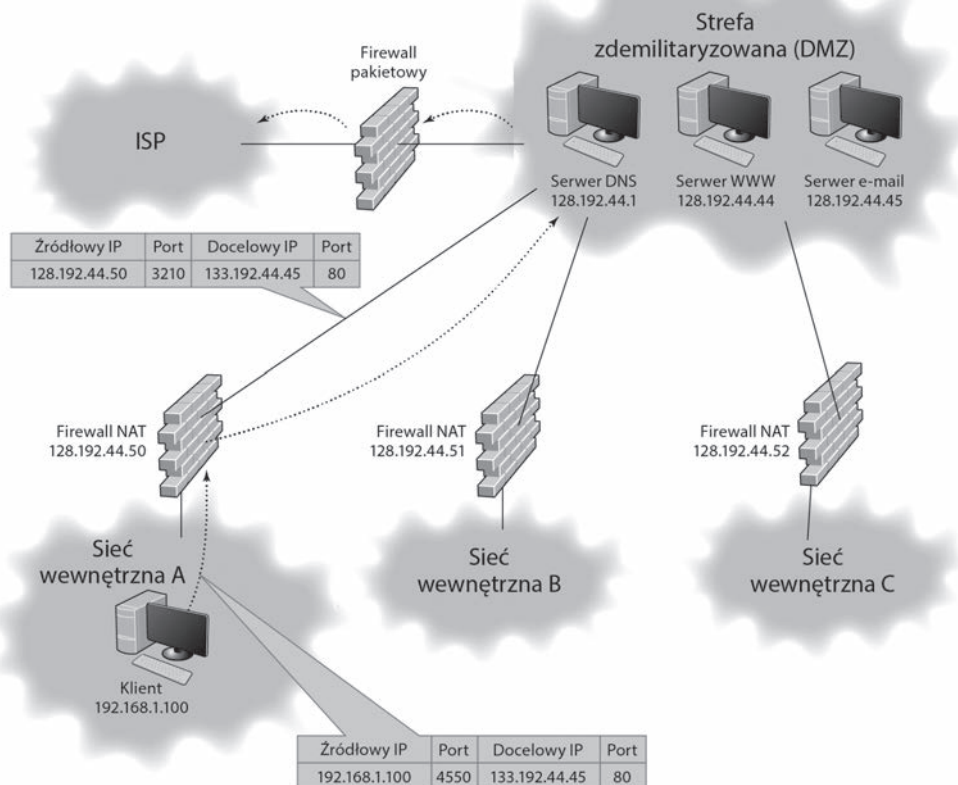
Zanim przejdziemy do opisu tej grupy firewalli, musimy pokrótce wyjaśnić, co kryje się pod akronimem NAT. Akronim ten jest skrótem od ang. *Network Address Translation*, co w dosłownym tłumaczeniu oznacza **tłumaczenie (translacje) adresów sieciowych**, czyli konwersję między adresami IP w publicznej przestrzeni internetu a **prywatnymi adresami IP**, nieinterpretowanymi na zewnątrz sieci (patrz rysunek 5.10). Mechanizm NAT jest przezroczysty; poza realizującym go firewallem inne komputery nic nie wiedzą o jego istnieniu. Głównym jego przeznaczeniem było oszczędne wykorzystywanie adresów IPv4, ale niejako przy okazji zauważono, że jest on wspaniałym środkiem zapewnienia bezpieczeństwa sieci — jeżeli komputery wewnątrz firmy opatrzone będą prywatnymi adresami IP, nie będą widoczne dla zewnętrznych intruzów, którzy tym samym nie będą w stanie ich atakować. Obecnie NAT wbudowywany jest w większość produkowanych firewalli i routerów, nawet tanich routerów do użytku domowego.

Firewall NAT jest czymś w rodzaju „okna na świat” dla komputerów wewnętrznych sieci firmowej. Gdy któryś z tych komputerów wysła pakiet przeznaczony dla komputera (serwera) gdzieś w internecie, pakiet ten trafia w pierwszej kolejności do firewalla, który rejestruje go w swoich tablicach i podmienia źródłowy adres IP na swój własny; podmieniany jest także numer portu źródłowego w enkapsulowanym segmencie TCP, na unikalną wartość, wybraną przez firewall. Oczywiście oryginalny adres źródłowy oraz oryginalny numer portu źródłowego zostają zapamiętane w tablicach firewalla. W ten oto sposób firewall staje się pełnomocnikiem (proxy) komputera wysyłającego żądanie, występując w jego imieniu; z punktu widzenia docelowego komputera w internecie jest to zgoła obojętne — wysła on odpowiedź na adres IP firewalla. Firewall odnajduje w swych tablicach (na podstawie podmienionego numeru portu źródłowego) adres IP wewnętrznego komputera oraz oryginalny numer portu źródłowego, i na ich podstawie przekazuje odpowiedź wewnętrznemu komputerowi.

Z perspektywy zewnętrznego komputera w internecie wszystko wygląda tak, jakby cała sieć firmowa składała się z jednego komputera (którym jest firewall). Założmy na przykład, że rejestrator przydzielił organizacji pulę adresów 128.192.55.X. Organizacja przydziela firewallowi NAT adres 128.192.55.1, lecz wewnętrzne komputery otrzymują adresy z puli prywatnej, na przykład 10.3.3.1, 10.3.3.2, 10.3.3.3 itd. (jak wynika z rysunku 5.10, adresy IPv4 o postaci 10.X.X.X są adresami prywatnymi), nie z puli 128.192.55.X. Nawet jeśli adresy wewnętrznych komputerów znane są hakerowi, nie jest on w stanie ich w żaden sposób użyć: nie są one trasowane w internecie, a jedyną do nich drogę zagraża firewall, niczym bodyguard całej sieci, nie do pokonania. Najciekawsze jest to, że cały mechanizm NAT jest całkowicie obojętny dla użytkowników wewnętrznych komputerów.

11.4.2.4. Architektura firewalli

W większości organizacji firewallle zorganizowane są w architekturę warstwową, której reprezentatywny przykład pokazaliśmy na rysunku 11.14. Pierwszym bastionem obronnym jest firewall pakietowy, oddzielający internet od sieci zawierającej wyłącznie serwery (głównie WWW, DNS oraz e-mail) przeznaczone do publicznego użytku. Sieć ta nazywana jest często **strefą zdemilitaryzowaną** (ang. DMZ — *DeMilitarized Zone*), ponieważ nie zapewnia swoim serwerom kompletnej ochrony. Firewall pakietowy zapewnia dostęp do DMZ pakietom HTTP, DNS i SMTP,



RYSUNEK 11.14. Przykładowy projekt sieci używającej firewalli

lecz blokuje ruch dla pakietów FTP z internetu do wspomnianych serwerów, uniemożliwiając modyfikowanie ich zawartości z zewnątrz sieci. Każda większa porcja sieci firmowej chroniona jest przez oddzielny firewall NAT, regulujący dostęp wewnętrznych użytkowników zgodnie z ustalonymi regułami.

Na rysunku 11.9 pokazaliśmy także drogę pakietu IP wysyłanego z komputera klienckiego w jednej z sieci wewnętrznych, chronionej przez firewall NAT. Pakiet ten w polu adresu źródłowego posiada adres komputera klienckiego (192.168.1.100), a w polu portu docelowego enkapsulowanego segmentu TCP znajduje się liczba 80, co wskazuje na protokół HTTP wysyłany do serwera WWW. Gdy pakiet ten dotrze do firewalla NAT, firewall zmienia źródłowy adres IP na swój własny (128.192.44.50), podobnie oryginalny numer portu źródłowego (4550) zostaje zastąpiony przez unikalny numer wygenerowany przez firewall (3210). W efekcie pakiet wysłany z gniazda 192.168.1.100:4550 po wysłaniu przez firewall do internetu wygląda jakby wysłano go z gniazda 128.192.44.50:3210. Jednocześnie firewall zapamiętuje w swoich tablicach, że numer portu źródłowego 3210 powiązany jest z oryginalnym adresem źródłowym 192.168.1.100. Oczywiście docelowy adres IP (128.192.44.44) i docelowy numer portu (80) pozostają niezmienione.

Pakiet dociera do serwera WWW, który wysyła odpowiedź do gniazda 128.192.44.50:3210, czyli na port 3210 firewalla. Firewall odnajduje w swoich tablicach, że port 3210 związany jest z oryginalnym gniazdem 192.168.1.100:4550. Zmienia więc docelowy adres IP na 192.168.1.100 i docelowy numer portu na 4550, po czym wysyła pakiet do sieci, czyli ostatecznie do komputera klienckiego. Serwer WWW nie wie nic o tym, na rzecz jakiego komputera rzeczywiście realizował żądanie.

11.4.2.5. Bezpieczeństwo fizyczne

Jednym z istotnych elementów zapobiegania nieautoryzowanemu dostępowi do wewnętrznych sieci firmy jest **bezpieczeństwo fizyczne**, czyli uniemożliwienie niepowołanym osobom dostępu do biur firmy, serwerowni, punktów dystrybucyjnych i poszczególnych urządzeń. Wszystkie punkty dystrybucyjne — główny oraz zdalne — powinny być adekwatnie zabezpieczone, a dostęp do nich powinien być monitorowany. Zaimplementowane systemy kontroli powinny ograniczać dostęp do zamkniętych pomieszczeń wyłącznie do autoryzowanego personelu. Komputery użytkowników powinny być wyposażone w zamki uniemożliwiające niepowołanym osobom włączanie zasilania; mogą być także chronione hasłem wymaganym do uruchomienia systemu.

Zwracaliśmy już uwagę na znaczenie kopii zapasowych i przechowywania ich w odległych lokalizacjach, a także na rozproszenie geograficzne ważnych serwerów, bo — w założeniu — ma to chronić cenne zasoby przed skutkami katastrof. Powstaje w tym momencie pytanie, czy takie rozproszenie faktycznie zwiększa bezpieczeństwo, czy może jednak zwiększa *liczbę lokalizacji narażonych na niebezpieczeństwo*? Liczbę miejsc narażonych na ataki? Liczbę koniecznych aktualizacji i zabezpieczeń? Bardzo często owe rozproszone serwery są częściami tej samej domeny logicznej; dla hakera oznacza to, że udane włamanie do jednego z nich daje mu dostęp do wszystkich innych. Wynika stąd ważna konieczność zdecydowania przez organizację, co — według własnego odczucia — daje jej większe poczucie bezpieczeństwa: pojedyncze centrum danych, z często sporządzanymi kopiami zapasowymi, czy proliferacja lokalizacji kolejnych serwerów.

Czasami zagrożenie zjawia się ze strony, z której — paradoksalnie — można by się go najmniej spodziewać. Metodą uzyskania nieautoryzowanego dostępu do sieci, znacznie prostszą od wyrafinowanych technik hakerskich, jest... zatrudnienie się w firmie w charakterze portiera lub ochroniarza, który w nocy, nie będąc przez nikogo niepokojonym, zainstaluje urządzenie podsłuchujące i rejestrujące komunikaty. Elementem wzmacniającym bezpieczeństwo fizyczne jest więc w tym przypadku należyta edukacja personelu odpowiedzialnego za zatrudnianie nowych pracowników. Nieuczciwy portier nie będzie miał szans na ukrycie swojego występkę, jeśli dostęp do newralgicznych pomieszczeń będzie ściśle ewidencjonowany i monitorowany.

Trzy obszary wrażliwe na **podśluchiwanie transmisji** (ang. *eavesdropping*) to bezprzewodowe sieci LAN, okablowanie i urządzenia sieciowe. Najprościej podsłuchiwać bezprzewodowe sieci LAN (WLAN), bo sygnał — radiowy lub podczerwieni — z punktów dostępowych przenika przez ściany pomieszczeń i budynków. Problematykę bezpieczeństwa sieci WLAN dyskutowaliśmy w rozdziale 7. i w tym miejscu nie będziemy do niej wracać.

Kable sieciowe są równie wdzięcznym obiektem amatorów podsłuchiwania transmisji, ponieważ często rozciągają się na dużych odległościach i zwykle nie są regularnie sprawdzane pod kątem obecności śladów „wgrzyzania się” do przewodów (ang. *wiretaping*). Kable stanowiące własność organizacji, instalowane w jej siedzibie, są na tę okoliczność najbardziej podatne:

znacznie łatwiej dostać się do żył skrętki biegnącej z MDF-u do konkretnego pomieszczenia, niż wyszukać kanał związany z organizacją w multipleksowanych obwodach dalekozasięgowego operatora telekomunikacyjnego. W związku z tym lokalne okablowanie powinno zostać ukryte między ścianami lub we wnętrzu stropu, a wyposażenie telefoniczne — wraz z szafkami — zamknięte w pomieszczeniach zabezpieczonych alarmem na wypadek próby włamania. Głównym tego celem jest ściśle kontrolowanie fizycznego dostępu do obwodów ze strony pracowników, dostawców i serwisantów.

DLA

11.3. Bezpieczeństwo danych wymaga fizycznego bezpieczeństwa

INŻYNIERA

Nie ma wątpliwości co do tego, że jeśli ktokolwiek może uzyskać fizyczny dostęp do Twojego serwera przez pewien odcinek czasu, to wszystkie informacje z Twoich komputerów — być może z wyjątkiem tych solidnie zaszyfrowanych — są dostępne dla hakerów.

W przypadku serwera Windows haker po prostu przeprowadza reboot systemu z płyty CD zawierającej dystrybucję systemu Knoppix (to odmiana Linuksa), być może po uprzedniej zmianie ustawień BIOS-u, odblokowującej możliwość bootowania z CD. Knoppix wyszukuje wszystkie napędy dyskowe w komputerze i umożliwia wygodne odczytywanie, na linuksowym pulpicie, wszystkich partycji NTFS i FAT32.

No dobrze, ale co z windowsowymi hasłami dostępu? Nic, po prostu nic — Knoppix zwyczajnie ignoruje ich istnienie. Haker może swobodnie odczytać, skopiować lub wysłać dowolny plik. Podobne, choć nieco bardziej skomplikowane ataki możliwe są także w serwerach linuksowych i w innych odmianach Uniksa.

Kabel kablowi nierówny, jedne są bardziej bezpieczne, inne mniej. Generalnie światłowody są zdecydowanie bardziej bezpieczne, bo wyprowadzenie sygnału świetlnego z włókna jest znacznie trudniejsze od uzyskania połączenia galwanicznego z żyłami skrętki. Skrętka też może stać się bardziej bezpieczna, jeśli zamknie się ją w trudnym do przebicia pancerzu. Istnieją także systemy alarmowe uruchamiające się w przypadku próby przebicia zewnętrznego płaszcza kabla, na przykład w U.S. Air Force używają kabli, w których pod szczelnym pancerzem płynie sprężony gaz. Jakiegokolwiek przebicie pancerza spowoduje spadek ciśnienia gazu i uruchomienie alarmu przez czujnik.

Urządzenia sieciowe — przełączniki i routery — powinny być umieszczane w zamykanych szafkach. Jak wyjaśnialiśmy w rozdziale 7., w sieci WLAN wszystkie komunikaty w sieci bezprzewodowej trafiają do wszystkich komputerów, i choć każdy komputer ignoruje komunikaty adresowane do innych, to przecież można w dowolnym komputerze zainstalować **program-szperacz** (ang. *sniffer*), kolekcjonujący wszystkie przepływające przez ten komputer komunikaty w celu późniejszej (nieautoryzowanej) ich analizy. Taki komputer można w sposób niezauważalny podłączyć do przełącznika kopiującego cały ruch sieciowy do ustalonej lokalizacji. **Bezpieczne przełączniki** (ang. *secure switches*) utrudniają nieco to zadanie, wymagając wprowadzenia specjalnego kodu autoryzującego w przypadku dodania nowego komputera.

11.4.3. Ochrona serwerów i klientów

11.4.3.1. Luki bezpieczeństwa

Nawet przy zapewnieniu absolutnego bezpieczeństwa fizycznego i uzbrojenia sieci w firewall, serwery i komputery klienckie nadal są zagrożone za sprawą luk bezpieczeństwa. Luka bezpieczeństwa (ang. *security hole*) to po prostu błąd oprogramowania (lub skutek niedopatrzania programistów) umożliwiający lub ułatwiający intruzom uzyskanie nieautoryzowanego dostępu. W powszechnie używanych systemach operacyjnych tkwią rozmaite luki bezpieczeństwa, bardzo dobrze znane hakerom — i producentom tychże systemów, którzy regularnie publikują stosowne **łaty** (ang. *patches*); sęk w tym, że menedżerowie sieci nie zawsze są świadomi nowo odkrywanych luk i nie zawsze w porę instalują publikowane aktualizacje.

Luki bezpieczeństwa to temat zasługujący na osobną książkę, a może nawet na całą bibliotekę, więc choćby pobieżne jego omówienie wykraczałoby poza ramy tej publikacji. Wiele luk bezpieczeństwa ma naturę czysto techniczną, czego przykładem jest choćby **atak przepełnienia bufora** (ang. *buffer overflow*), zmierzający do umieszczenia w specyficznym obszarze pamięci niewielkiego złośliwego kodu. Wiele innych ataków opiera się na prostych sztuczkach, które, choć skuteczne, nie są wcale oczywiste — jak choćby wysyłanie do serwera pakietów, wykorzystujących adres tego serwera jako zarówno docelowy, jak i źródłowy, w efekcie czego serwer sam siebie zaczyna zasypywać komunikatami i wkrótce kapitułuje.

Gdy tylko w aplikacji lub systemie wykryta zostanie nowa luka bezpieczeństwa, wieść o tym lotem błyskawicy ogarnia internet — i zaczyna się morderczy wyścig między hakerami a zespołami od zabezpieczeń, i współpraca po każdej ze stron barykady. Rolę centralnego banku informacji po stronie defensywnej pełni organizacja CERT, niezwłocznie rozpowszechniająca wieść o nowym zagrożeniu i sposobach przeciwdziałania mu, w sieci Web i w formie e-maili rozsyłanych do subskrybentów. Równie szybko reagują producenci dziurawego oprogramowania, publikując wkrótce stosowną łatę, neutralizującą możliwość wykorzystywania (eksploatowania) luki, dostępną do pobrania i zainstalowania. Bywa i tak, że odkrywca luki ujawnia swoją wiedzę jedynie producentowi, dając mu pewien czas na opracowanie stosownej łaty.

Niekiedy jednak sprawy przyjmują zupełnie inny obrót. Odkrywca nowej luki utrzymuje swą wiedzę w tajemnicy, sprzedając ją jednak cyberprzestępcom. Zainteresowane podmioty po drugiej stronie barykady (producenci i zespoły od zabezpieczeń) dowiadują się o luce dopiero wtedy, gdy jej przestępcza eksploatacja zdążyła już poczynić zauważalne spustoszenia. Takie ataki, realizowane w ukryciu jeszcze przed wynalezieniem antidotum, nazywamy **atakami dnia zero** (ang. *zero-day attacks*). Oczywiście firmy z branży bezpieczeństwa nie próżnują, starając się skupować interesujące informacje; w końcu wiele z aspektów zjawiska jest po prostu kwestią ceny.

Nawet jednak błyskawicznie udostępniane łatę nie spełnią pokładanych w nich nadziei, jeśli użytkownicy komputerów (w tym i menedżerowie sieci) nie wyrobią w sobie nawyku niezwłocznego instalowania publikowanych aktualizacji dla aplikacji czy systemów. Jego brak jest szczególnie groźny w przypadku, gdy między pojawieniem się luki a publikacją stosownej łatę upływa kilka tygodni czy nawet miesięcy. A propos — czy Ty, użytkowniku peceta z Windows czy MacIntosha, regularnie instalujesz dostępne aktualizacje?

DLA

11.5. Fałszywe antywirusy

MENEDŻERA

Świat wirusów komputerowych wciąż ewoluuje ku coraz bardziej wymyślnym konstrukcjom. Przed pojawieniem się internetu wirusy były przeważnie zabawkami do płatania figli — wyświetlania dziwnych tekstów, odgrywania muzyczki czy odwracania zawartości ekranu do góry nogami. Dzisiejsze wirusy to narzędzia do wykradania pieniędzy i wartościowych tajemnic. Gdy tylko wirus zainstaluje się na komputerze, może połączyć się z innym komputerem i przysyłać do niego wrażliwe dane. W celu przeciwdziałania takim zjawiskom tworzone są różne narzędzia antywirusowe. Jednak antywirus antywirusowi nierówny.

Wiele programów antywirusowych oferuje przeskanowanie komputera za darmo. Tak, za darmo! Zgodnie ze starym porzekadłem, jeżeli coś wydaje się zbyt piękne, by mogło być prawdziwe, faktycznie się tylko prawdziwym wydaje — i darmowe antywirusy nie są w tym względzie wyjątkiem. Chester Wisniewski z SophosLabs wyjaśnia, że gdy pobrałeś darmowy antywirus na swój komputer, to w istocie pobrałeś malware. Gdy taki program uruchomisz, jego wygląd i zachowanie stwarzać będą pozory, że masz do czynienia z profesjonalnym narzędziem — przyjazny interfejs GUI, często w wielu wersjach językowych, wiele opcji do wyboru. Ale nie wszystko złoto, co się świeci: gdy tylko uruchomisz skanowanie, zostaniesz zasypany komunikatami, z których wynikać będzie, że Twoje pliki to w większości trojany i robaki, i że Twój komputer w ogóle jest zainfekowany. Nieświadomy istoty rzeczy użytkownik przystaje na propozycję programu oferującego usunięcie „zainfekowanych” plików. Co więcej, skanowanie faktycznie ma miejsce, ale jego celem nie jest bynajmniej uwolnienie komputera od wirusów, tylko wyszukanie wrażliwych informacji i przesłanie ich na komputer hakerka.

Zamiast więc ulegać ułudzie darmowej usługi, lepiej wydać trochę pieniędzy na profesjonalną ochronę antywirusową, oferowaną przez produkty firm takich jak Sophos, Symantec, McAfee czy Kaspersky. Popularne czasopisma, między innymi „PC Magazine”, dostarczają corocznych raportów na temat różnych produktów antywirusowych, także tych darmowych. Pierwszym krokiem zmierzającym do ochrony przed opisanymi pułapkami jest bowiem edukacja.

Na podstawie: *Which Antivirus Is the Best* (www.pcantivirusreviews.com), Cassie Bodnar *Fake Antivirus: What Are They and How Do You Avoid Them* (blog.kaspersky.com).

Co ciekawe, niektóre luki bezpieczeństwa nie są lukami w ścisłym tego słowa znaczeniu, lecz rezultatami określonej polityki producentów, na przykład udostępniania systemów z preinstalowanymi kontami domyślnych użytkowników. Związane z tymi kontami domyślne hasła są powszechnie znane potencjalnym hakerom (admin/admin — skąd to znamy?), jeżeli więc menedżer sieci nie zatroszczy się o ich zmianę — szkoda gadać...

11.4.3.2. Systemy operacyjne

Agencje rządu USA wymagają określonych poziomów zabezpieczeń⁶ w używanych przez siebie systemach operacyjnych oraz aplikacjach sieciowych. Minimalny poziom zabezpieczeń, wymagany między innymi dla Windows, to C2. W stosunku do powszechnie wykorzystywanych systemów dąży się do zapewnienia bezpieczeństwa w wyższych kategoriach, na przykład B2. Niektóre krytyczne systemy wymagają zabezpieczenia na najwyższym poziomie — A1 lub A2.

Od niepamiętnych czasów toczą się debaty o tym, czy Windows jest systemem mniej bezpiecznym od innych systemów operacyjnych, na przykład Linuksa. Z każdym nowym atakiem na Windows odżywa na nowo spór między krytykami systemów Microsoftu („A nie mówiłem?”) a ich obrońcami, twierdzącymi, że systemy Windows, jako bardziej popularne, są częściej atakowane, i to w dużej mierze przez samych przeciwników.

DLA

11.4. Eksploatowanie luk bezpieczeństwa

INŻYNIERA

Aby haker mógł zrobić użytek z luki bezpieczeństwa, najpierw musi się o niej dowiedzieć. Jak? To oczywiście: przecież żyjemy w erze zautomatyzowanych narzędzi.

Haker rozpoczyna swoją akcję od zidentyfikowania serwerów znajdujących się w atakowanej sieci, za pomocą oprogramowania wypróbowującego kolejno wszystkie możliwe adresy IP w tej sieci. Ze zbioru znalezionych serwerów haker wybiera sobie kilka jako cel potencjalnych ataków.

Krok drugi to zidentyfikowanie usług świadczonych przez poszczególne serwery. W tym celu w każdym serwerze haker dokonuje systematycznego skanowania możliwych numerów portów TCP, wykrywając te, które reagują na żądania. Jeżeli na przykład serwer odpowiada na żądanie skierowane na port 80., to jest to niechybnie serwer WWW, a gdy reaguje na żądanie skierowane na port 25., jest serwerem e-mail.

W trzecim kroku haker próbuje zidentyfikować oprogramowanie — jego wersję i producenta — odpowiedzialne za świadczenie określonej usługi na określonym serwerze. Tak się składa, że różne aplikacje w różny sposób reagują na pewne specyficzne komunikaty; wiedza na ten temat wykorzystana została do konstrukcji zaawansowanych skanerów, które doświadczony haker pewnością posiada w swoim przyborniku.

Mając już wiedzę na temat profilu swojej ofiary, haker wykorzystuje kolejne skanery, tym razem wyszukujące znane luki w danym oprogramowaniu. Przykładowo, starsze serwery SMTP nie wymagały uwierzytelniania od użytkownika wysyłającego wiadomość pocztową; haker, mając do czynienia z takim serwerem, może preparować pakiety SMTP z fałszywym adresem źródłowym i zalewać adresatów poczty powodzią spamu. Z kolei niektóre wersje popularnych pakietów e-handlu umożliwiają wymuszanie na serwerach uniksowych wykonywanie pewnych poleceń, w wyniku umiejętnego użycia operatora potoku („|”) po którym następuje nazwa uploadowanego pliku. Po zakończeniu uploadu system automatycznie otworzy plik, uruchamiając scenariusz zaprogramowany w nim przez hakera.

⁶ Wyjaśnienie używanych poniżej symboli bezpieczeństwa znajdują Czytelnicy pod adresem https://pl.wikipedia.org/wiki/Kategoria:Standardy_bezpiecze%C5%84stwa_teleinformatycznego — przyp. tłum.

Zasadnicza różnica między tym, co potrafią aplikacje dla Windows, a tym, co potrafią aplikacje dla Linuksa, wynika wprost z odmiennej genealogii tych systemów. Protoplasta Linuksa — Unix — oryginalnie zaprojektowany został jako system dla wielu użytkowników, którzy mieli różne uprawnienia. Dostęp do krytycznych obszarów systemu mieli wyłącznie uprzywilejowani użytkownicy, będący jego administratorami, pozostali użytkownicy byli pozbawieni takiej możliwości.

Dla odróżnienia, system DOS, z którego począł się Windows — początkowo jako nakładka graficzna, później jako samodzielny system operacyjny — zaprojektowany został z myślą o pojedynczym użytkowniku, sprawującym niepodzielną kontrolę na swym komputerem, nomen omen — osobistym. Aplikacje dla Windows miały więc dostęp do wszystkich jego zakamarków, dzięki czemu mogły dokonywać wielu wspaniałych rzeczy, przed oczami użytkownika nieświadomego w dużej mierze tego, co dzieje się pod podszewką. Bogactwu aplikacji towarzyszyła przyjazność ich obsługi. To wszystko dostępne było dzięki właściwościom samego Windows, bez konieczności doinstalowywania czegokolwiek. Położenie kresu tej omnipotencji użytkownika wymagałoby głębokiej rekonstrukcji Windows od podszewki, wskutek czego wiele aplikacji zaprojektowanych dla dotychczasowych wersji przestałoby w nowej wersji funkcjonować. Dla większości użytkowników byłoby to zbyt wysoką ceną za realizację jakiejś bliżej nieokreślonej zachcianki o nazwie „bezpieczeństwo”.

Ta przyjazność ma — oczywiście — swoją cenę. Skoro aplikacje mogą robić niemal wszystko, i to bez wiedzy użytkownika, to mogą działać także na jego szkodę. Bo tak już jest, że dobry system operacyjny jest wynikiem nieuchronnego kompromisu między przyjaznością obsługi a restrykcjami, a bezpieczeństwo wymaga poświęcenia pewnej dozy atrakcyjności.

11.4.3.3. Konie trojańskie

Koń trojański (zwany też popularnie **trojanem**) to kolejny typ złośliwego oprogramowania, zawdzięczający nazwę analogii swego działania do podstępny, który rozstrzygnął losy niemal 10-letniej wojny między Grecją a Troją — historycy umiejscawiają ten konflikt w XII wieku przed Chrystusem. Trojanie lekkomyślnie przyjęli w darze od Greków olbrzymią drewnianą figurę konia, w której wnętrzu ukryci byli greccy wojownicy; ci pod osłoną nocy wyszli z ukrycia i otwarli bramy miasta, zapoczątkowując jego całkowite zniszczenie. Czy historia ta jest autentyczna, czy też stanowi jedynie alegorię jakiegoś kataklizmu, nie ma w tej chwili znaczenia, faktem jest natomiast, że destrukcyjne działanie wirusa z kategorii trojanów rozpoczyna się od wpuszczenia go do komputera, najczęściej przez zainstalowanie programu, w którym został ukryty. Programy ukrywające trojany — filmy, utwory muzyczne, gry, rozmaite programy narzędziowe — na pierwszy rzut oka nie wzbudzają podejrzeń. Gdy użytkownik ogląda ulubiony film czy też słucha ulubionej muzyki, nie jest świadomy tego, że równocześnie wirus, który właśnie opuścił swoją kryjówkę, przystąpił do działania, otwierając porty komputera i umożliwiając hakerowi całkowite przejście kontroli nad nim.

Do najbardziej spektakularnych trojanów w historii należy *Back Orifice*, który w 1998 roku agresywnie zaatakował serwery Windows i który jeszcze dziś nawiedza administratorów w ich koszmarnych snach. Intruz, posiadający możliwości operowania komputerem na równi z użytkownikiem siedzącym przed klawiaturą, mógł dowolnie manipulować zestawem i zawartością plików, ustawieniami rejestru, konfiguracją sieci i przekierowywaniem komunikatów.

Młodszy celebryci tego gatunku sztuki to *Optix Pro* (2004) i *MoSucker* (2009). Haker operujący na konsoli pierwszego z nich (rysunek 11.15) może jednym kliknięciem wyłączyć firewall i narzędzia antywirusowe w komputerze ofiary, podglądać i podsłuchiwać jego otoczenie (za pomocą podłączonych do niego kamery i mikrofonu, nawet gdy te wydają się aktualnie nie działać), a także wywoływać różne dodatkowe efekty, jak wyświetlanie napisów, generowanie dźwięków, zamiana funkcji przycisków myszy czy też manipulowanie tacką napędu CD/DVD.



RYSUNEK 11.15. Centrum sterowania wirusa Optix Pro

Narzędzia te są nie tylko skuteczne, ale i proste w obsłudze. Znacznie trudniejszą sprawą jest tworzenie programów chroniących przed ich destruktywnymi działaniami. Co w takim razie czeka nasze komputery w najbliższej przyszłości? Łatwo sobie wyobrazić, że rezydujący w komputerze trojan aktywuje się (powiedzmy) o 2:00 w nocy, wybiera losowo port, wysyła użytkownikowi e-mail „Byłem tu, wszedłem przez port #NNNNN, trochę nabałagałem, dobranoc”, wykonuje zaplanowaną destrukcję, uruchamia skrypt zacierający wszelkie ślady swej bytności, być może nawet usuwa swe pliki z dysku. Przerażające? Tak, ale możliwe.

Najczęstsze typy trojanów to **oprogramowanie szpiegujące** (ang. *spyware*), **oprogramowanie reklamowe** (ang. *adware*) i agenty DDos. Oprogramowanie szpiegujące, zgodnie z nazwą, rejestruje zdarzenia zachodzące w związku z aktywnością użytkownika — najczęściej jest to nagrywanie naciskanych klawiszy (taki trojan nazywany jest z ang. *keyloggerem*), dzięki czemu haker rozpoznaje wprowadzane przez użytkownika loginy i hasła, za pomocą których może następnie opróżnić konto bankowe ofiary. Oprogramowanie reklamowe śledzi aktywność użytkownika i w różnych momentach wyświetla (w wyskakujących oknach) ogłoszenia reklamowe — jeżeli na przykład użytkownik kliknie odsyłacz prowadzący do sklepu internetowego, spyware może wyświetlić okno z reklamą konkurencyjnego sprzedawcy, lub — co gorsza — przekierować użytkownika bezpośrednio do konkurencyjnego sklepu. O agentach DDos pisaliśmy już wcześniej w tym rozdziale.

Po drugiej stronie barykady mamy rozmaite narzędzia przeciwdziałające spełnianiu się tej ponurej wizji, uniwersalne lub ukierunkowane na konkretną kategorię trojanów (na przykład Spybot). Niektórzy producenci firewalli wyposażają swoje produkty w logikę antytrojanową, blokując możliwość przedostawania się trojanów do sieci i ich wypływanie z niej.

11.4.4. Szyfrowanie

Przechwycona przez hakera porcja danych będzie bezużyteczna, jeśli nie będzie on potrafił zinterpretować jej znaczenia. Temu celowi służy **szyfrowanie** (ang. *encryption*) danych, czyli przekształcenie ich za pomocą algorytmów realizujących odwracalne funkcje matematyczne. Szyfry znane były już w starożytności, obecnie jednak związana z nimi wiedza rozrosła się do potężnej dziedziny nauki, zwanej **kryptologią** (praktyczne wykorzystywanie jej metod nosi nazwę **kryptografii**)⁷.

Szyfrowanie informacji, której oryginalna postać nosi nazwę **tekstu jawnego** (ang. *plaintext*), przekształca ją do postaci **szyfrogramu** (ang. *ciphertext*). Przekształcenie to zależne jest od parametrów (jednego lub więcej) zwanych **kluczami** (ang. *keys*). Proces odwrotny, czyli odtwarzanie tekstu jawnego na podstawie szyfrogramu, nosi nazwę **deszyfracji** (ang. *decryption*) i jest niewykonalny bez znajomości odpowiednich kluczy.

Szyfrowane mogą być zarówno pliki przechowywane w pamięciach masowych (dyskach), jak i strumienie przesyłanych danych. Wiele firm wykorzystuje obie te możliwości, ponieważ szyfrowanie plików da się łatwo realizować przez odpowiednie ustawienia systemu. W Windows szyfrowanie pliku jest jednym z elementów jego właściwości, podobnie wiele witryn sieci Web przesyła dane w postaci zaszyfrowanej, przy użyciu protokołów w rodzaju HTTPS, użytkownik może też szyfrować całość komunikacji swojego komputera ze światem zewnętrznym, wykorzystując mechanizm VPN (patrz sekcja 9.4).

Metody szyfrowania dzielą się zasadniczo na dwie grupy. W **szyfrowaniu symetrycznym** zarówno szyfrowanie, jak i deszyfracja odbywają się przy użyciu tego samego klucza, w **szyfrowaniu asymetrycznym** szyfrowanie i deszyfracja odbywają się przy użyciu różnych kluczy.

11.4.4.1. Kryptografia symetryczna

Szyfrowanie symetryczne obejmuje dwa elementy: **algorytm** i **klucz**. Ten sam tekst jawny, zaszyfrowany za pomocą tego samego algorytmu, lecz przy użyciu różnych kluczy, daje różne szyfrogramy. By możliwa była szyfrowana komunikacja między dwoma użytkownikami, każdy musi znać wspólny klucz. Dobry algorytm szyfrowania nie musi być utrzymywany w tajemnicy, natomiast konieczne jest utajnienie klucza przed osobami niepowołanymi. Ponieważ szyfrowanie i deszyfracja są przekształceniami matematycznymi, klucz ma postać liczby całkowitej, pełniąc rolę parametru tych przekształceń. Oczywistym podejściem do próby złamania szyfru jest wypróbowanie w roli klucza wszystkich możliwych liczb składających się na przestrzeń kluczy; podejście takie nosi nazwę **ataku siłowego** (ang. *brute force*) i ma tym mniejsze szanse powodzenia, im większa jest liczebność tej przestrzeni, czyli im większa jest długość klucza (liczona w bitach). Przy dostatecznie długim kluczu czas i koszty przeprowadzenia takiego ataku byłyby niewspółmiernie duże w porównaniu z wartością chronionej informacji.

⁷ Wyczerpujący opis algorytmów i technik szyfrowania znajdują Czytelnicy w książce Williama Stallinsa *Kryptografia i bezpieczeństwo sieci komputerowych. Matematyka szyfrów i techniki kryptologii*, wyd. Helion 2011, <https://helion.pl/ksiazki/krybez.htm> — przyp. red.

DLA

1.6. Ransomware się opłaca — niekiedy dość słono

MENEDŻERA

Ransomware to specjalny typ złośliwego oprogramowania, którego zakładnikiem stają się dane ofiary; warunkiem uwolnienia tych danych jest wpłacenie okupu. Gdy użytkownik uruchomi plik zawierający takie oprogramowanie bądź wykona celowo spreparowany skrypt na stronie WWW, *ransomware* zaszyfruje wszystkie pliki na dysku komputera, przy użyciu klucza nieznanego ofierze, przez co straci ona dostęp do zawartości tych plików. Jednocześnie wyświetlone zostanie żądanie wpłacenia w ciągu dnia lub dwóch określonej kwoty na wskazane konto, w przeciwnym razie wszystkie pliki zostaną usunięte z komputera, gdy natomiast kwota zostanie wpłacona, ofiara otrzyma klucz umożliwiający odszyfrowanie plików.

Kwota żądanego okupu jest zwykle umiarkowana (na przykład 750 dolarów), więc użytkownikowi-ofierze bardziej opłaca się zapłacić okup niż samodzielnie próbować rozwiązać problem. Atakujący często żądają zapłaty w najbardziej umiędzynarodowionej walucie — bitcoinach, na konta ulokowane w krajach, w których cyberprzestępstwa nie są ścigane.

Ofiarami *ransomware* są przeważnie użytkownicy indywidualni. Są oni zwykle mniej wyedukowani informatycznie niż profesjonalne zespoły zatrudnione w firmach, no i łatwiej jest wyłudzić pieniądze od osoby fizycznej niż od firmy. Hakerzy jednak dokonują ataków *ransomware* na serwery małych i średnich firm, a żądane kwoty są wtedy znacząco wyższe niż w przypadku użytkowników indywidualnych. Ofiarami padają najczęściej firmy świadczące specjalistyczne usługi — lekarze, dentyści, prawnicy — po pierwsze dlatego, że mają szczególnie dużo do stracenia, po wtóre dlatego, że stać je na płacenie większych kwot.

Według raportu FBI, amerykańskie firmy tracą z tytułu *ransomware* około miliarda dolarów rocznie, włączając w to koszty przywrócenia danych i utracone zyski z powodu przerwy w funkcjonowaniu. Raport ten nie obejmuje jednak użytkowników indywidualnych, poza tym nie wszystkie firmy zgłaszają incydenty *ransomware* do FBI. Wymienioną kwotę należy więc raczej traktować jako wierzchołek góry lodowej.

Ponieważ obie strony szyfrowanej komunikacji współdzielił ten sam klucz, kryptografia symetryczna wiąże się z problemem **zarządzania kluczami**. Obie komunikujące się strony — komputery lub ich użytkownicy — muszą znać wspólny klucz, który tym samym musi zostać uprzednio przez nich uzgodniony. Uzgodnienie to musi przebiegać w taki sposób, by nawet przechwycenie wszystkich danych wymienianych w jego ramach nie umożliwiałoby odgadnięcia wynikowego klucza. To, wydawałoby się, niemożliwe zadanie jest wykonalne, dzięki pewnym własnościom matematycznym, a algorytm jego realizacji nosi nazwę **protokołu Diffiego-Hellmana**, od nazwisk jego wynalazców, Witfielda Diffiego i Martina Hellmana. Szczegółowy opis tego algorytmu, wraz z przykładem zastosowania, znajduje się pod adresem https://pl.wikipedia.org/wiki/Protok%C3%B3%C5%82_Diffiego-Hellmana, a także w przywoływanej już książce Williama Stallingsa. Jego zasadnicza idea — niemożność efektywnego obliczania logarytmów dyskretnych — wykorzystana została w wielu nowoczesnych algorytmach kryptograficznych, między

innymi w **kryptografii krzywych eliptycznych** (ECC — patrz https://pl.wikipedia.org/wiki/Kryptografia_krzywych_eliptycznych).

Oczywiście możliwe jest uzgadnianie kluczy w inny sposób — kiedyś przecież radzono sobie z tym bez uciekania się do subtelności matematycznych — ważne jest jednak utrzymywanie kluczy w absolutnej tajemnicy: ponieważ algorytmy szyfrowania i deszyfracji są jawne, więc ujawnienie klucza oznaczałoby skompromitowanie całej ochrony informacji.

11.4.4.1.1. DES

Jedną z podstawowych technik kryptografii symetrycznej był **szyfr DES** (ang. *Data Encryption Standard*), zaprojektowany przez firmę IBM w 1975 roku na zamówienie agencji USA noszącej obecnie nazwę Narodowego Instytutu Standaryzacji i Technologii (NIST). Najczęściej stosowana odmiana tego szyfru wykorzystuje klucz 56-bitowy; eksperci udowodnili, że wiadomość zaszyfrowana tą odmianą szyfru można przy użyciu odpowiednich narzędzi odtworzyć z szyfrogramu, bez znajomości klucza, w czasie krótszym niż 24 godziny. Nie może być więc uznany za zbyt bezpieczny i dlatego w 2001 roku wyparty został z użycia w USA przez nowszy kryptosystem AES.

11.4.4.1.2. Triple DES (3DES)

Triple DES — „potrójny DES” — oznaczany także skrótem 3DES, zgodnie z nazwą wykonuje szyfrowanie wg algorytmu równoważnego trzykrotnemu szyfrowaniu kluczem 56-bitowym, co można porównać do jednokrotnego szyfrowania z kluczem 168-bitowym. Algorytm szyfrowania 3DES wykorzystuje *dwa* klucze w trzech krokach, z których każdy jest równoważny szyfrowi DES: pierwszym krokiem jest zaszyfrowanie tekstu jawnego pierwszym kluczem, w drugim kroku rezultat pierwszego kroku *rozszyfrowywany* jest za pomocą drugiego klucza, w trzecim kroku wynik kroku drugiego szyfrowany jest pierwszym kluczem:

$$m_1 = E(K_1, M)$$

$$m_2 = D(K_2, m_1)$$

$$C = E(K_1, m_2)$$

(E oznacza szyfrowanie DES, D oznacza deszyfrację DES, M jest tekstem jawnym, C jest ostatecznym szyfrogramem).

Użycie dwóch identycznych kluczy równoważne jest pojedynczemu szyfrowaniu DES.

$$m_1 = E(K, M)$$

$$m_2 = D(K, m_1) = M$$

$$C = E(K, M)$$

Deszyfracja 3DES jest odwróceniem szyfrowania:

$$c_1 = D(K_1, C)$$

$$c_2 = E(K_2, c_1)$$

$$M = D(K_1, c_2)$$

11.4.4.1.3. AES

Nowym standardem NIST jest od 2001 roku kryptosystem **AES** (ang. *Advanced Encryption Standard*), pierwotnie noszący nazwę *Rijndael*, od nazwisk wynalazców: *Vincenta Rijmena* i *Joan Daemen*. Występuje w trzech odmianach, różniących się długością klucza: 128, 192 i 256-bitowej. Według oceny NIST, złamanie tego szyfru metodą *brute force* zajęłoby, przy użyciu najnowocześniejszych dostępnych dziś komputerów i technik, około 150 bilionów (150×10^{12}) lat. Oczywiście wobec żywiłowo rozwijających się technologii oszacowanie to trzeba nieustannie korygować w dół, ale i tak AES wydaje się nie do złamania w dającej się przewidzieć przyszłości. Od momentu wynalezienia DES do ostatecznego zastąpienia go przez AES minęło 26 lat, więc AES-owi można wróżyć podobnie długi żywot.

11.4.4.1.4. RC4

Szyfr RC4 — to skrót od *Rivest Cipher 4* — wynaleziony został przez Rona Rivesta z RSA Security w 1987 roku. Długość używanego klucza może sięgać 256 bitów, lecz najczęściej stosowany jest klucz 40-bitowy. Jest szybszy w działaniu niż DES, lecz tak samo jak on jest podatny na ataki *brute force*, za pomocą których można go złamać w ciągu dwóch dni.

11.4.4.2. Kryptografia z kluczami publicznymi

W kryptografii asymetrycznej szyfrowanie i deszyfracja odbywają się przy użyciu różnych kluczy. Najbardziej spektakularną odmianą kryptografii asymetrycznej jest **kryptografia z kluczami publicznymi**. Jej podstawy matematyczne sformułowane zostały w MIT w 1977 roku przez Rona Rivesta, Adiego Shamira i Leonarda Adlemana, a technika kryptograficzna nazwana została RSA, od pierwszych liter nazwisk wynalazców. Podobnie jak protokół Diffiego-Hellmana opierał się na trudności obliczania logarytmów dyskretnych, tak metoda RSA opiera się na praktycznej niemożności wykonania innej operacji matematycznej — **faktoryzacji**, rozkładu dużej liczby na czynniki pierwsze. Szczegóły tego algorytmu opisane są pod adresem https://pl.wikipedia.org/wiki/Kryptografia_klucza_publicznego oraz w cytowanej książce Wiliama Stallingsa.

Praktycznym zastosowaniem algorytmu RSA jest **infrastruktura klucza publicznego** (PKI — *Public Key Infrastructure*), zapewniająca nie tylko poufność informacji, lecz także cyfrowe podpisywanie dokumentów, stanowiące podstawę podpisu elektronicznego i profili zaufanych.

Z każdym podmiotem (firmą lub osobą fizyczną) uczestniczącym w PKI związane są dwa zależne od siebie klucze: **publiczny** i **prywatny**. Klucz publiczny jest ogólnie dostępny i potwierdzony jest certyfikatem; klucz prywatny znany jest wyłącznie swemu właścicielowi. Długość kluczy może wynosić 512, 1024, 2048 lub 4096 bitów.

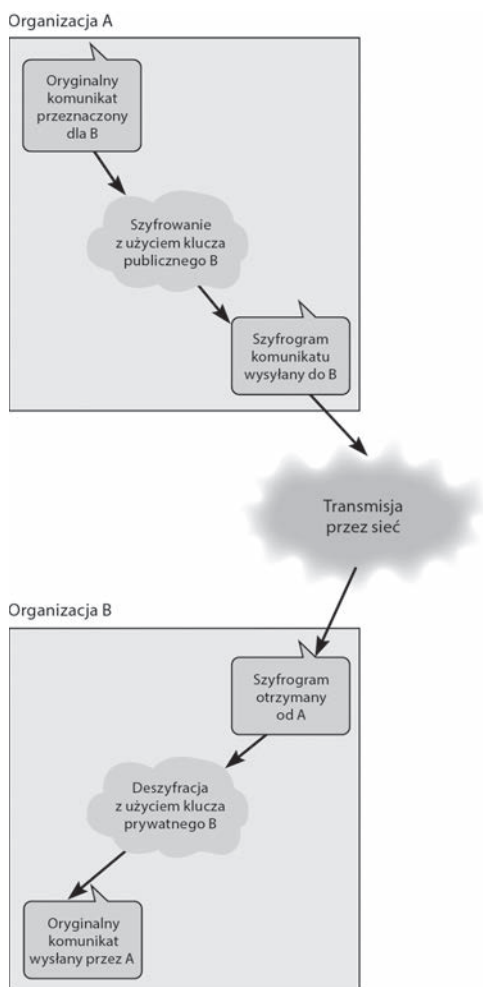
Działanie kryptografii z kluczami publicznymi opiera się na matematycznej funkcji jednokierunkowej. W celu ukrycia informacji nadawca szyfruje ją kluczem publicznym adresata i wynikowy szyfrogram wysyła adresatowi. Mimo iż szyfrogram ten jest dostępny dla osób niepowołanych, jest dla nich bezużyteczny, bo do odtworzenia tekstu jawnego niezbędny jest klucz prywatny, znany wyłącznie adresatowi. Na tym właśnie polega wspomniana jednokierunkowość: zaszyfrowanie tekstu jawnego w celu wysłania szyfrogramu adresowego jest sprawą prostą, zrealizowanie operacji odwrotnej — jest niewykonalne.

Nie można w tym momencie nie postawić pytania: skoro klucze publiczny i prywatny są względem siebie komplementarne, to czy nie można wydedukować klucza prywatnego podmiotu na podstawie powszechnie znanego jego klucza publicznego? Odpowiedź przecząca na to pytanie wynika ze wspomnianej wcześniej niewykonalności rozkładu na czynniki pierwsze dużej liczby,

która jest jednym z elementów klucza. Kryptografia z kluczami publicznymi jest więc jednym z najbardziej bezpiecznych systemów szyfrowania, o zastosowaniu uniwersalnym, z wyjątkiem może strategicznych systemów agencji rządowych, wymagających specjalnych zabezpieczeń.

Infrastruktura kluczy publicznych w znacznym stopniu ułatwia zarządzanie kluczami. Gdy dwoje uczestników chce nawiązać ze sobą szyfrowaną komunikację, każdy z nich odnajduje opublikowany klucz publiczny partnera i żadna procedura uzgadniania wspólnego klucza nie jest potrzebna. Jedynym problemem „zarządzania kluczami” jest w tym przypadku utrzymywanie w tajemnicy własnego klucza prywatnego.

Zasadę tę zilustrowaliśmy na rysunku 11.16. Gdy organizacja A zamierza w sposób bezpieczny przesłać komunikat do organizacji B, odnajduje w publicznym katalogu klucz publiczny tej ostatniej, szyfruje przy użyciu tego klucza wspomniany komunikat i uzyskany szyfrogram przesyła do organizacji B. Organizacja B deszyfruje otrzymany szyfrogram przy użyciu swego klucza prywatnego.



RYСУNEK 11.16. Bezpieczna transmisja danych z zastosowaniem kryptografii z kluczami publicznymi

11.4.4.2.1. Uwierzytelnianie tożsamości

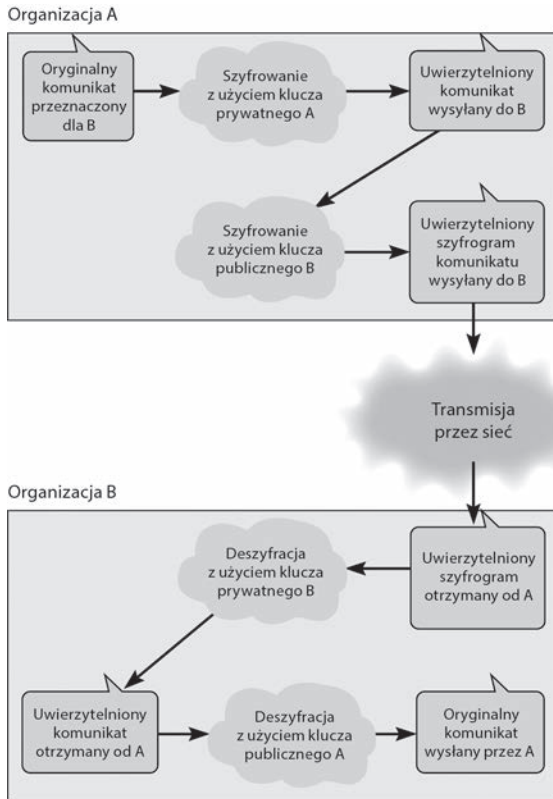
Drugą ważną funkcją kryptografii z kluczami publicznymi jest **uwierzytelnianie podmiotu** (ang. *authentication*), ściśle związane z funkcją **podpisu cyfrowego** (ang. *digital signature*). Natura niektórych komunikatów wymienianych między dwoma podmiotami wymaga, aby nadawca udowodnił przed adresatem swoją tożsamość, czyli **uwierzytelniał się** wobec niego (ang. *authenticate*); bez tego komunikaty te nie będą przedstawiać żadnej wartości. Transakcje bankowe, zakupy online, umowy prawne, zlecenia transakcji giełdowych — to tylko niektóre przykłady.

Kryptografia z kluczami publicznymi ma tę specyficzną cechę, że szyfrowanie i deszyfrowanie muszą odbywać się przy użyciu pary *komplementarnych* kluczy. Gdy nadawca, w celu utajnienia komunikatu, zaszyfruje go kluczem publicznym odbiorcy, do jego odszyfrowania należy użyć klucza prywatnego odbiorcy, w przypadku użycia innego klucza otrzymamy bowiem zwykle chaotyczną mieszankę bajtów zamiast oryginalnego komunikatu. Uwierzytelnianie za pomocą pary komplementarnych kluczy tym różni się od utajniania komunikatu, że szyfrowanie odbywa się przy użyciu klucza prywatnego nadawcy, zaś rozszyfrowywanie przy użyciu klucza publicznego nadawcy. Dla dużych objętościowo komunikatów, ze względów praktycznych szyfrowaniu podlega nie cały komunikat w oryginalnej postaci, ale ciąg charakterystycznych dla niego kluczowych elementów — nazwy nadawcy i odbiorcy, kwoty będące przedmiotem umowy itd. — lub (częściej) **wyciąg**, czyli wynik **funkcji mieszającej** zastosowanej do treści komunikatu (patrz https://pl.wikipedia.org/wiki/Funkcja_skr%C3%B3tu). Odbiorca odszyfrowuje otrzymany szyfrogram za pomocą klucza publicznego odbiorcy i porównuje odtworzone elementy kluczowe z treścią komunikatu (lub oblicza wyciąg komunikatu i porównuje go z tym otrzymanym w wyniku deszyfracji). Zgodność porównania oznacza, że kluczem użytym do zaszyfrowania komunikatu musiał być klucz komplementarny do klucza publicznego nadawcy, czyli jego klucza prywatny. Jest to dowodem na autentyczność nadawcy, ponieważ nikt poza nim nie zna jego klucza prywatnego.

Uwierzytelnianie może zostać połączone z utajnianiem, tak by dane uwierzytelniające zostały bezpiecznie dostarczone do odbiorcy. Schemat takiego przesyłania od organizacji A do organizacji B przedstawiliśmy na rysunku 11.17; dla prostoty przyjęliśmy, że uwierzytelnianiu podlega komunikat jako całość, nie tylko jego kluczowe elementy czy wyciąg. Rezultat zaszyfrowania oryginalnego komunikatu jest szyfrowany kluczem prywatnym A, a otrzymany szyfrogram zaszyfrowany jest kluczem publicznym B; wynik tego drugiego szyfrowania przesyłany jest do B. Po otrzymaniu przesyłki, B rozszyfrowuje ją przy użyciu swojego klucza prywatnego i wynik tej deszyfracji rozszyfrowuje przy użyciu klucza publicznego A. Jeśli wynikiem tej drugiej deszyfracji jest oryginalna treść komunikatu, mamy dowód na autentyczność nadawcy.

Możliwość udowodnienia przez nadawcę swojej tożsamości ma jeszcze jeden niebagatelny aspekt: zapobiega mianowicie wyparciu się faktu nadania wiadomości, gdy z jakichś powodów fakt ten stanie się dla nadawcy niewygodny. Jak wcześniej wyjaśnialiśmy, pozytywny wynik weryfikacji przeprowadzanych przez odbiorcę jest dowodem na autentyczność nadawcy. Ta własność kryptografii z kluczami publicznymi nosi nazwę **niezaprzeczalności** (ang. *nonrepudiation*).

W całym tym schemacie istnieje jednak bardzo poważne niedopatrzenie. Otóż odbiorca komunikatu nadal nie ma pewności, że nadawca *posługujący się* kluczem prywatnym podmiotu X, czyli kluczem komplementarnym do opublikowanego klucza publicznego podmiotu X, w rzeczywistości *jest* podmiotem X. Każdy może przecież opublikować swój klucz publiczny jako klucz organizacji XYZ i potem bez przeszkód odczytywać korespondencję kierowaną przez różnych



RYSUNEK 11.17. Bezpieczne przesyłanie danych uwierzytelniających za pomocą kryptografii z kluczami publicznymi

nadawców do firmy XYZ, odszyfrowując ją swym kluczem prywatnym. Ergo — kryptografia z kluczami publicznymi byłaby w skali globalnej jedynie niewiele znaczącą ciekawostką, gdyby jej wykorzystywanie nie było uregulowane w sposób formalno-prawny.

Taka właśnie jest rola **infrastruktury klucza publicznego** (PKI). Infrastruktura ta to system sprzętu, oprogramowania, organizacji i polityki, zaprojektowany w celu możliwości praktycznego wykorzystywania kryptografii z kluczami publicznymi w internecie. Na szczycie hierarchii tego systemu znajdują się **urzędy certyfikacyjne** (ang. CA — *Certificate Authorities*) — są to organizacje zaufania publicznego, które mogą zaświadczyć o tożsamości osoby lub organizacji dokonującej uwierzytelniania. Jedną z takich organizacji jest amerykańska firma Verisign. Osoba chcąc uzyskać certyfikat autentyczności swojego klucza publicznego rejestruje się w CA, przedstawiając jednocześnie jakieś dowody swojej tożsamości. Weryfikacja tożsamości może przebiegać na różnych poziomach, od prostego stwierdzenia autentyczności adresu e-mail do wnikliwego policyjnego śledztwa, sprawdzającego prawdziwość informacji przekazanych przez zainteresowanego w czasie osobistego wywiadu. Po pomyślnym zweryfikowaniu tożsamości zainteresowanego CA wydaje mu certyfikat, który pod względem matematycznym jest wynikiem zaszyfrowania klucza publicznego zainteresowanego przez klucz prywatny CA. Taki certyfikat dołączany jest przez nadawcę komunikatu do informacji uwierzytelniających lub publikowany na jego

stronie WWW. Odbiorca, chcąc upewnić się co do tożsamości nadawcy, a konkretnie — prawdziwości jego klucza publicznego, porównuje ów klucz z kluczem, na który wystawiony został certyfikat (sam certyfikat zaszyfrowany jest kluczem prywatnym CA, więc odbiorca musi go najpierw rozszyfrować kluczem publicznym CA). Odbiorca powinien następnie skontaktować się z CA w celu potwierdzenia, że certyfikat nadawcy nie został wcześniej unieważniony.

Dla wiadomości i certyfikatów o szczególnym znaczeniu CA wymaga, by posiadacz certyfikatu potwierdzał jego ważność przy każdej wysyłanej wiadomości. Nadawca wiadomości przesyła ją więc do CA, który — po sprawdzeniu ważności certyfikatu — sporządza tzw. **odcisk palca** (ang. *fingerprint*), stanowiący kombinację klucza prywatnego CA oraz kluczowych informacji (lub wyciągu) oryginalnego komunikatu. Nadawca wysyła następnie komunikat do odbiorcy (w sposób pokazany na rysunku 11.17), dołączając wspomniany odcisk palca. Odbiorca, wykorzystując klucz publiczny CA, weryfikuje następnie autentyczność odcisku palca — autentyczność ta jest dowodem na to, że certyfikat nadawcy był ważny w momencie wysyłania przez niego wiadomości.

11.4.4.3. Oprogramowanie kryptograficzne

PGP (ang. *Pretty Good Privacy*, dosł. „całkiem niezła prywatność”) to darmowa aplikacja, stworzona w 1991 roku przez Philipa Zimmermanna i rozwijana przez programistów na całym świecie; w 1994 roku ukazała się jej pierwsza polska wersja. Funkcjonuje na bazie algorytmu RSA i najczęściej wykorzystywana jest do szyfrowania poczty elektronicznej, łącznie z cyfrowym podpisywaniem wiadomości. Użytkownik chcący otrzymywać szyfrowane wiadomości publikuje swój klucz publiczny na stronie WWW lub przesyła go e-mailem do potencjalnych nadawców. Faktyczny nadawca może łatwo przekopiować klucz do oprogramowania PGP.

Pod względem praw własności PGP przechodziło różne koleje losu, obecnie dostępne jest w wersjach darmowej i komercyjnej. Od kwietnia 2010 roku właścicielem praw autorskich do programu jest firma Symantec.

SSL/TLS (ang. *Secure Sockets Layer* — warstwa bezpiecznych gniazd; ang. *Transport Layer Security* — bezpieczeństwo warstwy transportowej) to protokół szyfrowania powszechnie wykorzystywany w sieci Web, opracowany w 1994 roku przez firmę Netscape; w ciągu roku pojawiły się jego trzy wersje. W 1996 roku IETF powołała grupę roboczą Transport Layer Security, której zadaniem było rozwijanie protokołu SSL. Kolejna wersja protokołu ukazała się w 1999 roku, pod zmienioną nazwą TLS 1.0, choć chronologicznie był to protokół SSL 3.1. Protokół TLS funkcjonalnie umiejscowiony jest na styku warstwy aplikacyjnej i transportowej w pięciowarstwowym modelu internetowym, w siedmiowarstwowym modelu OSI zlokalizowany jest w warstwie prezentacji. Komunikaty wychodzące z warstwy aplikacyjnej są szyfrowane przez TLS przed osiągnięciem warstwy transportowej, podobnie komunikaty przechodzące z warstwy transportowej wzwyż hierarchii są deszyfrowane przed osiągnięciem warstwy aplikacyjnej.

Konwersację w ramach TLS rozpoczyna klient, wysyłając do serwera identyfikator sesji oraz propozycję algorytmów szyfrowania (zazwyczaj RC4, DES, 3DES i AES) i algorytmów kompresji do wyboru. Serwer odpowiada komunikatem zawierającym konkretny wybór, po czym wysyła klientowi swój certyfikat i klucz publiczny. W następnym kroku klient wysyła serwerowi wstępny klucz sesji, zaszyfrowany za pomocą klucza publicznego serwera, po czym klient i serwer ustalają wspólny klucz sesji, używany do szyfrowania symetrycznego. Od tego momentu klient i serwer mogą przysyłać między sobą szyfrowane komunikaty.

IPSec (ang. *IP Security* — bezpieczeństwo protokołu IP) to inny szeroko stosowany protokół szyfrowania. W przeciwieństwie do protokołu TLS, używanego głównie dla aplikacji webowych, IPSec może być używany przez znacznie większą liczbę protokołów warstwy aplikacyjnej. Funkcjonalnie IPSec plasuje się pomiędzy warstwą sieciową a protokołami TCP i UDP w warstwie transportowej. Może wykorzystywać różnorodne techniki szyfrowania, więc pierwszym krokiem w komunikacji między nadawcą a odbiorcą jest ustalenie konkretnej techniki i klucza — zadanie to spełnia protokół IKE (ang. *Internet Key Exchange* — internetowa wymiana kluczy). Każda ze stron generuje pseudolosowy klucz i przesyła go drugiej stronie za pomocą bezpiecznej komunikacji, uwierzytelnionej w ramach PKI, po czym obie strony obliczają ten sam klucz współdzielony. Między stronami jest także uzgadniany algorytm szyfrowania, najczęściej jest nim 3DES. Po dokonaniu tych uzgodnień można już rozpoczynać transmisję.

IPSec może działać w dwóch trybach. W **trybie transportowym** nagłówek pakietu IP pozostaje niezmieniony, natomiast zaszyfrowany zostaje jego ładunek użyteczny, czyli na przykład enkapsulowany w nim segment TCP. Jednocześnie między nagłówek pakietu IP a zaszyfrowany ładunek użyteczny wstawiany jest dodatkowy nagłówek — AH (ang. *Authentication Header* — nagłówek uwierzytelniający) lub ESP (ang. *Encapsulating Security Payload*, dosł. „zabezpieczająca enkapsulacja ładunku użytecznego”), zawierający niezbędne informacje dla odbiorcy, dotyczące szyfrowania.

W **trybie tunelowym** zaszyfrowaniu podlega cały pakiet IP, wynik szyfrowania poprzedzony zostaje nagłówkiem ESP, a całość zostaje enkapsulowana w pakiecie warstwy transportowej (opcjonalnie)⁸ i pakiecie IP, w którym docelowy adres IP identyfikuje agenta deszyfracji, nie miejsce przeznaczenia pierwotnego pakietu IP. Agent deszyfracji odtwarza oryginalny pakiet IP i przekazuje go do miejsca przeznaczenia. Tak właśnie funkcjonuje tunelowanie w ramach wirtualnej sieci prywatnej (patrz sekcja 9.4.3), agentem deszyfracji jest brama VPN. Ten tryb jest bezpieczniejszy od trybu transportowego, bo wewnątrz tunelu pakiet jest całkowicie ukryty przed potencjalnymi atakami.

11.4.4.4 Polityczne aspekty kryptografii

Obecnie rząd USA traktuje szyfrowanie jako rodzaj broni, więc eksport wykorzystujących je technologii podlega tym samym regulacjom prawnym, co eksport karabinów maszynowych i bomb. W związku z tym algorytmy szyfrowania oparte na kluczach dłuższych niż 64 bity objęte były przez długie lata embargiem eksportowym. Embargo to zniesione zostało prawie całkowicie na początku roku 2000 i utrzymane jedynie w stosunku do krajów uważanych za „wspierające terroryzm”. Amerykańskie embargo mogło być skuteczne wyłącznie do produktów amerykańskich, tymczasem poza kontynentem amerykańskim powstawało — i wciąż powstaje — znacznie lepsze oprogramowanie kryptograficzne, a jego twórcy wolni byli od konkurencji amerykańskich, bo tym postawiono tamę na granicach USA. W związku z tym lobbyści z kręgów amerykańskiego przemysłu IT zaczęli wywierać presję na swój rząd, by ten złagodził restrykcje.

⁸ Zobacz przypis na ten temat w sekcji 9.4.3 — *przyp. tłum.*

11.4.5. Uwierzytelnianie użytkowników

Gdy granice sieci i jej wnętrze zostaną zabezpieczone, kolejnym zabiegiem związanym z bezpieczeństwem jest zapewnienie, że dostęp do sieci i jej poszczególnych zasobów będą mieć wyłącznie upoważnieni do tego (autoryzowani) użytkownicy. Cel ten realizuje się w postaci uwierzytelniania użytkowników (ang. *users authentication*).

Podstawą uwierzytelniania użytkownika jest jego **profil** (ang. *profile*), przypisany do jego **konta** (ang. *account*) przez menedżera sieci. Profil ten określa szczegółowo, do jakich danych i zasobów ma dostęp użytkownik i jakiego rodzaju operacje (odczytywanie, modyfikowanie, tworzenie, usuwanie) wolno mu wykonywać na poszczególnych zasobach. W profilu mogą być także zdefiniowane inne szczegóły uprawnień użytkownika, między innymi ograniczenie logowania się użytkownika do określonych dni tygodnia lub określonych pór dnia, maksymalna liczba nieudanych prób logowania, po przekroczeniu której następuje zablokowanie konta, czy też limit czasu bezczynności, po przekroczeniu którego użytkownik zostaje automatycznie wylogowany (bo na przykład poszedł na obiad, zapomniawszy się uprzednio wylogować). Menedżer może też zablokować konto użytkownika, gdy zaobserwowana zostanie jego aktywność wzbudzająca podejrzenie naruszenia zasad bądź też gdy wyczerpie się jego budżet uprawniający do korzystania z konta.

Utworzenie konta i profilu dla nowego pracownika nie jest czynnością skomplikowaną, prawdziwym problemem jest natomiast dezaktywacja kont i profili pracowników odchodzących z firmy, bowiem menedżerowie sieci często nie są na bieżąco informowani o takich przypadkach. Przegląd kont w Uniwersytecie Georgia ujawnił, że 30% czynnych kont to pozostałość po byłych pracownikach. Takie „wiszące” konta są o tyle niebezpieczne, że mogą byłemu już pracownikowi posłużyć do wyłudzenia osobistych korzyści lub wręcz dokonania ataku na zasoby firmy, zwłaszcza gdy rozstanie pracownika z nią nastąpiło nie po jego myśli. Wiele systemów pozwala menedżerom sieci definiować okres ważności konta, co w pewnym stopniu automatycznie zmniejsza opisane ryzyko, niemniej jednak menedżerowie powinni być na bieżąco informowani o konieczności usuwania niepotrzebnych już kont.

Ogólnie rzecz biorąc, kryteria Twojego dostępu do chronionego zasobu bazować mogą na trzech kategoriach:

- wiedzy — czymś, o czym *wiesz* Ty, a nie wiedzą inni;
- posiadaniu — *dysponujesz* pewnym unikalnym przedmiotem („obiektem”), do którego nikt inny nie ma dostępu;
- byciu — identyfikują Cię niepowtarzalne cechy anatomiczne, odróżniające Cię od innych osób i mierzalne w sposób powtarzalny, nazywane cechami biometrycznymi. Potwierdzają one, że *jesteś* tym, za kogo się podajesz.

DLA

11.5. Łamanie haseł

INŻYNIERA

Na wstępie należy zaznaczyć, że łamanie haseł niekoniecznie musi mieć oblicze przestępcze. Zdarza się, że legalny użytkownik zapomina swoje hasło, bo loguje się do systemu automatycznie i gdy przychodzi konieczność zalogowania się *explicite*, staje przed barierą nie do pokonania. Próby łamania haseł mogą być ponadto weryfikacją ich siły, czyli trudności w przełamaniu.

Wszelkie metody łamania haseł w Windows i Linuksie opierają się na fakcie, że hasła te przechowywane są w bazach systemów w postaci *zahaszowanej*, czyli w postaci skrótów stanowiących wyniki ich przekształcenia przez funkcje mieszające. Z natury funkcji mieszającej wynika, że nieskończenie wiele różnych tekstów przekształcanych jest na ten sam skrót, jeżeli chcemy więc złamać hasło reprezentowane w postaci danego skrótu, musimy odnaleźć *dowolny* tekst przekształcany przez tę funkcję na ten właśnie skrót. Jednocześnie fundament każdej solidnej funkcji mieszającej jest nieodwracalność skrótu, czyli niemożność wykonania wspomnianej operacji w sposób bezpośredni. Próba taka musi mieć więc postać mniej lub bardziej inteligentnej formy ataku *brute force*.

Jednym z narzędzi do przeprowadzenia takiego ataku jest program L0phtCrack, utworzony w roku 1997 i wciąż rozwijany. Czas oczekiwania na rezultat może wynosić od kilku minut do kilku tygodni.

Znacznie bardziej efektywnym sposobem łamania (odzyskiwania) haseł jest wykorzystywanie tzw. **tęczowych tablic** (ang. *rainbow tables*). Program o nazwie Ophrack, autorstwa Philippe Oechslina z Politechniki w Lozannie, zdolny jest do efektywnego łamania haseł dla Windows, Linuksa i MacOS. Wspomniany wcześniej L0phtCrack także wykorzystuje tęczowe tablice, lecz Ophrack robi to znacznie bardziej inteligentnie, bardziej efektywnie wykorzystując pamięć RAM i szybki procesor, dzięki czemu (podobno) oczekiwanie na rezultat skraca się dziesięciokrotnie, a nawet stukrotnie. Dobry artykuł, wyjaśniający koncepcję tęczowych tablic i zawierający przykłady wykorzystania programu Ophrack, znajduje się pod adresem <https://nfsec.pl/hakin9/bruteforce.pdf>.

Skutecznym utrudnieniem dla metod opartych na dopasowywaniu skrótów haseł jest dodawanie do haseł tzw. **soli** (ang. *salt*). Jest to losowo wygenerowana liczba, która kombinowana jest z oryginalnym hasłem i dopiero wynik tej kombinacji przekształcany jest przez funkcję skrótu. Sól zapisywana jest w jawnej postaci obok skrótu i wykorzystywana w procesie weryfikacji hasła wprowadzanego przez użytkownika. Czyni to bezużytecznymi wszelkie „gotowce” w postaci baz skrótów dla wyrazów pobieranych ze słownika, jednak dla Ophracka nie stanowi poważnego utrudnienia.

Wniosek z powyższego jest oczywisty — hasła *nie są* absolutnym zabezpieczeniem systemu i wymagają uzupełnienia o inne metody, by można uznać chroniony przez nie system za rzeczywiście bezpieczny.

DLA

11.7. Definiowanie silnych haseł

MENEDŻERA

Hasło jest kluczem do konta użytkownika — o tym wszyscy wiemy. Większość użytkowników bardziej jednak obawia się zapomnienia własnego hasła niż jego złamania przez intruza, bo ta pierwsza ewentualność wydaje im się bardziej prawdopodobna — wszak nasza pamięć czasem płata figle, a perspektywa złamania hasła wydaje się czysto teoretyczna każdemu, kto nie doświadczył jego skutków na własnej skórze.

Wielu użytkowników dąży więc do ustanawiania haseł łatwych do zapamiętania, a niektórzy wręcz traktują konieczność zdefiniowania hasła jako wymysł producentów systemów i podchodzą do tego obowiązku dość niedbale. Z raportów hakerów, którym udało się bez trudu włamywać do źle zabezpieczonych systemów, wynika, że użytkownicy są sami sobie są winni, używając w roli haseł:

- słowo „hasło” lub „password”,
- wzorce wynikające z układu klawiatury, na przykład „qwerty” lub „12345678”,
- nazwy ulubionych rozrywek, na przykład „football”, „baseball”, „jazz” czy „poker”,
- imiona bohaterów filmowych — „gandalf”, „leja”, „batman”
(długo można by wyliczać).

Inne często spotykane hasła związane są z osobistymi atrybutami użytkowników: datami urodzin, numerami telefonów, imionami krewnych, partnerów lub ulubionych zwierząt, nazwami ulubionych miejsc itp. Zdecydowanie niezalecaną praktyką jest też używanie identycznych haseł do wielu różnych serwisów.

Im hasło silniejsze, tym trudniejsze do złamania, ale co w tym kontekście jest miarą siły? Wiele organizacji stawia określone wymagania i ograniczenia w kwestii haseł chroniących dostęp do zasobów, usług i witryn sieci Web, na przykład ustalając minimalną długość hasła. Właściwość ta wydaje się najważniejszym czynnikiem utrudniającym jego złamanie, ale dodatkowym utrudnieniem, często obligatoryjnym, jest użycie w hasle przynajmniej jednego wystąpienia znaku określonej kategorii — wielkiej litery, cyfry i znaku specjalnego. Utworzone w ten sposób hasło o długości 20 lub więcej znaków może być zarówno silne, jak i łatwe do zapamiętania.

Dodatkowymi wskazówkami, ułatwiającymi tworzenie silnych haseł, mogą być następujące:

- używaj *fraz hasłowych*, nie *wyrazów* ze słownika, wybierz trzy-cztery znane wyrazy uzupełnione o cyfry i znaki specjalne;
- unikaj znanych sentencji („to be or not to be”);
- nie używaj tych samych fraz do różnych witryn, zamiast tego wymyśl generalną frazę, przystosowywaną do specyfiki poszczególnych stron; takim przystosowaniem może być dołączenie do frazy generalnej przyrostka w postaci liczby wystąpień litery „a” w URL-u chronionej witryny;
- zawsze używaj osobnych haseł dla zasobów szczególnie wrażliwych, na przykład kont bankowości internetowej czy profili zaufanych.

11.4.5.1. Hasła

Podstawowym kryterium uwierzytelniania opartym na wiedzy są hasła — logujący się użytkownik musi wykazać się ich znajomością. Hasła bywają jednak wybierane niefortunnie, a przez to są łatwe do odgadnięcia przez intruzów. Niektóre organizacje wymagają haseł złożonych, czyli kilku następujących po sobie fraz hasłowych, rozdzielonych ciągami spacji. W ramce „Dla menedżera 11.7” znajduje się kilka wskazówek i sugestii dotyczących definiowania haseł silnych i jednocześnie łatwych do zapamiętania.

Obecnie coraz większą popularność zdobywają **menedżery haseł** (ang. *password managers*), na przykład Dashlane czy LastPass. Menedżer haseł to aplikacja (lub dodatek do przeglądarki) odpowiedzialna za bezpieczne przechowywanie haseł do odwiedzanych stron w sieci Web. Każdorazowo, gdy definiowane jest hasło dla nowej witryny, menedżer haseł oferuje jego zapamiętanie. Gdy witryna ta zostanie odwiedzona i wyświetlony zostanie dialog logowania, obok pola do wprowadzenia hasła pojawi się niewielki przycisk, którego kliknięcie spowoduje automatyczne wypełnienie pól loginu i hasła danymi dostarczonymi przez narzędzie. Menedżer przechowuje powierzone mu hasła na swojej prywatnej stronie, nie na jakimś konkretnym urządzeniu, choć można skonfigurować dane urządzenie tak, by menedżer automatycznie przeprowadzał logowanie z jego poziomu, pomijając dialog z użytkownikiem, i tym samym oszczędzał mu konieczności logowania się za każdym razem. I gdyby skradziono Ci Twoje urządzenie, na którym korzystasz z menedżera haseł, możesz — używając innego urządzenia — zalogować się na stronie menedżera i zmienić główne hasło dostępu dla skradzionego urządzenia. Gdy złodziej będzie chciał użyć skradzionego urządzenia, menedżer haseł każe mu podać Twój login i nowe hasło główne. Oczywiście rozwiązanie to ma sens jedynie w odniesieniu do Twoich *osobistych* urządzeń, nie do tych, z których korzystasz wspólnie z innymi użytkownikami.

11.4.5.2. Uwierzytelnianie dwuczynnikowe

System haseł daje raczej średnie zabezpieczenie; podobnie jak zamknięcie na klucz drzwi wejściowych do domu zniechęci amatora, lecz nie powstrzyma profesjonalnego włamywacza. Mimo to wiele organizacji ogranicza swoją ochronę wyłącznie do haseł, lecz mniej więcej jedna trzecia wymaga od uwierzytelniającego się użytkownika — oprócz podania hasła — także wykazania się posiadaniem pewnego obiektu. Takie uwierzytelnianie nazywamy **dwuczynnikowym** lub **dwuetapowym** (ang. *two-factor authentication*) — pierwszym czynnikiem uwierzytelnienia jest bowiem hasło, a drugim wspomniany obiekt. Z dwuczynnikowym uwierzytelnieniem mamy do czynienia podczas wypłaty z bankomatu: pierwszym czynnikiem uwierzytelnienia jest kod PIN, drugim fizyczny obiekt w postaci karty.

Coraz częściej w roli drugiego czynnika uwierzytelnienia występuje określony sposób użycia telefonu komórkowego lub smartfona. Użytkownik instaluje na swym telefonie aplikację mobilną (na przykład *Duo*) i integruje ją ze swym kontem bankowym. Gdy użytkownik loguje się do konta bankowego, oprogramowanie bankowe wysła alert do aplikacji mobilnej, za pomocą której może potwierdzić logowanie lub nakazać je odrzucić; przykład użycia *Duo* w tej roli przedstawiamy na rysunku 11.18. Rozwiązanie takie zwiększa bezpieczeństwo konta, bo intruz, mimo znajomości hasła, nie będzie mógł się zalogować, nie posiadając wspomnianego telefonu z aplikacją.



RYСУNEK 11.18. Uwierzytelnianie dwuczynnikowe za pomocą aplikacji mobilnej Duo

Inną odmianą uwierzytelniania dwuczynnikowego są **hasła jednorazowe** (ang. *one-time passwords*). Gdy użytkownik zaloguje się w tradycyjny sposób i jego hasło zostanie zaakceptowane, system generuje drugie hasło (często nazywane „kodem dostępu” — *access code*), przesyła je użytkownikowi za pomocą e-maila lub SMS-a, zależnie od profilu logowania związanego z kontem, i nakazuje jego wprowadzenie; w przeciwnym razie proces logowania nie zostanie ukończony. Wykradzione hasło główne nie wystarczy intruzowi do zalogowania, bo nie będzie on miał dostępu do hasła jednorazowego.

Uwierzytelnianie dwuczynnikowe jest więc dobrym zabezpieczeniem przed intruzami, łatwiej bowiem niepostrzeżenie wykraść użytkownikowi hasło lub PIN (czyli niematerialną informację) niż fizyczny obiekt (smartfon lub kartę bankomatową).

11.4.5.3. Uwierzytelnianie biometryczne

W przypadku szczególnie krytycznych aplikacji, uwierzytelnienie użytkownika następuje w rezultacie udowodnienia, że *jest on właśnie tym, a nie innym osobnikiem*, czyli posiada określoną postać niepowtarzalnych cech anatomicznych, przeważnie linii papilarnych lub siatkówki oka, stwierdzonych przez urządzenia skanujące **systemu biometrycznego**. Uwierzytelnianie biometryczne stosowane jest w ok. 15% organizacji na całym świecie. Mimo iż systemy biometryczne zaprojektowane zostały dla zastosowań, w których bezpieczeństwo jest czynnikiem na miarę „być albo nie być”, na rynku dostępne są ich tanie implementacje — przecież wiele laptopów i urządzeń mobilnych posiada wbudowanie czytniki linii papilarnych.

11.4.5.4. Uwierzytelnianie scentralizowane

Jednym z problemów związanych z uwierzytelnianiem użytkowników jest fakt posiadania przez danego użytkownika kont na wielu serwerach. Jest to nieco uciążliwe dla użytkownika, zmuszonego do zapamiętywania licznych loginów i haseł, jest to równocześnie prawdziwy koszmar dla menedżera sieci, zarządzającego kontami wszystkich użytkowników na każdym serwerze.

Rozwiązanie tej kłopotliwej sytuacji stanowi, stosowane przez coraz więcej firm, **uwierzytelnianie scentralizowane** (ang. *central authentication*), zwane także **uwierzytelnianiem sieciowym** (ang. *network authentication*) lub **pojedynczym logowaniem** (ang. *single sign-on*). Jego istotą jest logowanie się użytkownika nie do konkretnego serwera usługowego czy serwera plików, lecz do centralnego **serwera uwierzytelniającego** (ang. *authentication server*). Serwer uwierzytelniający sprawdza w swojej bazie podane przez użytkownika login i hasło i — po pomyślnej weryfikacji — wystawia użytkownikowi **certyfikat**, zwany także **potwierdzeniem tożsamości** (ang. *credentials*). Gdy użytkownik chce uzyskać dostęp do konkretnego serwera usługowego, nie musi podawać ponownie identyfikatora i hasła, bo ich rolę spełnia teraz wspomniany certyfikat; ewentualne „logowanie” ograniczone jest do podania hasła chroniącego dany serwer usługowy. Komunikacja między klientem a serwerem uwierzytelniającym oraz serwerami usługowymi intensywnie wykorzystuje szyfrowanie, zarówno na potrzeby ukrycia informacji, jak i wzajemnego uwierzytelniania klienta i serwerów usługowych. W ten sposób wykluczony jest zarówno dostęp do usługi nieuprawnionego do niej użytkownika, jak i dostęp uprawnionego użytkownika do niewłaściwego (na przykład spreparowanego) serwera usługowego.

W ramce „Dla menedżera 11.6” opisujemy szczegóły najpopularniejszej implementacji tego mechanizmu, noszącej nazwę Kerberos, opracowanej w Massachusetts Institute of Technology (MIT). Mimo iż w wielu systemach istnieje tylko jeden serwer uwierzytelniający, możliwe jest uruchomienie kilku takich serwerów w różnych częściach organizacji. Każdy serwer uwierzytelnia użytkowników ze swojej domeny, lecz przekazuje certyfikaty uwierzytelnienia także serwerom w innych domenach.

11.4.6. Obrona przed socjotechniką

Ogniwiem w łańcuchu bezpieczeństwa systemu komputerowego są pracujący w nim ludzie — niestety, często ogniwiem najsłabszym. To zadziwiające, jak często, zamiast wyszukiwać luki w zabezpieczeniach czy realizować czasochłonne ataki *brute force*, wystarczy — zwyczajnie poprosić. Ot, zwykły telefon do nic niepodejrzewającego użytkownika, że właśnie wydarzyła się awaria i w związku z tym ja — starszy menedżer, technik do zabezpieczenia itp. — potrzebuję Twojego hasła, żeby Twoje konto znów było aktywne. Użytkownicy chcą być pomocni i chętnie przekazują informację, o którą prosi przełożony lub kolega, i nawet jeżeli w pierwszej chwili może wydawać się śmieszne, że ktoś skłonny jest przekazywać wrażliwe informacje zupełnie nieznaney osobie, to wszystko jest kwestią motywacji, a właściwie manipulacji. Zespół technik, chwytów, podstępów itd. służący osiągnięciu określonych zachowań jednostki — lub nawet całych społeczeństw — nazywany jest **socjotechniką** lub **inżynierią społeczną** (ang. *social engineering*). Wyspecjalizowany w tej dziedzinie haker, jako bezceremonialny naciągacz, może — jak pokazują przykre doświadczenia — uzyskać każdą potrzebną informację, wszystko jest kwestią metody.

Specjaliści od bezpieczeństwa nie mają złudzeń: socjotechnika zbiera swoje żniwo, każda organizacja narażona jest na udane ataki tej kategorii i w każdej organizacji hakerzy próbują wyszukiwać i oszukiwać potencjalne ofiary, najczęściej nic niepodważających szeregowych pracowników. Odpowiednie szkolenia, których celem jest uwrażliwienie pracowników na opisywane zjawisko i przekonanie do absolutnej dyskrecji w kwestii ważnych informacji, nie zlikwidują co prawda problemu, ale mogą sprawić, że wobec uporu i nieustępliwości potencjalnych ofiar w jednej instytucji, haker postanowi poszukać sobie celów w innej.

Najpopularniejszą bodaj praktyką socjotechniczną jest podszywanie się nadawcy pod kogoś innego, zwane z angielska **phishingiem**. Nieprzypadkowe jest tu fonetyczne podobieństwo do słowa *fishing*, oznaczającego połów ryb (w postaci wędkarstwa lub rybołówstwa), bo phishing jest niczym innym jak próbą złowienia dużej liczby ofiar na podstępą przynętę. Haker może na przykład wysłać do milionów adresatów e-maile z informacją, że oto konto Pana/Pani musiało zostać czasowo zamknięte, by nie paść ofiarą masowego ataku, który — na szczęście — udało się powstrzymać; do ponownego aktywowania konta potrzebne są login i hasło, które uprzejmie prosimy podać po wejściu na stronę *www.cośtam.cośtam*. I wśród tych milionów adresatów na pewno znajdują się łatwowierni, którzy polecenie to wykonają — a po kilku sekundach haker, logując się do ich kont bankowych, pozbawi ich wszystkich środków, przelewając je na swoje konto. W sprytniejszym wariacie tego podstępu adresat może wyczytać w treści e-maila, że do jego konta w PayPal-u został dodany nowy użytkownik, w związku z czym urząd skarbowy przyznaje dodatkową ulgę podatkową, tylko konieczna jest weryfikacja numeru ubezpieczenia społecznego. Albo że bank XXX oferuje kredyt hipoteczny z niewiarygodnie niskim oprocentowaniem, którego przyznanie nastąpi natychmiast po podaniu numeru ubezpieczenia społecznego i numeru karty kredytowej.

DLA

11.6. Wewnątrz Kerberos

INŻYNIERA

Kerberos, najbardziej popularny protokół scentralizowanego uwierzytelniania, wykorzystywany jest przez wiele usług wymagających uwierzytelniania, między innymi przez *Active Directory* w Windows. Do szyfrowania informacji używany jest algorytm symetryczny, najczęściej DES.

Gdy użytkownik loguje się do systemu opartego na Kerberosie, wpisuje na swoim komputerze swój identyfikator i hasło. Oprogramowanie klienckie Kerberos (to działające w komputerze użytkownika) przesyła wspomniany identyfikator (ale nie hasło!) do serwera uwierzytelniającego, który przekazuje go do **centrum dystrybucji kluczy** (ang. KDC — *Key Distribution Center*).

KDC przeszukuje bazę użytkowników na obecność identyfikatora użytkownika, po czym wykonuje dwie rzeczy. Po pierwsze, generuje **bilet usługi** (ans. ST — *Service Ticket*), zawierający znacznik czasowy oraz (co ważniejsze) unikatowy **klucz sesji** (SK1), wykorzystywany w przyszłej komunikacji oprogramowania klienckiego z KDC, aż do wylogowania się użytkownika; unikatowość SK1 gwarantowana jest jego zależnością zarówno od identyfikatora użytkownika, jak i wskazania czasu w momencie rozpoczęcia jego generowania.

Teraz następuje najciekawszy element scenariusza: ST zostaje zaszyfrowany przy użyciu klucza zależnego od hasła użytkownika zapisanego w bazie przy jego identyfikatorze. Odszyfrowanie ST wykonalne jest więc tylko na komputerze, na którym istnieje informacja o tymże hasle; gdy użytkownik wprowadzi nieprawidłowe hasło, Kerberos zażąda ponownego logowania. Zauważmy, że hasło użytkownika *nigdy nie jest przesyłane przez sieć*.

Po drugie, KDC tworzy **paszport** (ang. TGT — *Ticket-Granting Ticket*), który jest unikatowym identyfikatorem, utworzonym przez zaszyfrowanie informacji o komputerze klienta wraz ze znacznikiem czasowym, przy użyciu sekretnego klucza (oznaczanego TGS — *Ticket Granting Service*), znanego tylko KDC i innym upoważnionym serwerom. TGT zostaje zaszyfrowany przy użyciu klucza SK1 i wynik tego szyfrowania przesyłany jest do oprogramowania klienckiego (a więc jego odtworzenie nie jest możliwe bez znajomości hasła użytkownika, od którego przecież zależny jest SK1). Oprogramowanie klienckie odszyfrowuje TGT, nie jest jednak w stanie zidentyfikować zawartych w nim informacji, bo nie zna klucza TGS. Od tego momentu użytkownik, aż do wylogowania, nie musi logować się ponownie, ponieważ dla każdej usługi wymagającej hasła oprogramowanie klienckie wysyła TGT w roli danych uwierzytelniających.

Gdy użytkownik po raz pierwszy próbuje uzyskać dostęp do serwera usługowego chronionego hasłem, serwer ten instruuje oprogramowanie klienckie, by pobrało z KDC związany z nim bilet usługi ST. Oprogramowanie to wysyła więc do KDC paszport (TGT) wraz z informacją o tym, do którego serwera chce uzyskać dostęp (jak pamiętamy, wszelka komunikacja między tym oprogramowaniem a KDC szyfrowana jest przy użyciu klucza SK1). KDC upewnia się, że użytkownik wciąż jest zalogowany, a następnie, po weryfikacji TGT, przesyła oprogramowaniu klienckiemu stosowny bilet ST oraz nowy klucz sesji (SK2), który w połączeniu z kluczem SK1 będzie używany do szyfrowania komunikacji między oprogramowaniem klienckim a serwerem usługowym. ST zawiera informacje uwierzytelniające oraz klucz SK2, w postaci zaszyfrowanej tajnym kluczem, znanym tylko KDC i serwerowi usługowemu.

Oprogramowanie klienckie przekazuje serwerowi usługowemu żądanie zalogowania (zawierające identyfikator użytkownika, znacznik czasowy i kilka innych informacji) w postaci zaszyfrowanej kluczem SK2, oraz bilet ST. Serwer usługowy odszyfrowuje ST (używając tajnego klucza znanego tylko jemu i KDC) i odczytuje dane uwierzytelniające oraz klucz SK2. Klucz SK2 jest następnie używany do deszyfracji żądania zalogowania. Po jego zweryfikowaniu serwer usługowy przesyła oprogramowaniu klienckiemu swoją wizytówkę, zaszyfrowaną kluczem SK2. W ten oto sposób klient i serwer zostają nawzajem uwierzytelnione wobec siebie i dlatego mogą bezpiecznie komunikować się ze sobą, za pomocą transmisji szyfrowanej przy użyciu połączonych kluczy SK2 i SK1.

Zauważmy na koniec, że hasło użytkownika pozostaje cały czas nieznane dla serwera usługowego.

11.4.7. Systemy zapobiegania włamaniom

Systemy zapobiegania włamaniom (ang. IPSs — *Intrusion Prevention Systems*) zaprojektowane zostały w celu wykrywania prób włamywania się do systemów i unieszkodliwiania ich. Istnieją dwa typy systemów IPS, większość menedżerów sieci instaluje oba.

Działanie **IPS opartego na sieci** (ang. *Network-Based IPS*) opiera się na **czujnikach IPS** (ang. *IPS sensor*), instalowanych w kluczowych obwodach sieci. Czujnik IPS to urządzenie sterowane specjalnym systemem operacyjnym, monitorującym wszystkie pakiety przepływające przez obwód i wysyłającym raporty o podejrzeniach włamań na **konsolę zarządzania IPS** (ang. *IPS management console*).

Drugi typ IPS to **IPS oparty na hoście** (ang. *Host-Based IPS*), fizycznie będący pakietem oprogramowania działającego w stacji roboczej lub na serwerze, monitorującego pakiety przychodzące i raportującego podejrzone przypadki na konsolę zarządzania IPS.

Oba typy IPS działają w oparciu o dwie podstawowe techniki detekcji włamań. Pierwszą z nich jest **wykrywanie nadużyć** (ang. *misuse detection*), polegająca na porównywaniu wyników monitorowania aktywności z sygnaturami znanych ataków. Gdy rozpoznana zostaje któraś z tych sygnatur, IPS wysyła alert i odrzuca podejrzone pakiety. Warunkiem powodzenia tej techniki jest utrzymywanie bazy wspomnianych sygnatur w aktualnym stanie, czyli jej aktualizowaniu w przypadku stwierdzenia nowych rodzajów ataku.

DLA

11.8. Socjotechnika znowu górą...

MENEDŻERA

Danny rozpoczął od odrobienia lekcji. Wkrótce zebrał wystarczająco dużo informacji, aby móc wcielić się w pracownika firmy. Znał jego nazwisko, wydział, numer telefonu i numer pracownika, a także nazwisko i numer telefonu jego szefa.

Nastąpiła cisza przed burzą. Dosłownie. Zgodnie z obmyślonym planem, Danny potrzebował teraz jeszcze jednej rzeczy, zanim wykona następny krok, i było to coś, nad czym nie miał kontroli: potrzebował burzy śnieżnej. Czekał na tak złą pogodę, która uniemożliwi pracownikom dojazd do pracy.

W czasie zimy w Południowej Dakocie — a tam właśnie miała siedzibę rzeczona firma — każdy, kto miał nadzieję na złą pogodę, nie musiał czekać zbyt długo. W piątkową noc nadeszła burza. Śnieg szybko przeszedł w marznący deszcz i do rana drogi zdążyły się zamienić w lodowiska.

Dla Danny'ego była to idealna okazja. Zadzwoił do firmy i poprosił o połączenie z jednym z informatyków. Podając nazwisko istniejącego pracownika, na temat którego zrobił wcześniej wywiad, powiedział:

— Tu Bob Billings. Pracuję dla Secure Communications Group. Jestem teraz w domu i nie mogę dojechać z powodu burzy. Problem polega na tym, że muszę dostać się z domu do mojego konta na serwerze, a zostawiłem token na biurku. Czy mógłby pan po niego pójść? Albo kogoś wysłać? A potem odczytać mój kod, kiedy będę chciał wejść? Nasz zespół dostał pilny termin i nie będę mógł skończyć mojej pracy. Nie mogę się dostać do biura, bo drogi są teraz zbyt niebezpieczne.

— Nie mogę wyjść z mojego biura — powiedział informatyk.

Danny zadziałał szybko:

— A ma pan może swój identyfikator?

— W centrum komputerowym jest jeden — stwierdził rozmówca — dla operatorów, do użytku w razie nagłych przypadków.

— Mam prośbę — powiedział Danny. — Wyświadczyłby mi pan przysługę? Mógłbym skorzystać z pańskiego identyfikatora, kiedy będę wchodził na konto? Do czasu, aż pogoda się poprawi.

— Mogę jeszcze raz prosić pana nazwisko? — zapytał informatyk.

— Bob Billings.

— Dla kogo pan pracuje?

— Dla Eda Trentona.

— A, tak. Znam go.

— Pracuję na drugim piętrze — ciągnął Danny. — Obok Roya Tuckera. Łatwiej byłoby po prostu pójść do mojego pokoju i przynieść mój identyfikator. Jest w górnej lewej szufladzie.

Danny był w miarę pewny, że jego rozmówca nie da się na to namówić. Po pierwsze, nie opuściłby swojego stanowiska w środku zmiany, by włączyć się gdzieś po odległych korytarzach budynku. Poza tym nie miał ochoty grzebać w czyimś biurku. Można się było założyć, że tego nie zrobi.

A informatyk nie chciał powiedzieć „nie” człowiekowi, który potrzebuje pomocy, ale nie miał też zamiaru powiedzieć „tak”. Dlatego postanowił przerzucić decyzję na kogoś innego:

— Moment, zapytam szefa.

Położył słuchawkę na biurku i Danny słyszał, jak podnosi drugą, łączy się i wyjaśnia sprawę. W pewnym momencie zrobił coś dziwnego: poświadczył za dzwoniącego, używając jego domniemanego nazwiska — Bob Billings.

— Znam go — powiedział szefowi. — Pracuje dla Eda Trentona. Możemy mu udostępnić identyfikator z centrum komputerowego?

Danny, trzymając słuchawkę, był zaskoczony tą niezwykłą i niespodziewaną formą pomocy, jakiej mu udzielono. Nie wierzył własnym uszom. Po paru kolejnych chwilach Kowalski wrócił do telefonu i powiedział:

— Mój szef chce z panem sam porozmawiać — po czym podał nazwisko i numer komórki szefa.

Danny zadzwonił do niego i opowiedział wszystko jeszcze raz, dodając parę szczegółów o projekcie, nad którym pracował, i o tym, dlaczego jego grupa musi koniecznie dotrzymać terminu.

— Prościej by było, gdyby ktoś po prostu poszedł i przyniósł mój identyfikator — powiedział. — Biurko nie powinno być zamknięte, a karta będzie chyba w górnej lewej szufladzie.

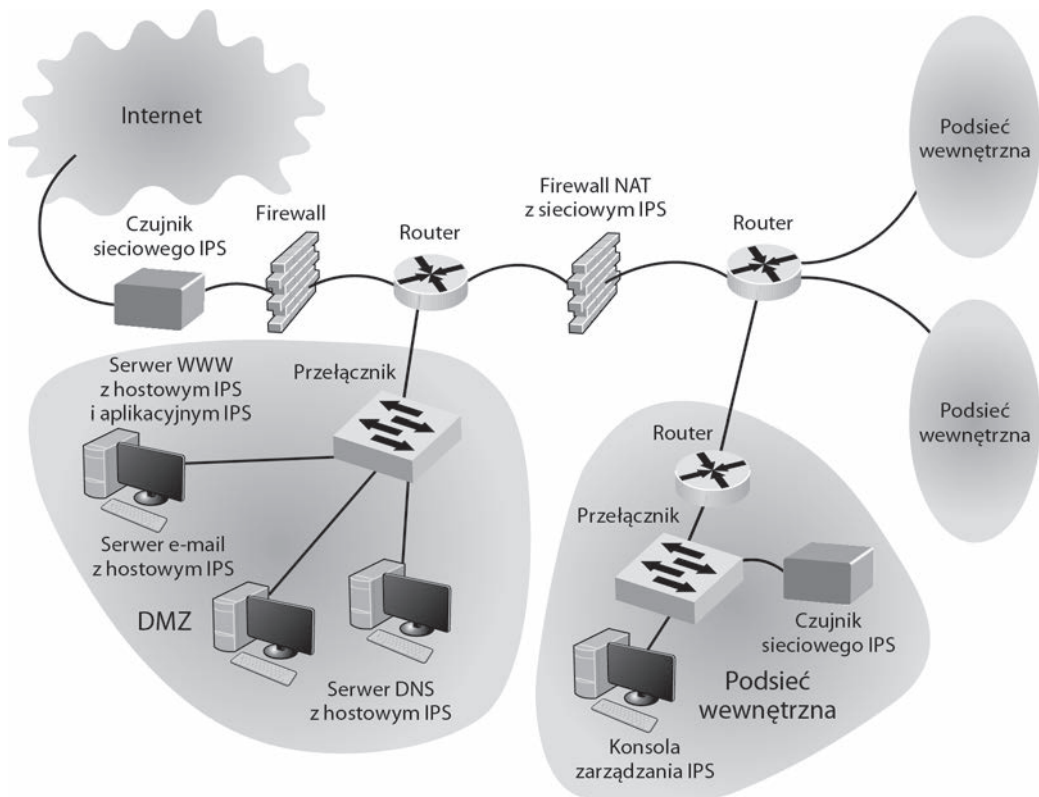
— Myślę, że tylko na weekend możemy pozwolić panu korzystać z identyfikatora awaryjnego. Powiem ludziom, którzy mają wtedy zmiany, żeby odczytywali kod, gdy będzie pan dzwonił — powiedział szef, po czym podał kod PIN, jakiego ma używać wraz z identyfikatorem.

Przez cały weekend, za każdym razem, gdy Danny chciał dostać się do systemu komputerowego firmy, musiał jedynie zadzwonić do centrum komputerowego i poprosić o odczytanie sześciu cyfr wyświetlanych w okienku identyfikatora.

Na podstawie: Kevin Mitnick, William L. Simon, Steve Wozniak, *Sztuka podstępu. Łamałem ludzi, nie hasła*. Wydanie II, Helion 2010.

Drugą ze wspomnianych technik jest **wykrywanie anomalii** (ang. *anomaly detection*), polegające na ciągłym konfrontowaniu monitorowanych aktywności ze wzorcami „normalnych” zachowań. Gdy pojawi się wyraźne odstępstwo od tych wzorców, na przykład nagły wzrost częstotliwości pakietów pingujących ICMP lub nienaturalnie duża częstotliwość nieudanych logowań na konto administratora lub menedżera, IPS wysyła alert i odrzuca podejrzone pakiety. Problemem związanym z tą techniką są fałszywe alerty, kwestionujące anomalne sytuacje spowodowane legalnymi przyczynami, na przykład nagłym wzrostem natężenia ruchu wskutek nerwowych nastrojów na Wall Street.

Systemy IPS funkcjonują najczęściej w połączeniu z innymi zabezpieczeniami, na przykład z firewallami, co pokazaliśmy na rysunku 11.19 — faktycznie, niektóre nowsze firewalle mają wbudowaną funkcjonalność IPS. Problem w tym, że czujniki IPS i konsole zarządzania IPS, jako narzędzia defensywy przeciwwłamaniowej, same stanowią częsty cel hakerów, dlatego muszą być przed tymi atakami szczególnie dobrze zabezpieczone. By utrudnić hakerom dokonywanie wspomnianych ataków, wiele organizacji instaluje redundantną strukturę IPS, na przykład instalując IPS na poziomie sieci od jednego dostawcy i IPS na poziomie hostów od innego.



RYСУNEK 11.19. System zapobiegania włamaniom (IPS)

DMZ = *DeMilitarized Zone* (strefa zdemilitaryzowana)

DNS = *Domain Name Service* (usługa nazw domenowych)

NAT = *Network Address Translation* (translacja adresów sieciowych)

Choć monitorowanie wykonywane przez systemy IPS jest istotne dla bezpieczeństwa sieci, to nie na wiele się ono zda, jeśli organizacja nie będzie dysponowała precyzyjnym atakiem odpowiedzi, czyli reakcji na trwającą właśnie próbę ataku. Większość dużych organizacji utrzymuje w tym celu zespoły SWAT (ang. *Special Weapons And Tactics* — specjalne wyposażenie i taktyka), które w takiej sytuacji natychmiast rozpoczynają działania obronne. Wzorcowym tego przykładem jest CERT, posiadająca takie zespoły na wypadek zagrożeń internetowych i pomagająca innym organizacjom w tworzeniu takowych.

Reagowanie na wykrytą próbę włamania jest trudniejsze niż mogłoby się zrazu wydawać. Załóżmy na przykład, że IPS wykrył próbę ataku DoS wychodzącego z pewnego adresu IP. Rytowną reakcją w takim przypadku powinno być odrzucanie wszystkich pakietów przychodzących spod tego adresu — i rzeczywiście byłoby ono rozsądne pod warunkiem, że źródłowy adres IP w tych pakietach *byłby autentyczny*. Hakerzy, spodziewając się takiej reakcji, złośliwie podmieniają źródłowe adresy IP w wysyłanych pakietach na adresy IP serwerów *znanych klientów* firmy, w rezultacie czego autentyczne pakiety, napływające ze strony tych klientów, są odrzucane. A klienci z zaskoczeniem konstatują, że ich sprawdzony dostawca nagle blokuje im próbę kontaktu...

11.4.8. Odtwarzanie po włamaniu

Pierwszymi krokami po wykryciu włamania powinny być zidentyfikowanie intruza i uniemożliwienie innym intruzom włamania w taki sam sposób. Większość firm ogranicza się do takiego właśnie zamknięcia drzwi dla intruzów, inne (ok. 30%) idą o krok dalej, śledząc aktywność intruza i zbierając dowody jego przestępczej działalności, po czym angażują policję, doprowadzając do jego aresztowania, wreszcie wytaczają mu proces cywilny (niezależnie od czekającej go odpowiedzialności karnej). Prawodawstwo niektórych stanów i prowincji zobowiązuje wręcz firmy do raportowania włamań i kradzieży danych klientów, więc wymieniony odsetek jest na ich obszarze wyraźnie wyższy.

Ponieważ różnorodne przestępstwa dokonywane z udziałem komputerów i internetu stały się już trwałym elementem świata IT, a także dlatego, że każdej akcji towarzyszy reakcja, incydentalne niegdyś metody obrony przed takimi przestępstwami przekształciły się z biegiem czasu w dyscyplinę naukową, zwaną **informatyką śledczą** (ang. *computer forensics*), zajmującą się technikami namierzania przestępców i gromadzenia materiału dowodowego, stanowiącego podstawę do wytaczania procesów karnych oraz ewentualnych cywilnych procesów odszkodowawczych. Podstawowe założenia informatyki śledczej są zasadniczo takie same jak w tradycyjnej kryminologii, jednakże wykorzystywane techniki śledcze są diametralnie różne. Po pierwsze — zidentyfikowanie potencjalnych dowodów. Po drugie — zabezpieczenie materiału dowodowego w postaci kopii i wykorzystanie ich do analiz. Po trzecie — sama analiza. Wynikiem analizy jest raport śledczy, przekazywany prokuratorowi.

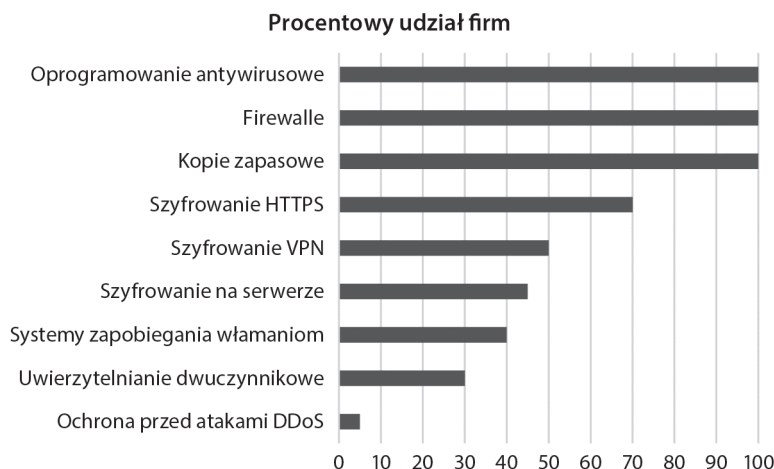
Należy dodać, że zaatakowane firmy często podejmują działania odwetowe przeciwko hackerom, same wcielając się w ich rolę. Może brzmieć to interesująco, jednak jest nielegalne.

Niektóre organizacje wyręczają w pewnym stopniu techników śledczych i we własnym zakresie zastawiają pułapki na intruzów. Pułapka ma zwykle postać serwera naszpikowanego niezwykle atrakcyjnymi informacjami — tyle że fikcyjnymi. Z jednej strony działa on jak przynęta, odwracająca uwagę hakera od prawdziwej sieci, z drugiej (właśnie) stanowi pułapkę: haker delektuje się obfitością wykradanych informacji, nie przypuszczając, że jego aktywność jest

szczegółowo monitorowana i zbierana w charakterze materiału dowodowego. To, że zawartość serwera nie przedstawia dla firmy żadnej wartości. Nie zmienia to faktu, że włamanie do tego serwera jest przestępstwem. Takie serwery-przynęty nazywane są popularnie *honeypotami*: **honeypot** — „garnek z miodem” to nazwa nawiązująca do przygód sympatycznego niedźwiadka Kubusia Puchatka, który wsadził łypek do garnka z miodem i nie potrafił go wyjąć, łakomstwo zwabiło go w pułapkę.

11.5. ZALECENIA PRAKTYCZNE

W tym rozdziale przedstawiamy liczne sugestie dotyczące zapewnienia ciągłości funkcjonowania biznesu oraz zapobiegania włamaniom do sieci. Solidna strategia zabezpieczeń rozpoczyna się od czytelnego planu odtwarzania po katastrofie i solidnej polityki bezpieczeństwa. Najlepszą inwestycją w bezpieczeństwo są szkolenia użytkowników na temat ochrony i odzyskiwania danych oraz czujności wobec podstępów socjotechnicznych. Nie oznacza to jednak, że można zapomnieć o technologicznej stronie zabezpieczeń. Na rysunku 11.20 przedstawiliśmy ranking popularności poszczególnych technologii bezpieczeństwa; w najbliższej przyszłości można się spodziewać, że uwierzytelnianie dwuczynnikowe i szyfrowanie zyskiwać będą na znaczeniu.



RYСУNEK 11.20. Powszechnie wykorzystywane kontrole bezpieczeństwa

Rzadko zdarza się tydzień, w którym nie zostaje wykryta luka w zabezpieczeniach, średnio w ciągu tygodnia pojawia się pięć nowych łat i innych aktualizacji.

Pojawiają się więc zewsząd pytania: czy to się nigdy nie skończy? Czy (nie)bezpieczeństwo stanie się trwałym elementem krajobrazu systemów informatycznych? W pewnym sensie — tak. Systemy informatyczne wciąż zyskują na znaczeniu, stając się nośnikami nowych możliwości z jednej strony i obiektem coraz większych zagrożeń z drugiej, a to stwarza nowe pole do popisu wszelkiego rodzaju cyberprzestępcom. Wysoka opłacalność wykradania (na sprzedaż) cennych informacji, w połączeniu ze wciąż zbyt małą skutecznością ścigania przestępców i wynikającym z niej poczuciem bezkarności, każą spodziewać się coraz większej liczby ataków w najbliższej przyszłości.

Stąd kolejne pytanie: czy musi być tak źle? Bez wątpienia, potrafimy się coraz lepiej zabezpieczać, coraz skuteczniej egzekwować politykę bezpieczeństwa i chronić dostęp do cennych zasobów. To jednak kosztuje i staje się coraz trudniejsze. Hakerzy tworzą i rozpowszechniają nowej generacji narzędzia do przypuszczania ataków, coraz lepsze, coraz trudniejsze do wykrywania i łatwe w obsłudze — łatwiejsze niż defensywne narzędzia do przeciwdziałania im.

Atakujący mają jeszcze jedną przewagę. Żeby ochronić sieć, należy załatać w niej *wszystkie luki na dłuższy czas*; żeby haker mógł dokonać ataku, wystarczy mu *jedna luka chwilowo niezaląatana*.

Dzień z życia menedżera bezpieczeństwa sieci

„Zarządzanie bezpieczeństwem sieci to kombinacja przenikliwości detektywa z umiejętnością jasnowidzenia”.

Codzienna praca menedżera bezpieczeństwa koncentruje się wokół trzech aktywności. Pierwsza z nich dotyczy tego, co dzieje się na zewnątrz firmy, czyli informacji o nowo wykrytych lukach zabezpieczeń i atakach, bo tym ostatnim sprzyjają nowe technologie i nowe sposobności. Chcąc chronić własne sieci, trzeba być na bieżąco z najnowszymi wirusami, trojanami i innym złośliwym oprogramowaniem, a także z narzędziami do tworzenia tychże. Obfitą porcję wiedzy na ten temat można znaleźć na dedykowanych stronach w sieci Web, w pierwszym rzędzie w witrynach CERT (www.cert.org) i SANS (www.sans.org). Na podstawie tej wiedzy kreowane są nowe wzorce standardowych (fabrycznych) ustawień komputerów, w założeniu bardziej odpornych na zewnętrzne ataki, a informacje o nowych zagrożeniach stają się podstawą rekomendacji w zakresie instalowania nowych uaktualnień i łat do znanych aplikacji. Dla menedżera stanowią one impuls do zweryfikowania polityki bezpieczeństwa w jego firmie i poinformowania użytkowników o ewentualnych zmianach.

Druga domena działalności menedżera bezpieczeństwa to nieustanne przyglądanie się własnej sieci — troskliwym okiem opiekuna, lecz także ze sprytem wyimaginowanego hakera, poszukującego słabych punktów i okazji do ataku, czyli w praktyce skanującego otwarte porty, poszukującego nieaktualizowanych komputerów i słabo chronionych miejsc w sieci. Analizie podatności na ataki towarzyszy obserwowanie sieci w celu wykrywania symptomów trwających już być może ataków, a także analizowanie zachowania poszczególnych komputerów, mogącego świadczyć o ich skompromitowaniu — takimi zachowaniami mogą być nienaturalne wzorce aktywności lub uruchamianie nieznanych usług.

Trzeci rodzaj aktywności menedżera bezpieczeństwa to czynna defensywa, czyli szybkie reagowanie na incydenty naruszenia bezpieczeństwa: identyfikowanie ich źródeł i przyczyn, kolekcjonowanie materiału dowodowego na użytek przyszłego śledztwa i procesu, no i oczywiście przywracanie zaatakowanych komputerów i aplikacji do normalnego stanu.

Źródło: dzięki uprzejmości Kenna Crooka

Czego więc możemy się spodziewać w przyszłym „bezpiecznym” środowisku firmy? Przede wszystkim ograniczeń, między innymi w postaci scentralizowanego **zarządzania pulpitami** (ang. *desktop management*), w ramach którego użytkownikom ograniczać się będzie możliwości operowania ustawieniami ich komputerów, wymuszać instalowanie na bieżąco dostępnych aktualizacji. Ponadto, w celu definitywnego pozbycia się ukrywających się (ewentualnie) wirusów i trojanów, co jakiś czas przeprowadzany będzie *reimaging* komputerów, czyli przywracanie zawartości dysków do stanu początkowego (fabrycznego), co oczywiście powodować będzie konieczność instalowania wszystkich aplikacji i aktualizacji od nowa, ale jednocześnie dawać będzie komputerom nowe życie, wolne od brzemienia zaszłości. Niekontrolowane pobieranie zewnętrznego oprogramowania będzie prawdopodobnie całkowicie zabronione. Upowszechniać się także będzie model aplikacji „cienkiego klienta” (patrz sekcja 2.2.3), jako skupiającego zabiegi konfiguracyjne głównie na serwerach.

Bezpieczeństwo ma swoją cenę, także w aspekcie wydajności. Ciągłe monitorowanie przychodzących i wychodzących pakietów WWW i e-mail zauważalnie spowalnia pracę sieci z perspektywy jej użytkowników. Szyfrowanie komunikacji między klientami i serwerami oraz szyfrowanie samej zawartości serwerów stanowi dodatkowy czynnik spowalniający. Rygorystycznie egzekwowana polityka bezpieczeństwa, z dotkliwymi konsekwencjami w razie jej naruszenia, to dodatkowe utrudnienia dla użytkowników. Znakomicie sprawdza się w tym przypadku znana zasada „wolniej jedziesz, dalej dojedziesz”; bezpieczeństwo ma swoją cenę, ale jest ona w pełni akceptowalna.

I tylko czasami spoglądamy z nostalgią wstecz, wspominając Złote Dni początków internetu, gdy beztrudno mogliśmy się zachwycać jego atrakcjami.

11.6. IMPLIKACJE DLA TWOJEGO CYBERBEZPIECZEŃSTWA

Cyberbezpieczeństwo było niegdyś gałęzią ezoteryki, kultywowaną przez niewielkich profesjonalistów. Brutalne koleje losu spowodowały jej zstąpienie z wyżyn na twardy grunt systemów informatycznych, w postaci wciąż rozwijającej się dyscypliny naukowej.

Specjalistyczna dziedzina wiedzy wymaga wykwalifikowanych specjalistów: według prognoz ekspertów, w 2020 roku ponad milion miejsc pracy w sektorze bezpieczeństwa systemów komputerowych będzie nieobsadzonych, ponieważ firmy nie nadążą z rekrutowaniem kandydatów stosownie do swego zapotrzebowania. Od większość z nich oczekiwać się będzie ponadstandardowych kwalifikacji i wieloletniego doświadczenia, choć w przypadku jednej trzeciej wystarczające będą podstawowe umiejętności. Dla Ciebie, jako potencjalnego przyszłego pracownika, jest to zdecydowanie dobra wiadomość.

Zdecydowanie złą wiadomością dla Ciebie jako użytkownika komputera jest konieczność skoncentrowania się na bezpieczeństwie w stopniu większym niż kiedykolwiek wcześniej, i to nie tylko na polu zawodowym, lecz także w życiu osobistym, którego istotne aspekty powierzylesz przecież komputerom, internetowi i urządzeniom mobilnym. W związku z tym mamy dla Ciebie kilka praktycznych porad, które powinny ułatwić Ci życie w tym względzie:

- Zasubskrybuj usługę regularnego archiwizowania zawartości Twojego komputera, by w razie jego awarii móc ją bezpiecznie odzyskać. W internecie znajdziesz oferty wielu firm, które profesjonalnie świadczą taką usługę.
- Upewnij się, że w ustawieniach systemu operacyjnego i oprogramowania antywirusowego na Twoim komputerze włączone są opcje automatycznej aktualizacji.
- Upewnij się, że Twoje urządzenia podłączone do internetu — smartfony i inteligentne telewizory smart TV — również są na bieżąco aktualizowane.
- Właściwie dobieraj hasła chroniące Twoje zasoby.
- Jeśli to możliwe, w serwisach, z których korzystasz, wybierz uwierzytelnianie dwuczynnikowe.
- Zainstaluj menedżer haseł i używaj go.
- Ochroń dostęp do swojego komputera, laptopa i smartfonu za pomocą hasła, PIN-u lub czytnika linii papilarnych. Statystyczny użytkownik urządzeń mobilnych traci średnio dwa-trzy urządzenia w ciągu swojej aktywności zawodowej, wskutek zgubienia lub kradzieży.
- Zainstaluj na swoim urządzeniu mobilnym oprogramowanie umożliwiające, w razie kradzieży lub zgubienia, usunięcie jego zawartości i unieruchomienie w sposób zdalny. Pamiętaj, że koszty użycia przez nieuprawnioną osobę danych zgromadzonych w Twoim urządzeniu mogą znacznie przekraczać jego cenę.
- *Nigdy* nie klikaj odsyłaczy znajdujących się w wiadomościach e-mail. 90% takich wiadomości to oszustwa phishingowe, więc każdy e-mail zawierający odsyłacz może być próbą oszustwa. Jeśli chcesz sprawdzić dany odsyłacz, wyszukaj go za pomocą Google.
- *Nigdy* nie klikaj wyskakujących okien z ostrzeżeniami, zwłaszcza tymi alarmującymi (na przykład o zawirusowaniu komputera); kliknięcie elementu, niezależnie od prezentowanych przez niego treści, może doprowadzić do zainstalowania wirusa lub ransomware przez stronę WWW. Najbezpieczniej w takiej sytuacji zamknąć przeglądarkę (lub „zabić” ją za pomocą menedżera zadań) czy nawet uruchomić ponownie komputer.
- *Nigdy* nie ujawniaj nikomu swoich haseł, czy to w trakcie rozmowy, czy też za pośrednictwem e-maili. Osoby, który pytają Cię o hasła, to w 99% hakerzy. Gdy ktoś telefonuje do Ciebie z pytaniem o hasło, natychmiast zakończ rozmowę. Gdy wpisujesz hasło, kod dostępu lub inną poufną informację, upewnij się uprzednio, że nikt nie podgląda ekranu Twojego komputera czy wyświetlacza urządzenia.

PODSUMOWANIE

Typy zagrożeń bezpieczeństwa. Ogólnie zagrożenia dla bezpieczeństwa sieci podzielić można na dwie kategorie: zagrożenia dla ciągłości funkcjonowania oraz włamania, czyli nieuprawniony dostęp. Ciągłość funkcjonowania może zostać zakłócona krótkoterminowo, lecz konsekwencją niektórych zakłóceń może być zniszczenie danych.

Katastrofy — te naturalne i te powodowane przez człowieka — mogą dokonywać fizycznych zniszczeń sprzętu na wielką skalę. Włamania, dokonywane z zewnątrz przez hakerów lub z wnętrza sieci przez pracowników, mają na celu nieuprawnione uzyskiwanie dostępu do cennych danych. Włamania wykorzystywane są do wykradania cennych informacji, do dokonywania oszustw w celu osiągnięcia osobistych korzyści oraz do niszczenia firmowych zasobów.

Ocena ryzyka. Bezpieczna sieć to taka, w której zaimplementowano system kontroli redukujących lub eliminujących zagrożenia. Zadaniem tego systemu jest wykrywanie i powstrzymywanie zagrożeń oraz niwelowanie ich skutków, gdy się urzeczywistnią. Pierwszym krokiem tworzenia bezpiecznej sieci jest przeprowadzenie oceny ryzyka dla bezpieczeństwa jej zasobów. Podstawą takiej oceny jest zidentyfikowanie cennych zasobów i uszeregowanie ich według znaczenia dla firmy, a następnie sporządzenie wykazu (również uszeregowanego według priorytetów) zagrożeń czyhających na poszczególne zasoby. Zagrożenia te konfrontowane są następnie z efektywnością kontroli zaimplementowanych w celu przeciwdziałania im. Dostępnych jest kilka frameworków umożliwiających przeprowadzanie takiej oceny według profesjonalnych, sformalizowanych zasad.

Ciągłość funkcjonowania. Podstawowymi zagrożeniami dla ciągłości funkcjonowania firmy są wirusy, kradzieże danych, ataki DoS/DDoS, awarie urządzeń i katastrofy. Zainstalowanie i bieżące aktualizowanie narzędzi antywirusowych to najważniejsza i najczęściej stosowana kontrola bezpieczeństwa. Zapobieganie atakom DoS/DDoS jest bardziej poważnym wyzwaniem i wymaga wsparcia sprzętowego. Wykradanie cennych, poufnych informacji jest zagrożeniem często niedocenianym, można je minimalizować przede wszystkim przez odpowiednie fizyczne zabezpieczenie sprzętu, szczególnie laptopów. Każde urządzenie prędzej czy później się zepsuje, więc projektanci sieci powinni przewidywać odpowiednią rezerwę sprzętową (przede wszystkim routerów i przełączników) na takie przypadki w newralgicznych miejscach sieci (na przykład na połączeniu z internetem i siecią szkieletową). Ponieważ katastrofy są w większości nieprzewidywalne i niewiele można zrobić w kierunku zapobiegania im, dlatego dla szybkiego przywrócenia sprawności po zaistnieniu takiej czy innej katastrofy konieczne jest opracowanie i symulacyjne przetestowanie planu wychodzenia z katastrofy oraz regularne sporządzanie kopii zapasowych ważnych danych i oprogramowania, przechowywanych w odległej lokalizacji.

Zapobieganie włamaniom. Intruzami mogą być zarówno zewnątrzni hakerzy, jak i pracownicy firmy działający na jej szkodę, a celem włamania kradzież cennych danych (na przykład numerów kart kredytowych klientów) bądź złośliwe niszczenie danych. Polityka bezpieczeństwa definiuje prawa i obowiązki poszczególnych użytkowników w zakresie dostępu do sieci i jej poszczególnych zasobów. Zadaniem firewalli jest powstrzymywanie intruzów na granicach sieci

poprzez odrzucanie niedozwolonych pakietów na podstawie wiedzy o znanych metodach ataków oraz ukrywanie wewnętrznych komputerów przed internetem, przez nadawanie im prywatnych adresów IP w ramach mechanizmu NAT. Konieczne jest także fizyczne zabezpieczenie urządzeń i instalacji, sieciowych i telefonicznych, przed niepożądanym dostępem. Instalowanie bieżących aktualizacji systemów operacyjnych i aplikacji, w celu łatania rozpoznanych luk w ich bezpieczeństwie, to kolejny warunek konieczny skutecznej obrony przed atakami hakerskimi. Szyfrowanie danych — tych przesyłanych i tych magazynowanych na dyskach — zwiększa ich bezpieczeństwo, bo czyni je nieczytelnymi w przypadku ich kradzieży. Uwierzytelnianie użytkowników uniemożliwia dostęp do systemu osobom nieupoważnionym i opiera się na posiadanej wiedzy (hasłach), posiadanych obiektach (uwierzytelnianie dwuczynnikowe) i posiadanych indywidualnych cechach anatomicznych (biometria). Ochrona przed atakami socjotechnicznymi, polegającymi na podstępny skłanianiu personelu sieci do ujawniania wrażliwych informacji, jest dość trudna i sprowadza się do należytego edukowania personelu w tej dziedzinie. Systemy przeciwdziałania włamaniom opierają się na sygnaturach znanych ataków oraz na monitorowaniu ruchu sieciowego w celu wykrywania jego anomalii, i poprzez generowanie alertów pozwalają menedżerom sieci na natychmiastowe przeciwstawianie się trwającemu atakom. Reagowanie na włamanie polega na rozpoznawaniu ich specyfiki w celu zapobiegania im w przyszłości oraz monitorowaniu aktywności intruza w celu zebrania materiału dowodowego i przekazania go organom ścigania. Jeżeli efektem włamania jest zniszczenie danych zagrażające działaniu systemu, należy przywrócić tym danym oryginalną postać.

KLUCZOWE TERMINY

3DES (<i>Triple DES</i>)	<i>Brute force</i>	DoS (<i>Denial of Service</i>)	Honeypot
ACL (<i>Access Control List</i>)	CA (<i>Certificate Authority</i>)	Dostępność	Hostowy IPS
AES (<i>Advanced Encryption Standard</i>)	Całkiem niezła prywatność (PGP)	DRP (<i>Disaster Recovery Plan</i>)	IKE (<i>Internet Key Exchange</i>)
Agent DDoS	CDP (<i>Continuous Data Protection</i>)	Dzielenie ryzyka	Informatyka śledcza
Algorytm	Certyfikat	Filtrowanie ruchu	Infrastruktura kluczy publicznych (PKI)
Analiza ruchu	Ciągła ochrona danych (CDP)	Finansowe konsekwencje incydentu	Integralność
Analizator anomalii ruchu	Ciągłość funkcjonowania	Firewall	Internetowa wymiana kluczy (IKE)
Atak dnia zero	Czujnik IPS	Firewall aplikacyjny	IPS (<i>Intrusion Prevention System</i>)
Atak siłowy	DDoS (<i>Distributed Denial of Service</i>)	Firewall NAT	IPSec (<i>IP Security Protocol</i>)
Bezpieczeństwo fizyczne	DES (<i>Data Encryption Standard</i>)	Firewall pakietowy	Kerberos
Bezpieczeństwo warstwy transportowej (TLS)	Deszyfracja	Framework oceny ryzyka	Klucz
Bezpieczny przełącznik	Detektor anomalii ruchu	Funkcja mieszająca	Klucz prywatny
		Haker	Klucz publiczny
		Hasło	Konsola zarządzania
		Hasło jednorazowe	IPS

Konto użytkownika	Outsourcing	Rozproszona odmowa	UPS (<i>Uninterruptible</i>
Kontrola detekcyjna	odtworzenia	usługi (DDoS)	<i>Power Supply</i>)
Kontrola korekcyjna	po katastrofie	RSA	Urząd certyfikacji (CA)
Kontrola odtwarzania	PGP (<i>Pretty Good</i>	Scenariusz zagrożenia	Uwierzytelnianie
z kopii zapasowej	<i>Privacy</i>)	Serwer odporny	Uwierzytelnianie
Kontrola opcji	<i>Phishing</i>	na awarie	biometryczne
zapasowej	PKI (<i>Public Key</i>	Serwer	Uwierzytelnianie
Kontrola prewencyjna	<i>Infrastructure</i>)	uwierzytelniający	dwuczynnikiowe
Koń trojański	Plan odtwarzania	Sieciowy IPS	Uwierzytelnianie
Koordynator DDoS	po katastrofie (DRP)	Socjotechnika	oparte na posiadaniu
Kopia zapasowa online	Podśluchiwanie	<i>Spoofing</i> IP	Uwierzytelnianie
Kraker	transmisji	SSL (<i>Secure Sockets</i>	oparte na wiedzy
Kryptografia	Polityka	<i>Layer</i>)	Uwierzytelnianie
z kluczami	bezpieczeństwa	Strategia kontroli	scentralizowane
publicznymi	Potrójny DES (3DES)	ryzyka	Uwierzytelnianie
Krytyczne aplikacje	Poufność	Symulowane	użytkownika
Lista kontroli dostępu	Prawne konsekwencje	odtworzenie	Warstwa bezpiecznych
(ACL)	incydentu	po katastrofie	gniazd (SSL)
Luka bezpieczeństwa	Produktywność	System biometryczny	Wirus
Łata	Profil użytkownika	System zapobiegania	Wojna informacyjna
Menedżer haseł	Program-szperacz	włamaniom (IPS)	Wykrywanie anomalii
<i>Mirroring</i>	Protokół	Szyfrogram	Wykrywanie nadużyć
Narzędzia	Diffiego-Hellmana	Szyfrowanie	Zaakceptowanie ryzyka
antywirusowe	Przeciwdziałanie	Szyfrowanie	Zagrożenie
NAT (<i>Network Address</i>	ryzyku	asymetryczne	Zapobieganie
<i>Translation</i>)	Pułapka	Szyfrowanie	nieautoryzowanemu
Ocena ryzyka	RAID (<i>Redundant</i>	symetryczne	dostępowi
Ocena wpływu	<i>Array of Independent</i>	Tekst jawny	Zarządzanie kluczami
Odbicie lustrzane	<i>Disks</i>)	Tęczowe tablice	Zarządzanie pulpitem
Odłożenie ryzyka	<i>Ransomware</i>	TLS (<i>Transport</i>	Zasilacz awaryjny
Odmowa usługi (DoS)	RC4	<i>Layer Security</i>)	(UPS)
Ograniczanie ruchu	Redundancja	Translacja adresów	Zasób
Oprogramowanie	Redundantna macierz	sieciowych (NAT)	Złośliwe
reklamowe	dyskowa (RAID)	Trojan	oprogramowanie
Oprogramowanie	Reputacja firmy	Tryb transportowy	(<i>malware</i>)
szpiegujące	Robak	IPSec	
		Tryb tunelowy IPSec	

PYTANIA

1. Jakie czynniki sprawiają, że należy zwracać szczególną uwagę na bezpieczeństwo sieci?
2. Opisz główne kroki procedury oceny ryzyka.
3. Nazwij i opisz główne obszary zagrożone łamaniem zabezpieczeń. Do kogo należy ocena stopnia zagrożenia (małe, średnie albo duże) każdego z tych obszarów? Uzasadnij odpowiedź.
4. Wymień kilka kryteriów różnicowania stopnia zagrożenia.
5. Jakie są najbardziej prawdopodobne zagrożenia dla bezpieczeństwa? Które są najbardziej krytyczne i dlaczego?
6. Wyjaśnij znaczenie budowania możliwych scenariuszy urzeczywistniania zagrożeń. Jakie są poszczególne kroki takiego scenariusza?
7. Jakie znaczenie ma ocena ryzyka i jak się ją oblicza?
8. Jakie są możliwe strategie kontrolowania ryzyka? Jakie są kryteria wyboru najlepszej strategii?
9. Jakie jest znaczenie poszukiwania usprawnień w obszarach przeciwdziałania ryzyku?
10. Jakie znaczenie ma plan odtwarzania po katastrofie? Wymień pięć typowych elementów takiego planu.
11. Co to jest wirus komputerowy? Jakie zagrożenie niesie ze sobą *ransomware*?
12. Wyjaśnij istotę ataku DoS.
13. Czym klasyczna postać ataku DoS różni się od jego wersji rozproszonej (DDoS)?
14. Na czym polega outsourcing odtwarzania po katastrofie? W jakiej sytuacji i dlaczego warto zawrzeć umowę o jego świadczenie?
15. Na czym polega sporządzenie kopii zapasowej online?
16. Opisz poszczególne kategorie intruzów próbujących włamywać się do sieci.
17. Opisz trzy najważniejsze elementy typowej polityki bezpieczeństwa.
18. Jakie są podstawowe aspekty zapobiegania włamaniom?
19. Jak zabezpiecza się sieć na jej granicach?
20. Co to jest bezpieczeństwo fizyczne i dlaczego jest tak istotne?
21. Jakie jest znaczenie podsłuchiwania transmisji z punktu widzenia bezpieczeństwa sieci?
22. Co to jest program-szperacz?
23. Co to jest firewall?
24. Jak działają poszczególne typy firewalli?
25. Co to jest *spoofing* IP?
26. Jak działa firewall NAT?
27. Co to jest luka bezpieczeństwa i jak można ją zneutralizować?
28. Wyjaśnij działanie trojanów.
29. Porównaj szyfrowanie symetryczne z asymetrycznym.
30. Opisz funkcjonowanie szyfrowania i deszyfracji w wariancie symetrycznym.

31. Opisz funkcjonowanie szyfrowania i deszyfracji w wariancie asymetrycznym.
32. Na czym polega zarządzanie kluczami?
33. Czym DES różni się od 3DES? Czym różni się AES od ES i 3DES?
34. Porównaj szyfr DES i kryptografię z kluczami publicznymi.
35. Jak przebiega uwierzytelnianie użytkownika?
36. Czemu służy infrastruktura kluczy publicznych (PKI)?
37. Jakie są zadania urzędu certyfikacji (CA)?
38. Czym PGP różni się od SSL/TLS?
39. Czym SSL/TLS różni się od IPsec?
40. Jakie są trzy podstawowe rodzaje uwierzytelniania użytkowników? Jakie są zalety i niedostatki każdego z nich?
41. Jakie są odmiany uwierzytelniania dwuczynnikowego i czym różnią się od siebie?
42. W jaki sposób systemy biometryczne mogą przyczyniać się do zwiększania bezpieczeństwa i jakie są podstawowe problemy związane z ich stosowaniem?
43. Dlaczego zarządzanie profilami użytkowników jest istotnym aspektem polityki bezpieczeństwa?
44. Co to jest socjotechnika i dlaczego stanowi tak poważne zagrożenie dla bezpieczeństwa?
45. Jakie są sposoby przeciwdziałania skuteczności zabiegów socjotechnicznych?
46. Co to są systemy zapobiegania włamaniom (IPS)?
47. Czym — w systemie IPS — różni się wykrywanie nadużyć od wykrywania anomalii?
48. Co to jest informatyka śledcza?
49. Do czego służą *honeypoty*?
50. Tylko niewielu konsultantów od bezpieczeństwa twierdzi, że szybki internet i technologie bezprzewodowe są ich sprzymierzeńcami. Jak myślisz, dlaczego?
51. Większość hakerów rozpoczynała nabywanie i doskonalenie swoich kwalifikacji jako nastolatkiem. Jak Ty — jako członek społeczności profesjonalnych użytkowników komputerów — widzisz możliwości zmniejszenia zainteresowania hakerstwem wśród osób w młodym wieku?
52. Niektórzy eksperci twierdzą, że publikowanie przez CERT informacji o atakach i zagrożeniach przynosi więcej szkody niż pożytku, bo w większym stopniu ułatwia ataki i do nich zachęca, niż im zapobiega. Jakie argumenty przemawiają na rzecz tej tezy, a jakie jej zaprzeczają? Czy uważasz, że CERT powinien nadal prowadzić swoją politykę informacyjną w obecnym kształcie?
53. Jakie jest główne ryzyko — oprócz możliwej odpowiedzialności karnej i cywilnej — nieautoryzowanego pobierania kopii utworów muzycznych z internetu?
54. Mimo iż względy bezpieczeństwa przemawiają za ochroną wszystkich serwerów, to niektóre serwery wymagają tej ochrony w stopniu większym niż inne. Które z nich wymagają jej najbardziej i dlaczego?

ĆWICZENIA

- A. Przeprowadź ocenę ryzyka sieci komputerowej w Twojej organizacji (na podstawie informacji, do których masz legalny dostęp).
- B. Przedstaw przykładowy przebieg śledztwa w związku z jakimś (być może fikcyjnym) atakiem hakerskim oraz przykładowy raport w tej sprawie do CERT.
- C. Oceń spodziewane korzyści i koszty zaabonowania przez Twoją organizację usługi outsourcingu odtwarzania po katastrofie.
- D. Porównaj korzyści i koszty używania firewalli w Twojej organizacji.
- E. Porównaj korzyści i koszty używania systemu IPS w Twojej organizacji.
- F. Porównaj korzyści i koszty szyfrowania danych (przechowywanych i transmitowanych) w Twojej organizacji.
- G. Porównaj korzyści i koszty usługi automatycznego wykonywania kopii zapasowych online w Twojej organizacji.

MINIANALIZY PRZYPADKÓW

I. Belmont State Bank to ogromny bank z setkami oddziałów połączonych z centralnym systemem komputerowym; niektóre oddziały używają w tym celu dedykowanych obwodów, inne korzystają z MPLS (patrz sekcja 9.3.3). Każdy oddział posiada zbiór różnorodnych komputerów i bankomatów przyłączonych do serwera. Serwer ten przechowuje dane o bieżących transakcjach i kilka razy na dobę przesyła je do wspomnianego centralnego systemu. Kasjerzy w każdym oddziale używają 4-cyfrowych haseł numerycznych, a każdy komputer kasjera jest skonfigurowany tak, by wyłącznie akceptować transakcje autoryzowane przez jego właściciela. Oceń ryzyko zagrożenia dla bezpieczeństwa w tej konfiguracji.

II. Western Bank jest małym rodzinnym bankiem z sześcioma oddziałami w całym kraju. Zamierza on udostępnić swoim klientom witrynę internetową, za pomocą której będą oni mogli kontrolować stan swoich kont i dokonywać zdalnych transakcji. Zaprojektuj kluczowe komponenty sprzętowe i programowe, niezbędne do zapewnienia bezpieczeństwa tej witryny.

III. Classic Catalog Company, część pierwsza. Classic Catalog Company to niewielka, lecz prężnie rozwijająca się firma sprzedaży wysyłkowej artykułów oferowanych za pośrednictwem drukowanych katalogów. Witryna firmy była dotąd hostowana przez lokalnego dostawcę internetu (ISP), lecz gdy sprzedaż poprzez internet zaczęła stanowić znaczący udział w ogólnym bilansie firmy, postanowiono przenieść tę witrynę do własnego wewnętrznego systemu komputerowego, dokonując jednocześnie modernizacji jego sieci. Firma mieści się w dwóch budynkach: w jednym znajdują się biura, w drugim hurtownia. W dwukondygnacyjnym biurowcu pracuje 60 komputerów: na parterze znajduje się 40, z czego 30 przeznaczonych jest do obsługi zamówień telefonicznych. Parterowy budynek hurtowni, znajdujący się 120 metrów od biurowca, po drugiej stronie firmowego parkingu, ma powierzchnię ok. 10 000 metrów kwadratowych.

Działa w nim 15 komputerów w dziale sprzedaży skupionym w jednym pomieszczeniu. Tytułem eksperymentu firma postanowiła wyposażyć pracowników w mobilne terminale, co ma usprawnić ich lokalizowanie i finalizowanie dostaw. Bazując na prognozie natężenia ruchu w nadchodzącym roku, firma planuje połączyć swe biura z dostawcą internetu za pomocą dedykowanych obwodów T1. W biurówcu działają trzy serwery: główny serwer WWW, serwer e-mail i wewnętrzny serwer wykorzystywany przez aplikacje obsługi zamówień i płatności. Dokonaj oceny ryzyka zagrożeń dla bezpieczeństwa w opisanej konfiguracji.

IV. Classic Catalog Company, część druga. Dla konfiguracji firmy opisanej w części pierwszej przedstaw założenia planu zapewnienia ciągłości funkcjonowania biznesu, uwzględniającego kontrole zmierzające do minimalizowania ryzyka oraz scenariusze odtwarzania po ewentualnej katastrofie.

V. Classic Catalog Company, część trzecia. Dla konfiguracji firmy opisanej w części pierwszej przedstaw założenia polityki bezpieczeństwa oraz projekt systemu kontroli, zapobiegający nieautoryzowanym dostępom do sieci.

VI. Classic Catalog Company, część czwarta. Jaką politykę aktualizacji zabezpieczeń proponowałbyś dla konfiguracji firmy opisanej w części pierwszej?

VII. Menedżery haseł. Aby uwolnić użytkowników od konieczności pamiętania wielu haseł, opracowano wiele różnych menedżerów haseł. Znajdź pięć najbardziej popularnych i porównaj ich możliwości oraz ceny.

ĆWICZENIA PRAKTYCZNE 11A

Zabezpiecz swój komputer

W tym rozdziale omawialiśmy tematykę bezpieczeństwa komputerowego, jego znaczenia dla oceny ryzyka zagrożeń (i minimalizowania tegoż ryzyka), a także ciągłości funkcjonowania i zapobiegania włamaniom. Przeciętnemu użytkownikowi komputera może się wydawać, że zagadnienia te istotne są raczej dla sieci korporacyjnych niż dla niego, jako indywidualnego użytkownika. Gdy jednak staje się on ofiarą ataku hakerskiego, rychło przekonuje się, że nawet w małej, domowej sieci LAN, do której przyłączony jest tylko jeden komputer stacjonarny lub laptop, bezpieczeństwo także ma niebagatelne znaczenie.

Z perspektywy użytkownika komputera domowego „ciągłość funkcjonowania” oznacza niezakłócony dostęp do edukacji, niezakłócone wykonywanie pracy zdalnej (istotne zwłaszcza przy pracy zespołowej) oraz nieprzerwany dostęp do informacji finansowych i rachunków bankowych. Znaczenie poufności i integralności danych związanych z wymienionymi aktywnościami nie wymaga komentarza. Wobec realności zagrożeń także i indywidualny użytkownik powinien dokonać oceny ryzyka ich urzeczywistnienia. W tym miejscu omówimy więc kilka zabiegów — prostych, lecz zdolnych w znaczącym stopniu poprawić bezpieczeństwo użytkowania komputera. Ograniczymy się do komputerów z systemem Windows, jako najpowszechniej używanym, jednak w podobny sposób (choć używając innych poleceń) można zwiększyć bezpieczeństwo komputera Apple.

A więc do dzieła — nie odkładaj niczego na później.

Ciągłość funkcjonowania

Wyjaśniliśmy już, co w kontekście użytkowania domowego komputera oznacza pojęcie ciągłości funkcjonowania, w tej chwili proponujemy wyobrazić sobie sytuację, gdy nagle ulega awarii dysk komputera, a jutro lub pojutrze trzeba oddać kolejny rozdział książki lub kolejny moduł oprogramowania. Perspektywa nieciekawa, jednak w zupełności realna. Dla nieprzygotowanego użytkownika — być może tragedia. A więc:

1. Rozpocznij od włączenia opcji automatycznego pobierania i instalowania uaktualnień Windows. Zagwarantuje to Twojemu komputerowi obecność krytycznych najnowszych łat i uaktualnień.
2. Zakup i zainstaluj profesjonalny pakiet antywirusowy znanego producenta (na przykład firmy Symantec), nie zapomnij włączyć w nim opcji automatycznego uaktualniania. Po zainstalowaniu pakiet antywirusowy prawdopodobnie zaproponuje Ci przeskanowanie komputera — choć operacja ta może trwać dość długo, nie rezygnuj z niej! Nie zapominaj także o regularnym skanowaniu komputera; większość pakietów umożliwia zdefiniowanie harmonogramu skanowania, niektóre oferują skanowanie „w tle”, podczas normalnej pracy komputera.
3. Nie wszystkie pakiety antywirusowe są wrażliwe na oprogramowanie szpiegujące (ang. *spyware*), być może będziesz więc musiał takowe doinstalować; ze swojej strony polecamy Spybot. Ponownie — nie zapomnij o włączeniu automatycznej aktualizacji.
4. Głównym nośnikiem wirusów, oprogramowania szpiegującego i oprogramowania reklamowego są pliki graficzne i muzyczne, dostępne za darmo w internecie. Nie pobieraj plików o nieznanym pochodzeniu, z nieznanego źródła. Pakiet antywirusowy chroniący Twój komputer powinien co prawda reagować na zagrożenia i wykrywać podejrzaną zawartość, ale jest tylko programem i ma swoje ograniczenia.
5. Opracuj własny plan odtwarzania po awarii — już teraz powinieneś założyć, że kiedyś się wydarzy. Które pliki są dla Ciebie najważniejsze? Bez których nie wyobrażasz sobie pracy? Których nie chciałbyś stracić, bo kosztowały Cię mnóstwo pracy? Sporządzaj często (najlepiej po każdej zmianie) ich kopie na zewnętrznych nośnikach — DVD, pendrive'ach lub na dyskach innych komputerów w Twojej sieci. Niestety, nośniki optyczne i pamięci *flash* bywają zawodne, a inne komputery w sieci też mogą ulec zniszczeniu, na przykład w wyniku pożaru czy powodzi. Najlepszym zabezpieczeniem będzie kopia zapasowa w chmurze, z automatyczną synchronizacją zmienionych danych. W internecie znajdziesz firmy z całego świata, oferujące taką usługę; prześcigają się one w ofertach, wystarczy wybrać tę najlepiej dopasowaną — funkcjonalnie i cenowo — do indywidualnych potrzeb.

Sprawdzian

1. Dokonaj analizy ryzyka związanego z Twoją siecią domową.
2. Sporządź plan odtwarzania po (ewentualnej) katastrofie dla Twojej sieci domowej.
3. Wybierz pakiet ochrony przed złośliwym oprogramowaniem najbardziej — Twoim zdaniem — dopasowany do Twojej sieci.

ĆWICZENIA PRAKTYCZNE 11B

Szyfrowanie na Twoim komputerze

Szyfrowanie danych jest dobrą metodą ich ochrony. Szyfrowanie jest szeroko wykorzystywane w internecie, nawet jeśli nie jesteś tego świadom — gdy dokonujesz zakupów przez internet, Twój komputer szyfruje wrażliwe informacje, między innymi numer Twojej karty kredytowej.

Czy jednak ma sens szyfrowanie w tak małej skali, jak komputer domowy? Zdecydowanie *tak*. Okaże się nieocenione, gdy Twój komputer zostanie skradziony. Że niby Twój komputer chroniony jest hasłem i to wystarczy? Nie wystarczy, włamanie do komputera chronionego hasłem jest bardzo łatwe.

Co innego jednak, gdy hasło to wykorzystywane jest do wygenerowania klucza, przy użyciu którego szyfrujesz dane zapisane na dysku. Masz do wyboru dwie opcje: szyfrowanie plików albo całego dysku. Gdy zaszyfrujesz pliki, złodziej, który nie zna hasła, nie zobaczy ich zawartości. Będzie mógł jednak oglądać zawartość niezasyfrowanych plików, będzie mógł doinstalowywać własne pliki i programy, i używać komputera bez przeszkód. Gdy jednak zaszyfrujesz cały dysk, dla osoby nieznającej Twojego hasła będzie on bezużyteczny, nadający się co najwyżej do ponownego sformatowania. Ten kijek ma jednak dwie strony: jeżeli zapomnisz hasła bądź na dysku zdarzy się jakieś uszkodzenie, jego zawartość będzie dla Ciebie stracona na zawsze. Stąd wniosek, że — *summa summarum* — szyfrowanie pojedynczych plików jest lepszym rozwiązaniem.

Celem tego ćwiczenia jest zaprezentowanie programu TrueCrypt. Jest to darmowy program z kategorii *open source*, dostępny w wersjach dla Windows, Linuksa i MacOS. Dostępny jest do pobrania ze strony <http://www.truecrypt.org/downloads>. Po pobraniu i zainstalowaniu z domyślnymi ustawieniami TrueCrypt otworzy okno przewodnika dla początkujących, znajdujące się pod adresem <http://www.truecrypt.org/docs/tutorial> — zalecamy zapoznanie się z nim. Potem już możesz wykonać następujące kroki:

1. Uruchom TrueCrypt; w systemie Windows znajdziesz jego skrót w Menu *Start*.
2. Wybierz opcję *Create Volume*. Otworzy się okno kreatora, wybierz z niego pierwszą opcję: *Create an encrypted file container*.
3. Wybierz wolumin w formie pliku (*Volume within a File*) — TrueCrypt nazywa go kontenerem (*Container*).
4. Wybierz opcję tworzenia standardowego woluminu (*Standard TrueCrypt volume*).
5. Wskaż miejsce, w którym ma zostać utworzony wolumin — z punktu widzenia systemu plików będzie to zwykły plik, który będziesz mógł skasować jak każdy inny. Wybierz folder *Dokumenty* i nadaj plikowi nazwę *Volume1*. **Uwaga!** Nie wskazuj istniejącego pliku, bo zostanie on nadpisany, nie zaszyfrowany, stracisz jego zawartość.
6. Kliknij przycisk *Next* w oknie kreatora.
7. Wskaż metodę szyfrowania — zalecamy pozostawienie domyślnego wyboru *AES*.
8. Określ rozmiar kontenera. Sugerujemy pozostawienie domyślnej wartości 1 MB, chyba że planujesz szyfrowanie większej liczby plików.
9. Ten krok jest najważniejszy — zdefiniuj hasło szyfrowania i potwierdź je; TrueCrypt użyje go do wygenerowania klucza. Kliknij przycisk *Next*.

10. Aby nadać kluczowi bardziej losowy charakter, wykonuj chaotyczne ruchy kursorem myszy przez kilka sekund, po czym kliknij przycisk *Next*.
11. Pomyślnie utworzyłeś wolumin-kontener. Kliknij *Exit* — okno kreatora zamknie się samoczynnie.
12. Wybierz literę dysku (powiedzmy J:), pod którą ma zostać zamapowany kontener.
13. Kliknij przycisk *Select File*; w oknie wyboru plików wybierz kontener, który zapisałeś w punkcie 5. (*Volume1*).
14. W oknie programu TrueCrypt kliknij przycisk *Mount*. Pojawi się żądanie hasła — wpisz to, które zdefiniowałeś w punkcie 9. i kliknij *OK*.
15. Utworzyłeś właśnie wirtualny dysk J: (w całości zaszyfrowany) z punktu widzenia systemu plików zachowujący się jak zwykły dysk. Możesz kopiować na niego pliki — będą szyfrowane „w locie”.

Chociaż szyfrowanie nie ochroni Cię przed złośliwym oprogramowaniem ani dostępem do zaszyfrowanego kontenera, to jednak dodaje pewną warstwę bezpieczeństwa, gdy pozostawisz publicznie dostępny włączony komputer. Możesz jednak w ten sposób zaszyfrować swój folder programu DropBox lub wykorzystać zamontowany kontener w charakterze przynęty na łakomych, lecz lekkomyślnych hakerów. Powodzenia!

UWAGA!

Rozwój programu TrueCrypt został zakończony w maju 2014 roku po tym, jak Microsoft zakończył wspieranie Windows XP. Następne wersje Windows — od Vista wzwyż — oferują zintegrowaną obsługę szyfrowanych dysków i obrazów dysków wirtualnych, integracja taka dostępna jest także w Linuksie i MacOS. TrueCrypt nie jest uważany za bezpieczny w tych wersjach Windows, należy więc pliki znajdujące się w woluminach utworzonych przez niego skopiować na natywny dysk szyfrowany lub dysk wirtualny.

Sprawdzian

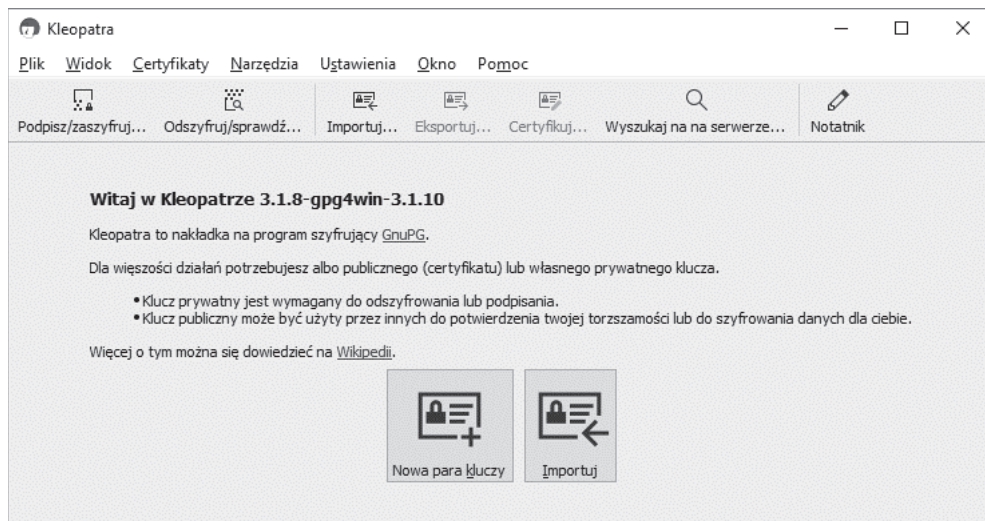
Zaszyfruj wybrany folder na swoim domowym komputerze. Zaprezentuj zrzut ekranu z wyświetlenia zawartości wynikowego woluminu.

ĆWICZENIA PRAKTYCZNE 11C

Laboratorium szyfrowania

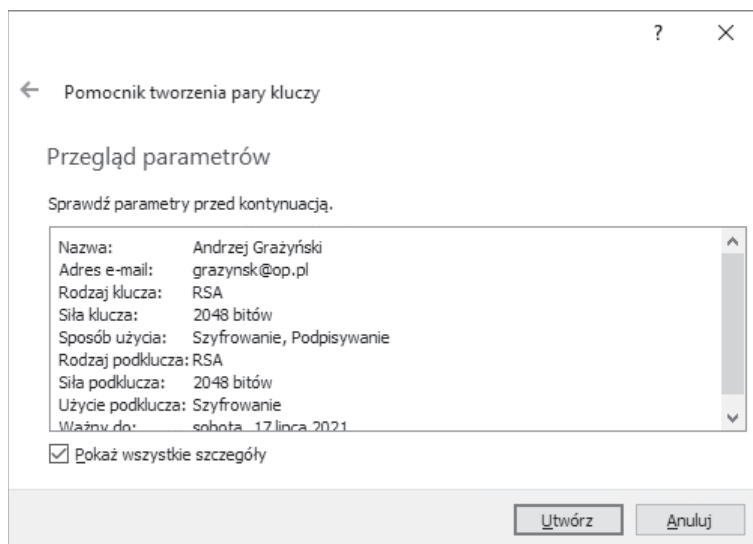
Przedmiotem niniejszego ćwiczenia jest zaszyfrowanie i deszyfracja wiadomości e-mail za pomocą standardu PGP, zaimplementowanego w postaci programu Gnu Privacy Guard (GPG), udostępnianego na licencji GPL. Do wykonania ćwiczenia potrzebny Ci będzie program Kleopatra. W wersji dla Windows jest to nakładka na program gpg4win, który możesz pobrać ze strony <https://www.gpg4win.org>. Dla użytkowników systemu MacOS program dostępny jest pod adresem <http://macgpg.sourceforge.net>.

1. Uruchom program Kleopatra. W oknie powitalnym (rysunek 11.21) wybierz opcję *Nowa para kluczy*.



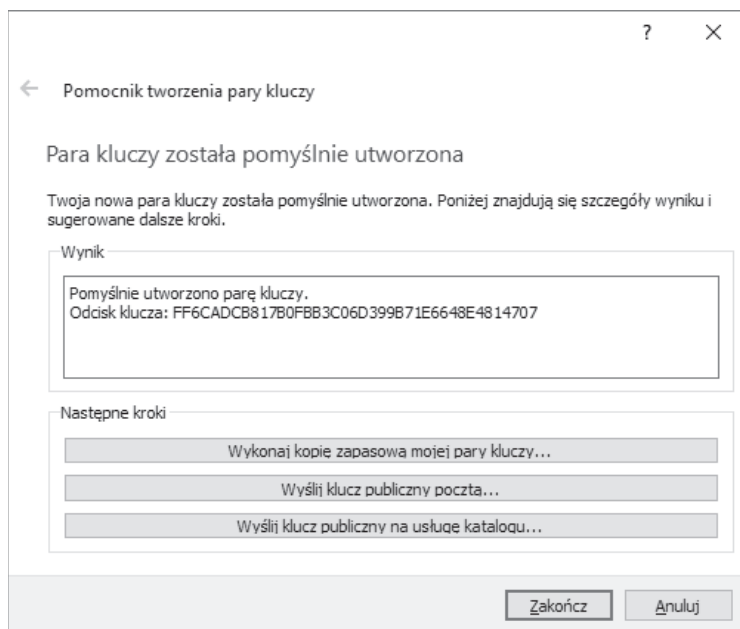
RYСУNEK 11.21. Okno powitalne programu Kleopatra

2. W wyświetlonym oknie wpisz nazwę klucza i adres e-mail. Kliknij przycisk *Dalej*.
3. Program wyświetli parametry tworzonej pary kluczy (rysunek 11.22). Zaakceptuj je, klikając przycisk *Utwórz*.



RYСУNEK 11.22. Parametry nowej pary kluczy

4. Program wyświetli żądanie zdefiniowania hasła, którym chroniona będzie nowa para kluczy. Hasło to będzie także używane w procesie szyfrowania i deszyfrowania wiadomości. Wskaźnik siły hasła powinien mieć zielony kolor i wskazywać 100%. Po wprowadzeniu hasła i kliknięciu przycisku *OK* program przystąpi do generowania kluczy. Zakończenie generowania oznajmione zostanie stosownym komunikatem (rysunek 11.23), zawierającym „odcisk palca” nowej pary kluczy oraz propozycje trzech opcjonalnych czynności. Kliknięcie przycisku *Zakończ* kończy proces generowania kluczy.



RYСУNEK 11.23. Pomyślne wygenerowanie pary kluczy

5. Kolejny krok to uczynienie klucza publicznego faktycznie publicznym, czyli oznajmienie go potencjalnym nadawcom, by mogli wysyłać do Ciebie szyfrowane wiadomości. Kliknij prawym przyciskiem myszy pozycję reprezentującą klucz w głównym oknie Kleopatry, wybierz opcję *Eksportuj* i wskaż lokalizację, w której klucz publiczny zostanie zapisany jako tekst ASCII, który można obejrzeć (na przykład) w Notatniku (rysunek 11.24). Teraz możesz już udostępnić swój klucz publiczny wszystkim osobom, od których chciałbyś otrzymywać szyfrowane wiadomości.
6. Oczywiście Kleopatra działa też w drugą stronę: możesz zaimportować klucz publiczny adresata, któremu chcesz wysłać zaszyfrowaną wiadomość (za pomocą opcji *Importuj* z paska narzędziowego), a następnie użyć tego klucza do zaszyfrowania wiadomości (za pomocą opcji *Zaszyfruj/podpisz* z paska narzędziowego). W wyświetlonym oknie (rysunek 11.25) wybierz żądane opcje i kliknij przycisk *Podpisz/zaszyfruj*. W podobny sposób można zaszyfrować wiadomość tekstową bez zapisywania jej do pliku: wybieramy opcję *Notatnik* z paska narzędziowego, wpisujemy (lub wklejamy ze schowka) tekst wiadomości, po czym wybieramy opcję *Podpisz/zaszyfruj*.

```

pubFF6CADCB817B0FBB3C06D399B71E664E4814707 — Notatnik
Plik Edycja Format Widok Pomoc
|----BEGIN PGP PUBLIC KEY BLOCK-----

mQENBF0vTy0BCADNHhPmg06uXYHPf3QaxJIcDz9/u1zrypWAHEIzohG5/cZOHgYq
iA98/frENyqGXGpEZTWybvVVEFP4UFnW+w722/Ejx9Kj1rhzyXAUHmAD3JWId06b
QWjv4bZ56PiJwH+sz5+d500h9WxNSJzh7RXYft32caAU8io9YL/b2+57TPjHbYaJ
g1AA+QLlj+i0wQeRuxFP1L11qrnopB8oTubtYZMoJVsmEdHMXJS/IUjJtzWrIAP
bv1zYKnhkhdwKywPVRTHTGqAD0ho0hbBTTu6ybfclAKQ1i63xaxsChR5H2WwOVbg
YRw0ij4/y+rqwuwlyv4x+8gg7x8pR0wqkoGBABEBAAg0DmdyYXp5bnNrQG9wLnBs
iQUBFBMBCAA+FiEE/2yty4F7D7s8Bt0Ztx5mSOSBRwcfAl0vTY0CGwMFCQPDxJMF
CwkIBwIGFQoJCAasCBBYCAwECGECFA4AACgkQtX5mSOSBRweE1Af7B18b0WIEWmY
1UtmzQPzgcQ1GhLPMpUy4sXEU2lUobgFLy2UIokryzUH5SrDpaVkyidVexfxal7
3g52eNantzxLOB284uQIkbyP5r6z8sN8DpKmDX9vkd1/WKTJwJnwKx1eGzHZ+RtW
d8skV8gEMSSeacmUA0dvS4TGzB5cRTnUK7IUt0yxJEO/HijPJuaoxm46k0DiHUG
QBPSGw4yBNAqnMHC2kA878Z5A6xkbK8uIBrGWPOx6UywrBTKiAazinClymKNXABK
jJ1iomc2utZnKvF/CbEOLSQN+M1T1dsQsxArnXS3IwKnT0Lw7c5o5K9G0DRxATS
7RTL0c8JerkBDQRdL02NAQgAwqPK1/+ryaiZa6rwoE3A74rJgtVTlrH19s5iaWNK
h6ha4kDra/k+6Pjr7BRAeYTNttZQsuR09sORWjJsh18F7BfTC5GPej9/lq71C+2x
4BrYgsiH9p670xcLePy7Suv0F68zC86j1447ZqatYyo/JX01Sx1VyqpzNFgSGsE0V
YmKzu9gaFuGCxRszXmVBnN1suu4ShSqivV+W1/WgmdwNoHPFH1sObImB5JS0DusN
ta2iN2bK1lWn6cFXA5Hb0hAIrIYz7dwSd4T+3G71J9ZCy5n1EaCaAwfGGjbxSIob
uyeG024i16yN9Ev7osvnVS6NNEH1EEIng7H0nC4ymotJ7QARAQABiQE8BBgBCAAm
FiEE/2yty4F7D7s8Bt0Ztx5mSOSBRwcfAl0vTY0CGwMFCQPDxJMACgkQtX5mSOSB
RwcmZwgAKK2+GYhAmUSztj67MPboeV7DZnwv8tmb7S1kFX0KYiG0QujMSRlH4I47
mV2Sg811uMsD/aeSSzPbYct77v6EMKwxJCD8KREmDyL6eQhQxc8PbbwFgwrSDtQ+
yd3Jh0uykxm6WlrInnsVBz0Aom2PT1iUMeb1br0LVFKG9PdEW5NgcjldQEldmnqBM
gD13FzbFp0aNBs48kIcFE/6cS0z+vPN71dii3nWz6owbvKFVrs/Jqs79JknMW4QN
18FfnFFJSZucuOXlmKeBEiVVB7Srvb4rUUVRLeit1jLANiQCuIZpZQHD4Q6v2SP3
D12KaXvT2nDcPeiqRCI3CrqUNMjmfQ==
=Kmi+
-----END PGP PUBLIC KEY BLOCK-----

```

RYSUNEK 11.24. Wygenerowany klucz publiczny

Podpisz/zaszyfruj pliki - Kleopatra

Podpisz / zaszyfruj pliki

Potwierdź autentyczność (podpisz)

☒ Podpisz jako: ☒ AndrewGrazynski <grazynski@op.pl> (certyfikowane, utworzoi)

Zaszyfruj

☒ Zaszyfruj dla mnie: ☒ Erago-243711 <gesa-3378@erago.fraunhofer.org> (certyfikow)

☒ Zaszyfruj dla innych:

☐ Zaszyfruj hasłem. Każdy który zna hasło będzie mógł odczytać dane.

Wynik

☐ Zaszyfruj / podpisz każdy plik osobno.

 C:/FOTKI/02.gif.gpg  

RYSUNEK 11.25. Szyfrowanie i (lub) podpisywanie wiadomości

7. Aby odszyfrować otrzymaną wiadomość, wklej jej treść do notatnika Kleopatry (musisz wkleić całą wiadomość, łącznie z nagłówkiem `BEGIN PGP MESSAGE` oraz stopką `END PGP MESSAGE`) i kliknij opcję *Odszyfruj/sprawdź* na pasku narzędziowym.

Sprawdzian

1. Wygeneruj parę kluczy „publiczny-prywatny” za pomocą Kleopatry. Prześlij plik `.asc`, zawierający klucz publiczny, na serwer swojej organizacji, zgodnie ze wskazaniami administratora.
2. Zaimportuj klucz publiczny Twojego kolegi do Kleopatry i wyślij mu zaszyfrowaną wiadomość na dowolny temat.
3. Twój kolega prześle Ci wiadomość zaszyfrowaną Twoim kluczem publicznym. Odszyfruj ją za pomocą Kleopatry.

ĆWICZENIA PRAKTYCZNE 11D

Projekt sieci dla rezydencji Apollo

W rezydencji Apollo gościmy ostatnio dość często, tym razem jako konsultanci od bezpieczeństwa komputerowego. Uniwersyteckie gremium odpowiedzialne za jakość usług i zadowolenie gości pragnie, by rezydencja zachowała ciągłość funkcjonowania, a jej sieć była solidnie chroniona, zarówno przed atakami z zewnątrz, jak i niecnymi działaniami ze strony niewdzięcznych studentów — to właśnie z ich strony pochodzi płowa wszystkich prób ataku, z czego gremium zarządzające doskonale zdaje sobie sprawę.

Sprawdzian

Zadaniem Twojej ekipy jest tym razem zaprojektowanie zabezpieczeń dla sieci funkcjonujących w rezydencji. Znając szczegóły konstrukcji tych sieci z poprzednich rozdziałów i kierując się ofertą przedstawioną na rysunku 11.26 zdecyduj, jakie należy w tym celu zakupić komponenty sprzętowe i programowe i jakie zasubskrybować usługi.

Urządzenia zabezpieczające	Cena
System ochrony przed atakami DDoS (detektor i analizator)	25 000
Firewall aplikacyjny	1 200
Brama VPN	700
System zapobiegania włamaniom (IPS)	15 000
Oprogramowanie zabezpieczające dla serwerów i routerów	Cena za egzemplarz dla jednego urządzenia
Oprogramowanie do ciągłej ochrony danych (CDP)	200
Oprogramowanie HTTP	150
E-mailowy system generowania haseł jednorazowych	1 000
Oprogramowanie do uwierzytelniania scentralizowanego	200
Oprogramowanie firewalla dla routera	100
Oprogramowanie firewalla NAT dla routera	200
Oprogramowanie zabezpieczające dla komputerów klienckich	Cena za egzemplarz dla jednego komputera
Oprogramowanie antywirusowe	20
Oprogramowanie VPN	10
Usługi bezpieczeństwa	Roczny abonament za jednego użytkownika
Uwierzytelnianie dwuczynnikowe z wykorzystaniem telefonu komórkowego	100

RYСУNEK 11.26. Oferta sprzętu, oprogramowania i usług przydatnych do zabezpieczenia sieci (ceny w dolarach)

SKOROWIDZ

5G, 396, 404

A

ACK, Acknowledgment, 178

ACL, Access Control Lists, 335, 446

administrowanie środowiskiem użytkownika, 510

ADPCM, Adaptive Differential Pulse-Code
Modulation, 127

adres

IPv4, 189, 191

IPv6, 239

MAC, 162, 188

URL, 70

adresowanie, 187, 218

bezklasowe, 191

podsieci, 193

rozgłoszeniowe, 192

adresy

dostępu do nośnika, 188

dynamiczne przydzielanie, 194

fizyczne, 188

IP prywatne, 448

przydzielanie, 188

rozwiązywanie, 195

typy, 187

warstwy łączy danych, 187, 198

ADSL, Asymetryczne DSL, 390

adware, 456

AES, Advanced Encryption Standard, 460

agent

DDoS, 432

transferu poczty, 75

użytkownika poczty, 75

agregowanie pasm, 396

algorytm, 457

AM, Amplitude Modulation, 120

analiza

pakietu HTTP, 169

poczty elektronicznej, 91

sieci

VPN, 372

WAN, 371

wymagań, 248

analizator anomalii ruchu, 434

ANSI, 46

antena

kierunkowa, 275

wielokierunkowa, 274

antywirusy, 453

AP, Access Point, 25, 273

aplikacje

krytyczne, 422

w sieci, 251

architektura

aplikacji, 56, 87

chmurowa, 58, 64

CORBA, 61

DSL, 389

dwuwarstwowa, 62

architektura

- e-mail, 75
- firewalli, 448
- FTTH, 393
- gwiazdy, 349
- kliencka, 59
- klient-serwer, 60, 63
- modemów kablowych, 391
- peer-to-peer, 67
- pierścienia, 348
- siatki, 350, 351
- sieci szkieletowej, 326
- sieci trasowanej, 323
- trójwarstwowa, 62
- trójwarstwowa e-mail, 76
- usługi modemów kablowych, 391
- wielowarstwowa, 62
- WiMax-a, 394
- z centralnym hostem, 58

ARP, Address Resolution Protocol, 198

ARP cache, 227

ARQ, Automatic Repeat reQuests, 153, 182

- z weryfikacją na bieżąco, 183

- z wycofywaniem do poprawności, 184

ASCII, 114, 155

asymetryczna DSL, 390

atak

- DDoS, 402, 433
- dnia zero, 452
- DoS, 432, 435
- ICMP, 435
- na tablicę procesów uniksowych, 435
- przepelnienia bufora, 452
- siłowy, 457
- UDP, 435

ataki ukierunkowane, 416

automatyczne żądania powtórzenia, 153, 183

awaria urządzeń, 437

B

badanie lokalizacji, 292

bajt, 114

baud rate, 122

baza informacji na potrzeby zarządzania, 502

bezpieczeństwo

- aplikacji, 85
- fizyczne, 450
- protokołu IP, 465
- sieci, 413

bezpieczna transmisja danych, 461

bezpieczne przesyłanie danych, 463

bezprowadowe sieci

- BYOD, 40
- LAN, 40, 271

bezprowadowy

- ethernet, 306
- punkt dostępowy, AP, 25

BGP, Border Gateway Protocol, 205

biegunowość, 116

bilet usługi, 472

bit

- parzystości, 150
- rate, 122

bity

- informacyjne, 159
- narzutu, 159

błąd kwantyzacji, 124

błędy

- kontrola, 153
- korekcja, 163
- progresywna korekcja, 151
- sieciowe, 145
- wykrywanie, 149, 163
- zapobieganie, 148, 163
- źródła, 146, 163

BN, Backbone Networks, 26, 45

bod, 122

bot, 432

brama

- IP, 209
- VPN, 359

broadcast, 190, 206, 331

BYOD, 40

C

CA, Carrier Aggregation, 396

CA, Certificate Authorities, 463

cache'owanie treści, 506

całkowity koszt posiadania, 526

CD, collision detection, 282

CDP, Continuous Data Protection, 440

centrum

- danych, 243, 295
- dystrybucji kluczy, 472
- dystrybucyjne, 392
- operacyjne sieci, 514

chmura Salesforce.com, 68

chmury
 obliczeniowe, 69
 prywatne, 57
 publiczne, 57
 wspólnotowe, 57
 ciało, 79
 odpowiedzi, 72
 żądania, 71
 ciągła ochrona danych, 440
 ciągłe ARQ, 183, 184
 ciągłość funkcjonowania, 430, 482, 489
 CIR, Committed Information Rate, 356
 CMTS, Cable Modem Termination System, 392
 CRC, Cyclic Redundancy Check, 151 154, 178
 CSU, Channel Service Units, 347
 cyberbezpieczeństwo, 44, 85, 128, 162, 217, 260, 305, 335, 367, 402, 480, 530
 cyfrowa linia abonencka, DSL, 104, 388
 cykliczna kontrola nadmiarowa, 151, 154, 178
 czterostopniowe uzgadnianie, 185

D

dane
 analogowe, 99, 124
 cyfrowe, 99, 113, 119
 datagram UDP, 362
 DCE, Distributed Computing Environment, 61
 DDoS, Distributed Denial-of-Service, 402, 434
 DE, Discard Eligible, 356
 definiowanie haseł, 468
 delineacja komunikatu, 163
 demodulator, 123
 DES, Data Encryption Standard, 459
 deszyfracja, 457
 detektor anomalii ruchu, 433
 dezorganizacja, 417
 diagram MDF, 321
 digitalizacja sygnału, 124
 DMZ, DeMilitarized Zone, 448, 476
 DNS, Domain Name Service, 197, 476
 cache, 229
 Request, 231
 Response, 231
 DOCSIS, 390
 dokumentowanie
 istniejących kontroli, 429
 konfiguracji, 510

domeny internetowe, 24
 najwyższego poziomu, 197
 DoS, Denial of Service, 432, 435
 dostarczanie treści, 507, 509
 dostawca
 chmur, 57
 internetu, 20, 358, 385
 usług internetowych, 382
 dostęp do internetu, 244, 397
 dostępność, 417
 sieci, 518
 DPCM, Differential Pulse-Code Modulation, 127
 DSL, Digital Subscriber Line, 20, 388, 403
 DSLAM, DSL access multiplexer, 390
 DSU, Data Service Units, 347
 działanie IPS, 474
 dzielenie ryzyka, 429

E

echo, 147
 efekt magistrali, 255
 efektywność transmisji, 160, 164
 e-handel, 298
 EIGRP, Enhanced Interior Gateway Routing Protocol, 205
 ekranowanie nośnika, 148
 eksploatacja luk bezpieczeństwa, 425, 454
 ekstranet, 27
 elektroniczne dostarczanie oprogramowania, 510
 elektrotechnika, 116
 e-mail, 74
 eMBB, 397
 enkapsulacja, 35
 ESP, Encapsulating Security Payload, 362
 ethernet, 118, 157, 282, 291, 305
 bezprowodowy, 283
 bezpieczeństwo, 288
 dostęp do nośnika, 284
 funkcja punktu koordynującego, 285
 rozproszona funkcja koordynująca, 284
 topologia, 284
 typy, 286
 przewodowy, 277
 bazujący na przełącznikach, 279
 oparty na koncentratorze, 278
 topologia, 278

F

faktoryzacja, 460
 fala nośna, 120
 fale dźwiękowe
 amplituda, 119
 częstotliwość, 119
 długość, 119
 faza, 119
 modulacja, 120
 fałszywa rekurencja DNS, 435
 fałszywe antywirusy, 453
 farma serwerowa, 45, 504
 FDM, Frequency Division Multiplexing, 103
 filtrowanie
 adresów MAC, 162, 289
 ruchu, 432
 Finger of Death, 435
 firewalle, 444
 aplikacyjne, 447
 NAT, 448
 pakietowe, 446
 flaga, 156
 floodowanie TCP SYN, 435
 FM, Frequency Modulation, 120
 Frame Relay, 356
 framework zarządzania siecią, 500
 FSK, Frequency-Shift Keying, 120
 FTTH, Fiber To The Home, 393, 403
 funkcje
 koordynujące, 284
 mieszające, 462
 punktu koordynującego, 285
 warstwy transportowej, 180

G

gigapop, 401
 główny punkt dystrybucyjny, 320, 390
 grupowanie warstw, 32
 GSM, 396

H

haker, 443
 hakywizm, 416
 hasło, 467–469
 HDLC, High-level Data Link Control, 154, 156

helpdesk, 530
 hermetyzacja, 35
 HFC, Hybrid Fiber Coax, 391
 honeypot, 478
 hotspoty Boingo, 109
 HTML, Hypertext Markup Language, 73
 HTTP, Hypertext Transfer Protocol, 32, 70, 77
 hybrydowe
 ARQ, 395
 sieci kablowe, 391

I

IaaS, Infrastructure as a Service, 67
 ICANN, 400
 ICMP, Internet Control Message Protocol, 205
 identyfikacja zagrożeń, 425
 IDF, Intermediate Distribution Points, 320
 IETF, 399
 IMAP, Internet Message Access Protocol, 75–77
 informatyka śledcza, 477
 infrastruktura
 jako usługa, IaaS, 67
 klucza publicznego, 460, 463
 integralność, 417
 interfejs, 199
 administrowania środowiskiem użytkownika, 510
 intermodulacja, 148
 internet, 381
 architektura, 383
 działanie, 383, 403
 rzeczy, IoT, 41
 technologie dostępu, 388
 zarządzanie, 399
 intranet, 27
 intruzja, 418
 inwentaryzacja zasobów IT, 421
 IoT, Internet of things, 42
 IP, Internet Protocol, 40, 177, 357
 IPS, Internet Protocol Suite, 173, 476
 oparty na goście, 474
 oparty na sieci, 474
 IPSec, IP Security, 465
 IPSs, Intrusion Prevention Systems, 474
 IPv4, 191
 IS-IS, Intermediate System to Intermediate System, 205
 ISO 8859, 114

ISP, Internet Service Provider, 20, 358, 382, 387,
 IXP, Internet Exchange Point, 384
 IXPs, Internet exchange points, 383

J

jakość usługi, 186
 jednostka danych, 32
 jednostki obsługi
 danych, 347
 kanałów, 347

K

kabel Cat 5, 135
 kabel
 koncentryczny, 107
 połączeniowy, 138
 światłowodowy, 105, 108
 karta sieciowa, 162, 270
 katastrofa, 439, 442, 482
 klasa
 usług, 186
 adresu, 237
 klastr, 504
 błędów, 146
 kliencka część NOS, 276
 klient, 25, 254
 cienki, 63
 e-mail, 75
 gruby, 63
 klient-serwer, 58
 klucz, 457
 prywatny, 460
 publiczny, 460, 494
 sesji, 472
 kluczowanie
 amplitudy, 120
 częstotliwości, FSK, 120
 fazy, PSK, 121
 kod, 114
 Hamminga, 152
 kodek, 99, 124
 kodowanie, 114
 dwuprądowe, 118
 z naprzemienną inwersją znaku, 118
 komórka telekomunikacyjna, 104
 komponenty sieci, 25, 242, 249, 305

komunikacja
 bezpołączeniowa, 182, 185
 danych, 19, 22
 połączeniowa, 182
 przez sieć zasilającą, 300
 komunikat HTTP, 33
 komunikatory internetowe, 80, 127
 komutacja pakietów, 354, 368
 koncentratory, hubs, 272
 konfiguracja
 komputerów klienckich, 509
 obwodu, 100
 sieci, 509
 TCP/IP, 210
 konie trojańskie, 455
 konto, 466
 kontrole
 bezpieczeństwa, 478
 błędów, 145, 153
 detekcyjne, 419
 korekcyjne, 419
 parzystości, 150
 prewencyjne, 419
 sieciowe, 419
 kontrolowanie awarii, 515
 konwerter elektryczno-optyczny, 392
 koordynator DDoS, 432
 kopie zapasowe, 440
 online, 441
 korekcja błędów, 151, 163
 koszty, 525
 kradzież poufnych informacji, 427
 kraker, 443
 krawędź
 e-handlu, 244
 przedsiębiorstwa, 244
 kryptografia, 457, 465
 krzywych eliptycznych, 459
 symetryczna, 457
 z kluczami publicznymi, 460–463
 kryptologia, 457
 kształtownik
 pasma, 505
 ruchu, 505
 kwadraturowe
 modulowanie amplitudowo-fazowe, QAM, 122
 kluczowanie fazy, 395

L

L2TP, 359, 370
 LAN, Local Area Network, 26, 40, 45, 75, 243, 267
 komponenty sieci, 269, 305
 polepszanie wydajności, 301
 sieci bezprzewodowe, 271
 LAP-M, Link Access Protocol for Modems, 184
 lista
 dystybucyjna, 74
 kontrolni dostępu, 335, 446
 LLC, Logical Link Control, 142
 logika
 aplikacyjna, 58
 biznesowa, 58
 dostępu do danych, 58
 prezentacji, 58
 lokalizowanie źródła przepięć, 149
 LTE, Long Term Evolution, 395, 404
 luki bezpieczeństwa, 403, 425, 452

Ł

łamanie haseł, 467
 łąta, 452
 łączniki VLAN, 330

M

MAC, Media Access Control, 142, 188
 MAC spoofing, 162
 macierze dyskowe, 298, 303
 magazynowanie danych, 58
 makrowirusy, 431
 malware, 425
 MAN, Metropolitan Area Networks, 26, 45, 365
 mapowanie sieci, 338
 MAR, Maximum Allowable Rate, 356
 maski podsieci, 194, 234
 masywna komunikacja pomiędzy maszynami, 397
 masywność online, 43
 MDF, Main Distribution Facility, 320, 390
 menedżer
 bezpieczeństwa sieci, 479
 haseł, 469
 operacyjny sieci, 324
 polityki sieciowej, 513
 sieci, 499, 523

MIB, Management Information Base, 501, 502
 mikrofae, 110
 mikser, 392
 MIME, Multipurpose Internet Mail Extension, 79
 MIMO, 395
 mirroring, 438
 MMOG, Massively Multiplayer Online Games, 43
 mMTC, 397
 moc, 116
 model
 chmur obliczeniowych, 57
 internetowy, 31
 obliczeń rozproszonych, 64
 referencyjny OSI, 28
 warstwy, 29
 sieci, 45
 warstwowy
 wady, 35
 zalety, 35
 modem, 99
 DSL, 389
 kablowy, 390–392, 403
 modulacja, 120
 amplitudy, AM, 120
 częstotliwości, FM, 120
 fazy, PM, 121
 impulsów
 kodowanych, PCM, 124, 125
 kodowanych różnicowo, DPCM, 127
 z adaptacyjnym kodowaniem różnic,
 ADPCM, 127
 modulator, 123
 monitorowanie sieci, 512, 514, 536, 540
 MPEG, Motion Picture Experts Group, 40
 MPLS, Multiprotocol Label Switching, 357
 MTBF, Mean Time Between Failures, 519
 MTTDiagnose, 519
 MTTFix, 519
 MTTRepair, 519
 MTTRespond, 519
 multicast, 190, 206
 multiplexer dostępowy, 390
 multipleksowanie, 102
 statystyczne z podziałem czasu, 103
 z podziałem czasu, 103
 z podziałem częstotliwości, 103
 z podziałem długości fali, 103

N

nadzorowany dostęp, 143
 nagłówek, 79
 odpowiedzi, 72
 pakietu SMTP, 92
 protokołu IPv4, 178
 protokołu IPv6, 178
 żądania, 71
 napięcie, 116
 NAT, Network Address Translation, 448, 476
 natężenie, 116
 nawałnica alarmów, 500
 nazwy serwerów, 196
 NCO, Network Cost of Ownership, 527
 NetView, 313
 NIC, Network Interface Card, 162, 270
 nieuprawniony dostęp, 425
 niezaprzeczalność, 462
 NOC, Network Operations Center, 514
 NOS, Network Operating System, 276
 część kliencka, 276
 część serwerowa, 276
 nośnik
 bezprzewodowy, 106
 informacji, 106
 kierowany, 106
 komunikacyjny, 129
 transmisyjny, 106
 wielomodowy, 108
 notacja z ukośnikiem, 191
 Nslookup, 228
 numerowanie pakietów, 185

O

obciążenia zwrotne, 527
 obciążenie, 254
 obliczenia gridowe, 67
 obrona przed socjotechniką, 471
 obserwowanie jednostek PDU, 50, 94
 obszar roboczy, 104
 obwody, 25, 98, , 129
 dedykowane, 100, 347
 dwupunktowe, 100
 dzierżawione, 367
 fizyczne, 98
 logiczne, 98
 multipleksowane, 102
 sieciowe, 270
 T-carrier, 352
 współdzielone, 100
 ocena ryzyka, 420, 482
 ochrona
 na granicy sieci, 444
 proaktywna, 444
 przed kradzieżą, 436
 przed skutkami katastrof, 439
 przed złośliwym oprogramowaniem, 430
 serwerów i klientów, 452
 ODBC, Open Database Connectivity, 62
 odbicie lustrzane, 438
 odcisk palca, 464
 odmowa usługi, 432, 433
 odpowiedzi HTTP, 70, 72
 odtwarzanie
 po katastrofie, 439
 po włamaniu, 477
 OFDM, 395
 ograniczniki pasma, 505
 okablowanie, 271
 poziome, 104
 strukturalne, 104
 szkieletowe, 104
 okno przesuwne, 156, 184
 ONU, Optical Network Unit, 393
 operatorzy telekomunikacyjni, 346
 opóźnienia, 366
 informacyjne, 20
 oprogramowanie
 jako usługa, SaaS, 65
 kryptograficzne, 464
 pośredniczące, 61
 reklamowe, 456
 szpiegujące, 456
 wspomagające zarządzanie siecią, 500
 zarządzania
 aplikacjami, 500
 punktami, 500
 systemem, 500
 urządzeniami, 500
 ortogonalne multipleksowanie z podziałem
 częstotliwości, 395
 OSI, 28
 OSPF, Open Shortest Path First, 205
 otwarte łącze bazodanowe, 62
 outsourcing, 441

P

- RIP, Routing Information Protocol, 205
- P2P, peer-to-peer, 87
- PaaS, Platform as a Service, 66
- PAD, packet assembly/disassembly device, 354
- pakiet, 32
 - ESP, 362
 - IP, 34, 177
 - SMTP, 79, 92
- pamięć masowa, 297
- panel krosowt, 299
- paszport, 473
- PCM, Pulse-Code Modulation, 124
- PDU, Protocol Data Unit, 32, 50
- pętla lokalna, 126, 190
- PGP, Pretty Good Privacy, 464
- phishing, 425, 472
- PKI, Public Key Infrastructure, 460, 463
- planowanie przepustowości, 254
- platforma jako usługa, PaaS, 66
- PLC, Power Line Communicaton, 300
- PM, Phase Modulation, 121
- poczta
 - elektroniczna, 74, 87
 - webowa, 76
- podpis cyfrowy, 462
- podpisywanie wiadomości, 494
- podsieci, 192
- podśluchiwanie transmisji, 450
- podszycanie się pod adresy MAC, 162
- podział sieci, 249
- pole, 154
 - Czas życia, 178
- polecenie
 - ipconfig, 224
- polityka bezpieczeństwa, 444
- połączenia
 - TCP, 214
 - telefoniczne, 126
- pomiar
 - prędkości, 410
 - ryzyka, 420
- pomieszczenie wyposażenia, 104
- POP, Post Office Protocol, 40, 75–77, 401
- POPs, Points Of Presence, 385
- port, 177, 180
 - docelowy, 180
 - źródłowy, 180
- pośrednie punkty dystrybucyjne, 320
- POTS, Plain Old Telephony Service, 42
- potwierdzenie tożsamości, 471
- poufność, 417
- PPP, Point-to-Point Protocol, 154, 158
- praktyki projektowe, 336
- prędkość
 - bitowa transmisji, 122
 - połączenia internetowego, 411
 - przesyłania, 100
- proces standaryzacyjny, 36
- profil, 466
 - sieciowy, 277
- program
 - do zarządzania przepustowością, 506
 - ipconfig, 224
 - Kleopatra, 492
 - NetView, 313
 - nslookup, 228
 - PuTTY, 80
 - SmartDraw, 265
 - SolarWinds, 537
 - TracePlus, 312
 - tracert, 230, 375
 - VisualRoute, 408
 - Wireshark, 372
- progresywna korekcja błędów, 151
- projektowanie sieci, 247, 342, 498
 - moduły, 246, 261
 - narzędzia, 257, 264
 - obwody, 254
 - szkieletowych, 332
 - technologiczne, 247
 - WAN, 363
 - wydajność, 532
- prośba o złożenie oferty, 258
- protokoły sieciowe, 32, 38
 - trasowania, 202, 205
 - warstwy łącza danych, 142, 154
 - zarządzania siecią, 501
- protokół
 - ARP, 198
 - BGP, 205
 - Diffiego-Hellmana, 458
 - EIGRP, 205
 - Ethernet, 157
 - ICMP, 205
 - IPv4, 177, 178, 214

IPv6, 178
 IMAP, 76
 IS-IS, 205
 LAP-M, 183
 OSPF, 205
 POP, 76, 95
 PPP, 158
 RFP, 258
 RIP, 205
 SMTP, 74, 78, 94, 174
 TCP, 176, 177, 214
 SDLC, 156
 SNMP, 502
 UDP, 177
 próbkowanie, 124
 przechwytywanie pakietów sieciowych, 168
 przeciwdziałanie ryzyku, 429
 przeglądarka WWW, 69
 przekazywanie żetonu, 144
 połączana sieć szkieletowa, 319
 przełączane

- obwody wirtualne, 355
- sieci szkieletowe, 322, 335

 przełącznik, 25, 126, 272, 318

- VLAN, 318
- w obudowie, 321
- warstwy 2., 280

 przepięcie, 147
 przepływ danych, 101
 przepróbkowanie, 125
 przepustowość, 161, 505

- obwodów, 122, 164, 304, 334, 366

 przepytanie, 143

- według listy, 144
- węzłowe, 144

 przesłuch, 147
 przewodowe sieci LAN, 270
 przezroczystość, 102
 przydzielanie adresów, 188
 PSK, Phase-Shift Keying, 121
 punkty

- dostępowe, AP, 273, 284, 290
- obecności, 385, 401
- wymiany ruchu internetowego, 383

 PuTTY, 80
 PVCs, permanent virtual circuits, 355

Q

QAM, Quadrature Amplitude Modulation, 122
 QPSK, Quadrature Phase Shift Keying, 395
 QoS, Quality of Service, 186

R

racjonalizacja, 513
 radio, 109
 RAID, 303, 438
 ramka, 154, 155

- bezprzewodowa, 286
- Ethernet II, 158
- ethernetowa, 34, 171
- ethernetowa 802.3ac, 157
- optymalny rozmiar, 162
- SDLC, 156

 ransomware, 425, 458
 raporty

- menedżerów, 520
- o problemach, 519
- raporty techniczne, 518

 RC4, Rivest Cipher 4, 460
 redukowanie

- kosztów, 528
- ruchu sieciowego, 505

 redundantna macierz dyskowa, 438
 regeneratory sygnału, 148, 273
 reguły biznesowe, 58
 rejestrowanie domen, 400
 retransmisja pakietu, 151
 RFID, Radio Frequency Identification, 81
 RFP, Request For Proposal, 258
 RMON, Remote Monitoring, 502
 robak, 431
 router, 199, 207, 318

- desygnowany, 203
- DSL, 389

 rozdzielanie pakietów, 354
 rozgałęźnik, 389
 rozgłaszanie, 190, 198, 331
 rozmiar ramki, 162
 rozmowy telefoniczne, 126
 rozproszona odmowa usługi, 402
 rozproszone środowisko obliczeniowe, 61

rozwiązywanie

- adresów, 218
- adresów warstwy łącza danych, 198
- nazw serwerów, 196
- problemów, 522

równoważenie obciążenia, 297, 504

RSVP, Resource Reservation Protocol, 186

RTP, Real-Time Transport Protocol, 186

RTSP, Real-Time Streaming Protocol, 186

rywalizacja, 143

ryzyko, 420, 429

S

SaaS, Software as a Service, 65

SAN, Storage Area Network, 297

satelity, 111

scalanie pakietów, 354

scenariusz kompromitacji bazy, 427, 428

schemat kodowania, 114

SDH, Synchronous Digital Hierarchy, 353

SDLC, Synchronous Data Link Control, 154, 156

segment TCP, 33, 177

segmentacja sieci, 181, 304

selektywne ARQ, 183

serwer, 25, 254

- aplikacji, 62
- autorytatywny, 197
- nazw, 196
- plików, 26
- pocztowy, 26, 75
- rozwiązujący, 197
- uwierzytelniający, 471
- WWW, 26

serwerowa część NOS, 276

serwery

- odporne na awarie, 438
- zwiększanie wydajności, 302

sesja, 182

siatka

- częściowa, 350
- pełna, 350

sieci

- aplikacje, 251
- kategoryzacja wymagań, 252
- komponenty architektoniczne, 242
- projektowanie, 241
- projektowanie technologiczne, 254
- użytkownicy, 251

sieciowy

- koszt posiadania, 527
- system operacyjny, NOS, 276

sieć

- 5G, 397
- bezprzewodowa, 267
- kampusowa, 243
- komunikacji danych, 22
- komutacji pakietów, 354
 - architektura, 354
 - prędkości, 356
 - przeplatanie się, 355
- lokalna, LAN, 26, 40, 45, 75, 243, 267
 - komponenty sieci, 269, 305
 - polepszanie wydajności, 301
 - sieci bezprzewodowe, 271
- miejska, MAN, 26, 45, 365
- nakładkowa, 291
- obwodów dedykowanych, 347
 - architektura gwiazdy, 349
 - architektura pierścienia, 347, 348
 - architektura siatki, 350, 351
- peer-to-peer, P2P, 68, 87
- przewodowa, 267
- rozległa, WAN, 45, 244, 345
 - analiza sieci, 371
 - projektowanie sieci, 363, 368
 - wydajność sieci, 365, 368
- szkieletowa, BN, 26, 45, 317, 387
 - projektowanie, 332
 - przełączane, 335
 - trasowane, 322, 335
 - wewnątrz budynku, 243
 - wydajność, 333, 336
- TCP/IP, 209, 219
- wirtualna LAN, VLAN, 325, 327, 335
 - działanie sieci, 328
- wirtualna sieć prywatna, VPN, 358, 368
 - architektura, 358
 - dostępowa, 360, 361
 - działanie, 360
 - ekstranetowa, 360
- Web, 69, 70, 87
- zarządzana, 499
- silnik cache'owania, 506
- skanowanie sieci, 340
- skrętka, 98, 106
- skrzynka odbiorcza, 75

- SLA, Service-Level Agreement, 522
- SMTP, Simple Mail Transfer Protocol, 74, 77, 174
- SNMP, Simple Network Management Protocol, 501
- SNR, Signal-to-Noise Ratio, 122
- socjotechnika, 471, 474
- SOHO, 299
- SolarWinds, 537
- SONET, 353
- sól, 467
- spyware, 456
- SQL injection, 86
- SSL/TLS, Secure Sockets Layer, 464
- stabilizator obciążenia, 504
- stacja czołowa, 392
- stałe obwody wirtualne, 355
- stan bezczynności, 155
- standard, 46
 - 802.11i, 289
 - ASCII, 114
 - IEEE 802.3ac, 157
 - ISO 8859, 114
 - Unicode, 115
- standardy
 - powszechne, 39
 - sieciowe, 36
 - zarządzania sieciami, 501
- status odpowiedzi, 72
- statystyki
 - awarii, 518
 - wydajności, 518
- STDM, Statistical Time Division Multiplexing, 103
- sterowanie
 - dostępem do nośnika, 142, 143, 163, 281
 - połączeniem logicznym, 142
 - przepływem, 185
 - transmisją, 177
 - wysokopoziomowe łączem danych, 156
- stopa błędów, 113, 146
- stos protokołów, 35
 - TCP/IP, 173, 176
- stosunek sygnału użytecznego do szumu, 122
- strategia
 - czysta, 57
 - mieszana, 57
- strefa zdemilitaryzowana, 448, 476
- struktura pakietu HTTP, 52
- sumy kontrolne, 150
- SVCs, switched virtual circuits, 355
- sygnał bezczynności, 155
- sygnał
 - radiowy, 109
 - zagłuszający, 282
- symbol rate, 122
- synchroniczne sterowanie łączem danych, 156
- system
 - autonomiczny, 385
 - DNS, 197
 - końcowy modemów kablowych, 392
 - zapobiegania włamaniom, 474, 476
- systemy operacyjne, 454
- szacowanie kosztów, 258
- szkolenia, 524
- szperacz pakietów, 168
- szum
 - biały, 146
 - gaussowski, 146
 - impulsowy, 147
- szybkość transmisji, 113
 - bitów informacyjnych, 166
 - danych, 122
- szyfr
 - AES, 460
 - DES, 459
 - RC4, 460
 - Triple DES, 459
- szyfrogram, 457
- szyfrowanie, 457, 490, 491, 494

Ś

- śledzenie
 - ethernetu, 311
 - problemów, 517
 - trasy pakietu, 375, 376, 377
- średni czas
 - pracy bezawaryjnej, 519
 - naprawy, 519
- światłowód, 107
 - do domu, 393, 403
 - gradientowy, 108
 - skokowy, 108

T

- tablica
 - przekazywania, 280
 - trasowania, 199

- TCO, Total Cost of Ownership, 526
 - TCP, Transmission Control Protocol, 33, 40, 177, 214
 - TCP/IP, 215
 - TDM, Time Division Multiplexing, 103
 - technologia Kinect, 42
 - technologie szerokopasmowe, 388
 - tekst jawny, 457
 - telefon VoIP, 128
 - telekomunikacja, 22
 - Telnet, 80
 - terminal sieci optycznej, 393
 - tęczowe tablice, 467
 - TLD, Top-Level Domain, 197
 - TracePlus, 312
 - tracert, 230, 375
 - translacja adresów sieciowych, 448, 476
 - transmisja
 - analogowa danych cyfrowych, 119, 130
 - asynchroniczna, 154, 155
 - cyfrowa, 116
 - danych analogowych, 124, 130
 - danych cyfrowych, 113, 130
 - danych przez ethernet, 118
 - komunikatu
 - HTTP, 33
 - przez warstwy, 32
 - SMTP, 78, 175
 - modemowa, 123
 - pełnodupleksowa, 102
 - półdupleksowa, 101
 - równoległa, 115
 - satelitarna, 111
 - sekwencji dwubitowych, 121
 - sympleksowa, 101
 - synchroniczna, 155
 - szeregowa, 115
 - trasowane sieci szkieletowe, 322, 335
 - trasowanie, 198, 200, 205, 218
 - dynamiczne, 201
 - scentralizowane, 200
 - statyczne, 200
 - w internecie, 203
 - TRIB, Transmission Rate of Information Bits, 166
 - Triple DES, 459
 - trojan, 455
 - tryb transmisji, 115
 - transportowy, 465
 - tunelowy, 465
 - TTL, Time To Live, 178
 - tunel VPN, 359
 - tworzenie plików MP3, 137
 - typy
 - adresów, 187
 - architektur aplikacji, 56, 87
 - DSL, 390
 - ethernetu, 282
 - FTTH, 393
 - modemów kablowych, 392
 - sieci, 26, 112
 - trasowania, 200
 - WiMax-a, 394
 - zagrożeń bezpieczeństwa, 417, 482
- ## U
- UDP, User Datagrams Protocol, 177
 - ujednolicony lokalizator zasobu, URL, 70
 - ultraniezawodna transmisja, 397
 - umowa SLA, 522
 - UMTS, 396
 - unicast, 331
 - Unicode, 115
 - unikanie katastrofy, 439
 - UPS, Uninterruptible Power Suppliers, 439
 - URL, Uniform Resource Locator, 70
 - URLLC, 397
 - urządzenia zarządzane, 499
 - urzędy certyfikacyjne, 463
 - usługa, 186
 - dedykowanego obwodu, 348
 - dostarczania treści, 508
 - modemów kablowych, 391
 - nazw domenowych, 476
 - usługi
 - ethernetowe, 357
 - Frame Relay, 356
 - IP, 357
 - SONET, 353
 - T-carrier, 352
 - telekomunikacyjne, 106
 - uwierzytelnianie
 - biometryczne, 470
 - dwuczynnikowe, 469
 - podmiotu, 462
 - scentralizowane, 471
 - tożsamości, 462
 - użytkowników, 466

uzgadnianie trójstopniowe, 182
użytkownicy sieci, 251

V

VisualRoute, 408
VLAN, Virtual LAN, 157, 325, 335
 działanie sieci, 328
VoIP, Voice over IP, 41, 127
VPN, Virtual Private Network, 358, 368
 architektura, 358
 dostępowe, 360, 361
 działanie, 360
 ekstranetowe, 360

W

WAN, Wide Area Network, 45, 244, 345
 analiza sieci, 371
 projektowanie sieci, 363, 368
 wydajność sieci, 365, 368
wardriving, 313
warstwa
 aplikacyjna, 32, 33, 55–96
 bezpiecznych gniazd, 464
 fizyczna, 29, 31, 34, 97–139
 intersieci, 32
 łącza danych, 29, 31, 34, 141–171
 prezentacyjna, 30
 sesji, 30
 sieciowa, 30, 31, 34, 173
 sprzętowa, 32
 transportowa, 30–33, 173, 217
warstwy modelu sieciowego, 215
warwalking, 313
wąskie gardła, 255
WDM, Wavelength Division Multiplexing, 103
webcasting, 84
wejście, 104
WEP, Wired Equivalent Privacy, 288
węzeł optyczny, 392
wideokonferencje, 82
 desktopowe, 83
wieloprotokołowe przełączanie etykiet, 357
wiersz żądania, 71
Wi-Fi, 291, 292
WiMax, 393, 404
Wireshark, 372

wirtualizacja, 296
wirtualne sieci LAN, VLAN, 157, 325, 335
 działanie sieci, 328
wirtualne sieci prywatne, VPN, 358, 368
 architektura, 358
 dostępowe, 360, 361
 działanie, 360
 ekstranetowe, 360
wirtualny
 pulpit, 86
 serwer, 504
włamania, 443
włókna jednomodowe, 108
worm, 431
WPA, Wi-Fi Protected Access, 289
wsparcie dla użytkowników, 522, 532
wstrzykiwanie kodu, SQL injection, 86
wtyczka Cat 5, 136
wybór
 architektury, 68
 nośnika, 112
wydajność
 sieci, 498
 szkieletowej, 336
 WAN, 365, 368
 urządzeń, 334, 366
wykrywanie
 anomalii, 476
 błędów, 149
 kolizji, 282
 nadużyć, 474
wyposażenie klienta, 389
wzmocniacze, 148

Z

zabezpieczenia, 416
 fizyczne, 436
zabiegi socjotechniczne, 415, 425
zagrożenia, 417, 425
załącznik, 74, 79
zanikanie sygnału, 147
zapewnienie ciągłości funkcjonowania, 417
zapobieganie
 błędom, 148
 włamaniom, 443, 482
zarządzane punkty dostępowe, 294

zarządzanie

- awariami, 512, 532
- internetem, 399
- kluczami, 458
- konfiguracją, 509, 532
- kosztami, 525, 533
- okablowaniem, 275
- przepustowością, 505
- pulpitami, 480
- ruchem sieciowym, 503
- sesjami, 182
- sieciami, 497, 501, 502, 505
- wydajnością, 512, 532

zasilacze awaryjne, 439

zasoby IT, 421

zautomatyzowane dostarczanie

- oprogramowania, 510

zdalne monitorowanie, 502

zestaw protokołów internetowych, 173

złośliwe oprogramowanie, 129, 425, 430

znakowanie, 330

Ż

źródła

- błędów, 146, 163
- kosztów, 525
- przepięć, 149

Ż

żądania

- HTTP, 70, 71
- powtórzenia, 182
- dostępu, 143

żądanie GET, 170

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 

Komunikacja w sieci: być albo nie być współczesnego biznesu!

Połączenie komputerów w sieć zapoczątkowało prawdziwy skok cywilizacyjny. Dziś zarządzanie przepływem danych między komunikującymi się urządzeniami stanowi jeden z fundamentalnych czynników rozwoju współczesnego biznesu. Wraz z upowszechnianiem się kolejnych nowinek, takich jak internet rzeczy, nasze życie coraz bardziej staje się życiem online. W konsekwencji zrozumienie zagadnień dotyczących mobilności, bezpieczeństwa, skalowalności i zarządzania siecią może mieć niebagatelne znaczenie dla pomyślnej realizacji celów biznesowych. Dotyczy to nie tylko inżynierów, lecz także twórczych jednostek, które dzięki kreatywnemu zagospodarowaniu technologii zwiększają konkurencyjność swoich firm.

To kolejne wydanie znakomitego podręcznika dla inżynierów i menedżerów, którzy chcą dogłębnie zrozumieć fundamentalne koncepcje związane z sieciami i komunikowaniem danych. Wyczerpująco omówiono tu podstawy funkcjonowania sieci komputerowych — szczególną uwagę zwrócono na różnorodne aspekty projektowania niezawodnej sieci i zarządzania nią. Książka została zaktualizowana i uzupełniona o najnowsze zalecenia w dziedzinie cyberbezpieczeństwa zarówno dla starszych rozwiązań, jak i dla architektur chmurowych czy sieci wykorzystujących urządzenia mobilne. Znakomitym uzupełnieniem treści są analizy przypadków oraz praktyczne zestawy ćwiczeń, dzięki którym można lepiej zrozumieć opisane koncepcje i techniki.

W tej książce między innymi:

- modele sieci i standardy sieciowe
- funkcje poszczególnych warstw sieci
- rodzaje sieci, w tym sieci szkieletowe, LAN i bezprzewodowe
- bezpieczeństwo sieci i zapobieganie włamaniom
- zarządzanie ruchem w sieci i kontrolowanie awarii

Dr Jerry FitzGerald był autorem wczesnych wydań tej książki, które ukazały się w latach 80. ubiegłego wieku. Obecnie jest dyrektorem założonej w 1977 roku firmy Jerry FitzGerald & Associates.

Dr Alan Dennis wykłada na Uniwersytecie Indiany. Specjalizuje się w tworzeniu oprogramowania wspomagającego pracę zespołową. Jest konsultantem Departamentu Obrony Stanów Zjednoczonych i Armii Australijskiej.

Dr Alexandra Durcikova wykłada w Price College of Business w Oklahomie. Interesuje się systemami zarządzania wiedzą oraz bezpieczeństwem sieci w kontekście psychologicznych aspektów zasad i ograniczeń.

  helion.pl  HELION SA ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl	<i>Sprawdź nasze szkolenia!</i>  AKADEMIA IT & BUSINESS WWW.SZKOLENIA.HELION.PL	KOD KORZYŚCI <i>Sięgnij po więcej!</i>   ISBN 978-83-283-5767-9  9 788328 357679 Cena: 129,00 zł
---	---	---

WILEY