

IDŹ DO

PRZYKŁADOWY ROZDZIAŁ



SPIS TREŚCI

KATALOG KSIĄŻEK

KATALOG ONLINE

ZAMÓW DRUKOWANY KATALOG

TWÓJ KOSZYK

DODAJ DO KOSZYKA

CENNIK I INFORMACJE

ZAMÓW INFORMACJE
O NOWOŚCIACH

ZAMÓW CENNIK

CZYTELNIA

FRAGMENTY KSIĄŻEK ONLINE

Podstawy ochrony komputerów

Autorzy: Rick Lehtinen, Deborah Russell, G T Gangemi

Tłumaczenie: Julia Szajkowska

ISBN: 978-83-246-0780-8

Tytuł oryginału: [Computer Security Basics](#)

Format: B5, stron: 304



Zadbaj o bezpieczeństwo swojego komputera

- Poznaj zagrożenia, na jakie narażony jest komputer
- Naucz się kontrolować dostęp do komputera
- Stosuj techniki zapewniające bezpieczeństwo w sieci

Czy mój komputer na pewno jest bezpieczny? Wiele osób zadaje sobie to pytanie dopiero w momencie, kiedy system zaczyna zachowywać się w podejrzany sposób. Okazuje się wówczas, że skaner wykrywa dziesiątki, a nawet setki wirusów, programy zaczynają działać nieprawidłowo, a z dysku giną ważne dane. Pół biedy, jeśli jest to tylko domowy komputer z prywatnymi plikami. Dużo gorsze skutki może mieć włamanie do firmowej bazy danych lub przechwycenie poufnej komunikacji.

Książka „Podstawy ochrony komputerów” to wszechstronne wprowadzenie do najważniejszych zagadnień dotyczących bezpieczeństwa danych i sprzętu. Czytając ją, poznasz zagrożenia, jakie czyhają na użytkowników komputerów, ale także skuteczne techniki ochrony. Nauczysz się kontrolować dostęp do danych, prowadzić efektywną politykę zabezpieczeń, wykrywać i usuwać wirusy oraz zapobiegać przenikaniu ich do systemu. Dowiesz się, jak zapewnić bezpieczeństwo komputera w sieci oraz jak używać szyfrowania do przesyłania poufnych informacji. Przeczytasz też o najnowszych technikach zabezpieczenia bazującego na danych biometrycznych (wzorce siatekówek czy odciskach palców) oraz ochronie sieci bezprzewodowych.

- Niebezpieczeństwa grożące użytkownikom komputerów
- Kontrolowanie dostępu do komputera
- Walka z wirusami
- Prowadzenie skutecznej polityki zabezpieczeń
- Bezpieczne korzystanie z sieci
- Szyfrowanie poufnych danych
- Komunikacja bez ryzyka
- Zabezpieczenia biometryczne
- Tworzenie bezpiecznych sieci bezprzewodowych

**Stosuj skuteczne zabezpieczenia i zapewnij
maksymalne bezpieczeństwo swojemu komputerowi!**

Wydawnictwo Helion
ul. Kościuszki 1c
44-100 Gliwice
tel. 032 230 98 63
e-mail: helion@helion.pl



Spis treści

Przedmowa	7
Część I Bezpieczeństwo dzisiaj	11
1. Wstęp	13
Nowe zagrożenie	13
Czym jest bezpieczeństwo komputera?	19
Zagrożenia	23
Dlaczego należy kupować zabezpieczenia	29
Co ma zrobić użytkownik?	31
Podsumowanie	32
2. Krótka historia zabezpieczeń	33
Informacja i kontrola nad nią	33
Ochrona komputera — dawniej i dziś	35
Wczesne próby ochrony komputerów	38
Krok w stronę standaryzacji	42
Ustawodawstwo i pełnomocnictwa dotyczące ochrony komputerów	48
Podsumowanie	56
Część II Ochrona komputera	57
3. Ochrona systemu komputerowego i kontrola dostępu	59
Po czym poznać bezpieczny system?	59
Dostęp do systemu — logowanie	60
Podsumowanie	87

4. Wirusy i inne przejawy dzikiego życia	89
Koszty ponoszone w związku z działaniem szkodliwego oprogramowania	89
Wirusy a zdrowie publiczne	90
Wirusy, robaki i konie trojańskie (ojej!)	90
Kto pisze wirusy?	100
Remedium	102
Rozdmuchana sprawa wirusów	103
Odrobina profilaktyki	104
Podsumowanie	104
5. Ustanowienie i utrzymanie polityki zabezpieczeń	107
Zabezpieczenia administracyjne	108
Ogólnie rozumiane planowanie i administracja	109
Codzienna administracja	114
Podział obowiązków	120
Podsumowanie	121
6. Ataki przez sieć i słabe punkty internetu	123
Internet	123
Czym są protokoły sieciowe?	127
Delikatna sieć	134
Podsumowanie	143

Część III Zabezpieczenia komunikacyjne 145

7. Szyfrowanie	147
Odrobina historii	148
Czym jest kodowanie?	150
Standard kodowania danych	161
Inne algorytmy kryptograficzne	171
Uwierzytelnianie wiadomości	176
Rządowe programy kryptograficzne	177
Ograniczenia eksportowe	179
Podsumowanie	179
8. Komunikacja i zabezpieczenia sieci	181
Kiedy łączność jest bezpieczna?	182
Modemy	185
Sieci	186
Ochrona sieci	194
Podsumowanie	206

Część IV Inne rodzaje zabezpieczeń	207
9. Zabezpieczenia fizyczne i biometryka	209
Zabezpieczenia fizyczne	210
Zamki i klucze — wczoraj i dziś	213
Biometryka	218
Delikatne przypomnienie	224
Podsumowanie	225
10. Zabezpieczenia sieci bezprzewodowej	227
Jak się tu dostaliśmy?	227
Infrastruktura bezprzewodowa w dzisiejszych czasach	228
Jak działa sieć bezprzewodowa?	232
Zabawa z polami	235
O co chodzi z tymi decybelami?	239
Dlaczego jest to tak ważne?	239
Zalecam odbiór zbiorczy	239
Bezprzewodowe ataki w warstwie fizycznej	240
Podsumowanie	251

Część V Dodatki	253
A Model OSI	255
B Tempest	259
C Orange Book, FIPS PUBS i Common Criteria	265
Skorowidz	289

Ochrona systemu komputerowego i kontrola dostępu

Ochrona komputera to pojęcie, które dotyczy zagadnień związanych z zapewnieniem właściwego zamknięcia pomieszczeń serwerowni i centrali telefonicznej, zabezpieczenia komputera przed niepowołanym dostępem, ochrony kont odpowiednio silnymi hasłami, zabezpieczania plików i trzymania się terminów wykonywania kopii zapasowej, która chroni dane przed unicestwieniem, szyfrowania wiadomości przesyłanych w sieci i wreszcie stosowania osłon, które zmniejszą prawdopodobieństwo przechwycenia danych poprzez odczyt zmian pola elektromagnetycznego, emitowanego przez urządzenia (system TEMPEST). Ale ludzie, mówiąc o ochronie komputera, mają z reguły na myśli tak zwaną **ochronę systemu komputerowego**, co jest synonimem ochrony danych.

Po czym poznać bezpieczny system?

Najbardziej intuicyjna definicja zabezpieczonego systemu komputerowego to stwierdzenie, że Twój komputer nie robi niczego, co do niego nie należy — nawet jeśli użytkownicy nie wypełniają swoich obowiązków w tym względzie. Bezpieczny system nie gubi danych w nim zapisanych, nie pozwala na złośliwe lub przypadkowe ich zmienianie ani nie zezwala na odczyt i modyfikacje tych danych przez ludzi do tego nieupoważnionych.

W jaki sposób działa ochrona systemu? Istnieją cztery główne sposoby zabezpieczania:

Kontrola dostępu do systemu

Metody kontroli dostępu zapewniają, że żaden nieupoważniony użytkownik nie dostanie się do systemu, a także zachęcają (lub wręcz zmuszają) uprawnionych użytkowników do świadomego zwiększania poziomu zabezpieczeń — jednym ze sposobów jest regularna zmiana haseł. Poza tym, system chroni hasła i zapisuje dane dotyczące działalności poszczególnych użytkowników w systemie, szczególnie jeśli ich działania wkraczają na pole zabezpieczeń (na przykład prowadzi dzienniki wejścia do systemu, otwarcia plików, korzystania ze specjalnych uprawnień). Kontrola dostępu do systemu to serce procesu uwierzytelniania.

Kolejny dział obejmuje podstawy zagadnienia kontroli dostępu. W dodatku C znajduje się opis wymogów odpowiedzialności, jakie zawarte są w dokumencie *Orange Book*. Określają one metody kontroli dostępu, właściwe dla danego poziomu bezpieczeństwa systemu. Dokument *Orange Book*, choć został już zastąpiony przez kolejny — *Common Criteria*, nadal jest ważnym źródłem danych dotyczących bezpieczeństwa.

Kontrola dostępu do danych

Zbiór tych metod określa, które osoby mają mieć dostęp do określonych danych i jakie mogą podejmować działania. Innym terminem określającym ten typ kontroli dostępu jest **autoryzacja**, która definiuje uprawnienia użytkownika po zakończeniu procesu uwierzytelniania. System, z którym pracujesz, może obsługiwać swobodną kontrolę dostępu (DAC — *przyp. tłum.*), zezwalającą Ci na określenie, czy inni użytkownicy mają prawo odczytywać lub zmieniać Twoje dane. Oczywiście system może korzystać z narzuconej kontroli dostępu¹. Wtedy to system określa reguły dostępu, bazując na poziomach bezpieczeństwa przypisanych poszczególnym użytkownikom, plikom i innym elementom systemu. Kontrola dostępu oparta o role² jest systemem hybrydowym, który rozszerza proces indywidualnej autoryzacji na grupy użytkowników.

Zarządzanie systemem i zabezpieczeniami

Ta grupa metod, niezwiązanych z pracą w sieci, w praktyce tworzy lub przełamuje system zabezpieczeń, jasno określając obowiązki administratora systemu, ucząc jego użytkowników odpowiednich zachowań i pilnując, aby stosowali się do polityki ochrony firmy. Do niej także zaliczamy ogólniej rozumiane zarządzanie zabezpieczeniami, jak choćby określenie zagrożeń, na jakie narażony jest Twój system, i kosztu stworzenia systemu, który Cię przed nimi uchroni.

Rozdział 5. poświęcony został podstawom planowania i administracji ochroną systemu. W dodatku C znajdziesz informacje dotyczące wymagań, jakie poszczególnym poziomom zabezpieczeń stawia dokumentacja *Orange Book*.

Projektowanie systemu

Te sposoby zabezpieczeń korzystają z podstawowych możliwości, jakie dają oprogramowanie i urządzenia, których używamy. Przykładem może tu być architektura systemu, dopuszczająca segmentację pamięci, co pozwala oddzielić procesy uprzywilejowane od tych nieuprzywilejowanych.

I chociaż dogłębna analiza projektu bezpiecznego systemu wykracza poza pole tematyczne tej książki, w dodatku C wskazaliśmy podstawowe wymagania *Orange Book* wobec systemów o różnych poziomach zabezpieczeń.

Dostęp do systemu — logowanie

System kontroluje dostęp do swoich zasobów, i jest to pierwszy etap zapewnienia ochrony komputera. Kto może wejść do systemu? Jak system zadecyduje, czy użytkownik ma prawo z niego korzystać? W jaki sposób wreszcie pilnować działania każdego z użytkowników?

¹ MAC — *przyp. tłum.*

² RBAC — *przyp. tłum.*

Próba wejścia do systemu jest swego rodzaju realizacją scenariusza wezwanie-odpowiedź. Ty mówisz systemowi, kim jesteś, a on żąda, byś potwierdził swoją tożsamość, dostarczając mu informacji, które będą odpowiadać tym zapisanym uprzednio w jego pamięci. Fachowe określenie tego dwustopniowego procesu to **identyfikacja i uwierzytelnianie**.

Identyfikacja i uwierzytelnianie

Identyfikacja to sposób na poinformowanie systemu, kim jesteś. **Uwierzytelnianie** ma z kolei dowiedzieć, że faktycznie jesteś tym, za kogo się podajesz. W każdym systemie wielodostępowym, w którym użytkownicy pracują w sieci lokalnej, oraz w większości komputerów osobistych i przenośnych musisz dokonać identyfikacji, a system musi sprawdzić wiarygodność Twojej deklaracji, zanim zezwoli Ci na pracę. Istnieją trzy klasyczne już metody postępowania w takim przypadku.

Co wiesz?

Najlepiej znanym z nich jest podanie hasła. Teoretycznie, jeżeli znasz tajne hasło do swojego konta, to jesteś jego właścicielem. W praktyce nie jest to tak proste, ponieważ mogłeś udostępnić komuś hasło albo ukradziono Ci je. Jeżeli gdzieś je zapisałeś, istnieje szansa, że ktoś je przeczytał. Jeżeli zdradzisz je komuś, to ta osoba może przekazać je dalej. Hasło zbyt proste można łatwo odgadnąć lub złamać je.

Co masz?

Przykładami są klucze, tokeny, odznaki i karty magnetyczne, bez których nie masz dostępu do swojego konta. Znow, zgodnie z teorią, posiadacz klucza jest jednoznaczny z właścicielem konta. Problemem jest to, że klucz może zostać skradziony, czy też możesz go pożyczyć komuś, kto zrobi jego kopię. Klucze elektroniczne, odznaki i karty magnetyczne są pewną formą uwierzytelniania i pozwalają uzyskać dostęp do budynków i pomieszczeń z komputerami. Najbardziej zaawansowane technologicznie tokeny to urządzenia, które stale obliczają nowe hasła, bazując na bieżącej dacie i korzystając z algorytmów zabezpieczających. Te same hasła są obliczane przez system. Hasło użytkownika starającego się o dostęp do systemu musi zgadzać się z hasłem obliczonym przez system.

Kim jesteś?

Formą potwierdzenia Twojej tożsamości mają być zachowania i pewne cechy fizjologiczne — odciski palców, dłoni, wzór siatkówki czy tęczówki, próbki głosu, podpis, kolejność wciskanych klawiszy. Systemy biometryczne porównują Twoje cechy osobnicze z danymi, które mają zapisane, i na tej podstawie weryfikują Twoją tożsamość. Czasami pojawiają się problemy z niepoprawną pozytywną i negatywną identyfikacją; zdarza się, że pełnoprawny użytkownik zostanie odrzucony, a osoba bez uprawnień dopuszczona do systemu. Istnieje jeszcze jeden poważny problem z tym rodzajem uwierzytelniania — ludzie nie czują się swobodnie, stając w obliczu wspomnianych metod weryfikacji.

Uwierzytelnianie wielostopniowe

Uwierzytelnianie wielostopniowe to połączenie wspomnianych powyżej metod. W ten sposób, jeśli osoba zagrażająca Twoim danym przedrze się przez pierwszy stopień zabezpieczeń, nadal musi pokonać kolejne. Hasła są nadal zdecydowanie najczęściej wybieraną formą zabezpieczeń. W wielostopniowym systemie zabezpieczeń podaniu nazwy użytkownika i hasła

towarzyszyłyby inne formy identyfikacji. W rzeczywistości tokeny i urządzenia biometryczne raczej nie zastępują konwencjonalnych nazw użytkownika i haseł, a bardziej są do nich dodatkiem. Więcej informacji na ten temat znajduje się w rozdziale 9.

Proces logowania

Większość systemów wymaga od użytkownika podania unikatowej nazwy użytkownika lub innego identyfikatora oraz hasła. Identyfikator to zwykle Twoje imię, inicjały, nazwa działu lub numer konta (bazujący na Twoim imieniu i (lub) grupie zaszerzegowania) przypisany przez administratora systemu. Hasło składa się przeważnie z ciągu liter i (lub) cyfr i powinno być znane tylko Tobie.

Konkretne rozwiązania uwierzytelniania, oparte o podawanie nazwy użytkownika i hasła, różnią się w zależności od systemu operacyjnego. Jednak zawsze w jego skład wchodzi przynajmniej jedna z podanych poniżej metod.

Szyfrowanie

Można zaszyfrować hasło tak, aby nikt niepowołany nie był w stanie go poznać, nawet jeżeli będzie śledził transmisję czy badał przechowywane dane. Przeciwnieństwem szyfrowania jest wysyłanie wiadomości **jawnym tekstem**. W takim razie hasło lub inne informacje są wysyłane bez żadnych modyfikacji.

Wyzwanie-odpowiedź

Ta metoda wymaga uwierzytelniania na początku transmisji, a żądanie to jest powtarzane podczas jej trwania w losowo określonych odstępach czasu.

Kolejne działy poświęciłem opisowi kilku modeli tych mechanizmów.

Password Authentication Protocol

Protokół uwierzytelniania hasłem (PAP) zobowiązuje użytkownika do podania jego nazwy i hasła, które są następnie porównywane z wartościami zapisanymi w tablicy. Protokół PAP w zasadzie nie różni się niczym od klasycznego logowania do systemu Unix. Informacje o hasle i nazwie użytkownika są przekazywane jawnie, bez szyfrowania.

Challenge Handshake Authentication Protocol (CHAP)

Protokół CHAP jest takim rodzajem uwierzytelniania, w którym urządzenie przeprowadzające tę procedurę (zwykle jest to serwer sieciowy) wysyła programowi klienckiemu numer identyfikacyjny i losowo wygenerowaną wartość. Zarówno nadawca, jak i odbiorca mają wcześniej ustalone tajne hasło. Klient łączy to hasło z wartością losową (lub zlepkiem liter) i numerem identyfikacyjnym i oblicza na tej podstawie nową wartość, korzystając z tak zwanej **funkcji mieszającej**. Tak obliczona wartość jest wysyłana do urządzenia uwierzytelniającego, który posiada ten sam łańcuch znaków i oblicza wartość. Teraz następuje porównanie obu wartości, i jeśli są one zgodne, program kliencki jest dopuszczany do serwera. Aby zwiększyć jeszcze bezpieczeństwo, urządzenie uwierzytelniające może dodatkowo wykonywać okresowo sprawdzanie typu wyzwanie-odpowiedź.

Uwierzytelnianie wzajemne

Uwierzytelnianie wzajemne jest procesem dwustronnym. Klient dokonuje uwierzytelnienia na serwerze, a serwer „przedstawia się” klientowi lub stacji roboczej. W ten sposób serwer sprawdza, czy użytkownik korzysta z autoryzowanego stanowiska pracy. Jeśli tak nie jest, serwer nie zezwala na dostęp. Uwierzytelnianie wzajemne zapobiega **atakom maskaradowym**, w których atakujący podszywa się pod uprawnionego użytkownika, aby dostać się do danych, a także **atakami ze spotkaniem w środku**, w których atakujący przedstawia się serwerowi jako uprawniony użytkownik, a użytkownikowi jako jego serwer. Mając już połączenie z obydwojema, czerpie z informacji przesyłanych przez obie strony lub wysyła im szkodliwe dane.

Hasło jednorazowe

Hasło jednorazowe (OTP — ang. *one-time password* — *przypp. tłum.*) jest tak naprawdę odmianą układu nazwa użytkownika-hasło. W tej formie uwierzytelniania użytkownik tworzy sobie hasło, a przy kolejnych próbach potwierdzania tożsamości użytkownika używane są jego wariacje. W ten sposób nigdy nie korzystamy z tego samego hasła. Nawet jeśli agresor pozna nasze hasło, nie będzie mógł z niego ponownie skorzystać.

Uwierzytelnianie przed każdą sesją

Używanie metody, która wymaga od użytkownika potwierdzenia swojej tożsamości przed każdą wymianą informacji, jest nużące, ale jest jedną z lepszych form ochrony. Jednym ze sposobów jej realizacji jest zwiększanie wartości licznika przy każdej transmisji danych. Ponieważ hasło ulega ciągłym zmianom, użytkownik jest zabezpieczony przed podsłuchującymi i podglądającymi go osobami niepowołanymi.

Tokeny

Token jest formą wzmocnienia ochrony, ponieważ dostarcza do naszych zabezpieczeń warstwę „co masz?”. Jest to z reguły niewielkie urządzenie, które dostarcza odpowiedzi na wyzwanie rzucone w momencie próby logowania.

Token może być rozmiarów karty kredytowej i mieć wbudowaną klawiaturę. Przy próbie wejścia do systemu serwer generuje wyzwanie, liczbę losową. Użytkownik wprowadza ją do karty tokenu, a on wyświetla odpowiedź, którą należy wprowadzić do systemu. Serwer ma już obliczoną swoją wartość i porównuje ją z wynikiem otrzymanym przez token. Jeśli liczby są identyczne, to użytkownik zostaje uwierzytelniony.

Inne tokeny stosują rozwiązania oparte o aktualny czas. Liczba, którą wyświetlają, zmienia się w regularnych odstępach czasu, z reguły kilka razy w ciągu godziny. Wejście do systemu następuje po podaniu nazwy użytkownika, hasła i wartości wygenerowanej przez token w danej chwili. Jeśli odpowiada ono temu wyliczonemu przez system oraz jeśli konto zostanie uwierzytelnione, użytkownik otrzymuje dostęp do danych.

Poważną wadą tokenów są ich niewielkie rozmiary i stosunkowo wysoka cena. W razie awarii czy zgubienia trzeba będzie go zastąpić nowym urządzeniem. W zależności od producenta i jakości wyrobu ceny tokenów utrzymują się w przedziale od 90 do 350 zł za sztukę.

Pewną alternatywą są tokeny programowe. Podczas próby podłączenia się do systemu użytkownik wprowadza swój PIN, dla którego token tworzy hasło jednorazowe. Kod PIN nigdy nie jest przesyłany. Token programowy przestaje działać, jeśli kilkakrotnie otrzyma nieprawidłowy kod.

Urządzenia biometryczne

Ta klasa urządzeń wykorzystuje indywidualne cechy użytkownika, odciski palców, kontur dłoni, wzór siatkówki lub tętnówki, próbki głosu, pisma czy sposób wciskania klawiszy. Używa się ich, jednak rzadko, jako jednego ze stopni zabezpieczeń. Ta forma uwierzytelniania jest najodpowiedniejsza, jeśli w systemie ochronnym, poza określeniem „co wiesz” i „co masz”, wprowadzamy warstwę „kim jesteś”. Urządzenia biometryczne mają skłonność do wykonywania fałszywych identyfikacji pozytywnych (stwierdzają, że jesteś kimś, kim nie jesteś) lub negatywnych (mówią, że nie jesteś tym, kim w rzeczywistości jesteś), więc większość systemów posiada wbudowany współczynnik wiarygodności, obliczany na podstawie tych dwóch wartości. Sprzęt tej klasy może sprawdzać się doskonale do sprawdzania nazwy użytkownika czy hasła, więc należy się spodziewać wzrostu zakresu jego użycia. Urządzenia biometryczne opiszę szerzej w rozdziale 9.

Dostęp zdalny (TACACS i RADIUS)

Pomimo ciągłego wzrostu popularności połączeń szerokopasmowych, takich jak sieci kablowe i DSL, dostęp do sieci typu *dial-up* nadal cieszy się sporą popularnością. Aby użytkownik zdalny mógł połączyć się z siecią, w wielu firmach pojawiają się banki modemów. Takie połączenia muszą być dobrze zabezpieczone i nadzorowane, a użytkownikom potrzebne są proste metody uwierzytelniania, żeby szybko mogli dostać się do potrzebnych im informacji. Ale zamiast setek, a nawet tysięcy potencjalnych klientów, każdy, kto posiada telefon i komputer, staje się potencjalnym użytkownikiem.

Naprzeciw tej potrzebie wychodzą specjalistyczne systemy uwierzytelniania. Oto dwa spośród nich:

- Remote Authentication Dial-In User Service (RADIUS).
- Terminal Access Controller Access Control System (TACACS).

Obsługa dużej liczby wejść do systemu może skutecznie zablokować każdy serwer. Podane protokoły potrafią przenieść procesy uwierzytelniania i autoryzacji z serwera sieciowego na serwer centralny. Przeciętna sieć w firmie ma serwer dostępowy połączony z pulą modemów, obsługujących połączenia przychodzące, dla których usługi uwierzytelniania świadczą serwery TACACS i RADIUS. Użytkownik zdalny łączy się z serwerem dostępowym, który wysyła żądanie uwierzytelnienia do jednego z serwerów TACACS lub RADIUS. One dokonują rozpoznania użytkownika i dają mu dostęp do wewnętrznych zasobów sieci. Użytkownicy zdalni są klientami serwerów dostępowych, a serwery dostępne są klientami serwera RADIUS.

Poza uwierzytelnianiem użytkownika, oba serwery mogą śledzić jego poczynania, sprawdzać czas trwania sesji, protokoły i porty, z jakich korzystał użytkownik podczas niej, adresy, z jakimi się łączył, a nawet przyczynę przerwania sesji. Opisane serwery tworzą dzienniki aktywności w sieci, więc można bez problemu sprawdzić, kto otwierał poszczególne sesje. Wszystkie potrzebne informacje dotyczące sesji można przedstawić w formie arkusza i dokładnie zbadać. To wszystko sprawia, że często procesy uwierzytelniania i autoryzacji są przekazywane do serwerów TACACS i RADIUS.

Dokumentacja TACACS jest dostępna w RFC 1492. RADIUS jest opisany w RFC 2139 (*RADIUS Accounting*, kwiecień 1997) i RFC2865 (*Remote Authentication Dial-In User Service (RADIUS)*, czerwiec 2000).

DIAMETER

DIAMETER to protokół, który dokonuje uwierzytelnienia użytkowników łączących się przez linię telefoniczną oraz dokonuje ich autoryzacji i pozwala rozliczać zasoby sieciowe. Protokół DIAMETER wywodzi się ze wspomnianego wcześniej protokołu RADIUS, który miał ograniczenia co do sposobu połączenia. RADIUS pracował tylko z protokołami modemowymi, takimi jak SLIP (ang. *Serial Line Interface Protocol*) i PPP (ang. *Point to Point Protocol*). Dawniej wystarczyło to w zupełności; dziś użytkownicy korzystają z telefonów komórkowych, które mogą łączyć się z internetem³, oraz z palmtopów i innych urządzeń kieszonkowych, które obsługują różne protokoły. Stąd wyniknęła potrzeba wprowadzenia bardziej elastycznego narzędzia.

Protokół DIAMETER wprowadza nowe możliwości do protokołu RADIUS — teraz można prosić o dodatkowe informacje dotyczące logowania, wykraczające poza podstawowe uwierzytelnianie. DIAMETER, poza przeprowadzeniem standardowej weryfikacji nazwy użytkownika i hasła, prosi o udzielenie dodatkowych informacji, które wspomagają proces uwierzytelniania albo dają użytkownikowi dostęp do innych części systemu.

Ten protokół obsługuje wymienione poniżej rozszerzenia usług sieciowych:

Operacje roamingowe (ROAMOPS — Roaming Operations)

To procedury, mechanizmy i protokoły, które zapewniają przełączanie użytkownika pomiędzy poszczególnymi grupami dostawców usług internetowych (ISP — ang. *Internet Service Provider*). Operacje o zasięgu globalnym, które wykonywane są z jednego konta, powinny opierać się o aplikacje obsługujące *roaming*.

Wymagania serwera dostępowego (NASREQ — Network Access Server Requirements)

Jest to rozszerzenie serwera dostępowego, pozwalające mu obsługiwać nie tylko zwykłe połączenia typu *dial-up*, ale również dostęp do prywatnych sieci wirtualnych (VPN — ang. *Virtual Private Network*), oraz poprawiające jakość uwierzytelniania i przełączania między sieciami (*roamingu*).

Routing adresu IP węzłów bezprzewodowych i przenośnych (MobileIP — IP Routing for Wireless/Mobile Hosts)

Rozwój jakże potrzebnej usługi routingu pozwala przełączać adresy IP dzięki IPv4 (ang. *Internet Protocol ver. 4 — przyp. tłum.*) lub IPv6 (ang. *Internet Protocol ver. 6 — przyp. tłum.*) między podsieciami i różnymi nośnikami.

Protokół DIAMETER został opisany w wersjach roboczych standardów IETF (ang. *Internet Engineering Task Force*):

RFC 3589: Kody sterujące w protokole Diameter dla 3GPP, wersja 5, wrzesień 2003 (ang. *Diameter Command Codes for 3GPP Release 5*).

RFC 3588: Protokół Diameter, wrzesień 2003 (ang. *Diameter Based Protocol*).

RFC 2486: Identyfikator dostępu do sieci, styczeń 1999 (ang. *Network Access Identifier*).

RFC 2607: Łączenie serwerów proxy i polityka roamingu, czerwiec 1999 (ang. *Proxy Chaining and Policy in Roaming*).

³ Ang. *smart phone* — przyp. tłum.

Kerberos

W 1980 roku w *Massachusetts Institute of Technology* (MIT) powstał jeden z najważniejszych i najbardziej złożonych schematów uwierzytelniania, który nazwano **Kerberos**. Dzięki niemu można bezpiecznie przesyłać dane w tak niebezpiecznym ośrodku, jakim jest internet. Nazwano go na cześć mitologicznego, trójąłowego psa, strzegącego bram Hadesu — Cerbera. Proces wejścia do systemu przebiega w trzech etapach:

1. Użytkownik dostarcza nazwy i hasła. Żądanie uzyskania dostępu do systemu zostaje przesłane do **serwera uwierzytelniania (AS — authentication server)**, który dokonuje weryfikacji tożsamości użytkownika. Kiedy serwer odbiera to żądanie, automatycznie tworzy klucz w dwóch kopiach, tzw. **klucz sesji**. Klucze sesji służą do wymiany informacji między użytkownikiem a zasobem, do którego próbuje on otrzymać dostęp.

Jedna kopia klucza sesji jest szyfrowana kluczem, który przechowuje użytkownik. Drugą koduje się przy pomocy klucza przechowywanego przez serwer. Klucz serwera nazywamy **biletem**. Serwer AS wysyła bilet razem z kluczem użytkownika z powrotem do użytkownika, który otwiera jedną z kopii hasłem, jakim jest jego własny klucz. W ten sposób otrzymuje klucz sesji, po czym może sprawdzić jego poprawność. Użytkownik ma dostęp jedynie do kopii zaszyfrowanej własnym kluczem. Druga została przecież zabezpieczona kluczem serwera.

2. Teraz użytkownik tworzy kolejną wiadomość, tzw. **wartość uwierzytelniającą**, do której dodaje aktualną godzinę i tworzy **sumę kontrolną** tego wyrażenia. Suma kontrolna to forma matematycznej weryfikacji dokonywana na serii znaków. Znaki są do siebie dodawane i łączone w szereg. Kolejni użytkownicy powtarzają proces obliczania sumy kontrolnej i porównują otrzymany wynik z pierwotną wartością. Jeśli sumy są różne, oznacza to, że niektóre znaki uległy zmianie podczas przechowywania lub transmisji, więc na wszelki wypadek odrzuca się całość danych. Użytkownik koduje sumę kluczem sesji wysyła bilet i wartość uwierzytelniającą na serwer, który ma dostarczyć mu żądanych zasobów.

3. Serwer odbiera obie zaszyfrowane kopie klucza sesji, korzysta z biletu (klucza serwerowego) do otwarcia kopii nim zaszyfrowanej, wypakowuje klucz sesji i sprawdza, czy pochodzi on od użytkownika, oceniając to na podstawie obecności klucza użytkownika. Druga wiadomość jest otwierana kluczem sesji. Teraz serwer dokonuje porównania oznaczenia czasu i sumy kontrolnej z bieżącym czasem. W ten sposób oceniana jest spójność odpowiedzi i jej autor. Jeśli wszystkie wartości są poprawne, użytkownik otrzymuje prawo dostępu do serwera.

Proste, prawda? W rzeczywistości system Kerberos jest trochę bardziej skomplikowany, ale ogólnie powyższe wyjaśnienia zamykają temat. Kerberos jest odpowiedzią na rozwój technologii. Większość sieci skupiała się na zagadnieniu komunikacji, pomijając kwestie bezpieczeństwa. Większość danych była przesyłana jawnie. Pojawienie się narzędzi takich jak pakiety dało operatorom możliwość obserwacji każdego bitu nieszyfrowanych danych, jaki przepływał po łączach, co było niedopuszczalne. Kerberos stara się rozwiązać ten problem, dając nam do ręki potężne narzędzie szyfrujące i pewne uwierzytelnianie. Pomimo swojej złożoności stał się wzorem, według którego powstawały kolejne rozwiązania z tej dziedziny.

Rzetelnych informacji na temat systemu Kerberos dostarczają strona internetowa instytutu MIT (<http://web.mit.edu/kerberos/www/>) oraz źródła RFC organizacji IETF (<http://www.ietf.org/rfc/rfc1510.txt>).

Hasła

Pomimo istnienia szerokiej gamy zabezpieczeń nadal największą popularnością cieszy się uwierzytelnianie hasłem i nazwą użytkownika, które pojawia się w większości systemów Unix, Linux i Windows. Na przykład Unix wyświetla zgłoszenie:

login:

po czym oczekuje imienia, nazwiska lub ustalonej wcześniej kombinacji ich obu, na przykład inicjału imienia, po którym nastąpi maksymalnie siedem liter nazwiska czy innego, dowolnego identyfikatora, którym Cię obdarzono.

Po podaniu nazwy użytkownika jesteś proszony o wpisanie hasła:

Password:

Wpisujesz hasło (które najczęściej w żaden sposób nie jest przedstawiane na ekranie, chyba że w postaci szeregu gwiazdek, których liczba może, ale nie musi odpowiadać liczbie znaków hasła). System sprawdza Twoją tożsamość na podstawie wprowadzonych danych. Praca z systemem stanie się możliwa dopiero wtedy, kiedy hasło będzie odpowiadało przechowywanemu przez system.

Wskazówki dotyczące ochrony hasła

Odpowiedzialność za zapewnienie właściwej ochrony hasła spoczywa zarówno na administratorsze sieci, jak i na użytkowniku. Pamiętaj, zabezpieczenie hasła leży w interesie wszystkich. Osoba niepowołana, która pozna Twoje hasło, może zniszczyć Twoje pliki, ale również zagrozić pozostałym danym w systemie.

- Nie zezwalaj na dostęp do systemu niezabezpieczony hasłem. Jeśli jesteś administratorem systemu, upewnij się, że każdemu kontu odpowiada hasło.
- Nie zachowuj domyślnych hasła systemowych. Zmieniaj wszystkie hasła testowe, instalacyjne oraz gości — na przykład *root*, *system*, *test*, *demo* i *guest* — zanim zezwolisz innym użytkownikom na pracę w nim.
- Nigdy nie ujawniaj nikomu swojego hasła. Jeśli nie będziesz miał wyjścia, bo na przykład będziesz w domu, a ktoś w pracy musi z niego skorzystać, zmień je potem tak szybko, jak będzie to możliwe, albo poproś administratora o utworzenie nowego, zanim nie zmienisz swojego.
- Nie zapisuj nigdzie hasła, a szczególnie nie w pobliżu komputera, terminalu czy swojego biurka. Jeśli je zanotujesz, nie pisz przy nim, że to hasło. Najlepiej dodaj do niego, na początku lub końcu, kilka zbędnych znaków, a obok zapisz kilka innych potencjalnych hasła. Dobrym pomysłem jest zapisanie go od końca.
- Nie wpisuj hasła, gdy ktoś patrzy.
- Nie zapisuj hasła w sieci ani nie wysyłaj go pocztą elektroniczną. Cliff Stool w swojej książce *The Cuckoo's Egg* opisuje, jak napastnik przeszukał wiadomości jego poczty pod kontem słowa „hasło”.
- Nie używaj ciągle tego samego hasła. Nawet jeśli nic mu nie zagraża, zmieniaj je regularnie.

Sieć USENET podaje bardzo dobre przysłowie na tę okazję: „Hasło powinno być jak szczoteczka do zębów. Korzystaj z niego codziennie, zmieniaj regularnie i nie dziel się nim z przyjaciółmi”.

Hasła stanowią pierwszą linię zabezpieczeń przed atakiem. Aby chronić swój system i dane, musisz wybrać sobie odpowiednio silne hasła i chronić je skutecznie.

Wszystkie powyższe uwagi tyczą się komputera, na którym pracuje tylko jedna osoba. Jeśli z komputera korzysta wielu użytkowników, na przykład pracownicy na różnych zmianach albo kierownicy, którzy posiadają własne konta na Twoim sprzęcie, dzięki czemu mogą wspomagać właściwe jego utrzymanie, to każde z tych haseł jest narażone na atak.

Wskazówki dotyczące wyboru hasła

Jeśli masz możliwość samodzielnie ustalić swoje hasło, wybierz takie, które niełatwo będzie odgadnąć. Oto kilka rad:

- Wybieraj hasła, które nie są słowami (w żadnym z języków) czy imionami (szczególnie Twoimi ani należącymi do postaci fikcyjnych, jak Hamlet czy Gandalf, ani też członków Twojej rodziny lub zwierzaków).
- Buduj hasła złożone ze zlepeków liter i cyfr. Nigdy nie używaj hasła złożonego z samych cyfr (szczególnie unikaj numeru telefonu czy numeru PESEL).
- Wybieraj długie hasła. Jeśli hasło ma tylko kilka znaków, atakujący z łatwością sprawdzi wszystkie kombinacje. Większość systemów wymaga haseł o długości przynajmniej od 6 do 8 znaków. Niektóre z nich dopuszczają użycie haseł długich nawet na 40 znaków.
- Stosuj różne hasła dla każdej maszyny i węzła dostępowego, z którymi pracujesz.
- Wprowadzenie do hasła znaków specjalnych, czyli & lub \$ oraz innych, zwiększa poziom bezpieczeństwa hasła, ponieważ atakujący ma do wyboru większą liczbę znaków. Uważaj jednak z użyciem cyfr, które przypominają litery, ponieważ agresor doskonale zna te nawyki. Z tego samego powodu omijaj mowę 31337⁴. Dowiedz się od administratora, których znaków wolno Ci używać.

Najlepsze hasła składają się z wielkich i małych liter, znaków specjalnych i cyfr. Hasło wcale nie musi być bezsensowne. Hasła bez znaczenia są często zapisywane, co niszczy efekt starannego doboru znaków je tworzących. Znów podsuwamy kilka pomysłów:

- Łącz krótkie słowa ze znakami specjalnymi i cyframi, na przykład Mam3psa.
- Używaj akronimów zdań, które zapamiętasz. Wybieraj zdania, których nie będzie można rozpoznać. Prosty przykładem jest: „O nie, zapomniałem tego zrobić”, czyli Onztz.
- Zwiększ bezpieczeństwo, dodając kilka znaków specjalnych lub cyfr: :0nz;tz czy 0n5zt0z.
- Wybierz bezsensowne słowo, które jednak da się wymówić, na przykład 8Bektag czy szmoaz12.

⁴ Hack mowa — wbrew nazwie nie jest to żargon używany przez hakerów, ale oparty na języku angielskim slang rozpowszechniony na różnego rodzaju chatach sieciowych, IRC, forach — popularny także wśród nastoletnich włamywaczy komputerowych, krakerów oprogramowania, piratów softwarowych i fanów FPS. W żargonie tym słowa zapisuje się przy pomocy cyfr i innych kodów ASCII (np. 1337 zamiast leet wymawiane jak el eat, tj. elite), fonetycznie, wykorzystując brzmienie zbliżone do angielskich słów (np. „b”, czytane jak nazwa litery zamiast angielskiego czasownika *be* — być), opuszczając niewymawiane litery, robiąc celowe błędy typograficzne (np. „pr0n” zamiast „porn”) i zapisując angielską liczbę mnogą z literą „z” zamiast „s” — *przypr. tłum.*

Ochrona haseł

Decyzje dotyczące dostępu to serce systemu zabezpieczeń, a są one podejmowane na podstawie podawanych haseł. Dlatego tak ważne jest, by system chronił swoje hasła i pozostałe informacje dotyczące logowania.

Większość administratorów chroni hasła na trzy podstawowe sposoby: wybierają mało oczywiste hasła, sprawiają, że ciężko jest złamać obsługę logowania, i bardzo mocno ochraniają plik, w którym hasła są przechowywane.

Tworzenie haseł niebanalnych wymaga odrobiny praktyki i okazjonalnego sprawdzania, czy system jest odporny na zwykły atak słownikowy. Takie testowanie pozwala wykryć użytkowników, których hasła są zbyt proste do odgadnięcia.

Ochrona haseł przed kradzieżą wymaga ukrycia pliku, w którym są przechowywane, a czasami wykonania jednokierunkowego szyfrowania, tzw. **skrótów wiadomości** lub **mieszania**, które przechowuje hasła w utajnionej postaci.

Ochrona nazwy użytkownika i hasła podczas wejścia do systemu

Większość dostawców oferuje cały zastęp metod obsługi logowania i zarządzania hasłami, które można dowolnie łączyć, aby zapewnić maksymalną ochronę dla systemu. Ponieważ są to przydatne usługi, które w dodatku łatwo wprowadzić do systemu, często są już w niego wbudowane. Tabela 3.1 zawiera krótkie zestawienie takich cech.

Tabela 3.1. Przykładowe elementy obsługi uwierzytelniania nazwa użytkownika-hasło

Cecha	Znaczenie
Komunikaty systemowe	Większość systemów wyświetla banery przed i (lub) po udanym wejściu do systemu. Dawniej banery miały postać przyjaznych identyfikatorów, ale liczne sprawy sądowe ukazały nierozważność tych działań. Niektóre systemy pozwalają administratorowi zablokować tego typu wiadomości, gdyż mogą one ułatwić obserwatorowi rozpoznanie systemu, do którego loguje się użytkownik. Jeśli atakujący połączy się przez protokół telnet i zorientuje się, że ma do czynienia z systemem Solaris, będzie to dla niego bardzo cenną wskazówką.
Ograniczona liczba prób	Po danej liczbie nieudanych logowań (którą określa administrator systemu), system blokuje Ci dostęp do danych z danego terminala. Niektóre systemy wykonują tę operację, nie informując Cię o tym. Pisane obecnie agresywne skrypty obchodzą to zabezpieczenie, próbując otworzyć kilka sesji naraz.
Ograniczony czas dostępu	Można ograniczyć uprawnienia niektórych użytkowników lub stanowisk pracy, tak aby logowanie z nich było możliwe tylko w godzinach pracy biura.
Zwiększenie czasu pomiędzy kolejnymi nieudanymi próbami logowania	Kolejne nieudane logowania do systemu oddzielane są od siebie coraz dłuższymi odstępami czasu. Po pierwszej próbie system potrzebuje jednej sekundy na ponowne uruchomienie, po drugiej próbie zabiera mu to już dwie sekundy, po trzeciej — cztery, po czwartej — osiem. To zapobiega powtarzającym się atakom na zasadzie pełnego przeglądu wszystkich kombinacji.
Komunikat ostatniego logowania	Kiedy wchodzisz do systemu, może on wyświetlić informacje o dacie i godzinie ostatniego logowania. Wiele systemów informuje również o liczbie nieudanych prób logowania od czasu poprzedniej udanej próby. Dzięki temu wiesz, czy ktoś inny próbował dostać się do Twojego konta, na przykład możesz w ten sposób wychwycić próbę, która miała miejsce w środku nocy, albo serię powtarzających się podejrzanych prób logowania. Jeśli nie rozpoznajesz w nich swojej działalności, natychmiast powiadom o nich administratora.

Tabela 3.1. Przykładowe elementy obsługi uwierzytelniania nazwa użytkownika-hasło — ciąg dalszy

Cecha	Znaczenie
Hasła zmieniane przez użytkownika	W wielu systemach użytkownik może zmieniać hasło dowolną liczbę razy, po pierwotnym utworzeniu go przez administratora. Może wręcz istnieć obowiązek zmieniania go po upływie określonego czasu.
Hasła generowane przez system	Niektóre systemy wymagają korzystania z haseł generowanych automatycznie przez system, nie polegając na pomysłowości użytkownika. Czasami udostępniają listę haseł do wyboru, więc możesz zdecydować się na takie, które będzie Ci łatwo zapamiętać. Niestety hasła tworzone przez system są często tak trudne do zapamiętania, że użytkownicy nagminnie je zapisują. Inne niebezpieczeństwo grozi ze strony ujawnienia algorytmu szyfrującego. W takim razie cały Twój system będzie narażony na niebezpieczeństwo.
Starzenie się i wygasanie hasła	Po upływie określonego czasu — na przykład na koniec każdego miesiąca — hasło może stracić swoją ważność. Przeważnie nowe hasła muszą być inne niż wszystkie poprzednie. System powinien wysłać odpowiedni komunikat, zanim zażąda zmiany hasła, ponieważ szybki wybór hasła sprawia, że z reguły jest ono słabe. Niektóre systemy dopuszczają ingerencję administratora, jeśli zostaną naruszone zasady bezpieczeństwa, na przykład poprzez natychmiastowe usunięcie ważności hasła. Przez jakiś czas system może przechować informacje o starych hasłach, aby uniknąć ich powtarzania.
Minimalna długość	Krótkie hasła są łatwiejsze do odgadnięcia, zatem niektóre systemy wymagają haseł o minimalnej długości, z reguły sześciu bądź ośmiu znaków, ale im hasło jest dłuższe, tym lepiej.
Blokowanie kont	Wprowadzenie ograniczenia daje administratorowi możliwość zablokowania dostępu do systemu niektórym użytkownikom albo zamknięcia kont, które nie były używane przez określony czas.

Bezpieczne przechowywanie hasła

Każdy system musi gdzieś przechowywać dane niezbędne do uwierzytelniania użytkowników. Z reguły poprawne hasła znajdują się w specjalnym pliku z hasłami. Dostęp do tego pliku jest możliwy tylko w ściśle określonych warunkach — podczas rejestracji nowego użytkownika, podczas zmiany istniejącego hasła oraz w trakcie procesu uwierzytelniania. Nawet administrator *nie ma* dostępu do przechowywanych haseł, co jest formą ochrony przed dostaniem się do pliku osób niepowołanych czy pojawieniem się w nim uszkodzonych bądź zagrożonych ujawnieniem haseł. Może on co najwyżej zmienić hasło na jakąś wartość przejściową, jak „let-mein” czy „hasło”, i żądać od użytkownika, aby przy kolejnym logowaniu zmienił je na własne.

Ochrona haseł jest jednym z najważniejszych elementów systemu zabezpieczeń. Aby chronić plik, w którym hasła są zapisane, system z reguły korzysta zarówno z szyfrowania danych, jak również kontroluje dostęp do danych.

Szyfrowanie

Większość systemów koduje dane przechowywane w systemowym pliku z hasłami. Proces szyfrowania (opisany dokładniej w rozdziale 7.) przekształca pierwotną informację do postaci przypadkowo zlepionych znaków. Szyfrowanie gwarantuje, że nawet w sytuacji, w której intruz przełamie zabezpieczenia systemu, nie będzie w stanie poznać haseł; będą dla niego niezrozumiałym połączeniem znaków.

Większość systemów korzysta z jednokierunkowego kodowania. Oznacza to tyle, że hasło nigdy nie powinno zostać odkodowane. Hasło dostarczone Ci przez administratora zostaje zaszyfrowane jeszcze przed umieszczeniem go w pliku haseł. Jego pierwotna postać

nie jest nigdzie zapisana, nawet w pamięci. Podczas każdego logowania i wpisywania hasła system dokonuje jego zaszyfrowania i w tej formie porównuje je z również zaszyfrowaną wersją z pliku z hasłami. Pamiętaj też, że hasło nigdy nie pojawia się na ekranie monitora.

Kontrola dostępu

Oczywiście uparty przeciwnik zdoła w końcu złamać szyfr naszych hasel. Wiele systemów przechowuje hasła w tak zwanych plikach **shadow password**. Są to pliki o najwyższym stopniu zabezpieczeń, do których dostęp mają jedynie administratorzy, ponieważ w ACL każdego pliku podane są identyfikatory należące wyłącznie do nich. (Rozważania dotyczące ACL znajdują się w części „Kontrola dostępu w praktyce”).

Ataki na hasła

Istnieją dwa podstawowe sposoby łamania hasel. Pierwszy sprowadza się do sprawdzenia wszystkich możliwych kombinacji znaków za pomocą kolejnych pojedynczych prób logowania (jest to tak zwany **atak na zasadzie pełnego przeglądu**). Oczywiście, jak wszystko inne w świecie komputerów, i ten proces został zautomatyzowany. Krakerzy używają komputerów, które próbują za nich odgadnąć hasła. Teoretycznie im dłuższe jest hasło, tym więcej czasu musi upłynąć, zanim zostanie ono złamane. Na przykład, aby znaleźć hasło zbudowane z ośmiu losowo wybranych znaków, trzeba najpierw sprawdzić 2,8 trylionu kombinacji. To zadanie na kilka tygodni, nawet dla szybkich komputerów.

Drugą metodą jest tak zwany atak **słownikowy**. Jest on skuteczny, ponieważ większość użytkowników nie korzysta z losowo wybranych hasel, ani nawet z takich, które są znośnie bezpieczne. Przeciętny użytkownik wybiera śmiesznie proste hasła — swoje inicjały, imiona dzieci, numery rejestracyjne samochodu itp. Badania wykazują, że znaczny odsetek hasel wybieranych przez użytkowników komputera można odgadnąć bez trudu. Krakerzy, mając do pomocy internetowe słowniki i spisy (ortograficzne, imion, zwierząt, samochodów, postaci książkowych i filmowych, miejsc i wiele innych), są w stanie w dość łatwy sposób odgadnąć większość hasel wybieranych przez ludzi. Ale jeśli rozsądnie wybierzesz swoje hasło (zobacz ramkę „Wskazówki dotyczące wybierania hasel”), atakujący nie powinien go odgadnąć, nawet jeśli korzysta ze słownika.

Autoryzacja

Po zakończeniu procesu uwierzytelniania system używa Twojego identyfikatora (i informacji dotyczących zabezpieczeń z nim związanych), aby określić, czy masz prawo pracować na danym komputerze bądź w sieci. Proces określania Twoich uprawnień nazywa się **autoryzacją**. Na przykład, jeśli próbujesz dokonywać zmian w ważnym pliku, system najpierw sprawdzi listę identyfikatorów osób uprawnionych do odczytu i zapisu danych w takim pliku i zweryfikuje, czy Twój identyfikator się na niej znajduje. Dostęp do pliku otrzymasz jedynie wtedy, gdy Twój identyfikator tam widnieje.

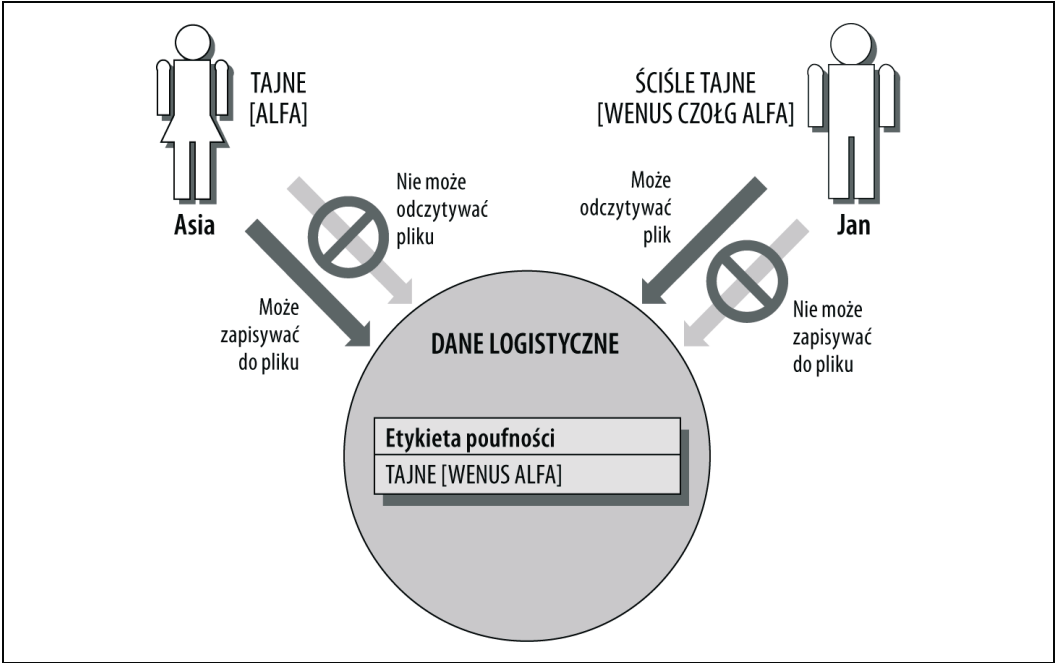
Systemy zazwyczaj przechowują plik, w którym zapisują uprawnienia poszczególnych użytkowników i ich charakterystykę. W zależności od systemu, taki plik nazywa się profilem zabezpieczeń, profilem uwierzytelniania lub listą użytkowników. Twój profil zawiera informacje o uprawnieniach, jakie posiadasz (na przykład TAJNE), mówi, czy wolno Ci zmieniać

własne hasło, logować się w dni wolne od pracy, czy możesz uruchamiać programy robiące kopie zapasowe i inne uprzywilejowane aplikacje itp. Istnieją sytuacje, w których profil jest tym samym plikiem, co lista haseł; informacje mogą być też przechowywane z dala od haseł. Zawsze jednak system chroni je, bo każde zagrożenie może negatywnie odbić się na bezpieczeństwie całego systemu.

Jedną z najważniejszych informacji zawartych w profilu uwierzytelniania i liście użytkowników jest typ użytkownika. Większość systemów obsługuje szereg kategorii użytkowników czy tak zwanych ról. Zwykle występują wśród nich: zwykły użytkownik, administrator systemu i jego operator. Systemy o podwyższonym stopniu zabezpieczeń definiują dodatkowo oficera bezpieczeństwa. Każda kategoria użytkowników charakteryzuje się swoimi przywilejami i obowiązkami — dotyczy to na przykład programów, które dany użytkownik może uruchamiać. I tak, administrator systemu może robić w zasadzie wszystko, łącznie z omijaniem, zmienianiem **wymogów** zabezpieczeń czy tworzeniem nowych, fałszywych kont na użytek atakujących system. (Mocą administratora jest jego szczególne uprawnienie w kwestii zabezpieczeń; odsyłam do opisu narzędzi administracyjnych i najmniejszego przywileju w rozdziale 5.).

Etykiety poufności

Wszystkie elementy systemu, który działa w oparciu o narzuconą kontrolę dostępu, podmioty i obiekty, posiadają przypisaną im etykietę poufności. Składa się ona z dwóch części: klasyfikacji i zestawu kategorii, nazywanych czasami **przedziałami**. Etykieta pokazana na rysunku 3.1 składa się właśnie z takich części.



Rysunek 3.1. Etykieta poufności, która składa się z dwóch części

Klasyfikacja ma pojedynczy, hierarchiczny poziom. W tak zwanym modelu zabezpieczeń wojskowych (bazującym na wielopoziomowej polityce zabezpieczeń Departamentu Obrony Stanów Zjednoczonych) istnieją cztery istotne poziomy:

ŚCIŚLE TAJNE
TAJNE
POUFNE
JAWNE

Każda z klasyfikacji jest obdarzona większymi prawami niż ta znajdująca się pod spodem.

Rzeczywista definicja klasyfikacji zależy od administratorów systemu lub oficera bezpieczeństwa. Jeśli Twoja organizacja przetwarza tajne informacje rządowe, model nadawania etykiet musi być zgodny z modelem wojskowym. Ale ogólnie etykiety mają reprezentować każdą klasyfikację i zestaw kategorii, jaki ma sens w danym przypadku. Na przykład ze stroną komercyjną może wiązać się tak zdefiniowana hierarchia firmy:

KORPORACJA
ODDZIAŁ
DZIAŁ

Możesz też zdefiniować hierarchię w oparciu o poziomy zaufania:

POUFNY
TYLKO DLA KIEROWNICTWA
ZASTRZEŻONE DLA FIRMY
PUBLICZNE

lub:

ZASTRZEŻONE
POUFNE
PUBLICZNE

Kategorie (przedziały) nie podlegają hierarchii i mają reprezentować poszczególne dziedziny danych w Twoim systemie. Razem składają się na zestaw kategorii (zestaw przedziałów). Zestaw może składać się z dowolnej liczby elementów.

W środowiskach związanych z wojskiem mogłyby wystąpić takie kategorie:

ANTYTERRORYZM
CZOŁG
OKRĘT PODWODNY
WENUS
STEALTH

W środowisku biznesowym kategorie częściej odpowiadają poszczególnym wydziałom firmy, nazwom produktów, kampaniom reklamowym, czy innemu dowolnie wybranemu zestawowi, związanemu z jej działalnością:

KSIEGOWOŚĆ
PUBLIC RELATIONS
MARKETING
SPRZEDAŻ
BADANIA I ROZWÓJ

Cały pomysł polega na tym, że nawet użytkownik o najwyższym poziomie klasyfikacji nie jest automatycznie upoważniony do przeglądania wszystkich informacji z tego poziomu. Aby móc poznać dane z kategorii ANTYTERRORYZM, musisz ich „potrzebować”.

Informacja podzielona na kategorie

Informacje w systemach wojskowych są często podzielone tak, jak stopniuje się je ze względu na poufność. Myśl o tym podziale, jak o odpowiedniku kryterium „**potrzebowania**” **informacji**. Nawet jeśli dana osoba ma wysoki poziom zaufania, to jeżeli nie ma żadnych przekonywujących argumentów za tym, aby znała to czy tamto zagadnienie, dane pozostają poza jej zasięgiem. Zasadę tę, w dość radykalny sposób, zastosował pewien sierżant piechoty morskiej armii Stanów Zjednoczonych, który dowodził załogą helikoptera wyposażonego w tajne urządzenie elektroniczne — jedną z „zabawek” wywiadu.

Pewnego dnia do helikoptera zająztał ciekawski kapitan. Sierżant natychmiast wypchnął go z maszyny, w wyniku czego kapitan spadł ze schodków prowadzących do helikoptera. W normalnych okolicznościach nie byłby to rozsądny ruch, sprzyjający karierze sierżanta. Jednak swoim zachowaniem sierżant dał doskonały przykład dobrych nawyków bezpieczeństwa. Przewaga rangi kapitańskiej została w tej sytuacji „przebita” przez potrzebę udzielenia, bądź nie, pewnej informacji. Sierżant doskonale zdawał sobie sprawę, że kapitan nie musi wiedzieć nic na temat urządzenia znajdującego się na pokładzie helikoptera, więc powziął natychmiast odpowiednie środki ostrożności.

Pomimo natury całego incydentu, sierżantowi udało się zachować godne tradycje jednostki, w której służył. Kapitan, spadając po schodach, słyszał wyraźnie wykrzyczane: „Przepraszam, **panie kapitanie!**”.

Modele dostępu

Istnieją dwa główne modele dostępu, z którymi zapewne się spotkasz. Model Bell-LaPadula skupia się przede wszystkim na zagadnieniach tajności danych, zatem najbardziej rozpowszechnił się on w kręgach obrony. Model Biba zajmuje się przede wszystkim spójnością przesyłanych informacji, więc cieszy się powodzeniem w kręgach zajmujących się transakcjami finansowymi i biznesem.

Model Bell-LaPadula. W roku 1973 panowie David Bell i Leonard LaPadula jako pierwsi opisali teoretycznie problem wielopoziomowych zabezpieczeń, które mogłyby pojawić się w Departamencie Obrony. Wykorzystali do tego formalizm matematyczny, o czym wspomniano w rozdziale 2. W opisie posłużyli się notacją matematyczną i teorią zbiorów, które pozwoliły im zdefiniować ideę stanu bezpiecznego, tryby dostępu i zasady jego udzielania. Nazwali go Teorią Podstaw Bezpieczeństwa. W modelu tym o udzieleniu podmiotowi (którym najczęściej jest użytkownik) dostępu do obiektu (z reguły pliku) decyduje wynik porównania klasyfikacji zabezpieczeń obiektu z poziomem zaufania, jaki posiada podmiot. Istnieją trzy podstawowe reguły:

- Zasada * (zasada gwiazdki).
- Prosta zasada bezpieczeństwa.
- Zasada spokoju.

Zasada *, czyli **zasada gwiazdki** stwierdza, że podmiot ma prawo zapisu w obiekcie (zwykle pliku) tylko wtedy, gdy poziom zabezpieczeń obiektu jest większy bądź równy poziomowi zaufania, jakim obdarzony jest podmiot. Dzięki temu podmiot o wysokim poziomie zaufania nie zapisze ściśle tajnych informacji w pliku o niedużym poziomie zabezpieczeń, dostępnym dla użytkowników o niższym poziomie zaufania. Zasada ta zapobiega sytuacji, w której użytkownik o wysokim poziomie zaufania mógłby skopiować poufne dane do dokumentu o niskim poziomie bezpieczeństwa — w ten sposób „poufne” dane **tracą swoją wartość** albo zmienia im się zaszeregowanie z „ściśle tajne” na „jawne”. Ten sposób klasyfikacji nazywa się czasami zasadą **zapisywania wzwyż** lub nie **zapisywania w dół**.

Prosta zasada bezpieczeństwa, nazwana tak na cześć swojej prostoty, mówi, że użytkownik (podmiot) może odczytywać pliki (obiekty) jedynie wtedy, gdy poziom jego zaufania jest wyższy lub równy poziomowi zabezpieczenia pliku. Oznacza to tyle, że użytkownik o poziomie „tajny” nie może odczytać pliku sklasyfikowanego jako „ściśle tajny”, ale z powodzeniem odczyta dane z plików „tajnych” i „poufnych”. Często zasadę tę określa się jako **zasadę odczytu w dół** lub **zakaz odczytu wzwyż**.

Zgodnie z **zasadą spokoju** poziom zabezpieczeń obiektu nie może zostać zmieniony podczas przetwarzania go przez system komputerowy. Dzięki niej żaden program ani atak nie może zmienić poziomu zabezpieczeń otwartego, a tym samym podatnego na zagrożenia pliku.

Model Biba. Model wojskowy nie sprawdza się we wszystkich warunkach. Innym ważnym modelem jest model Biba, zwany przez niektórych **odwróconym modelem Bell-LaPauda**. W sektorze bankowo-komercyjnym najważniejszym aspektem zabezpieczeń jest spójność danych (to znaczy niepopelnianie pomyłek przy oznaczaniu miejsc dziesiętnych); ich bezpieczeństwo jest mniej istotne. (Załóżmy, że ktoś kupuje 10 milionów akcji jakiejś firmy. W takim razie, prędzej czy później, znajdzie się w przeglądzie finansowym, więc raczej nie uda mu się zachować tajemnicy. Ale jeśli podczas transakcji ktoś zamieniłby wartość cyfry po przecinku i w ten sposób zmienił cenę każdej akcji o 10 groszy, to cena sprzedaży zmieniłaby się tak, że nawet Bill Gates odczułby różnicę).

W modelu Bell-LaPadula użytkownik z wysokimi prawami dostępu nie może zapisywać nic w dokumentach o niskim poziomie zabezpieczeń (ani tworzyć takich dokumentów), co zapobiega przepływowi tajnych informacji do osób nieposiadających uprawnień do ich odczytu. Model Biba odwraca tę sytuację, zatem użytkownik o niewielkich prawach dostępu nie ma prawa zapisywać informacji w dokumentach o wyższym poziomie poufności. Zakłada się, że dokładność podawanych informacji i ich wiarygodność rośnie wraz ze wzrostem poufności dokumentu. **Zapis wzwyż** narażałby dane o wyższym poziomie nienaruszalności na zmieszanie z danymi o niższym poziomie nienaruszalności. Tak samo odczyt plików o niższym poziomie zaufania przez użytkownika o wysokich prawach dostępu mógłby zanieczyścić dokumenty o wyższym poziomie nienaruszalności.

Kontrola dostępu w praktyce

Na szczęście przeciętny użytkownik, w tym także Ty, nie ma nic wspólnego z zawiłymi modelami matematycznymi, opisującymi kontrolę dostępu. Zasady te wprowadzono już w życie w kilku mechanizmach kontroli.

Po co kontrolować dostęp?

Jeśli jesteś jedynym użytkownikiem komputera, nie musisz martwić się ustawianiem praw dostępu. Jesteś właścicielem wszystkich plików. Jeśli chcesz jakiś udostępnić, wystarczy nagrać go na dysk, udostępnić w sieci folder, w którym się on znajduje, zapisać ten plik na wspólnej przestrzeni serwera czy wypalić go na płycie CD lub DVD.

Niestety sprawa nie jest tak prosta w przypadku komputerów dzielonych między różnymi użytkownikami. Rozpoczęcie pracy w systemie, który obsługuje dostęp wielu użytkowników, zobowiązuje Cię do dbania o ochronę danych i kontrolowania dostępu do nich. Przecież nie każdy użytkownik systemu powinien mieć dostęp do Twoich plików. A już na pewno nie chciałbyś, aby każdy mógł je zmieniać.

Nawet jeśli ufasz wszystkim, którzy mają dostęp do Twoich plików, i wiesz, że nie będą ich zmieniać, to powinieneś chronić dane przed wypadkami losowymi. Załóżmy, że Ty i Tomek pracujecie w różnych częściach jednego katalogu. Oboje pracujecie nad tym samym projektem i oboje wybieracie identyczne nazwy dla swoich plików. Kilkomu nieszczęsnymi uderzeniami w klawisze Tomek może zmienić katalog i usunąć Twój plik, myśląc, że to jego praca. Jeśli Twoje pliki nie są chronione, to Tomek nie napotka żadnych przeszkód na swojej drodze. Kontrola dostępu do pliku (którą można również odnieść do innych obiektów w systemie, na przykład katalogów i urządzeń) może zapobiec takim sytuacjom.

Można mówić o przynajmniej trzech podstawowych typach kontroli dostępu, które zapewniają różne poziomy zabezpieczeń Twoich plików:

- Swobodna kontrola dostępu (DAC).
- Narzucona kontrola dostępu (MAC).
- Kontrola dostępu oparta o role (RBAC).

Jeśli zdecydujesz się na wybór **swobodnej kontroli dostępu** (DAC), to będziesz mógł decydować o tym, w jaki sposób chcesz chronić swoje dane i komu je udostępnić. Bardziej złożoną formą ochrony jest **narzucona kontrola dostępu** (MAC), w której to system odpowiada za ochronę Twoich danych. W zabezpieczeniach typu MAC każdemu obiektowi w systemie przypisana jest odpowiednia etykieta. Korzystając z zależności systemu bezpieczeństwa zdefiniowanych dla Twojej firmy, system operacyjny decyduje, czy użytkownik powinien dostać dostęp do pliku, porównując etykietę użytkownika i etykietę pliku. W zabezpieczeniach opartych na rolach (RBAC), otrzymujesz pewne przywileje w oparciu o Twoją pozycję w hierarchii służbowej firmy. Jeśli jesteś księgowym, to uzyskujesz dostęp do danych, do których mają dostęp pozostali pracownicy tego działu. Poniżej znajdziesz szczegółowy opis wszystkich trzech typów kontroli dostępu.

Swobodna kontrola dostępu

Ten typ kontroli zapewnia użytkownikowi dostęp do pliku (i innych obiektów w systemie, jak na przykład urządzenia czy katalogi) w oparciu o jego tożsamość a także grupę, do której przynależy. O jakiej swobodzie można mówić w przypadku tej formy kontrolowania dostępu? W przeciwieństwie do narzuconej kontroli, gdzie to system decyduje o udzielaniu dostępu, model DAC opiera się na Twoim uznaniu czyichś praw. To właściciel pliku decyduje, czy udostępnić komukolwiek dane. Model MAC nie daje takich możliwości.

Model DAC informuje system operacyjny, kto ma prawo pracować z Twoimi plikami, oraz pozwala Ci sprecyzować rodzaj dostępu dla danego użytkownika. Możesz udostępnić do odczytu dany plik wszystkim pracownikom swojej firmy, ale uprawnienia do wprowadzania zmian pozostawić jedynie dla siebie i kierownika. Większość systemów obsługuje trzy podstawowe typy dostępu:

Odczyt

Umożliwia odczyt danych zawartych w pliku.

Zapis

Pozwala zapisywać dane w pliku (zmieniać dane bądź zastępować plik innym).

Wykonywanie

Ten typ dostępu ma odniesienie jedynie do plików będących programami. Prawo wykonywania daje Ci możliwość uruchomienia programu.

Własność. Istnieje wiele typów swobodnej kontroli dostępu. Jeden z nich uwzględnia prawa własności do plików, katalogów i urządzeń.

Jeśli stworzyłeś dany plik, jesteś jego właścicielem. Identyfikator Twojej tożsamości jest umieszczany w nagłówku pliku. System może opierać wszystkie decyzje dotyczące dostępu do pliku na prawach własności. Właściciel będzie mógł go odczytać i zmieniać jego zawartość. Żaden inny użytkownik nie będzie posiadał uprawnień do działań na tym pliku. To bardzo prosty schemat, ale zupełnie niepraktyczny. Przede wszystkim nie pozwala Ci udostępniać własnych plików.

W zasadzie każdy system przechowuje informacje o właścicielach plików i wiele swoich decyzji opiera właśnie o tę cechę (na przykład, niezależnie od innych mechanizmów, system może zezwolić Ci na usunięcie pliku jedynie wtedy, jeśli jesteś jego właścicielem).

Elementy sterujące własny/grupowy/publiczny. Wiele systemów reguluje dostęp do plików, dzieląc świat użytkowników na trzy kategorie i określając uprawnienia każdej z nich. W niektórych systemach nazywa się to kontrolą **własny/grupowy/publiczny** (ang. *self/group/public* — *przyp. tłum.*). W systemach Unix nazywa się ją kontrolą **użytkownik/grupowy/inny** (ang. *user/group/other* (UGO) — *przyp. tłum.*).

Własny

Grupa określająca twórcę i właściciela pliku, czyli Ciebie.

Grupowy

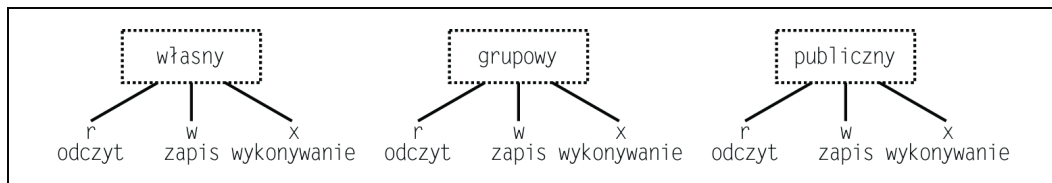
Opisuje pewną grupę użytkowników. Na przykład, wszyscy użytkownicy pracujący w określonym dziale mogą należeć do grupy R&D (ang. *Research and Development* — dział badań i rozwoju — *przyp. tłum.*).

Publiczny

Wszyscy inni — użytkownicy inni niż należący do Twojej grupy.

Prawa dostępu. Każdy plik ma wydzielone bity, nazywane **prawami dostępu**⁵. Ich znaczenie często pokrywa się ze schematem przedstawionym na rysunku 3.2.

⁵ Określają one, co możesz zrobić z danym plikiem — *przyp. tłum.*



Rysunek 3.2. Elementy sterujące własny/grupowy/publiczny

Jeśli wypiszesz listę plików (poleceniem `ls -l`), to w systemach z rodziny Unix bądź Linux otrzymasz spis praw dostępu w następującej postaci:

```
-rw-rw-r-- 1 franek r&d 81904 Nov 7 13:25 AKTUALIZACJE
```

Jeśli w miejscu symbolu określonego prawa dostępu wystąpi myślnik (-), to oznacza, że użytkownik **nie ma** prawa odczytywać, zapisywać bądź wykonywać danego pliku. Na przykład, z powyższego kodu wynika, że właściciel pliku (frank) może odczytywać i zapisywać dane w pliku AKTUALIZACJE (rw-), członkowie grupy, do której on należy (r&d) mają identyczne uprawnienia (rw-), a pozostali użytkownicy mogą jedynie odczytywać dane z pliku (r--). (Myślnik poprzedzający cały ciąg znaków ma specjalne, niezwiązane z uprawnieniami dostępu do pliku znaczenie w systemie Unix).

Oto kilka kolejnych przykładów.

Plik SZACHY to gra komputerowa. Jego prawa dostępu mają następującą postać:

```
-rwxrwxrwx 1 libr games 61799 May 10 10:11 SZACHY
```

Każdy ma prawo odczytać ten plik, zapisać w nim informacje, a także uruchomić go.

Plik SRC95 jest fragmentem kodu, nad którym pracują ludzie zrzeszeni w grupie r&d. Oto jego prawa dostępu:

```
-rw-rw---- 1 zosia r&d 55660 Dec 19 11:42 SRC95
```

Właściciel tego pliku oraz członkowie grupy mogą odczytywać i zmieniać go. Nikt inny nie ma do niego dostępu.

Elementy sterujące własny/grupowy/publiczny stanowią bardzo dobrą metodę ochrony danych zawartych w pliku. Ale co zrobić, jeśli zaistnieje potrzeba ochrony plików w inny sposób przed różnymi użytkownikami lub jeśli będziesz chciał ograniczyć dostęp do pliku jednemu użytkownikowi?

Jeśli Zosia jest właścicielem pliku ZNACZNIK i chce, aby Tomek (członek jej grupy) miał możliwość czytania i zmieniania tego pliku, ustali następujące prawa dostępu:

```
-rw-rw---- 1 zosia r&d 22975 Jan 10 10:14 ZNACZNIK
```

Jeśli Zosia będzie chciała, żeby Tomek mógł odczytywać dane z pliku ZNACZNIK, a jednocześnie będzie chciała udostępnić go do odczytu i zapisu dla Marii, może z niej uczynić właściciela pliku (o zezwoleniach r i w), a Tomka pozostawić w grupie, która posiada jedynie uprawnienia do odczytu (r). Ale co zrobić, jeśli do grupy Tomka należą inni użytkownicy, którzy nie powinni poznać zawartości pliku ZNACZNIK? W jaki sposób Zosia ma wykluczyć groźnego Stefana?

Swobodna kontrola dostępu wydaje się odrobinę uciążliwa, ale jest bardzo elastyczna. Pewne skomplikowane manewry pozwoliłyby osiągnąć opisane powyżej cele samymi tylko elementami sterującymi własny/grupowy/publiczny, ale im bardziej złożone staną się Twoje

potrzeby, tym bardziej nieporęczne staną się rozwiązania oferowane przez tę formę dostępu do plików. Kolejny opisany przez nas system oferuje trochę mniejszą elastyczność.

Narzucona kontrola dostępu

Ta forma dostępu do plików jest odpowiednia dla systemów, w których przechowywane są szczególnie poufne informacje (na przykład tajne informacje rządowe czy ważne dane korporacyjne). Systemy obsługujące narzuconą kontrolę dostępu muszą zaopatrzyć *wszystkie* swoje podmioty (np. użytkowników, programy) oraz obiekty (pliki, katalogi, urządzenia, okna, porty) w etykietę określającą poziom zaufania, jakim dany obiekt czy podmiot się cieszy. Etykieta użytkownika określa jego poziom zaufania i często bywa nazywana zgodą lub pozwoleniem dostępu. Etykieta pliku mówi, jaki poziom zaufania musi mieć użytkownik, aby się dostać do tego pliku. Kontrola typu MAC korzysta z etykiet, aby określić, kto ma mieć dostęp do danych informacji w Twoim systemie.

Wzięte razem, etykiety i kontrola MAC, zapewniają **wielopoziomą politykę ochronną** — politykę, która obsługuje wiele klasyfikacji informacji na różnych poziomach zabezpieczeń wewnątrz jednego systemu komputerowego.

W przeszłości systemy wojskowe potrafiły obsłużyć tylko jeden poziom zabezpieczeń. Systemy te, znane jako **systemy o wysokich zabezpieczeniach**, wymagały, aby każdy ich użytkownik legitymował się najwyższym poziomem zaufania, wymaganym przez dowolne dane w systemie. Przykładowo system, który obsługiwał dane o etykiecie TAJNE, nie pozwalał na pracę nikomu, kto nie posiadał poziomu zaufania TAJNE (niezależnie od tego, jak dobrze były chronione dane oznaczone tą etykietą).

Chociaż wiele źródeł rządowych nadal działa w trybie wysokich zabezpieczeń, stosowane są już systemy wielopoziomowe. Wielopoziomowość jest możliwa dzięki dzieleniu danych i wprowadzaniu ich do odpowiednich przedziałów. Takie systemy radzą sobie doskonale z jednoczesną obsługą użytkowników o wysokim i niskim poziomie zaufania (czy nawet bez określonej etykiety), dając symultaniczny dostęp do różnych typów plików SCI (ang. *sensitive compartmented intelligence*).

W tej części postaramy się przedstawić najważniejsze sprawy związane z nadawaniem etykiet, poziomami zabezpieczeń, i kontrolą MAC, ale są to skomplikowane zagadnienia, które posiadają głębokie korzenie historyczne. Dodatek C stara się krótko opisać wymagania, jakie *Orange Book* stawia tym dziedzinom zabezpieczeń.

Import i eksport danych. W systemach o narzuconej kontroli dostępu istotnymi zagadnieniami są: dopuszczenie importu informacji z innych systemów komputerowych oraz eksportowanie ich do zewnętrznych systemów. Kontrola MAC posiada mnóstwo zasad regulujących zagadnienia wysyłania i pobierania danych. Określają one, na jakie urządzenia można kopiować informacje, a na jakie można wypuszczać wydruki. Na przykład możesz nie uzyskać pozwolenia wydruku danych na drukarce, która znajduje się w publicznej części budynku. Istnieją również zasady nadawania etykiet poszczególnym urządzeniom i drukarkom (ze stronami transparentowymi, oraz nagłówkami stron i rozbiegówkami). Konkretnie przykłady zostały omówione w dodatku C.

Decydowanie o dostępie

W systemach o kontroli dostępu MAC wszystkie decyzje dotyczące wydania pozwolenia dostępu do pliku są podejmowane przez sam system. W przeciwieństwie do modelu DAC, który dawał użytkownikowi prawo do udzielania innym dostępu do pliku na własne ryzyko, model MAC ceduje wszelkie decyzje na system. Decyzja o przyznaniu dostępu do obiektu (na przykład pliku) wiąże się ze zbadaniem wszystkich trzech poniższych czynników:

- Etykieta podmiotu — na przykład Twój poziom zaufania:
SCISLE TAJNE [WENUS CZOLG ALFA]
- Etykieta obiektu — na przykład plik o nazwie LOGISTYKA o poziomie zabezpieczeń:
TAJNE [WENUS ALFA]
- Żądanie dostępu — na przykład Twoja próba odczytu z pliku LOGISTYKA.

Gdy próbujesz odczytać dane z pliku LOGISTYKA, system porównuje poziom zaufania, który został Ci nadany, z etykietą pliku, sprawdzając w ten sposób, czy masz prawo go odczytać.

Chociaż wielopoziomowe systemy zabezpieczeń dają nam wiele korzyści, są również źródłem frustracji. Zdarza się na przykład, że system pozwoli Ci zapisać dane do pliku, a następnie odmówi Ci prawa do odczytu!

Kontrola dostępu oparta o role

Ten typ regulowania dostępu do danych opiera się na zaszeregowaniu użytkownika według roli, jaką sprawuje. Weźmy dla przykładu kierowników działu finansowego. Oni mogą potrzebować dostępu do wszelkiego rodzaju danych księgowych: podatków, listy płac, należności, wpływów, salda. Natomiast urzędnik z sekcji należności będzie potrzebował dostępu jedynie do pewnej części danych księgowych, a inżynier z działu rozwoju nie będzie ich potrzebował prawie wcale. Rola, którą system przypisuje użytkownikowi, opiera się na **idei najmniejszego przywileju**. Rola jest definiowana w oparciu o minimalną liczbę zezwoleń, która pozwala wykonać stawiane przed użytkownikiem zadania. Jeżeli zmianie ulegną przywileje przypisane do danej roli, można usunąć lub dodać pozwolenia. Takie rozwiązanie daje większą elastyczność, bo zamiast zmieniać zezwolenia dla każdego użytkownika, zmienia się definicja samej roli.

Ponieważ użytkownik może pełnić więcej niż jedną rolę naraz, istnieje możliwość powstania konfliktu. Jedna z ról może zezwalać na dostęp do źródła, podczas gdy druga będzie go broniła. Konflikty tego typu muszą być rozstrzygane indywidualnie dla każdego użytkownika. Ogólnie rzecz biorąc, przyjmuje się te rozwiązania, które udzielają najmniejszych praw.

Listy kontroli dostępu

Nadszedł czas, by wyjaśnić, jak implementuje się modele przedstawione w poprzednim dziale. Listy kontroli dostępu (ACL — ang. *access control lists*) stanowią spis użytkowników i grup z przypisanymi im zezwoleniami. Dzięki nim można łatwiej przeprowadzać swobodną kontrolę dostępu. Implementacja listy ACL zależy w dużej mierze od systemu operacyjnego. Na przykład w bezpiecznym systemie opartym o technologię Unix plik WYPLATA byłby chroniony plikiem ACL w postaci:

```
<jan.ksieg, r>  
<asia.wypl, rw>
```

gdzie:

- `jan` i `asia` to identyfikatory użytkowników, którzy mają dostęp do pliku `WYPLATA`,
- `ksieg` i `wypl` to identyfikatory grup, do których należą,
- `r` i `w` określają rodzaj dostępu; `r` oznacza, że użytkownik ma prawo odczytywać dany plik, a `w` mówi, że może również wprowadzać w nim zmiany.

Jeśli `jan` został zaklasyfikowany do grupy `ksieg`, to może jedynie odczytywać plik. Jeśli należałby do jakiegokolwiek innej grupy, w ogóle nie miałby do niego dostępu. Podobnie ma się sprawa z użytkownikiem `asia`, która należąc do grupy `wypl`, może odczytywać i zmieniać zawartość pliku.

Listy kontroli dostępu najczęściej obsługują znaki specjalne, które pozwalają określić bardziej ogólne zasady dostępu do pliku. Można na przykład zapisać:

```
<*.*, r>
```

aby zażądać możliwości odczytu (`r`) pliku przez dowolnego (`*`) użytkownika każdej (`*`) z grup. Można również zastosować zapis:

```
<@.*,rw>
```

co jest równoznaczne ze stwierdzeniem, że tylko właściciel (`@`) danego pliku może go odczytać (`r`) i zmienić (`w`).

Niektóre systemy zezwalają na wykluczenie określonego użytkownika z grona osób, które mają dostęp do danego pliku, na przykład definiując znak zerowy czy pojęcia `none` lub `null`, które będzie można podać w miejsce znaków `r` i `w`.

```
<stefan.*,none >
```

Największe i najmniejsze zezwolenie

W przykładzie dotyczącym dostępu do pliku `WYPLATA`, jaki mają użytkownicy `jan` i `asia`, kryje się konflikt dostępu. Prawa dostępu, jakimi cieszy się użytkownik `jan`, zależą od grup, do jakich on należy. Dopóki pozostaje członkiem jedynie grupy `ksieg`, może tylko odczytywać wspomniany plik. Gdy zostanie dodatkowo dołączony do jakiegokolwiek innej grupy, automatycznie utraci do niego dostęp.

Konflikt powstaje ze względu na istnienie **zasady najmniejszego zezwolenia**. Zezwolenia skumulowane są różnie obsługiwane przez różne systemy, takie jak na przykład NetWare, wykorzystujący Novell Directory Service, a także WindowsNT i Windows 200x, używający Active Directory Service.

Na szczęście znajomość systemu Unix UGO⁶ przygotuje Cię do pracy z dowolnym innym systemem. Bez trudu będziesz rozróżniać reguły, więc jedynie, co będziesz musiał opanować, to ich nowe nazwy i obsługa w danym systemie. Jest to bardzo istotne, choć teoretyczne zagadnienie. W rzeczywistości administratorzy radzą sobie z tymi problemami, klonując użytkownika, który ma odpowiednie uprawnienia, i zmieniając jego nazwę. Nie omijaj zbyt chętnie tej kwestii: z listami dostępu spotkasz się jeszcze przy okazji projektowania systemu ochrony komputera, na przykład zapory ogniowej albo pakietu takiej zapory dla routerów czy systemów wykrywania intruzów.

⁶ *User/Group/Other* — kontrola typu użytkownik/grupowy/inny — *przypl. tłum.*

Usługi katalogowe

Wspominałem już poprzednio o systemach śledzenia nieupoważnionych użytkowników. Niektóre z nich, jak urządzenia biometryczne, nadal czekają na swoje pięć minut. Z kolei inne, jak uwierzytelnianie nazwa-hasło, są już całkiem przestarzałe. Jednym z usprawnień w technologii uzyskiwania dostępu do systemu jest integracja uwierzytelniania, autoryzacji i rozliczania (AAA — ang. *Authentication, Authorization, Accounting*). Te ulepszone systemy katalogowania przechowują informację o każdym użytkowniku, włączając w to przypisane mu atrybuty, na przykład informację o kontaktach, dane osobiste, włącznie z informacją, do jakich danych ma on dostęp. Tworzy to swoistą bazę danych usług katalogów. **Usługa katalogowa** jest w zasadzie ogromną bazą danych, która przechowuje informacje o oddziaływaniach między obiektami. Obiektami mogą być użytkownicy, grupy, sprzęt czy oprogramowanie.

Bardzo popularną usługą jest Active Directory, z której korzystają serwery i kontrolery domen pracujące na systemie Windows. Implementacja Active Directory pomaga zarządzać obiektami i zasobami środowiska sieciowego, działając jednocześnie jako centralny punkt zabezpieczeń. Active Directory dostarcza środków kontrolowania wszystkich zdarzeń, jakie mają miejsce w serwerach logowania (kontrolerach domen). Politykę audytu można skierować na monitorowanie konkretnych przejawów aktywności, tworzenie raportów i powiadamianie wybranych pracowników o zajściu tych zdarzeń. Zadania dotyczące grup można zdefiniować tak, aby pomagały zarządzać grupami użytkowników i komputerów. Można zastosować je do witryn, domen czy jednostek organizacyjnych, które są zdefiniowane wewnątrz struktury Active Directory.

Najświeższe usługi uwierzytelniania zależą od układu parametrów zdefiniowanych w normie ISO X.500, która określa standardy utrzymania łatwego dostępu do użytkownika i jego atrybutów. Norma X.500 jest tak obszernym systemem, że z reguły dostęp do katalogu X.500 uzyskuje się dzięki prostszemu narzędziu, Lightweight Directory Access Protocol (LDAP). Dalsza część rozdziału jest poświęcona normie X.500 i sposobom wykorzystania LDAP, jego aplikacji, problemów i możliwości, jakie daje.

Przykład poczty e-mail

Chyba najłatwiejszym przykładem, który pokaże uwierzytelnianie oparte o usługi katalogów jest sytuacja, z jaką na co dzień mieli do czynienia administratorzy poczty elektronicznej, zanim nastała era LDAP i X.500. Wyobraź sobie połączenie dwóch ogromnych przedsiębiorstw. Teraz wyobraź sobie, że obie te firmy korzystały z innych systemów poczty elektronicznej, które są kompatybilne jedynie w teorii. Użytkownik musi pojawić się w każdym z systemów (czyli mieć w nim konto pocztowe), aby móc się zalogować. Oznacza to, że dwie grupy administratorów muszą nagle przenieść pełną listę użytkowników z tej drugiej firmy do swojego systemu. Oczywiście oznacza to również, że w efekcie ich działań każdy użytkownik będzie posiadał dwóch klientów poczty, za pomocą których będzie musiał regularnie sprawdzać pocztę. Obie części nowej firmy przeżyją ciężki okres, podczas którego będą próbowały sprawnie się ze sobą komunikować. O problemach, jakie spotkają klientów tej firmy, lepiej nie wspominać.

Dość rozsądnym wyjściem z sytuacji wydaje się przeniesienie kont pocztowych jednej z firm do systemu tej drugiej. Niestety może to być zupełnie niewłaściwe podejście, ponieważ serwery pocztowe operują klastrami podczas wewnętrznej komunikacji, opierając się o lokalizację

użytkownika i załadowanie systemu. Przeniesienie jednego systemu do drugiego może skutkować drastycznymi zmianami infrastruktury, które mogą w rezultacie prowadzić to zakłóceń świadczonych usług. Może okazać się, że łatwiej jest powiększyć systemy pocztowe obu firm i uporządkować wszystkie sprawy w przyszłości, korzystając z serwerów bramkowych, aby oba systemy mogły się porozumieć.

Korzystanie z katalogu LDAP sprawia, że wszyscy użytkownicy znajdują się w jednej bazie, która jest automatycznie replikowana przez system na serwery logowania w całej sieci. Zasoby pocztowe mogą być przekazane użytkownikom, nawet jeśli przeniosą się z jednego miasta do drugiego. Zmiany danych osobowych, takich jak nazwiska, zmiany pozdrowienia używanego w listach czy tytułu użytkownika są dokonywane w jednym miejscu, skąd przenoszą się we wszystkie niezbędne lokalizacje.

X.500

ISO X.500 jest międzynarodowym standardem, który charakteryzuje się dwiema znakomitymi cechami:

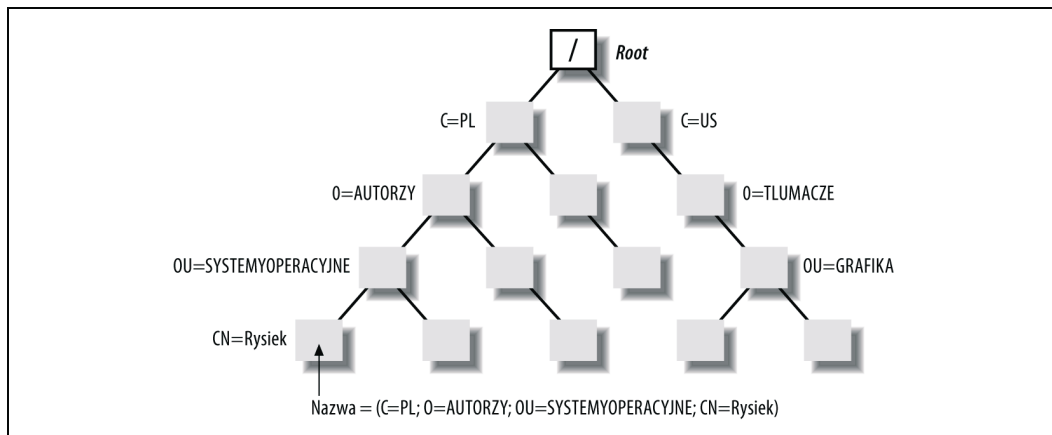
- jest obiektowy,
- korzysta ze struktury hierarchicznej.

ISO X.500 traktuje każdy zapisany w nim komputer i każdego użytkownika jak **obiekt**. Serwer „Kochanowskiego 12”, operator kopii zapasowych „tkotek” i administrator systemu „IMJasinska” są obiektami. Każdy obiekt posiada własne atrybuty. „Kochanowskiego 12” może mieć adres IP równy 198.168.212.12. Użytkownik tkotek ma przypisany atrybut opisujący grupę operatorów kopii bezpieczeństwa. Administrator IMJasinska będzie miała przypisane atrybuty **imię** o wartości Iza i **nazwisko** Jasińska. W sytuacji idealnej schemat bazy danych dla każdej klasy obiektów byłby spójny. Przypadek, kiedy zajdzie potrzeba zmiany tego schematu, aby dopasować do siebie elementy danych z obu firm, nie powinien stanowić zbyt trudnego wyzwania.

Drugą zadziwiającą cechą normy X.500 jest jej hierarchiczna struktura. Wszystkie osoby i komputery w systemie dostają **nazwy publiczne** (CN — ang. *common names*). Nazwa publiczna to nazwa, jaką Ty czy ja moglibyśmy przypisać dowolnie wybranej rzeczy. Obiekty z nadanymi nazwami publicznymi zostają zgromadzone w kontenerach określonych mianem **jednostek organizacyjnych** (OU — ang. *organizational units*). OU są zebrane w obiektach nazwanych **organizacjami** (O — ang. *organizations*). Z kolei organizacje umieszczono w obiektach o nazwie **kraje** (C — ang. *countries*). Kraje umieszczone są w Root, który stanowi początek całego drzewa systemu plików. Cała struktura hierarchiczna jest pokazana na rysunku 3.3.

Informacyjne drzewo katalogów (DIT — ang. *directory information tree*) nie ma być bazą danych ogólnego użytku. Zostało zoptymalizowane do częstego odczytu i rzadkiego zapisu. Swoją budową przypomina bazę danych, ale jego idea została rozwinięta po to, aby przechować informacje o obiektach związanych z dziedziną telekomunikacji.

Zaletą hierarchicznego schematu X.500 jest to, że wprowadził on standard przechowywania plików, uwzględniający atrybuty i zezwolenia przypisane danemu użytkownikowi. Taka struktura znacznie ułatwia dzielenie danych i ich kopiowanie. Niezależnie, czy kopiujemy dane jednego użytkownika, czy całego działu firmy, powielanie adresów sprowadza się po prostu do skierowania kopii danego obiektu do nowego drzewa katalogów; wszystkie obiekty składowe zostaną przeniesione automatycznie. Przeniesienie kopii jest równoznaczne z jej gotowością do użycia.



Rysunek 3.3. Hierarchiczna struktura X.500

Każde centrum X.500 posiada własny schemat katalogów. Jego administrator musi aktualizować tylko dane z własnego centrum. Użytkownicy mają dostęp jedynie do lokalnych systemów X.500. Jeśli poszukiwana osoba pracuje w lokalnym centrum, dane na jej temat będą dostępne w lokalnym systemie X.500. Jeżeli natomiast pracuje w punkcie zdalnym, to lokalny X.500 automatycznie kontaktuje się ze zdalnym X.500, aby pobrać jej dane. To przełączanie się między lokalnymi i zdalnymi systemami odbywa się bez wiedzy użytkownika. Jeśli okaże się, że dostęp X.500 do zdalnych centrów jest zbyt wolny, można tak zmienić konfigurację, aby wymiana danych między dwoma centrami polegała na kopiowaniu ich i ciągłej aktualizacji. Proces ten nazywamy **replikacją**.



Główną cechą tego systemu jest to, że skupia się on na potrzebach indywidualnego użytkownika, który chce dostać się do systemu, zamiast koncentrować się na statycznych elementach, takich jak listy haseł. Użytkownicy są mapowani w katalogu zgodnie z ich zaszerzowaniem w strukturze organizacji.

System katalogów X.500 może przechowywać ogromne ilości informacji: numerów telefonów i faksów, adresów, adresów komputerów, pozycji służbowej, zakresu obowiązków itd. Można go oczywiście rozszerzyć tak, aby przechowywał informacje potrzebne w danej organizacji. Te duże ilości danych i wrodzona złożoność systemu X.500 sprawiły, że bardzo rozpowszechnił się protokół LDAP.

Protokół LDAP

Protokół LDAP (ang. *Lightweight Directory Access Protocol*) jest standardem, który opisuje protokół sieciowy dostępu do informacji zawartej w katalogach. W momencie, w którym stało się oczywiste, że trzeba będzie wprowadzić standardy przechowywania informacji w katalogach, IBM, Microsoft, Lotus i Netscape postanowiły wprowadzić obsługę protokołu LDAP. Protokół ten został zaprojektowany w taki sposób, żeby dostosować złożony system katalogów, jakim jest X.500, do potrzeb współczesnych sieci, w tym również internetu. Serwery katalogów, korzystające z protokołu LDAP, są uruchamiane na komputerach głównych w sieci internet, co oznacza, że programy klienckie, które mają wprowadzoną obsługę LDAP, mogą połączyć się z takim serwerem i przejrzeć listę jego katalogów.

Historia X.500

Standard X.500 został opisany w normach ISO/IEC 9594 oraz ITU-T X.500 Recommendations. Jak dotąd pojawiły się cztery wersje tego standardu

Pierwsza edycja, 1988 rok

Pierwsza edycja pojawiła się jako wieloczęściowy standard oparty na normach ISO/IEC 9594:1990 i CCITT X.500 (1988) Series of Recommendations.

Druga edycja, 1993 rok

Druga edycja została oparta o normy ISO/IEC 9594:1995 i ITU-T X.500 (1993). Wprowadzono do niej kilka przydatnych funkcji, takich jak replikowanie informacji między katalogami, kontrola dostępu oraz rozszerzenie modelu informacji i możliwości zarządzania.

Trzecia edycja, 1997 rok

Trzecia edycja bazuje na ISO/IEC 9594:1998 i ITU-T X.500 (1997). Dodano do niej kilka małych oraz większych rozszerzeń. Dodano w niej cechę zwaną **contexts**, która pozwala na wyróżnienie informacji w kontekście, w jakim uzyskano do niej dostęp. Wprowadzono również model OSI zarządzania katalogiem. Pojawiły się też nowe, ważne elementy zabezpieczeń oraz rozszerzenia już istniejących.

Czwarta edycja, 2001 rok

Czwarta edycja jest oparta o ISO/IEC 9594:2001 i ITU-T X.500 (2001). Dodano w niej możliwości zarządzania usługami, dopasowanie oparte o mapowanie, rodziny wejść i obsługę stosu TCP/IP.

Piąta edycja, 2005 rok

Piąta edycja bazuje na normach ISO/IEC 9594:2005 i ITU-T X.500 (2005).

Serwery LDAP zaczynają pracę od indeksacji wszystkich danych, które pojawiają się w ich wpisach. Podczas przetwarzania **żądania** wszystkie dane muszą przejść przez odpowiednie **filtry**. Są to krótkie, wprowadzone przez użytkownika wyrażenia, które opisują, jakie cechy należy uwzględnić, a jakie odrzucić. Filtry pozbywają się niepotrzebnych wpisów, przedstawiając użytkownikowi tylko te, którymi był zainteresowany. Filtrem może być osoba, grupa, a nawet dość egzotyczne wyrażenie w stylu: „każdy użytkownik na północnym zachodzie wybrzeża Pacyfiku, który twierdzi, że posiada przenośne urządzenie 802.11”. Użytkownik może również określić, jaka liczba danych, takich jak nazwisko, imię, tytuł, adres e-mail, numer telefonu, ma odpowiadać jego zapytaniu.

Przestrzeń nazw protokołu LDAP. W protokole LDAP każdą listę w katalogu nazywamy wpisem. Każdy wpis może mieć jeden lub więcej atrybutów. Każdy atrybut ma określony typ i przynajmniej jedną wartość. Oto przykład wpisu:

```
cn = mojmalypistestowy  
objectclass = osoba
```

Skrót **cn** oznacza nazwę publiczną (**commonName**), której wartością jest **mojmalypistestowy**. Dzięki wykorzystaniu atrybutu **objectClass** wpisowi został przypisany typ.

System i jego użytkownicy wprowadzają wpis za wpisem, nadając im nazwy. Wpisy są nazywane w oparciu o jeden z ich atrybutów. W pokazanym przykładzie **cn=mojmalypistestowy** jest **względna wyróżniona nazwą** (RDN — ang. *relative distinguished name*).

Hierarchia. Przestrzeń nazw protokołu LDAP jest, tak samo jak w X.500, hierarchiczna. W praktyce sprowadza się to do tego, że pełna nazwa obiektu musi zawierać ścieżkę dostępu do tego obiektu. Tą pełną nazwę określa się mianem **nazwy wyróżnionej** (DN — ang. *distinguished name*). (Owo „wyróżnienie” oznacza tyle, co „całkowicie i jednoznacznie zidentyfikowany”).

Możliwości magazynowania protokołu LDAP. W katalogach X.500 można przechowywać w zasadzie każdy rodzaj danych. Bez problemu radzi on sobie z tekstem, fotografiami do dokumentów, informacjami biometrycznymi, niezbędnymi do identyfikacji i uwierzytelnienia, adresami WWW, adresami FTP oraz innymi wskaźnikami. Jednak wzrost ilości przechowywanych danych wiąże się nieodmiennie ze wzrostem potrzebnej pamięci.

Różne rodzaje danych są przechowywane w atrybutach różnych typów. Każdy typ atrybutu ma określoną składnię. Atrybuty zdefiniowane przez użytkownika, składnie i klasy obiektów dają administratorom zabezpieczeń możliwość dopasowania katalogów ściśle do ich potrzeb.

Protokół LDAP jest szczególnie przydatny we wszystkich tych miejscach, gdzie internet stał się częścią sieci albo jej schematu uwierzytelniania. W jego skład wchodzi protokoły, które pozwalają na samoreplikację danych między różnorakimi centrami. Takie aktualizacje krążą w sieci, jak każdy inny rodzaj przepływu informacji.

Zarządzanie tożsamością

Mogłoby się wydawać, że wraz z nastaniem X.500 i protokołu LDAP dostęp do systemu oparty na katalogach osiągnął szczyt swoich możliwości. To nieprawda. Czego doświadczy użytkownik, który w ciągu swojego dnia pracy musi kontaktować się z czterema czy z pięcioma sieciami? Brzemie zarządzania jest ciężkie, nawet jeśli mamy na myśli jedynie bieżące zarządzanie nazwami użytkownika oraz wszystkimi niezbędnymi zmianami i rotacjami. Spada ono na barki administratorów sieci, którzy muszą utrzymywać porządek i synchronizować ze sobą kilka baz danych zawierających dane logowania użytkowników. Sytuację znacznie uprościłoby rozwiązanie pozwalające użytkownikowi uzyskiwać uwierzytelnienie i autoryzację w ramach jednego procesu w obrębie jednej bazy. Na podstawie wyników tego procesu sieć mogłaby tworzyć certyfikaty lub tokeny, które służyłyby użytkownikowi podczas pracy. Byłby to mniej więcej odpowiednik systemu Microsoft Passport, który na podstawie jednego logowania daje użytkownikowi dostęp do wielu usług.

Ten postępujący trend do wprowadzenia połączonych procesów uwierzytelniania i autoryzacji nosi nazwę **zarządzania tożsamością** lub **sfederowanego zarządzania tożsamością**. Mówi się o jego sfederowaniu, ponieważ umożliwia unifikację — jedno logowanie zastępuje dostęp do kilku różnych aplikacji, na przykład logowania do sieci, rozliczenia rozmów międzymiastowych, dostępu do poczty elektronicznej, dostępu do bezpiecznych rejonów sieci bez podawania numeru PIN.

Podstawą idei zarządzania tożsamością jest nasze dążenie do zminimalizowania trudów związanych z kontrolą różnych haseł. Jest to potrzeba ekonomiczna, która wiąże się z możliwością zlikwidowania nadmiarowych baz danych, które przechowują informacje o tej samej grupie ludzi. Jest to również potrzeba prawna, ponieważ w ten sposób będziemy lepiej chronić prywatność użytkowników, którzy w katalogach firmy przechowują swoje osobiste dane, takie jak informacje medyczne, te o ich upośledzeniach lub o stanie zdrowia. Taka zmiana mogłaby na przykład zaowocować możliwością wysyłania wiadomości e-mail informacyjnych do osób niewidzących lub niedowidzących, pracujących w firmie, w formacie odpowiednim dla ich oprogramowania lektorskiego.

Nastąpił gwałtowny rozwój oprogramowania i standardów, umożliwiających tworzenie międzyaplikacyjnych systemów zarządzania tożsamością. Przykładem może tu być język Security Assertion Markup Language (SAML), który dostarcza jednego rozwiązania, pozwalającego ustalić tożsamość użytkownika końcowego. Korporacje Boeing i Southwest Airlines rozwijają wspólnie projekt oparty o integrację aplikacji, wykorzystującą SAML, który udostępnia mechanizm linii lotniczych dostęp do instrukcji napraw, które znajdują się na serwerach należących do Boeinga. Dostęp jest udzielany na podstawie pojedynczego logowania. Kolejnym protokołem, który niebawem zostanie zatwierdzony, jest SPML (ang. *Service Provisioning Markup Language*). Jest istotny z tego powodu, że zintegruje systemy świadczące usługi zaopatrywania kont użytkowników.

Presje prawne i finansowe

Wszystkie instytucje stają dziś przed obowiązkiem zapewnienia ochrony prywatności użytkowników ich systemów i dokładności danych finansowych firmy oraz nadzorowania i rejestrowania swoich wysiłków. Są to działania, które wymuszają na nich nowo wprowadzone prawa. Najważniejsze spośród nich⁷ to ustawa *Sarbanes-Oxley* (zgodność ze standardem przechowywania danych finansowych), ustawa *Gramm-Leach-Bliley* (ochrona prywatnych danych finansowych), ustawa *Health Insurance Portability and Accountability Act* (bezpieczne przesyłanie i przechowywanie danych medycznych), ustawa *Family Educational Rights and Privacy Act* (ochrona danych uczniów i studentów) oraz ustawa *U.S. Patriot* (antyterroryzm i egzekwowanie prawa). W Wielkiej Brytanii wprowadzono ustawę *U.K. Data Protection and Regulation of Investigatory Powers*⁸. Przedsiębiorcy próbują jak najszybciej wprowadzać w życie te często radykalne zmiany, mając na szczególnej uwadze wszystko, co wiąże się z zachowaniem prywatności. Znacznie łatwiej jest chronić dane, które znajdują się w jednej, dobrze chronionej bazie danych, niż takie, które są rozproszone w wielu bazach. Wizja kar pieniężnych, a nawet więzienia, za niedopatrzenia związane z przestrzeganiem nowego prawa skutecznie zachęca do zwrócenia baczniejszej uwagi na zaawansowane zarządzanie i kontrolę danych.

Polskie prawodawstwo

Zakres podmiotowy ustawy o ochronie danych osobowych, określony został możliwie szeroko, w celu zapewnienia właściwej kontroli przetwarzania danych osobowych. Ustawa ta obowiązuje szeroko pojętą sferę publiczną i prywatną. Oznacza to, że wymagania określone w ustawie i aktach wykonawczych muszą być spełnione przez przedsiębiorstwa i inne podmioty przetwarzające dane osobowe w ramach swojej działalności, jak również przez podmioty, które prowadzą jedynie zbiory danych osobowych pracowników (akta pracownicze, zbiory kadrowe).

Ponadto forma przetwarzania danych osobowych nie ma znaczenia w kontekście obowiązywania zakresu przedmiotowego ustawy - przepisy ustawy mają zastosowanie zarówno do podmiotów przetwarzających dane w formie tradycyjnej - jak i w formie elektronicznej (bazy danych, systemy informatyczne). Każdy, tj. osoby prawne, osoby fizyczne, jednostki nie posiadające osobowości prawnej oraz jednostki administracji, kto choćby przechowuje lub korzysta w inny sposób z danych osobowych bez względu na strukturę w jakiej są wykorzystywane oraz

⁷ W Stanach Zjednoczonych — *przyp. tłum.*

⁸ Ustawa dotycząca ochrony danych i regulacji uprawnień śledczych — *przyp. tłum.*

fazę przetwarzania danych musi mieć świadomość, że jego działanie i zaniechanie w tym zakresie podlega ocenie pod kątem przepisów ustawy o ochronie danych osobowych oraz aktów wykonawczych.

Podsumowanie

Początek rozdziału obejmował wprowadzenie do prostych systemów uwierzytelniania, polegających na podaniu nazwy użytkownika i hasła. Przedstawił także sposoby kontynuacji tego rozwiązania — różne systemy uwierzytelniania i autoryzacji użytkownika, w tym systemy RADIUS, TACACS, DIAMETER oraz Kerberos. Rozwój technik AAA jest kontynuowany przez usługi katalogowe, jak X.500 i protokół LDAP. Na końcu poświęciliśmy kilka słów idei przyszłości — zarządzaniu tożsamością.

Odpowiednio dobrane hasła są pierwszym buforem chroniącym firmę przed atakami. Stworzenie silnego hasła i częste jego zmienianie nie jest zadaniem trudnym, ale użytkownicy czują opór przed tymi działaniami albo niwelują pożytek z nich płynący, zapisując swoje hasła i zostawiając je w widocznych miejscach. Złożone systemy uwierzytelniania, również te uwzględniające elementy biometryczne, zdają się obiecywać solidną ochronę sieci i urządzeń.

Jeżeli zajdzie potrzeba rozszerzenia stosowanych już zabezpieczeń, z pomocą mogą nam przyjść tokeny, dodatkowy czynnik procesu logowania. Ta metoda uwierzytelniania tworzy kolejną warstwę ochronną w złożonym systemie zabezpieczeń. Inne wspomniane metody to hasła jednorazowe, system Kerberos i urządzenia biometryczne.

Ostatnio pojawiające się rozwiązania, oparte o logowanie do systemów katalogowych, pozwalają łączyć procesy uwierzytelniania, autoryzacji i rozliczania (AAA — *przyp. tłum.*). Włączanie baz danych wymagających wielokrotnego uwierzytelniania do federacyjnego systemu zarządzania tożsamością, który umożliwia dostęp do różnych źródeł informacji za pomocą jednego logowania, powoduje, że do tego procesu można włączyć coraz więcej prywatnych informacji dotyczących użytkownika.