

SYSTEMY OPERACYJNE

Architektura, funkcjonowanie i projektowanie



WILLIAM STALLINGS

WYDANIE IX

Helion

Tytuł oryginału: Operating Systems: Internals and Design Principles (9th Edition)

Tłumaczenie: Zdzisław Płoski

ISBN: 978-83-283-3759-6

Authorized translation from the English language edition, entitled: OPERATING SYSTEMS: INTERNALS AND DESIGN PRINCIPLES, Ninth Edition; ISBN 0134670957; by William Stallings; published by Pearson Education, Inc. Copyright © 2018, 2015, 2012, 2009 by Pearson Education, Inc., Hoboken, New Jersey 07030.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from Pearson Education, Inc. Polish language edition published by HELION S.A. Copyright © 2018.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiejkolwiek 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 Wydawnictwo HELION 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 Wydawnictwo HELION nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Wydawnictwo HELION
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/syopa9>
Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

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

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

SPIS TREŚCI

PRZEDMOWA	23
O AUTORZE	33
CZĘŚĆ I. PODSTAWY	35
ROZDZIAŁ 1. RZUT OKA NA SYSTEM KOMPUTEROWY	37
1.1. PODSTAWOWE SKŁADOWE	38
1.2. EWOLUCJA MIKROPROCESORÓW	40
1.3. WYKONANIE ROZKAZU	40
1.4. PRZERWANIA	43
Przerwania i cykl rozkazowy	45
Przetwarzanie przerwania	47
Przerwania wielokrotne	50
1.5. HIERARCHIA PAMIĘCI	53
1.6. PAMIĘĆ PODRĘCZNA	56
Motywy	56
Zasady działania pamięci podręcznej	56
Projektowanie pamięci podręcznej	58
1.7. BEZPOŚREDNI DOSTĘP DO PAMIĘCI	60
1.8. ORGANIZACJA WIELOPROCESOROWA I WIELORDZENIOWA	61
Wieloprocesory symetryczne	62
Komputery wielordzeniowe	64
1.9. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	66
Podstawowe pojęcia	66
Pytania sprawdzające	66
Zadania	67
DODATEK 1A. CHARAKTERYSTYKA WYDAJNOŚCI PAMIĘCI DWUPOZIOMOWYCH	69
Lokalność	70
Działanie pamięci dwupoziomowej	72
Wydaźność	73

ROZDZIAŁ 2. PRZEGLĄD SYSTEMÓW OPERACYJNYCH	77
2.1. CELE I FUNKCJE SYSTEMU OPERACYJNEGO	78
System operacyjny jako interfejs użytkownik – komputer	79
System operacyjny jako zarządca zasobów	81
Łatwość ewolucji systemu operacyjnego	82
2.2. ROZWÓJ SYSTEMÓW OPERACYJNYCH	83
Przetwarzanie seryjne	83
Proste systemy wsadowe	84
Wieloprogramowe systemy wsadowe	87
Systemy z podziałem czasu	90
2.3. NAJWAŻNIEJSZE OSIĄGNIĘCIA	92
Proces	93
Zarządzanie pamięcią	96
Ochrona informacji i bezpieczeństwo	99
Planowanie operowania zasobami	99
2.4. DROGA DO WSPÓŁCZESNYCH SYSTEMÓW OPERACYJNYCH	101
2.5. TOLEROWANIE AWARII	104
Podstawowe pojęcia	104
Wady	106
Mechanizmy systemu operacyjnego	106
2.6. PROBLEMY PROJEKTOWANIA SYSTEMÓW OPERACYJNYCH WIELOPROCESORÓW I KOMPUTERÓW WIELORDZENIOWYCH	107
Rozważania dotyczące wieloprocesorowego symetrycznego SO	107
Rozważania dotyczące wielordzeniowych SO	108
2.7. PRZEGLĄD SYSTEMU MICROSOFT WINDOWS	110
Rodowód	110
Architektura	111
Model klient-serwer	114
Wątki i SMP	115
Obiekty systemu Windows	116
2.8. TRADYCYJNE SYSTEMY UNIKSOWE	118
Historia	118
Opis	119
2.9. NOWOCZESNE SYSTEMY UNIKSOWE	120
System V Release 4 (SVR4)	122
System BSD	122
Solaris 11	122
2.10. LINUX	123
Historia	123
Struktura modularna	124
Składowe jądra	126
2.11. ANDROID	129
Architektura oprogramowania Androida	129
Środowisko wykonawcze Androida	132
Architektura systemu Android	135
Czynności	136
Zarządzanie zasilaniem	136

2.12. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	137
Podstawowe pojęcia	137
Pytania sprawdzające	138
Zadania	138

CZĘŚĆ II. PROCESY 141

ROZDZIAŁ 3. POJĘCIE PROCESU I PRZEBIEG STEROWANIA 143

3.1. CZYM JEST PROCES?	145
Podstawy	145
Procesy i bloki kontrolne procesów	146
3.2. STANY PROCESU	147
Dwustanowy model procesu	150
Tworzenie i likwidowanie procesów	151
Model pięciostanowy	153
Procesy zawieszone	157
3.3. OPIS PROCESU	162
Struktury sterowania systemu operacyjnego	163
Struktury sterowania procesu	165
3.4. STEROWANIE PROCESAMI	172
Tryby wykonywania	172
Tworzenie procesów	174
Przełączanie procesów	174
3.5. WYKONYWANIE SYSTEMU OPERACYJNEGO	178
Jądro nieprocesowe	178
Wykonywanie w procesach użytkownika	179
System operacyjny oparty na procesach	180
3.6. ZARZĄDZANIE PROCESAMI W SYSTEMIE UNIX SVR4	181
Stany procesu	181
Opis procesu	183
Sterowanie procesami	186
3.7. PODSUMOWANIE	186
3.8. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	187
Podstawowe pojęcia	187
Pytania sprawdzające	187
Zadania	188

ROZDZIAŁ 4. WĄTKI 193

4.1. PROCESY I WĄTKI	194
Wielowątkowość	195
Funkcjonowanie wątków	198
4.2. RODZAJE WĄTKÓW	200
Wątki poziomu użytkownika i wątki poziomu jądra	200
Inne organizacje	205
4.3. WIELORDZENIOWOŚĆ I WIELOWĄTKOWOŚĆ	207
Wydajność oprogramowania na wielu rdzeniach	207
Przykład zastosowania: oprogramowanie gier Valve	210

4.4. ZARZĄDZANIE PROCESAMI I WĄTKAMI W SYSTEMIE WINDOWS	212
Zarządzanie zadaniami drugoplanowymi i cyklami istnienia aplikacji	213
Proces w systemie Windows	214
Obiekty procesów i wątków	215
Wielowątkowość	217
Stany wątków	217
Zaplecze podsystemów SO	218
4.5. ZARZĄDZANIE WĄTKAMI I WIELOPRZETWARZANIEM SYMETRYCZNYM W SYSTEMIE SOLARIS	219
Architektura wielowątkowa	219
Uzasadnienie	220
Struktura procesu	220
Wykonanie wątku	222
Przerwania jako wątki	223
4.6. ZARZĄDZANIE PROCESAMI I WĄTKAMI W SYSTEMIE LINUX	223
Prace Linuxa	223
Wątki Linuxa	225
Przestrzenie nazw Linuxa	226
4.7. ZARZĄDZANIE PROCESAMI I WĄTKAMI W SYSTEMIE ANDROID	229
Aplikacje Androida	229
Czynności	230
Procesy i wątki	232
4.8. WIELKA CENTRALNA EKSPEDYCJA W SYSTEMIE MAC OS X	233
4.9. PODSUMOWANIE	235
4.10. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	236
Podstawowe pojęcia	236
Pytania sprawdzające	236
Zadania	236
ROZDZIAŁ 5. WSPÓLBIEŻNOŚĆ — WZAJEMNE WYKLUCZANIE I SYNCHRONIZACJA	243
5.1. WZAJEMNE WYKLUCZANIE — PODEJŚCIA PROGRAMOWE	246
Algorytm Dekkera	246
Algorytm Petersona	250
5.2. PODSTAWY WSPÓLBIEŻNOŚCI	251
Prosty przykład	252
Szkodliwa rywalizacja	254
Kwestie związane z systemem operacyjnym	255
Interakcja procesów	255
Wymagania dotyczące wzajemnego wykluczania	259
5.3. WZAJEMNE WYKLUCZANIE — ZAPLECZE SPRZĘTOWE	260
Blokowanie przerwań	260
Specjalne rozkazy maszynowe	261
5.4. SEMAFORY	263
Wzajemne wykluczanie	269
Problem producenta-konsumenta	269
Implementacja semaforów	275

5.5. MONITORY	277
Monitor z sygnałem	277
Alternatywny model monitora z powiadamianiem i rozgłaszaniem	281
5.6. PRZEKAZYWANIE KOMUNIKATÓW	283
Synchronizacja	284
Adresowanie	285
Format komunikatów	287
Dyscyplina kolejkowania	287
Wzajemne wykluczanie	288
5.7. PROBLEM CZYTELNIKÓW I PISARZY	290
Czytelnicy mają pierwszeństwo	291
Pisarze mają pierwszeństwo	291
5.8. PODSUMOWANIE	294
5.9. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	295
Podstawowe pojęcia	295
Pytania sprawdzające	295
Zadania	296

ROZDZIAŁ 6. WSPÓLBIEŻNOŚĆ — ZAKLESZCZENIA I GŁODZENIE 311

6.1. NA CZYM POLEGA ZAKLESZCZENIE	312
Zasoby nieużywalne	316
Zasoby zużywalne	317
Grafy przydziału zasobów	318
Warunki występowania zakleszczenia	320
6.2. ZAPOBIEGANIE ZAKLESZCZENIOM	321
Wzajemne wykluczanie	321
Przetrzykiwanie i oczekiwanie	321
Brak wyłączeń	322
Czekanie cykliczne	322
6.3. UNIKANIE ZAKLESZCZEŃ	322
Odmowa wszczynania procesu	323
Odmowa przydziału zasobu	324
6.4. WYKRYWANIE ZAKLESZCZEŃ	328
Algorytm wykrywania zakleszczenia	328
Rekonstrukcja (uzdrawianie)	329
6.5. POŁĄCZONE STRATEGIE POSTĘPOWANIA Z ZAKLESZCZENIAMI	330
6.6. PROBLEM OBIADUJĄCYCH FILOZOFÓW	331
Rozwiązanie z użyciem semaforów	332
Rozwiązanie z użyciem monitora	333
6.7. MECHANIZMY WSPÓLBIEŻNOŚCI W SYSTEMIE UNIX	334
Potoki	335
Komunikaty	335
Pamięć dzielona	335
Semafore	336
Sygnały	336

6.8. MECHANIZMY WSPÓŁBIEŻNOŚCI W JĄDRZE LINUXA	337
Operacje niepodzielne	338
Wirujące blokady	339
Semafore	342
Bariery	343
6.9. ELEMENTARNE OPERACJE SYNCHRONIZACJI W SYSTEMIE SOLARIS	345
Zamek wzajemnego wykluczania	346
Semafore	347
Blokada czytelnicy – pisarz	347
Zmienne warunków	348
6.10. MECHANIZMY SYNCHRONIZACJI W SYSTEMIE WINDOWS	348
Funkcje czekania	348
Obiekty dyspozytora	349
Sekcje krytyczne	350
Wąskie blokady czytelnicy – pisarze i zmienne warunków	351
Synchronizacja bez blokad	351
6.11. KOMUNIKACJA MIĘDZYPROCESOWA W SYSTEMIE ANDROID	352
6.12. PODSUMOWANIE	353
6.13. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	354
Pytania sprawdzające	354
Zadania	354

CZĘŚĆ III. PAMIĘĆ 363

ROZDZIAŁ 7. ZARZĄDZANIE PAMIĘCIĄ	365
7.1. WYMAGANIA DOTYCZĄCE ZARZĄDZANIA PAMIĘCIĄ	366
Przemieszczanie	367
Ochrona	368
Współużytkowanie	368
Organizacja logiczna	369
Organizacja fizyczna	369
7.2. PODZIAŁ PAMIĘCI	370
Podział stały	371
Podział dynamiczny	374
System kumpłowski	377
Przemieszczanie	380
7.3. STRONICOWANIE	381
7.4. SEGMENTACJA	385
7.5. PODSUMOWANIE	386
7.6. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	386
Podstawowe pojęcia	386
Pytania sprawdzające	387
Zadania	387
Dodatek 7A. ŁADOWANIE I KONSOLIDACJA	390
Ładowanie	390
Konsolidacja	394

ROZDZIAŁ 8. PAMIĘĆ WIRTUALNA	397
8.1. SPRZĘT I STRUKTURY STEROWANIA	398
Lokalność i pamięć wirtualna	400
Stronicowanie	402
Segmentacja	413
Połączenie stronicowania i segmentacji	414
Ochrona i dzielenie	415
8.2. OPROGRAMOWANIE SYSTEMU OPERACYJNEGO	417
Zasady pobierania	418
Zasady umieszczania	419
Zasady zastępowania	419
Zarządzanie zbiorem rezydującym	426
Zasady czyszczenia	433
Kontrola załadowania	434
8.3. ZARZĄDZANIE PAMIĘCIĄ W SYSTEMACH UNIX I SOLARIS	436
System stronicowania	436
Alokator pamięci jądra	439
8.4. ZARZĄDZANIE PAMIĘCIĄ W LINUXIE	440
Pamięć wirtualna Linuxa	441
Przydział pamięci w jądrze	443
8.5. ZARZĄDZANIE PAMIĘCIĄ W SYSTEMIE WINDOWS	445
Mapa adresów wirtualnych w systemie Windows	445
Stronicowanie w systemie Windows	445
Wymiana w systemie Windows	447
8.6. ZARZĄDZANIE PAMIĘCIĄ W ANDROIDZIE	447
8.7. PODSUMOWANIE	448
8.8. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	449
Podstawowe pojęcia	449
Pytania sprawdzające	449
Zadania	450
 CZĘŚĆ IV. PLANOWANIE	 455

ROZDZIAŁ 9. PLANOWANIE JEDNOPROCESOROWE	457
9.1. RODZAJE PLANOWANIA PROCESORA	459
Planowanie długoterminowe	461
Planowanie średnioterminowe	461
Planowanie krótkoterminowe	462
9.2. ALGORYTMY PLANOWANIA	462
Kryteria planowania krótkoterminowego	462
Zastosowanie priorytetów	463
Alternatywne zasady planowania	465
Porównanie efektywności	478
Planowanie uczciwych udziałów	482
9.3. TRADYCYJNE PLANOWANIE UNIKSOWE	486
9.4. PODSUMOWANIE	488

9.5. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	489
Podstawowe pojęcia	489
Pytania sprawdzające	489
Zadania	490

ROZDZIAŁ 10. PLANOWANIE WIELOPROCESOROWE, WIELORDZENIOWE I W CZASIE RZECZYWISTYM	495
10.1. PLANOWANIE WIELOPROCESOROWE I WIELORDZENIOWE	496
Ziarnistość	497
Zagadnienia projektowe	498
Planowanie procesów	500
Planowanie wątków	502
Wielordzeniowe planowanie wątków	508
10.2. PLANOWANIE CZASU RZECZYWISTEGO	509
Podstawy	509
Charakterystyka systemów czasu rzeczywistego	510
Planowanie czasu rzeczywistego	513
Planowanie terminów nieprzekraczalnych	515
Planowanie monotonicznego tempa	519
Odwroćenie priorytetów	522
10.3. PLANOWANIE W SYSTEMIE LINUX	525
Planowanie czasu rzeczywistego	525
Planowanie poza czasem rzeczywistym	526
10.4. PLANOWANIE W SYSTEMIE UNIX SVR4	528
10.5. PLANOWANIE W SYSTEMIE UNIX FREEBSD	530
Klasy priorytetów	530
SMP i zaplecze wielordzeniowości	531
10.6. PLANOWANIE W SYSTEMIE WINDOWS	533
Priorytety procesów i wątków	534
Planowanie wieloprocessorowe	536
10.7. PODSUMOWANIE	536
10.8. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	537
Podstawowe pojęcia	537
Pytania sprawdzające	537
Zadania	537

CZĘŚĆ V. WEJŚCIE-WYJŚCIE I PLIKI 541

ROZDZIAŁ 11. ZARZĄDZANIE WEJŚCIEM-WYJŚCIEM I PLANOWANIE DYSKÓW	543
11.1. URZĄDZENIA WEJŚCIA-WYJŚCIA	544
11.2. ORGANIZACJA FUNKCJI WEJŚCIA-WYJŚCIA	546
Ewolucja funkcji wejścia-wyjścia	547
Bezpośredni dostęp do pamięci	547
11.3. ZAGADNIENIA PROJEKTOWANIA SYSTEMU OPERACYJNEGO	549
Cele projektowe	549
Logiczna struktura funkcji wejścia-wyjścia	550

11.4. BUFOROWANIE WEJŚCIA-WYJŚCIA	552
Bufor pojedynczy	553
Bufor podwójny	555
Bufor cykliczny	555
Użyteczność buforowania	555
11.5. PLANOWANIE DYSKÓW	556
Parametry wydajnościowe dysku	556
Sposoby planowania dysku	559
11.6. RAID	564
RAID — poziom 0	565
RAID — poziom 1	570
RAID — poziom 2	571
RAID — poziom 3	571
RAID — poziom 4	572
RAID — poziom 5	573
RAID — poziom 6	573
11.7. PAMIĘĆ PODRĘCZNA DYSKU	574
Zagadnienia projektowe	574
Zagadnienia dotyczące wydajności	576
11.8. WEJŚCIE-WYJŚCIE W SYSTEMIE UNIX SVR4	578
Pamięć podręczna buforów	578
Kolejka znakowa	580
Niebuforowane wejście-wyjście	580
Urządzenia uniksowe	580
11.9. WEJŚCIE-WYJŚCIE W SYSTEMIE LINUX	581
Planowanie dysku	581
Linuksowa podręczna pamięć stron	584
11.10. WEJŚCIE-WYJŚCIE W SYSTEMIE WINDOWS	585
Podstawowe własności wejścia-wyjścia	585
Wejście-wyjście asynchroniczne i synchroniczne	586
Programowa realizacja RAID	587
Kopie-cienie tomów dyskowych	587
Szyfrowanie tomów	587
11.11. PODSUMOWANIE	587
11.12. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	588
Podstawowe pojęcia	588
Pytania sprawdzające	589
Zadania	589
ROZDZIAŁ 12. ZARZĄDZANIE PLIKAMI	593
12.1. W SKRÓCIE	594
Pliki i systemy plików	594
Struktura pliku	595
Systemy zarządzania plikami	597
12.2. ORGANIZACJA I DOSTĘP DO PLIKÓW	600
Serta	601
Plik sekwencyjny	603
Indeksowany plik sekwencyjny	603

Plik indeksowy	604
Plik bezpośredniego dostępu, czyli plik haszowany	605
12.3. B-DRZEWA	605
12.4. KATALOGI PLIKÓW	609
Zawartość	609
Struktura	610
Nazewnictwo	612
12.5. DZIELENIE PLIKÓW	612
Prawa dostępu	612
Dostęp jednoczesny	614
12.6. BLOKOWANIE REKORDÓW	615
12.7. ZARZĄDZANIE PAMIĘCIĄ DRUGORZĘDNĄ	616
Lokowanie plików	616
Zarządzanie wolną przestrzenią	622
Tomy, inaczej — woluminy	624
Niezawodność	624
12.8. ZARZĄDZANIE PLIKAMI W SYSTEMIE UNIX	625
I-węzły	626
Umieszczanie plików	628
Katalogi	629
Struktura tomu	630
12.9. WIRTUALNY SYSTEM PLIKÓW W LINUXIE	630
Obiekt superbloku	632
Obiekt i-węzła	633
Obiekt k-wpisu	633
Obiekt pliku	633
Pamięci podręczne	634
12.10. SYSTEM PLIKÓW WINDOWS	634
Podstawowe cechy NTFS	634
Wolumin NTFS i struktura pliku	635
Odtwarzalność	638
12.11. ZARZĄDZANIE PLIKAMI W ANDROIDZIE	639
System plików	639
SQLite	641
12.12. PODSUMOWANIE	641
12.13. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	642
Podstawowe pojęcia	642
Pytania sprawdzające	643
Zadania	643

CZĘŚĆ VI. SYSTEMY WBUDOWANE 645

ROZDZIAŁ 13. WBUDOWANE SYSTEMY OPERACYJNE 647

13.1. SYSTEMY WBUDOWANE	648
Pojęcie systemu wbudowanego	648
Procesory aplikacji a procesory do zadań specjalnych	650
Mikroprocesory	650

Mikrokontrolery	652
Systemy głęboko wbudowane	653
13.2. CHARAKTERYSTYKA WBUDOWANYCH SYSTEMÓW OPERACYJNYCH	653
Środowiska macierzyste i docelowe	655
Metody opracowywania	657
Adaptacja istniejących handlowych systemów operacyjnych	657
Wbudowany system operacyjny skonstruowany w określonym celu	657
13.3. WBUDOWANY LINUX	658
Charakterystyka wbudowanego Linuxa	658
Systemy plików wbudowanego Linuxa	661
Zalety wbudowanego Linuxa	661
μClinux	662
Android	665
13.4. TinyOS	665
Bezprzewodowe sieci czujników	665
Cele systemu TinyOS	666
Składowe TinyOS-a	668
Planista TinyOS-a	670
Przykładowa konfiguracja	671
Interfejs zasobów TinyOS-a	673
13.5. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	675
Podstawowe pojęcia	675
Pytania sprawdzające	675
Zadania	675
ROZDZIAŁ 14. MASZyny WIRTUALNE	679
14.1. KONCEPCJA MASZyny WIRTUALNEJ	680
14.2. HIPERWIZORY	683
Hiperwizory	684
Parawirtualizacja	686
Wirtualizacja z asystą sprzętu	687
Aplikacje wirtualne	688
14.3. WIRTUALIZACJA KONTENEROWA	688
Jądrowe grupy sterowania	689
Pojęcie kontenera	689
Kontenerowy system plików	693
Mikrouslugi	694
Docker	694
14.4. ZAGADNIENIA DOTYCZĄCE PROCESORÓW	695
14.5. ZARZĄDZANIE PAMIĘCIĄ	697
14.6. ZARZĄDZANIE WEJŚCIEM-WYJŚCIEM	699
14.7. HIPERWIZOR VMWARE ESXI	701
14.8. MICROSOFTOWY HYPER-V I ODMIANY XENA	703
14.9. MASZYNA WIRTUALNA JAVY	705
14.10. ARCHITEKTURA MASZyny WIRTUALNEJ LINUX VSERVER	706
Architektura	706
Planowanie procesów	707

14.11. PODSUMOWANIE	709
14.12. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	709
Podstawowe pojęcia	709
Pytania (problemy) sprawdzające	710
Zadania	710

ROZDZIAŁ 15. BEZPIECZEŃSTWO SYSTEMÓW OPERACYJNYCH 711

15.1. INTRUZI I ZŁOŚLIWE OPROGRAMOWANIE	712
Zagrożenia dostępu do systemu	712
Środki zaradcze	714
15.2. PRZEPEŁNIENIE BUFORA	716
Ataki z przepełnieniem bufora	717
Obrona w fazie kompilacji	720
Obrona w fazie wykonania	723
15.3. KONTROLOWANIE DOSTĘPU	725
Kontrolowanie dostępu w systemie plików	725
Zasady kontrolowania dostępu	727
15.4. KONTROLOWANIE DOSTĘPU W SYSTEMIE UNIX	733
Konwencjonalna kontrola dostępu do plików uniksowych	733
Listy kontroli dostępu w UNIX-ie	735
15.5. HARTOWANIE SYSTEMÓW OPERACYJNYCH	736
Instalowanie systemu operacyjnego — ustawienia początkowe i łatanie	737
Usuwanie zbędnych usług, aplikacji i protokołów	738
Konfigurowanie użytkowników, grup i uwierzytelniania	739
Kształtowanie kontroli zasobów	739
Wprowadzanie dodatkowych środków bezpieczeństwa	740
Testowanie bezpieczeństwa systemu	740
15.6. DBAŁOŚĆ O BEZPIECZEŃSTWO	741
Rejestrowanie zdarzeń	741
Składowanie i archiwizowanie danych	742
15.7. BEZPIECZEŃSTWO W SYSTEMIE WINDOWS	742
Schemat kontroli dostępu	743
Żeton dostępu	743
Deskryptory bezpieczeństwa	744
15.8. PODSUMOWANIE	748
15.9. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	748
Podstawowe pojęcia	748
Pytania sprawdzające	749
Zadania	749

ROZDZIAŁ 16. SYSTEMY OPERACYJNE CHMUR I INTERNETU RZECZY 753

16.1. OBLICZENIA W CHMURZE	754
Elementy chmury obliczeniowej	754
Modele usług chmurowych	756
Modele realizacyjne chmur	757
Wzorcowa architektura chmury obliczeniowej	760

16.2. CHMUROWE SYSTEMY OPERACYJNE	763
Infrastruktura jako usługa	763
Wymagania na chmurowy system operacyjny	765
Ogólna architektura chmurowego systemu operacyjnego	766
OpenStack	772
16.3. INTERNET RZECZY (IR)	780
Rzeczy w internecie rzeczy	781
Rozwój IR	781
Elementy urządzeń podłączonych do IR	782
IR i kontekst chmury	782
16.4. SYSTEMY OPERACYJNE IR	785
Urządzenia ograniczone	785
Wymagania dotyczące SOIR	787
Architektura SOIR	789
RIOT	790
16.5. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	793
Podstawowe pojęcia	793
Pytania sprawdzające	793

DODATKI 795

DODATEK A. ZAGADNIENIA WSPÓLBIEŻNOŚCI 797

A.1. SZKODLIWA RYWALIZACJA I SEMAFORY	798
Sformułowanie problemu	798
Pierwsza próba	798
Druga próba	800
Trzecia próba	801
Czwarta próba	802
Dobra próba	803
A.2. PROBLEM ZAKŁADU FRYZJERSKIEGO	805
Nieuczciwy zakład fryzjerski	805
Uczciwy zakład fryzjerski	808
A.3. Zadania	810

DODATEK B. PROJEKTY I ZADANIA Z PROGRAMOWANIA SYSTEMÓW OPERACYJNYCH 811

B.1. PROJEKTY DOTYCZĄCE SEMAFORÓW	812
B.2. PROJEKT DOTYCZĄCY SYSTEMÓW PLIKÓW	813
B.3. OS/161	813
B.4. SYMULACJE	814
B.5. PROJEKTY PROGRAMISTYCZNE	815
Projekty zdefiniowane w podręczniku	815
Dodatkowe większe projekty programistyczne	815
Małe projekty programistyczne	816
B.6. PROJEKTY BADAWCZE	816
B.7. ZADANIA TYPU „PRZECZYTAJ I ZDAJ SPRAWĘ”	817

B.8. PRACE DO NAPISANIA 817

B.9. TEMATY DO DYSKUSJI 817

B.10. BACI 818

LITERATURA 819

ŹRÓDŁA I ZASADY ICH UDOSTĘPNIENIA 831

SKOROWIDZ 833

SPIS TREŚCI ROZDZIAŁÓW ONLINE

ROZDZIAŁ 17. PROTOKOŁY SIECIOWE	1005
17.1. ZAPOTRZEBOWANIE NA ARCHITEKTURĘ PROTOKOŁÓW	1007
17.2. ARCHITEKTURA PROTOKOŁÓW TCP/IP	1010
Warstwy TCP/IP	1010
Protokoły TCP i UDP	1011
Protokoły IP i IPv6	1012
Działanie TCP/IP	1014
Zastosowania TCP/IP	1016
17.3. GNIAZDA	1017
Pojęcie gniazda	1017
Wywołania interfejsu gniazd	1018
17.4. PRACA SIECIOWA W SYSTEMIE LINUX	1019
Wysyłanie danych	1022
Odbieranie danych	1022
17.5. PODSUMOWANIE	1023
17.6. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	1023
Podstawowe pojęcia	1023
Pytania sprawdzające	1024
Zadania	1024
Dodatek 17A. TFTP — BANALNY PROTOKÓŁ PRZESYŁANIA PLIKÓW	1027
Wprowadzenie do TFTP	1027
Pakiety TFTP	1027
Rzut oka na przesyłanie	1029
Błędy i opóźnienia	1030
Składnia, semantyka i koordynacja w czasie	1031
 ROZDZIAŁ 18. PRZETWARZANIE ROZPROSZONE, KLIENT-SERWER I GRONA	1033
18.1. OBLICZENIA W UKŁADZIE KLIENT-SERWER	1034
Co to są obliczenia klient-serwer?	1034
Aplikacje klient-serwer	1036
Warstwa pośrednia	1043
18.2. ROZPROSZONE PRZEKAZYWANIE KOMUNIKATÓW	1045
Niezawodność a zawodność	1047
Blokowanie a nieblokowanie	1048

18.3. ZDALNE WYWOŁANIA PROCEDUR	1048
Przekazywanie parametrów	1050
Reprezentowanie parametrów	1050
Wiązanie klienta z serwerem	1050
Synchroniczne czy niesynchroniczne	1051
Mechanizmy obiektowe	1051
18.4. GRONA, CZYLI KLASTRY	1052
Konfigurowanie grom	1053
Zagadnienia projektowe systemów operacyjnych	1055
Architektura grona komputerów	1057
Grona w porównaniu z SMP	1058
18.5. SERWER GRONA W SYSTEMIE WINDOWS	1058
18.6. BEOWULF I GRONA LINUKSOWE	1060
Właściwości Beowulfa	1060
Oprogramowanie Beowulfa	1061
18.7. PODSUMOWANIE	1062
18.8. LITERATURA	1063
18.9. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	1063
Podstawowe pojęcia	1063
Pytania sprawdzające	1064
Zadania	1064
ROZDZIAŁ 19. ROZPROSZONE ZARZĄDZANIE PROCESAMI	1067
19.1. WĘDRÓWKA PROCESÓW	1068
Uzasadnienie	1068
Mechanizmy wędrówki procesów	1069
Negocjowanie wędrówki	1073
Eksmisja	1074
Przeniesienia z wywłaszczaniem lub bez wywłaszczania	1075
19.2. ROZPROSZONE STANY GLOBALNE	1075
Stany globalne i migawki rozproszone	1075
Algorytm migawki rozproszonej	1078
19.3. ROZPROSZONE WZAJEMNE WYKLUCZANIE	1080
Koncepcje rozproszonego wzajemnego wykluczania	1080
Porządkowanie zdarzeń w systemie rozproszonym	1083
Kolejka rozproszona	1086
Metoda przekazywania żetonu	1089
19.4. ZAKLESZCZENIE ROZPROSZONE	1091
Zakleszczenie w przydziale zasobów	1091
Zakleszczenie w przekazywaniu komunikatów	1097
19.5. PODSUMOWANIE	1102
19.6. LITERATURA	1102
19.7. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA	1104
Podstawowe pojęcia	1104
Pytania sprawdzające	1104
Zadania	1105

ROZDZIAŁ 20. PRAWDOPODOBIENSTWO I PROCESY STOCHASTYCZNE W ZARYSIE	1107
20.1. PRAWDOPODOBIENSTWO	1108
Definicje prawdopodobieństwa	1108
Prawdopodobieństwo warunkowe i niezależność	1111
Twierdzenie Bayesa	1112
20.2. ZMIENNE LOSOWE	1113
Funkcje rozkładu i gęstości	1113
Ważne rozkłady	1114
Wiele zmiennych losowych	1117
20.3. ELEMENTARNE KONCEPCJE PROCESÓW STOCHASTYCZNYCH	1118
Statystyka pierwszego i drugiego rzędu	1119
Stacjonarne procesy stochastyczne	1120
Gęstość widmowa	1121
Przyrosty niezależne	1122
Ergodyczność	1126
20.4. ZADANIA	1127
 ROZDZIAŁ 21. ANALIZA KOLEJEK	 1131
21.1. ZACHOWANIE KOLEJEK — PROSTY PRZYKŁAD	1133
21.2. PO CO ANALIZOWAĆ KOLEJKI?	1138
21.3. MODELE KOLEJEK	1140
Kolejka jednoserwerowa	1140
Kolejka wieloserwerowa	1144
Podstawowe zależności obsługi masowej	1145
Założenia	1146
21.4. KOLEJKI JEDNOSERWEROWE	1147
21.5. KOLEJKI WIELOSERWEROWE	1150
21.6. PRZYKŁADY	1150
Serwer bazy danych	1150
Obliczanie percentyli	1152
Wieloprocessor ściśle powiązany	1153
Problem wieloserwera	1154
21.7. KOLEJKI Z PRIORYTETAMI	1156
21.8. SIECI KOLEJEK	1157
Dzielenie i łączenie strumieni ruchu	1158
Kolejki posobne (tandemowe)	1159
Twierdzenie Jacksona	1159
Zastosowanie w sieci komutacji pakietów	1160
21.9. INNE MODELE KOLEJEK	1161
21.10. SZACOWANIE PARAMETRÓW MODELU	1162
Próbkowanie	1162
Błędy próbkowania	1164
21.11. LITERATURA	1165
21.12. ZADANIA	1165

PROJEKT PROGRAMISTYCZNY NR 1. OPRACOWANIE POWŁOKI	1169
PROJEKT PROGRAMISTYCZNY NR 2. POWŁOKA DYSPOZYTORA HOST	1173
DODATEK C. PROBLEMATYKA WSPÓLBIEŻNOŚCI	1181
DODATEK D. PROJEKTOWANIE OBIEKTOWE	1191
DODATEK E. PRAWO AMDAHLA	1203
DODATEK F. TABLICE HASZOWANIA	1207
DODATEK G. CZAS ODPOWIEDZI	1211
DODATEK H. POJĘCIA SYSTEMÓW KOLEJKOWANIA	1217
DODATEK I. ZŁOŻONOŚĆ ALGORYTMÓW	1223
DODATEK J. URZĄDZENIA PAMIĘCI DYSKOWEJ	1227
DODATEK K. ALGORYTMY KRYPTOGRAFICZNE	1239
DODATEK L. INSTYTUCJE NORMALIZACYJNE	1251
DODATEK M. GNIAZDA — WPROWADZENIE DLA OSÓB PROGRAMUJĄCYCH	1263
DODATEK N. MIĘDZYNARODOWY ALFABET WZORCOWY (IRA)	1291
DODATEK O. BACI — SYSTEM WSPÓLBIEŻNEGO PROGRAMOWANIA BEN-ARIEGO	1295
DODATEK P. STEROWANIE PROCEDURAMI	1307
DODATEK Q. ECOS	1313
SŁOWNIK	1329
SKOROWIDZ	1343

Rozdział 2

Przegląd systemów operacyjnych

2.1. CELE I FUNKCJE SYSTEMU OPERACYJNEGO

- System operacyjny jako interfejs użytkownik – komputer
- System operacyjny jako zarządca zasobów
- Łatwość ewolucji systemu operacyjnego

2.2. ROZWÓJ SYSTEMÓW OPERACYJNYCH

- Przetwarzanie seryjne
- Proste systemy wsadowe
- Wieloprogramowe systemy wsadowe
- Systemy z podziałem czasu

2.3. NAJWAŻNIEJSZE OSIĄGNIĘCIA

- Proces
- Zarządzanie pamięcią
- Ochrona informacji i bezpieczeństwo
- Planowanie operowania zasobami

2.4. DROGA DO WSPÓŁCZESNYCH SYSTEMÓW OPERACYJNYCH

2.5. TOLEROWANIE AWARII

- Podstawowe pojęcia
- Wady
- Mechanizmy systemu operacyjnego

2.6. PROBLEMY PROJEKTOWANIA SYSTEMÓW OPERACYJNYCH WIELOPROCESORÓW I KOMPUTERÓW WIELORDZENIOWYCH

- Rozważania dotyczące wieloprocessorowego symetrycznego SO
- Rozważania dotyczące wielordzeniowych SO

2.7. PRZEGLĄD SYSTEMU MICROSOFT WINDOWS

- Rodowód
- Architektura
- Model klient-serwer
- Wątki i SMP
- Obiekty systemu Windows

2.8. TRADYCYJNE SYSTEMY UNIKSOWE

- Historia
- Opis

2.9. NOWOCZESNE SYSTEMY UNIKSOWE

- System V Release 4 (SVR4)
- System BSD
- Solaris 11

2.10. LINUX

- Historia
- Struktura modułarna
- Składowe jądra

2.11. ANDROID

- Architektura oprogramowania Androida
- Środowisko wykonawcze Androida
- Architektura systemu Android
- Czynności
- Zarządzanie zasilaniem

2.12. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA

W TYM ROZDZIALE POZNASZ I ZROZUMIESZ:

- podsumowanie, na najogólniejszym poziomie, podstawowych funkcji SO;
- omówienie ewolucji systemów operacyjnych od wczesnych, prostych systemów wsadowych aż po nowoczesne, złożone systemy;
- krótkie wyjaśnienie każdego z głównych osiągnięć w badaniach SO, jak określono w podrozdziale 2.3;
- omówienie kluczowych obszarów projektowania, które przyczyniły się do rozwoju nowoczesnych systemów operacyjnych;
- zdefiniowanie i omówienie maszyn wirtualnych i wirtualizacji;
- zagadnienia projektowania SO wynikłe z pojawienia się organizacji wieloprocessorowych i wielordzeniowych;
- podstawową strukturę systemu Windows;
- istotne elementy tradycyjnego systemu UNIX;
- wyjaśnienie nowych właściwości występujących w nowoczesnych systemach uniksowych;
- omówienie Linuxa i jego związków z UNIX-em¹.

Nasz kurs systemów operacyjnych rozpoczynamy od krótkiego przeglądu historycznego. Jest to opowieść sama w sobie ciekawa, a ponadto służy jako przegląd podstaw SO. W pierwszym podrozdziale zajmiemy się przeznaczeniem i funkcjami systemów operacyjnych. Dalej przyjrzymy się ich ewolucji od prymitywnych systemów wsadowych do wyrafinowanych, wielozadaniowych systemów obsługujących jednocześnie wielu użytkowników. W pozostałej części rozdziału prześledzimy historię i przedstawimy ogólną charakterystykę dwóch systemów operacyjnych, które służą jako przykładowe w całej książce.

2.1. CELE I FUNKCJE SYSTEMU OPERACYJNEGO

System operacyjny (SO) jest programem, który nadzoruje wykonywanie aplikacji (programów użytkowych) i działa jak interfejs między aplikacjami a sprzętem komputerowym. Można go traktować jako mający do spełnienia trzy cele:

- **wygoda** — SO sprawia, że komputer jest wygodniejszy w użyciu;
- **sprawność** — SO umożliwia efektywne użytkowanie zasobów komputerowych;

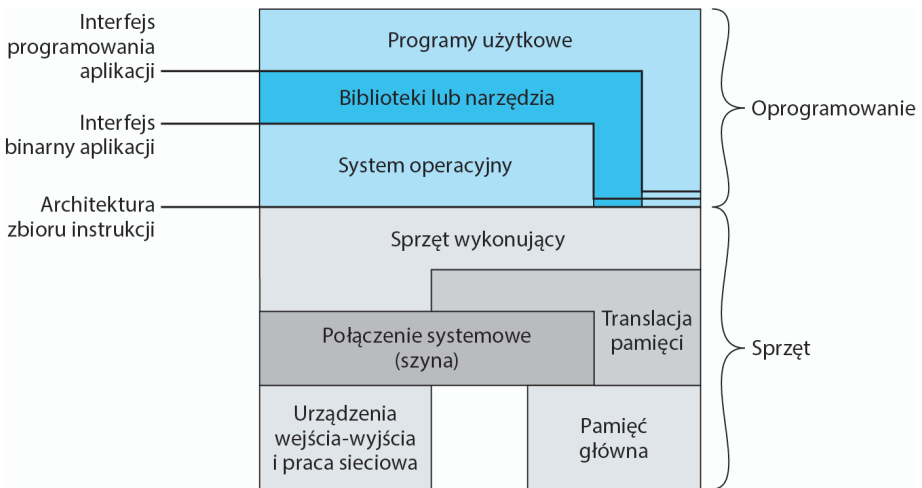
¹ Oraz podstawowych cech systemu Android, o czym tu nie wspomniano — *przyp. tłum.*

- **zdolność do ewolucji** — SO powinien być skonstruowany w taki sposób, aby można go było skutecznie rozwijać, testować i rozszerzać o nowe funkcje systemowe, nie zaburzając już świadczonych usług.

Prześledźmy te trzy aspekty SO po kolei.

System operacyjny jako interfejs użytkownik – komputer

Sprzęt i oprogramowanie służące do udostępniania użytkownikowi aplikacji można rozpatrywać warstwowo, jak pokazano na rysunku 2.1. Użytkownika tych aplikacji (użytkownika docelowego) na ogół nie obchodzą szczegóły sprzętu komputerowego. Postrzega on system komputerowy w kategoriach zbioru aplikacji. Aplikację można wyrazić w języku programowania, opracowuje ją programista aplikacji². Gdyby komuś przyszło napisać program użytkowy w formie zbioru instrukcji maszynowych, w pełni odpowiadający za sterowanie sprzętem komputera, okazałoby się to przedsięwzięciem złożonym ponad wszelką miarę. Aby ułatwić tę robotę, tworzy się zbiór programów systemowych. Niektóre z nich są nazywane programami narzędziowymi lub bibliotecznymi. Realizują one często używane funkcje, pomagające w budowaniu programów, zarządzaniu plikami i sterowaniu urządzeniami wejścia-wyjścia. Programista posłuży się nimi, tworząc aplikację, natomiast aplikacja w trakcie swojego działania będzie się uciekać do tych narzędzi, aby wykonać pewne funkcje. W skład SO wchodzi zestaw najważniejszych programów systemowych. SO ukrywa przed osobą programującą szczegóły sprzętowe, dostarczając jej wygodnego interfejsu do użytkowania systemu komputerowego. Działa jak mediator, ułatwiając programiście i programom użytkowym dostęp do narzędzi i usług oraz ich wykorzystanie.



Rysunek 2.1. Sprzęt komputerowy i struktura oprogramowania

² On lub ona — podkreśliły to tutaj, zaznaczając, że uwaga ta odnosi się do wszelkich zawężonych do rodzaju męskiego form w tym przekładzie — *przyp. tłum.*

W skrócie — SO zazwyczaj dostarcza usług w następujących obszarach:

- **Opracowywanie programów.** SO udostępnia różnorodność narzędzi i usług, takich jak edytory i debugery (programy uruchomieniowe) pomagające programiście w tworzeniu programów. Na ogół usługi te przybierają postać programów narzędziowych, które — choć nie są nieodłączną częścią rdzenia SO — są dostarczane wraz z SO i określane jako narzędzia opracowywania programów użytkowych.
- **Wykonywanie programów.** Aby nadać bieg gotowemu programowi, trzeba wykonać kilka kroków. Rozkazy i dane należy załadować do pamięci, urządzenia zewnętrzne i pliki muszą być zainicjowane, a także trzeba przygotować inne zasoby. SO wywiązuje się wobec użytkownika z tych planowych obowiązków.
- **Dostęp do urządzeń wejścia-wyjścia.** Każde urządzenie zewnętrzne wymaga do działania specyficznego zbioru instrukcji lub sygnałów. SO dostarcza jednolitego interfejsu ukrywającego te szczegóły, dzięki czemu programiści mogą kontaktować się z takimi urządzeniami, używając po prostu operacji pisania i czytania.
- **Kontrolowany dostęp do plików.** Aby umożliwić dostęp do plików, SO musi szczegółowo rozróżniać nie tylko rodzaj urządzenia wejścia-wyjścia (napęd dysku czy taśmy), lecz także strukturę danych zawartych w plikach zmagazynowanych na nośniku danych. W przypadku systemu z wieloma użytkownikami SO może tworzyć mechanizmy ochronne w celu kontrolowania dostępu do plików.
- **Dostęp do systemu.** W przypadku systemów do wspólnego czy publicznego użytku SO kontroluje dostęp do systemu jako całości oraz do poszczególnych zasobów systemu. Funkcja dostępu musi zapewniać ochronę zasobów i danych przed nieupoważnionymi użytkownikami, a także rozwiązywać konflikty powstające na tle rywalizacji o zasoby.
- **Wykrywanie błędów i reagowanie.** Podczas działania systemu komputerowego mogą występować różnego rodzaju błędy. Są wśród nich wewnętrzne i zewnętrzne błędy sprzętowe (jak błąd pamięci lub wadliwe albo nie całkiem sprawne działanie urządzenia) oraz rozmaite błędy programowe (jak dzielenie przez zero, próba dostępu do zakazanego miejsca pamięci lub niezdolność SO do spełnienia żądania aplikacji). W każdym przypadku SO musi odpowiednio zareagować, usuwając skutki błędu w sposób możliwie najmniej inwazyjny dla wykonywanych aplikacji. Reakcja taka może być różnego rodzaju: od zakończenia programu, który spowodował błąd, do ponowienia operacji lub zwykłego powiadomienia aplikacji o błędzie.
- **Rozliczanie.** Dobry SO będzie gromadził statystykę wykorzystania różnych urządzeń i monitorował parametry działania, na przykład czas odpowiedzi. Tego rodzaju informacje są przydatne w każdym systemie, gdyż służą przewidywaniu konieczności przyszłych ulepszeń i strojeniu systemu w celu poprawienia jego wydajności. W systemie z wieloma użytkownikami informacje te mogą być użyte do wystawiania rachunków.

Na rysunku 2.1 zaznaczono również trzy podstawowe interfejsy typowego systemu komputerowego. Są to:

- **Architektura zbioru rozkazów** (ang. *instruction set architecture* — ISA). ISA definiuje repertuar instrukcji maszynowych, którym dany komputer jest zdolny dać posłuch. Ten interfejs stanowi rozgraniczenie między sprzętem a oprogramowaniem. Zauważmy, że zarówno programy użytkowe, jak i narzędziowe mogą mieć bezpośredni dostęp do ISA. Dla tych programów wydziela

się podzbiór repertuaru instrukcji (ISA użytkownika). System operacyjny rozporządza dodatkowymi instrukcjami języka maszynowego, służącymi do zarządzania zasobami systemowymi (ISA systemu).

- **Interfejs binarny aplikacji** (ang. *application binary interface* — ABI). ABI definiuje standard binarnej przenośności między programami. Określa też interfejs wywołań systemowych łączących z systemem operacyjnym oraz z zasobami sprzętowymi i usługami dostępnymi w systemie za pośrednictwem ISA użytkownika.
- **Interfejs programów użytkowych** (interfejs programowy aplikacji, ang. *application programming interface* — API³). API daje programowy dostęp do zasobów sprzętowych i usług osiągalnych w systemie za pośrednictwem ISA użytkownika, uzupełniony o wywołania biblioteczne w języku wysokiego poziomu (ang. *high-level language* — HLL). Dowolne wywołanie systemowe jest zwykle wykonywane przez biblioteki. Używanie API umożliwia łatwe przenoszenie oprogramowania użytkowego (przez rekompilację) do innych systemów udostępniających ten sam API.

System operacyjny jako zarządca zasobów

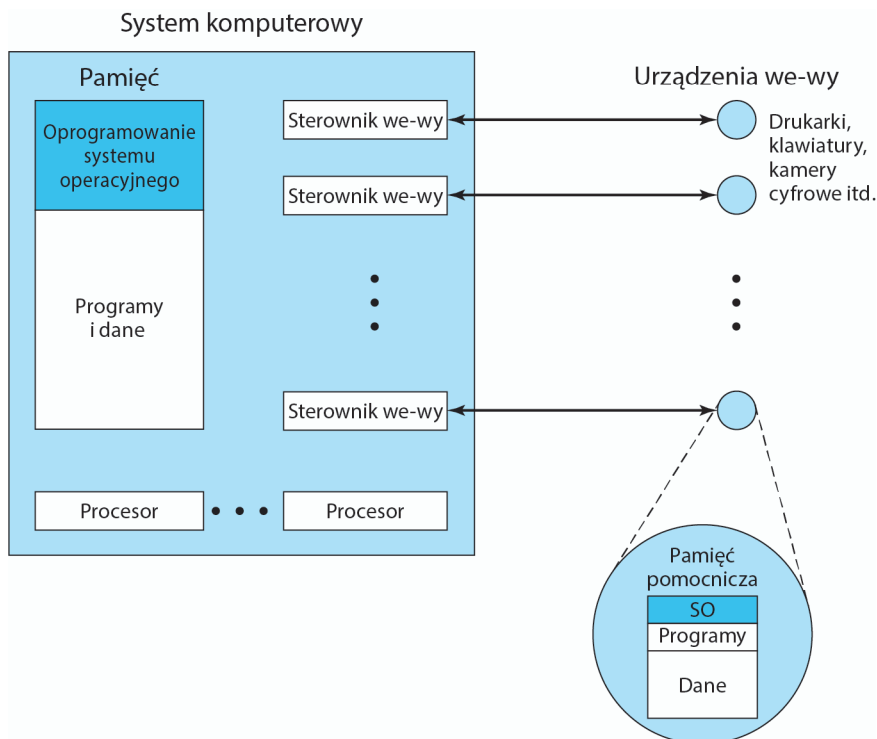
SO odpowiada za kontrolowanie użycia zasobów komputera, takich jak wejście-wyjście, pamięć główna i pomocnicza oraz czas działania procesora. Wszelako ten nadzór jest sprawowany w osobliwy sposób. Zwykle jesteśmy skłonni uważać mechanizm kontrolny za coś zewnętrznego w stosunku do tego, co jest kontrolowane, lub przynajmniej coś, co jest odrębną i wydzieloną częścią tego, co podlega kontroli. (Przykładem jest system centralnego ogrzewania, sterowany przez termostat — przyrząd odrębny względem źródła ciepła i kaloryferów). Rzecz ma się inaczej w przypadku SO, który jako mechanizm sterujący jest niezwykły z dwu powodów:

- SO działa w ten sam sposób jak zwykle oprogramowanie komputera, to znaczy jest programem lub zestawem programów wykonywanych przez procesor.
- SO często wyzbywa się kontroli i musi polegać na procesorze, który umożliwia mu jej odzyskanie.

Jak inne programy komputerowe SO składa się z rozkazów wykonywanych przez procesor. W czasie działania SO decyduje o sposobie przydziału czasu procesora i o tym, które zasoby są dostępne do użycia. Aby jednak procesor mógł te decyzje realizować, musi zaprzestawać wykonywania programu SO i wykonywać inne programy. Tak więc SO zrzeka się sterowania, aby procesor mógł wykonać jakąś „pożyteczną” pracę, po czym wznawia kontrolę na czas potrzebny do przygotowania procesora do wykonania następnej pracy. Służące temu mechanizmy zostaną wyjaśnione w dalszej części rozdziału.

Na rysunku 2.2 naszkicowano główne zasoby zarządzane przez SO. Część systemu rezyduje w pamięci głównej. Obejmuje ona **jądro** (ang. *kernel* albo *nucleus*), które zawiera najczęściej używane w systemie funkcje i — w danym czasie — inne części SO, będące akurat w użyciu. Pozostała część pamięci głównej zawiera programy użytkowe i narzędziowe oraz dane. SO i sprzętowa jednostka zarządzania pamięcią w procesorze sprawują razem kontrolę nad przydzielaniem pamięci głównej, co jeszcze zobaczymy. SO decyduje, kiedy urządzenie zewnętrzne może być użyte przez wykonywany program, i kontroluje dostęp do plików i ich użycie. Sam procesor jest jednym z zasobów i SO musi określać, ile czasu procesora będzie poświęcone na wykonanie programu danego użytkownika.

³ Akronim rozwijany również jako *application programmer interface*, czyli interfejs programisty aplikacji — przyp. tłum.



Rysunek 2.2. System operacyjny jako zarządca zasobów

Łatwość ewolucji systemu operacyjnego

Z biegiem czasu ważniejsze systemy operacyjne będą ewoluowały z kilku powodów:

- **Unowocześnienia sprzętowe i nowe rodzaje sprzętu.** Na przykład w pierwszych wersjach UNIX-a i macintoshowego Mac OS-a nie wykorzystywano mechanizmu stronicowania, ponieważ systemy te działały na procesorach bez sprzętowego zaplecza stronicowania⁴. Kolejne wersje tych systemów operacyjnych zostały zmodyfikowane tak, aby mogły wykorzystywać stronicowanie. Również zastosowanie terminali graficznych i terminali pracujących w trybie stronicowym zamiast działających w trybie przewijania wiersz po wierszu wpłynęło na projekt SO. Na przykład terminal graficzny z reguły umożliwia użytkownikowi oglądanie na ekranie kilku aplikacji w tym samym czasie za pośrednictwem „okien”. To wymaga bardziej wyrafinowanego wsparcia ze strony SO.
- **Nowe usługi.** W odpowiedzi na zapotrzebowania użytkowników lub wymagania administratorów SO jest rozszerzany o nowe usługi. Jeśli na przykład okaże się, że za pomocą istniejących narzędzi trudno jest zapewnić użytkownikom dobrą wydajność, do SO mogą być dołączone nowe narzędzia pomiarowe i sterujące.
- **Poprawki.** Każdy SO ma wady. Z upływem czasu są one wykrywane i poprawiane. Oczywiście poprawka może spowodować powstanie nowych wad.

⁴ O stronicowaniu powiemy pokrótce w dalszej części rozdziału, a dokładniej zajmiemy się nim w rozdziale 7.

Z koniecznością regularnego uaktualniania SO wiążą się pewne wymagania projektowe. Oczywiście postulatem jest zapewnienie modularności budowy systemu, z jasno zdefiniowanymi interfejsami między modułami, oraz to, że powinien on być dobrze udokumentowany. W wielkich programach — a takimi są typowe współczesne systemy operacyjne — to, co można by nazwać prostą modularyzacją, jest niewystarczające [DENN80a]. To znaczy trzeba zrobić znacznie więcej niż dokonać zwykłego podziału programu na moduły. Powróćmy do tego tematu w dalszej części rozdziału.

2.2. ROZWÓJ SYSTEMÓW OPERACYJNYCH

Próbując zrozumieć kluczowe wymagania dotyczące SO i znaczenie ważniejszych cech współczesnego SO, dobrze jest spojrzeć, jak wyglądał rozwój systemów operacyjnych na przestrzeni lat.

Przetwarzanie seryjne

W pierwszych komputerach z lat 40. i pierwszej połowy lat 50. XX wieku programista miał bezpośredni kontakt ze sprzętem komputerowym — nie było żadnego systemu operacyjnego. Komputery były uruchamiane z konsoli zawierającej lampki kontrolne i przełączniki stanowiące rodzaj urządzenia wejściowego, do tego dochodziła drukarka. Program w kodzie maszynowym był ładowany z urządzenia wejściowego (np. z czytnika kart). Jeśli program zatrzymał się z powodu błędu, powód błędu był wyświetlany za pomocą lampek. Gdy program zakończył działanie w normalny sposób, wyniki pojawiały się na drukarce. W tych pierwszych systemach występowały dwa główne problemy:

- **Planowanie.** W większości instalacji posługiwano się papierowymi arkuszami rezerwacji czasu komputera. Użytkownik zazwyczaj mógł zarezerwować porcję czasu wielkości pół godziny lub jej wielokrotność. Użytkownik mógł na przykład zapisać się na godzinę i skończyć pracę po trzech kwadransach, co skutkowało tym, że czas pracy komputera był marnowany. Z drugiej strony, zdarzało się, że użytkownik popadał w kłopoty, nie zdążył zakończyć obliczeń w przydzielonym czasie i musiał je przerwać przed uzyskaniem rozwiązania.
- **Czas przygotowania do pracy** (ang. *setup time*). Pojedynczy program, zwany **zadaniem** (ang. *job*), mógł wymagać załadowania do pamięci kompilatora oraz programu w języku wysokiego poziomu (programu źródłowego), przechowania skompilowanego programu (programu wynikowego), a następnie załadowania i skonsolidowania programu wynikowego z typowymi funkcjami. Każdy z tych kroków mógł obejmować zakładanie i zdejmowanie taśm lub ładowanie pakietów kart do czytnika. Jeśli przytrafił się błąd, pechowy użytkownik musiał na ogół ciąg tych przygotowań zaczynać od nowa. Powodowało to znaczne straty czasu na samo przygotowanie programu do działania.

Ten tryb pracy można by nazwać *przetwarzaniem seryjnym*, zważywszy na fakt, że użytkownicy podchodzili do komputera raz po razie podczas jednej sesji. Z biegiem czasu opracowano różne systemowe narzędzia programowe w dążeniu do usprawnienia przetwarzania seryjnego. Należały do nich biblioteki typowych funkcji, konsolidatory, ładowacze, debugery i moduły (programy) sterujące urządzeniami zewnętrznymi, udostępniane jako typowy software wszystkim użytkownikom.

Proste systemy wsadowe

Pierwsze komputery były bardzo drogie, toteż maksymalizowanie wykorzystania procesora było niezwykle istotną kwestią. Marnowanie czasu wskutek godzinowych rezerwacji i czasu zużywanego na przygotowanie do pracy było nie do przyjęcia.

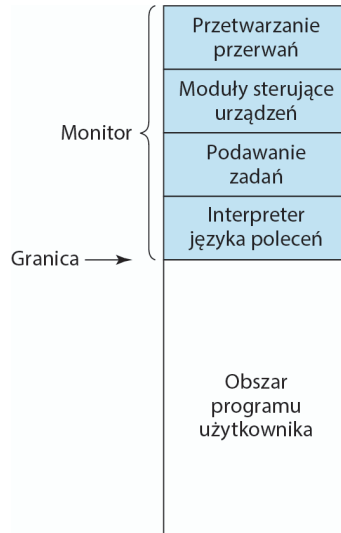
Aby poprawić wykorzystanie, opracowano koncepcję **wsadowego systemu operacyjnego** (ang. *batch OS*). Wydaje się, że pierwszy wsadowy SO (a zarazem pierwszy SO jakiegokolwiek typu) powstał w połowie lat 50. XX wieku w General Motors, opracowany na użytek komputera IBM 701 [WEIZ81]. Pomysł ten został następnie ulepszony i zrealizowany na komputerze IBM 704 przez wielu nabywców IBM-ów. We wczesnych latach 60. grupa dostawców opracowała wsadowe systemy operacyjne dla swoich systemów komputerowych. Szczególnie godny uwagi okazał się IBSYS, system operacyjny IBM-owskich komputerów 7090/7094, ze względu na jego szerokie oddziaływanie na inne systemy.

Głównym pomysłem w prostym schemacie przetwarzania wsadowego jest zastosowanie oprogramowania zwanego **monitorem**. Użytkownik w systemie operacyjnym tego typu nie ma już bezpośredniego dostępu do procesora. Zamiast tego przekazuje zadanie na kartach lub taśmie operatorowi komputera, który grupuje zadania jedno po drugim w pakiet (wsad) i umieszcza go na urządzeniu wejściowym do użycia przez monitor. Każdy program jest skonstruowany w taki sposób, aby po zakończeniu działania następowało przejście z powrotem do monitora, a wówczas monitor automatycznie rozpoczyna ładowanie następnego programu.

Żeby zrozumieć działanie tego schematu, spójrzmy na niego z dwóch perspektyw: monitora i procesora.

- **Perspektywa monitora.** Monitor nadzoruje ciąg zdarzeń. Aby to było możliwe, większa część monitora musi stale przebywać w pamięci głównej gotowa do działania (rysunek 2.3). Tę część nazywa się **monitorem rezydentnym** (ang. *resident monitor*). Reszta monitora składa się z narzędzi i typowych funkcji, ładowanych jako podprogramy do programu użytkownika na początku każdego zadania, które ich wymaga. Monitor czyta po jednym zadaniu z urządzenia wejściowego (zazwyczaj z czytnika kart lub taśmy magnetycznej). Przeczytane zadanie jest umieszczane w obszarze programu użytkownika i następuje przekazanie do niego sterowania. Gdy zadanie jest ukończone, zwraca sterowanie do monitora, który natychmiast czyta następne zadanie. Wyniki każdego zadania są przesyłane na urządzenie wyjściowe, takie jak drukarka, aby mogły być dostarczone użytkownikowi.
- **Perspektywa procesora.** W którymś momencie procesor wykonuje instrukcje z obszaru pamięci głównej zawierającego monitor. Powodują one przeczytanie następnego zadania do innego obszaru pamięci głównej. Po przeczytaniu zadania procesor napotyka w monitorze rozkaz skoku, polecający mu kontynuowanie działania, zaczynając od początku programu użytkownika. Procesor wykonuje więc instrukcje w programie użytkownika, aż napotka zakończenie lub błąd. Każde z tych zdarzeń powoduje, że następny rozkaz procesor pobiera z programu monitora. Tak więc wyrażenie „sterowanie jest przekazywane do zadania” oznacza po prostu, że procesor obecnie pobiera i wykonuje rozkazy w programie użytkownika, a „przekazanie sterowania z powrotem do monitora” oznacza, że procesor pobiera teraz i wykonuje rozkazy z programu monitora.

Monitor wykonuje funkcję planowania: pakiet zadań tworzy kolejkę, a zadania są wykonywane tak szybko jak to możliwe, bez przestojów między jednym a drugim. Monitor poprawia również **czas przygotowania zadania** (czas instalowania zadania, ang. *job setup time*). Każde zadanie



Rysunek 2.3. Rozplanowanie pamięci monitora rezydentnego

jest zaopatrzone w instrukcje wyrażone w dość prostym **języku sterowania zadaniami** (ang. *job control language* — JCL). Jest to specjalny rodzaj języka programowania używany do instruowania monitora. Prosty przykład to ten, w którym użytkownik przedkłada do wykonania program napisany w języku programowania Fortran i trochę danych do użycia w tym programie. Wszystkie instrukcje Fortranu i dane znajdują się na osobnych kartach perforowanych lub tworzą osobne rekordy na taśmie. Oprócz wierszy fortranowych i wierszy danych zadanie zawiera instrukcje sterujące zadaniem, oznaczane za pomocą początkowego znaku \$. Ogólna postać takiego zadania wygląda następująco:

```

$JOB
$FTN
.
.
.
$LOAD
$RUN
.
.
.
$END

```

} Instrukcje w Fortranie

} Dane

Aby wykonać to zadanie, monitor czyta wiersz \$FTN i ładuje kompilator odpowiedniego języka ze swojej pamięci masowej (zwykle jest to taśma). Kompilator tłumaczy program użytkownika na kod wynikowy, który zostaje przechowany w pamięci głównej lub w pamięci masowej. Jeśli jest przechowany w pamięci głównej, operacja jest określana jako „kompiluj, ładuj i wykonaj”. Jeśli jest przechowany na taśmie, wymagana jest instrukcja \$LOAD. Ta instrukcja jest czytana przez monitor, który odzyskuje sterowanie po operacji kompilowania. Monitor uruchamia ładowacz, a ten wprowadza program wynikowy do pamięci (na miejsce kompilatora) i przekazuje mu sterowanie. W ten sposób duży segment pamięci głównej może być dzielony między kilka różnych podsystemów, aczkolwiek tylko jeden podsystem może być wykonywany w danym czasie.

Podczas wykonywania programu użytkownika każda instrukcja wejścia powoduje przeczytanie jednego wiersza danych. Instrukcja wejścia w programie użytkownika powoduje wywołanie procedury wejścia będącej częścią SO. Procedura wejścia upewnia się, że program nie czyta przez przypadek wiersza JCL. Gdyby się tak zdarzyło, wystąpi błąd i sterowanie wraca do monitora. Po zakończeniu zadania użytkownika monitor przegląda wiersze wejściowe dopóty, dopóki nie napotka następnej instrukcji JCL. Tak więc system jest chroniony przed programem ze zbyt dużą lub zbyt małą liczbą wierszy danych⁵.

Monitor, czyli wsadowy SO, jest po prostu programem komputerowym. Funkcjonuje dzięki zdolności procesora do pobierania rozkazów z różnych fragmentów pamięci głównej, aby na przemian przejmować i oddawać sterowanie. Pożądane są również pewne inne cechy sprzętu:

- **Ochrona pamięci.** Gdy działa program użytkownika, nie wolno mu zmienić obszaru pamięci zawierającej monitor. Jeśli dojdzie do takiej próby, sprzęt procesora powinien wykryć błąd i przekazać sterowanie monitorowi. Monitor może wtedy zaniechać zadania, wydrukować komunikat błędu i załadować następne zadanie.
- **Czasomierz** (ang. *timer*). Czasomierz jest używany do ochrony przed zmonopolizowaniem systemu przez jedno zadanie. Jest ustawiany na początku każdego zadania. Jeśli czas się wyczerpie, program użytkownika jest zatrzymywany i sterowanie wraca do monitora.
- **Instrukcje uprzywilejowane.** Pewne instrukcje z poziomu maszynowego są oznaczone jako uprzywilejowane i mogą być wykonywane tylko przez monitor. Jeśli procesor napotka taki rozkaz podczas wykonywania programu użytkownika, wystąpi błąd powodujący przekazanie sterowania do monitora. Do instrukcji uprzywilejowanych należą instrukcje wejścia-wyjścia, tak więc monitor zachowuje kontrolę nad wszystkimi urządzeniami wejścia-wyjścia. Chroni to na przykład przed niezamierzonym przeczytaniem przez program użytkownika instrukcji sterujących z następnego zadania. Jeśli program użytkownika chce wykonać operację wejścia-wyjścia, musi zgłosić zapotrzebowanie na jej wykonanie w monitorze.
- **Przerwania.** Wczesne modele komputerów nie miały tej zdolności. Ta cecha sprawia, że SO staje się bardziej elastyczny w przekazywaniu sterowania do programów użytkowych i jego (od nich) odzyskiwaniu.

Rozważania nad ochroną pamięci i instrukcjami uprzywilejowanymi naprowadziły na pomysł trybów działania. Program użytkownika pracuje w **trybie użytkownika** (ang. *user mode*), w którym pewne obszary pamięci są chronione przed wykorzystaniem przez użytkownika, a pewne instrukcje są zabronione. Monitor działa w trybie systemu lub — jak to się zwykle nazywać — w **trybie jądra** (ang. *kernel mode*), w którym instrukcje uprzywilejowane mogą być wykonywane, a dostęp do chronionych obszarów pamięci jest dozwolony.

Oczywiście system operacyjny można zbudować bez tych właściwości. Jednak dostawcy komputerów szybko pojęli, że w rezultacie powstawał chaos, więc nawet stosunkowo proste systemy wsadowe były wyposażane w te cechy sprzętowe.

We wsadowym SO czas procesora jest naprzemiennie przydzielany na wykonywanie programów użytkowników i działanie monitora. Odbyna się to za cenę dwóch wyrzeczeń. Pewna partia

⁵ Chyba że użytkownik lub operator zapomniał wstawić kartę \$JCL. Aby ustrzec się przed „połykaniem” następnego zadania w charakterze „danych” przez poprzedni program, w używanym w Polsce na maszynach ODRA 1305 ICL-owskim systemie GEORGE 3 stosowano dodatkową kartę z tak zwanym ogranicznikiem uniwersalnym, której dodanie na końcu każdego zadania było rutynowym obowiązkiem operatora — *przyp. tłum.*

pamięci głównej jest obecnie przekazana we władanie monitorowi i część czasu procesora jest konsumowana przez monitor. I jedno, i drugie stanowi swoistą daninę. Mimo to proste systemy wsadowe poprawiają wykorzystanie komputera.

Wieloprogramowe systemy wsadowe

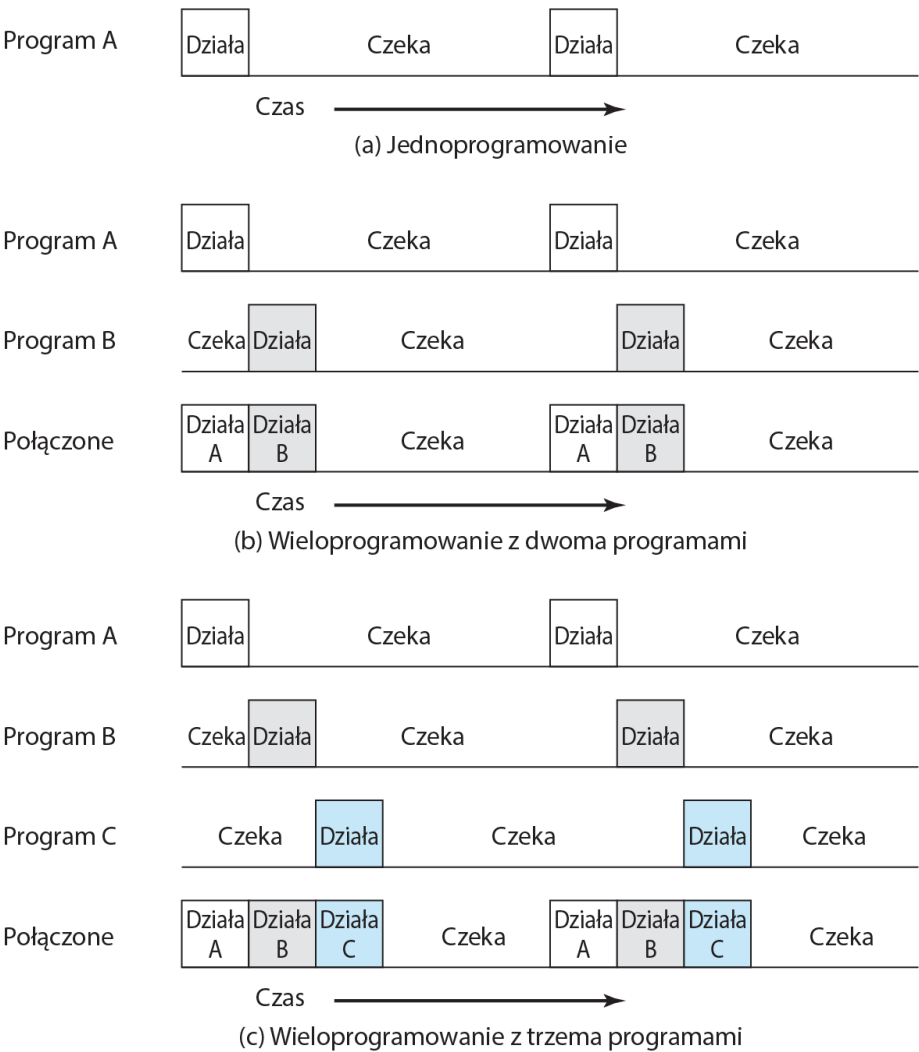
Nawet z automatycznym podawaniem zadań umożliwiającym przez prosty, wsadowy SO procesor jest często bezczynny. Problemem są urządzenia wejścia-wyjścia, które są wolniejsze w porównaniu z procesorem. Na rysunku 2.4 przedstawiono szczegóły reprezentatywnego obliczenia. Dotyczy ono programu, który przetwarza plik rekordów i wykonuje średnio 100 instrukcji maszynowych na rekord. W tym przykładzie komputer spędza ponad 96% czasu w oczekiwaniu na zakończenie przez urządzenia wejścia-wyjścia przesyłania danych do i z pliku. Na rysunku 2.5a przedstawiono tę sytuację: mamy tu do czynienia z jednym programem, co jest określane jako **jednoprogramowość** (ang. *uniprogramming*). Procesor działa przez pewien czas, aż do napotkania operacji wejścia-wyjścia. Wówczas musi zaczekać, aż instrukcja wejścia-wyjścia zakończy działanie, i dopiero wtedy może kontynuować pracę.

Czytanie jednego rekordu z pliku	15 μs
Wykonanie 100 instrukcji	1 μs
Zapisanie jednego rekordu do pliku	15 μs
Ogółem	31 μs
Procent wykorzystania CPU = $\frac{1}{31} = 0,032 = 3,2\%$	

Rysunek 2.4. Przykład wykorzystania systemu

Tej nieefektywności można uniknąć. Wiemy, że pamięci musi wystarczyć na utrzymywanie SO (monitora rezydentnego) i jednego programu użytkownika. Załóżmy, że mamy miejsce na SO i dwa programy użytkownika. Kiedy jedno zadanie musi czekać na wejście-wyjście, procesor może się przełączyć do drugiego zadania, które być może nie czeka na wejście-wyjście (zob. rysunek 2.5b). Mało tego — możemy powiększyć pamięć, aby pomieściła trzy, cztery lub więcej programów, i dokonywać przełączeń między nimi wszystkimi (zob. rysunek 2.5c). To podejście jest nazywane **wieloprogramowością** (ang. *multiprogramming*) lub **wielozadaniowością** (ang. *multitasking*). Stanowi ono centralny motyw nowoczesnych systemów operacyjnych.

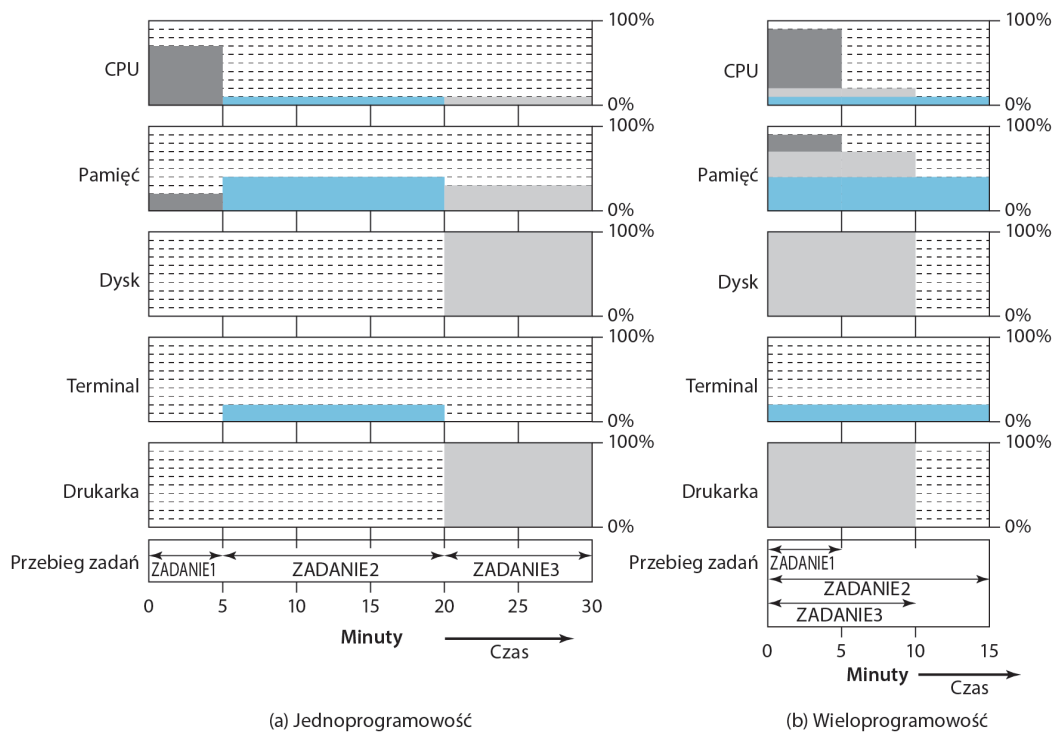
Aby zilustrować korzyści płynące z wieloprogramowości, podamy prosty przykład. Rozważmy komputer mający 250 MB dostępnej pamięci (nieużywanej przez system operacyjny), dysk, terminal i drukarkę. W tej samej chwili do wykonania trafiają trzy programy: ZADANIE1, ZADANIE2 i ZADANIE3 o charakterystykach jak w tabeli 2.1. Zakładamy minimalne zapotrzebowanie na procesor w zadaniach ZADANIE1 i ZADANIE3 i stałe używanie dysku i drukarki przez ZADANIE3. W prostym środowisku wsadowym zadania te będą wykonywane po kolei. Wobec tego ZADANIE1 zakończy się po 5 minutach, ZADANIE2 musi czekać przez 5 minut i potem będzie działał przez 15 minut. ZADANIE3 rozpocznie się po 20 minutach i zakończy w 30 minucie od chwili przedłożenia do wykonania. Średnie wykorzystanie zasobów, przepustowość i czas odpowiedzi pokazano w kolumnie jednoprogramowości tabeli 2.2. Wykorzystanie urządzeń przedstawiono na rysunku 2.6a. Widać wyraźnie duże niewykorzystanie wszystkich zasobów uśrednione w wymaganym 30-minutowym odcinku czasu.



Rysunek 2.5. Przykład wieloprogramowości

Tabela 2.1. Przykładowe charakterystyki wykonania programów

	ZADANIE1	ZADANIE2	ZADANIE3
Rodzaj zadania	Dużo obliczeń	Dużo we-wy	Dużo we-wy
Czas trwania	5 minut	15 minut	10 minut
Zapotrzebowanie na pamięć	50 MB	100 MB	75 MB
Dysk potrzebny?	Nie	Nie	Tak
Terminal potrzebny?	Nie	Tak	Nie
Drukarka potrzebna?	Nie	Nie	Tak



Rysunek 2.6. Histogramy wykorzystania

Tabela 2.2. Wpływ wieloprogramowości na wykorzystanie zasobów

	Jednoprogramowość	Wieloprogramowość
Użycie procesora	20%	40%
Użycie pamięci	33%	67%
Użycie dysku	33%	67%
Użycie drukarki	33%	67%
Zużyty czas	30 minut	15 minut
Przepustowość	6 zadań/godzinę	12 zadań/godzinę
Średni czas odpowiedzi	18 minut	10 minut

Załóżmy teraz, że zadania są wykonywane współbieżnie w systemie wieloprogramowym. Ponieważ między zadaniami jest mało rywalizacji o zasoby, wszystkie trzy mogą działać prawie w minimalnym czasie, koegzystując w komputerze z innymi (przyjmując, że ZADANIE2 i ZADANIE3 mają przydzielone dość czasu procesora, aby utrzymać w ruchu swoje operacje wejścia i wyjścia). ZADANIE1 nadal będzie wymagało 5 minut do zakończenia, lecz pod koniec tego czasu ZADANIE2 będzie już w jednej trzeciej wykonane, a ZADANIE3 będzie wykonane do połowy. Wszystkie trzy zadania zakończą się po 15 minutach. Poprawa jest ewidentna, gdy przyjrzymy się kolumnie wieloprogramowości w tabeli 2.2, otrzymanej z histogramu przedstawionego na rysunku 2.6b.

Podobnie jak w prostym systemie wsadowym, wieloprogramowy system wsadowy musi mieć wsparcie ze strony pewnych właściwości sprzętu. Najważniejszą dodatkową cechą przydatną w wieloprogramowości jest sprzęt umożliwiający przerwania pochodzące z wejścia-wyjścia i od DMA (bezpośredniego dostępu do pamięci). Mając sterowane przerwaniami operacje wejścia-wyjścia lub DMA, procesor może wydać polecenie wejścia-wyjścia w jednym zadaniu i kontynuować wykonywanie innego zadania w czasie wykonywania wejścia-wyjścia przez sterownik urządzenia. Gdy operacja wejścia-wyjścia dobiegnie końca, działanie procesora jest przerywane i sterowanie przekazuje się do programu obsługi przerwania w SO. Po obsłużeniu przerwania SO przekaże sterowanie innemu zadaniu.

Wieloprogramowe systemy operacyjne są dosyć skomplikowane w porównaniu z systemami wykonującymi jeden program, czyli **jednoprogramowymi**. Aby mieć pewną liczbę zadań gotowych do działania, trzeba je utrzymywać w pamięci, co wymaga jakiejś formy **zarządzania pamięcią**. Ponadto jeśli kilka zadań jest gotowych do działania, procesor musi decydować, które ma być wykonywane, a ta decyzja wymaga jakiegoś algorytmu planowania. Te zagadnienia będą omówione dalej w tym rozdziale.

Systemy z podziałem czasu

Z zastosowaniem wieloprogramowości **przetwarzanie wsadowe** (ang. *batch processing*) może być całkiem sprawne. Jednak w wielu zadaniach jest pożądaný tryb, w którym użytkownik pozostaje w bezpośredniej interakcji z komputerem. Rzeczywiście, w niektórych zadaniach, jak na przykład przetwarzanie transakcji, tryb interakcyjny ma duże znaczenie.

Obecnie wymaganie na obliczenia interaktywne można zaspokoić — i często tak się dzieje — stosując wydzielony komputer osobisty lub stację roboczą. W latach 60. nie było tej możliwości. Komputery były wówczas w większości wielkie (gabarytowo — *przyp. tłum.*) i kosztowne. Wymyślono więc podział czasu.

Skoro wieloprogramowość umożliwia obsługiwanie przez procesor wielu zadań wsadowych, postępując analogicznie, można ją zaprząć do obsługiwania wielu zadań interakcyjnych. W tym drugim przypadku nazywa się tę technikę **podziałem czasu** (ang. *time sharing*), ponieważ czas procesora jest dzielony między wielu użytkowników. Do systemu z podziałem czasu ma jednoczesny dostęp wielu użytkowników za pośrednictwem terminali, a SO przeplata wykonywanie programów każdego z nich w krótkotrwałych dostępach do procesora, czyli kwantach obliczeń. Jeśli więc w danym czasie n użytkowników domaga się obsługi, to każdy z nich będzie miał do dyspozycji średnio tylko $1/n$ efektywnej mocy obliczeniowej komputera, nie wliczając nakładów wymaganych przez SO. Zważywszy jednak na stosunkowo wolny czas reakcji człowieka, czas odpowiedzi w dobrze zaprojektowanym systemie powinien być podobny jak w przypadku użycia wydzielonego komputera.

Zarówno w przetwarzaniu wsadowym, jak i w podziale czasu jest wykorzystywana wieloprogramowość. Zasadnicze różnice zestawiono w tabeli 2.3.

Tabela 2.3. Porównanie wieloprogramowania wsadowego z podziałem czasu

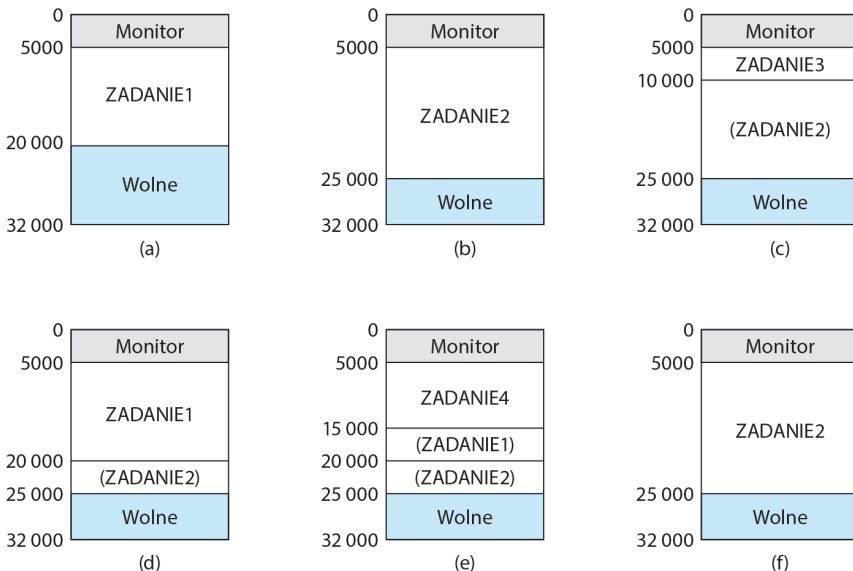
	Wieloprogramowość wsadowa	Podział czasu
Główny cel	Maksymalizacja wykorzystania procesora	Minimalizacja czasu odpowiedzi
Źródło dyrektyw dla systemu operacyjnego	Polecenia w języku sterowania zadaniami podane w opisie zadania	Polecenia wprowadzane z terminala

Jednym z pierwszych skonstruowanych systemów operacyjnych z podziałem czasu był Compatible Time-Sharing System (CTSS) [CORB62], opracowany w MIT przez grupę o nazwie Project MAC (Machine-Aided Cognition, z ang. percepcja wspomagana maszynowo, albo Multiple-Access Computers, z ang. komputery z wielodostępem). Początkowo, w 1961 roku, system ten został utworzony dla komputerów IBM 709, a potem przeniesiono go na maszyny IBM 7094.

W porównaniu z późniejszymi systemami CTSS był prymitywny. Działał na komputerze z 32 000 słów 36-bitowych w pamięci głównej, z rezydentnym monitorem zajmującym 5000 z nich. Kiedy miało nastąpić przekazanie sterowania do interakcyjnego użytkownika, program i dane użytkownika były ładowane do pozostałych 27 000 słów pamięci głównej. Program był zawsze ładowany tak, aby zaczynał się od słowa 5000. Upraszczało to zarówno monitor, jak i zarządzanie pamięcią. Zegar systemowy generował przerwania w przybliżeniu co 0,2 sekundy. Przy każdym przerwaniu zegarowym SO przejmował sterowanie i mógł przydzielić procesor innemu użytkownikowi. To postępowanie nazywa się **kwantowaniem czasu** (dosł. plastrowaniem czasu, ang. *time slicing*). Dzięki temu w regularnych odstępach czasu bieżący użytkownik mógł być wywłaszczony, a inny użytkownik załadowany. Aby zachować stan programu starego użytkownika do późniejszego wznowienia, przed wczytaniem programów i danych nowego użytkownika programy i dane starego użytkownika były zapisywane na dysku. Gdy później ten program otrzymywał następną rundę, kod i dane starego użytkownika sprowadzano z powrotem do pamięci głównej.

Aby zminimalizować ruch na dysku, pamięć użytkownika była wypisywana tylko wówczas, gdy wchodzący program mógł ją zapisać na nowo. Zasadę tę ilustruje rysunek 2.7. Załóżmy, że mamy czterech użytkowników interakcyjnych z następującymi zapotrzebowaniami na pamięć (w słowach):

- ZADANIE1 — 15 000;
- ZADANIE2 — 20 000;
- ZADANIE3 — 5000;
- ZADANIE4 — 10 000.



Rysunek 2.7. Działanie CTSS

Na początku monitor ładuje ZD1⁶ i przekazuje mu sterowanie (rysunek 2.7a). Później monitor postanawia przekazać sterowanie do ZD2. Ponieważ ZD2 wymaga więcej pamięci niż ZD1, ZD1 musi być najpierw wypisane, dopiero potem można załadować ZD2 (rysunek 2.7b). Dalej jest ładowane do wykonania ZD3. Ponieważ jednak ZD3 jest mniejsze niż ZD2, część ZD2 może pozostać w pamięci, zmniejszając czas zapisu na dysku (rysunek 2.7c). Następnie monitor decyduje się oddać sterowanie z powrotem ZD1. Na okoliczność ponownego ładowania ZD1 trzeba uprzednio wypisać dodatkową porcję ZD2 (rysunek 2.7d). Przy ładowaniu ZD4 część pozostającego w pamięci ZD1 i fragment ZD2 pozostają zachowane (rysunek 2.7e). Jeśli teraz zostanie uaktywnione ZD1 lub ZD2, będzie potrzebne tylko częściowe ładowanie. W tym przykładzie jako następne działa ZD2. Wymaga to przepisania na dysk ZD4 i pozostałej porcji ZD1, po czym będzie można wczytać brakującą porcję ZD2 (rysunek 2.7f).

Podjęcie CTSS jest prymitywne w porównaniu ze współczesnym podziałem czasu, lecz było skuteczne. Dzięki jego skrajnej prostocie można było zminimalizować rozmiar monitora. Ponieważ zadanie ładowano zawsze na to samo miejsce w pamięci, nie trzeba było stosować technik relokacji podczas ładowania (co jest omówione dalej). Technika wypisywania tylko tego, co konieczne, minimalizowała aktywność dysku. Wykonywany na komputerze 7094 system CTSS obsługiwał maksymalnie 32 użytkowników.

Podział czasu i wieloprogramowość rodzą w SO mnóstwo nowych problemów. Jeśli wiele zadań znajduje się w pamięci, to trzeba je chronić przed wzajemnym oddziaływaniem, na przykład przed zmianą danych jednego przez drugie. W obecności wielu użytkowników należało chronić system plików, aby dostęp do konkretnego pliku mieli tylko upoważnieni użytkownicy. Należało uwzględnić rywalizację o zasoby, takie jak drukarki i urządzenia pamięci masowej. Te i inne problemy z możliwymi rozwiązaniami będą się pojawiać w tym podręczniku.

2.3. NAJWAŻNIEJSZE OSIĄGNIĘCIA

Systemy operacyjne są jednymi z najbardziej skomplikowanych rodzajów oprogramowania, jakie kiedykolwiek zbudowano. Odzwierciedla to dążenie do zaspokojenia trudnych i niekiedy przeciwstawnych celów wydoby, wydajności i zdolności ewolucyjnej. W pracy [DENN80a] wysunięto cztery zasadnicze, teoretyczne kierunki rozwoju systemów operacyjnych:

- procesy,
- zarządzanie pamięcią,
- ochrona informacji i bezpieczeństwo,
- planowanie i zarządzanie zasobami.

Każdy kierunek jest scharakteryzowany przez wytyczne, lub abstrakcje, wykształcone w celu rozwiązania trudnych problemów praktycznych. Te cztery obszary rozpatrywane łącznie obejmują wiele kluczowych zagadnień projektowych i implementacyjnych nowoczesnych systemów operacyjnych. Przedstawiony w tym podrozdziale krótki przegląd tych czterech obszarów zarysowuje znaczną część zagadnień poruszonych w podręczniku.

⁶ Dla wygody skracamy dalej ZADANIE do ZD — *przyp. tłum.*

Proces

Centralną koncepcją projektowania systemów operacyjnych jest **proces** (ang. *process*). Po raz pierwszy terminu tego użyli projektanci systemu Multics w latach 60. XX wieku [DALE68]. Jest on nieco ogólniejszy niż *zadanie* (ang. *job*). Sformułowano wiele definicji terminu *proces*, między innymi takie:

- program w trakcie wykonywania;
- egzemplarz (ang. *instance*) programu działającego na komputerze;
- obiekt (pewna całość, ang. *entity*), któremu można przydzielić procesor w celu wykonywania;
- jednostka (ang. *unit*) aktywności charakteryzująca się pojedynczym, sekwencyjnym wątkiem wykonania, bieżącym stanem i przydzielonym zbiorem zasobów systemu.

Pojęcie to powinno stać się jaśniejsze w miarę kontynuowania przez nas tematu.

Trzy główne fronty rozwoju systemów komputerowych doprowadziły do powstania problemów koordynacji i synchronizacji, które przyczyniły się do ukształtowania koncepcji procesu: wieloprogramowość działań wsadowych, podział czasu i systemy transakcji przebiegających w czasie rzeczywistym. Jak widzieliśmy, wieloprogramowość zaprojektowano po to, aby utrzymywać procesor i urządzenia wejścia-wyjścia (w tym urządzenia pamięci masowej) w jednoczesnym działaniu w celu uzyskania maksymalnej wydajności. Zasadniczy mechanizm przedstawia się następująco: w reakcji na sygnały zakończenia operacji wejścia-wyjścia procesor jest przełączany między różnymi programami przebiegającymi w pamięci głównej.

Drugim frontem rozwoju był uniwersalny podział czasu. Tutaj podstawowym celem projektowym jest żywe reagowanie na potrzeby indywidualnego użytkownika i dodatkowo — ze względów ekonomicznych — zdolność jednoczesnego obsługiwanie wielu użytkowników. Te cele są ze sobą zgodne dzięki stosunkowo długim czasom reakcji użytkownika. Typowy użytkownik potrzebuje na przykład średnio 2 sekundy czasu przetwarzania na minutę, więc blisko 30 takich użytkowników powinno móc wspólnie korzystać z tego samego systemu bez zauważalnych wzajemnych utrudnień. Oczywiście w takich kalkulacjach trzeba uwzględnić nakład pracy ze strony SO⁷.

Trzecią ważną linią rozwoju były systemy przetwarzania transakcji w czasie rzeczywistym. W tym przypadku pewna liczba użytkowników wprowadza zapytania lub uaktualnienia do bazy danych. Przykładem jest system rezerwacji linii lotniczych. Zasadniczą różnicą między systemem przetwarzania transakcji a systemem z podziałem czasu jest ograniczenie tego pierwszego do jednego lub kilku zastosowań, podczas gdy użytkownicy systemu z podziałem czasu mogą się angażować w opracowywanie programów, wykonywanie zadań i korzystanie z różnych aplikacji. W obu przypadkach czas reakcji systemu ma kapitalne znaczenie.

Podstawowym narzędziem programistów systemowych opracowujących wczesne interakcyjne, wieloprogramowe i dostępne dla wielu użytkowników systemy było przerwanie. Działanie każdego zadania mogło być zawieszone wskutek wystąpienia określonego zdarzenia w rodzaju zakończenia operacji wejścia-wyjścia. Procesor mógł przechować pewien rodzaj kontekstu (np. licznik programu i inne rejestry) i przejść do procedury obsługi przerwania, która potrafiła określić naturę przerwania i przetworzyć je, po czym następowało wznowienie przerwane zadania użytkownika lub podjęcie innego zadania.

⁷ I są to kalkulacje w aspekcie liczbowym historyczne, na miarę mocy obliczeniowej komputerów z lat 60. — *przyj. tłum.*

Zaprojektowanie oprogramowania systemowego koordynującego te różne czynności okazało się nadzwyczaj trudne. Wobec wielu zadań w toku realizacji, z których każde wymagało wykonania wielu kroków po kolei, przeanalizowanie wszystkich potencjalnych kombinacji ciągów zdarzeń stało się niemożliwe. Wskutek braku systematycznych środków koordynacji i kooperacji aktywności programiści uciekali się do doraźnych metod opartych na ich rozumieniu środowiska, które SO miał kontrolować. Wysiłki te były narażone na subtelne błędy programowania, których skutki dawały się zauważyć tylko w stosunkowo rzadkich układach zdarzeń. Były to błędy trudne do diagnozowania, gdyż należało je odróżnić od błędów w oprogramowaniu aplikacji i błędów sprzętowych. Nawet w wypadku wykrycia błędu niełatwo było ustalić jego przyczynę z powodu wielkich trudności w precyzyjnym odtworzeniu warunków jego powstania. Najogólniej biorąc, są cztery główne przyczyny takich błędów [DENN80a]:

- **Niewłaściwa synchronizacja.** Niejednokrotnie zdarza się, że procedura musi być zawieszona, aby zaczekać na zdarzenie w innej części systemu. Na przykład program, który zapoczątkowuje na wejściu-wyjściu operację czytania, musi czekać przed podjęciem dalszego działania, aż dane staną się dostępne w buforze. W takich razach jest potrzebny sygnał z innej procedury. Niewłaściwe zaprojektowanie mechanizmu sygnalizacji może prowadzić do utraty sygnałów lub ich podwajania.
- **Wadliwe wzajemne wykluczanie.** Często się zdarza, że więcej niż jeden użytkownik lub program będzie usiłował skorzystać ze wspólnego zasobu w tym samym czasie. Na przykład dwóch użytkowników mogłoby próbować jednocześnie redagować ten sam plik. Musi istnieć jakiś rodzaj mechanizmu wzajemnego wykluczania, który pozwala w danej chwili uaktualniać plik tylko jednemu programowi. Poprawność realizacji takiego wzajemnego wykluczania jest trudna do zweryfikowania z uwzględnieniem wszystkich możliwych ciągów zdarzeń.
- **Nieokreślone działanie programu.** Wyniki konkretnego programu powinny w normalnych okolicznościach zależeć tylko od danych wejściowych tego programu, nie zaś od działań innych programów we wspólnie użytkowanym systemie. Jednak gdy programy wspólnie korzystają z pamięci i ich wykonanie jest przeplatane przez procesor, mogłyby oddziaływać na siebie, zapisując wspólne obszary pamięci w nieokreślony sposób. Tak więc kolejność planowania poszczególnych programów mogłaby wpłynąć na wyniki konkretnego programu.
- **Zakleszczenia.** Może się zdarzyć, że dwa lub więcej programów zawiesi się, oczekując na siebie wzajemnie. Na przykład dwa programy mogą potrzebować dwóch urządzeń zewnętrznych do wykonania pewnej operacji (np. do kopiowania z dysku na taśmę). Jeden z programów przejmuje kontrolę nad jednym z urządzeń, a drugi program rozporządza drugim urządzeniem. Każdy z programów czeka, aby drugi zwolnił potrzebne urządzenie. Takie zakleszczenie może zależeć od przypadkowej kolejności przydziału i zwalniania zasobów.

Aby poradzić sobie z tymi problemami, należy systematycznie nadzorować poszczególne programy wykonywane przez procesor i odpowiednio nimi sterować. Podstawą takich metod jest pojęcie procesu. Możemy przyjąć, że na proces składają się trzy elementy:

1. Wykonywalny program.
2. Związane z programem potrzebne mu dane (zmienne, obszar roboczy, bufory itp.).
3. Kontekst wykonywania programu.

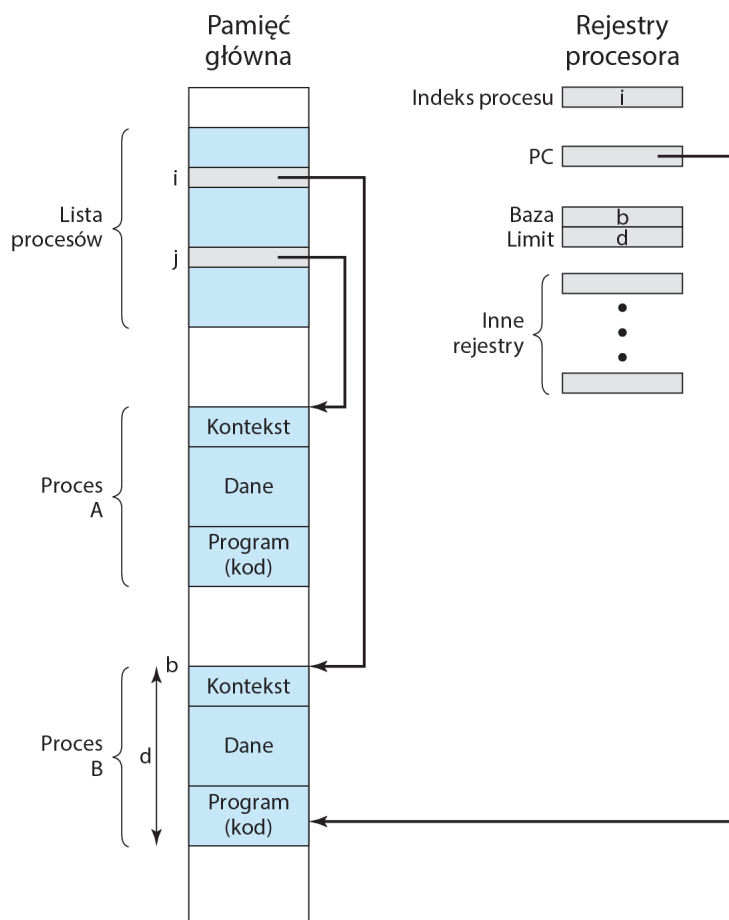
Ostatni element ma istotne znaczenie. **Kontekst wykonywania** (ang. *execution context*), inaczej: **stan procesu** (ang. *process state*), stanowią wewnętrzne dane, dzięki którym SO potrafi nadzorować proces i nim sterować. Te wewnętrzne informacje są oddzielone od procesu, ponieważ jako informacje systemowe są mu zabronione. Kontekst obejmuje wszystkie informacje potrzebne systemowi do zarządzania procesem, a procesorowi do właściwego wykonywania procesu. Do kontekstu należy zawartość różnych rejestrów procesora, takich jak licznik programu i rejestry danych. Zawiera on również informacje wykorzystywane przez SO, jak priorytet procesu oraz to, czy proces czeka na zakończenie jakiejś operacji wejścia-wyjścia.

Na rysunku 2.8 pokazano sposób, w jaki można zarządzać procesami. Dwa procesy, A i B, rezydują w dwu obszarach pamięci głównej, to znaczy każdemu z nich jest przydzielony blok pamięci, który zawiera program, dane i informacje dotyczące kontekstu. Każdy proces jest zapisany na liście procesów budowanej i utrzymywanej przez SO. Lista procesów ma po jednym wpisie dla każdego procesu, zawierającym wskaźnik do położenia bloku pamięci z procesem. Wpis ten może również zawierać część lub całość kontekstu wykonywania procesu. Pozostała część kontekstu wykonywania jest pamiętana gdzie indziej, być może w samym procesie (jak pokazano na rysunku 2.8) lub często w osobnym regionie pamięci. Rejestr indeksu procesu zawiera indeks położenia na liście procesów wpisu procesu aktualnie sterującego procesorem. Licznik programu (PC) wskazuje następny rozkaz do wykonania w danym procesie. Rejestry bazowy i graniczny określają obszar pamięci zajmowany przez proces. Rejestr bazowy zawiera adres początkowy tego obszaru, a rejestr graniczny jego rozmiar (w bajtach lub słowach). Licznik programu i wszystkie odniesienia do danych są interpretowane względem rejestru bazowego i nie mogą przekroczyć wartości w rejestrze granicznym. Zapobiega to wzajemnemu zaburzaniu się procesów.

Na rysunku 2.8 indeks procesu (a bezpośrednio licznik programu, PC) wskazuje, że wykonywany jest proces B. Proces A był wykonywany wcześniej, lecz chwilowo został przerwany. Zawartości wszystkich rejestrów w momencie przerwania A zostały zapisane w jego kontekście wykonywania. System operacyjny może później przełączyć procesy i wznowić wykonywanie procesu A. Przełączenie procesów składa się z przechowania kontekstu B i odtworzenia kontekstu A. Gdy w liczniku programu zostanie umieszczona wartość wskazująca miejsce w obszarze programu A, proces A automatycznie podejmie działanie.

Tak więc proces jest urzeczywistniany w postaci struktury danych. Proces może być wykonywany albo może czekać na wykonanie. W każdym momencie cały **stan** procesu jest zawarty w jego kontekście. Ta struktura umożliwia opracowywanie skutecznych sposobów zapewniania koordynacji i kooperacji między procesami. Można projektować i wprowadzać do systemu nowe właściwości (np. priorytety), rozszerzając kontekst o dowolne nowe informacje niezbędne do ich uzyskania. Wielokrotnie w tej książce napotkamy przykłady wykorzystania struktury procesu do rozwiązywania problemów wynikających z wieloprogramowości i dzielenia zasobów.

Na koniec zajmiemy się tutaj krótko koncepcją **wątku** (ang. *thread*). Pojedynczy proces, który ma przydzielone pewne zasoby, można w istocie podzielić na wiele współbieżnych wątków, które kooperatywnie wykonują pracę procesu. To prowadzi do nowego poziomu równoległych działań, którymi trzeba zarządzać za pomocą metod sprzętowych i programowych.



Rysunek 2.8. Typowa implementacja procesu

Zarządzanie pamięcią

Potrzeby użytkowników mogą być zaspokajane najlepiej przez środowisko obliczeniowe umożliwiające programowanie modularne i elastyczne użytkowanie danych. Administratorzy systemów potrzebują sprawnego i uporządkowanego zarządzania przydzielaniem pamięci. Aby spełnić te wymagania, na systemie operacyjnym spoczywa pięć podstawowych powinności dotyczących zarządzania pamięcią:

1. **Izolowanie procesów.** System operacyjny musi powstrzymywać niezależne procesy przed ingerencją w obszary pamięci należące do innych procesów, zarówno w odniesieniu do instrukcji, jak i danych.
2. **Automatyczny przydział i zarządzanie.** Programy powinny być lokowane w hierarchii pamięci dynamicznie, stosownie do potrzeb. Przydział powinien być przezroczysty dla osoby programującej. Programista jest więc uwalniany od problemów związanych z ograniczeniami pamięci, a SO może zyskiwać na efektywności dzięki przydzielaniu pamięci zadaniom tylko stosownie do zapotrzebowań.

3. **Umożliwianie programowania modularnego.** Programiści powinni móc definiować moduły programu, dynamicznie je tworzyć i likwidować oraz zmieniać ich rozmiary.
4. **Ochrona i kontrolowanie dostępu.** Współużytkowanie pamięci na każdym poziomie hierarchii pamięci daje możliwość zaadresowania przez program pamięci w obrębie innego programu. Jest to pożądane, gdy pewne aplikacje wymagają dzielenia pamięci w sensie wspólnego korzystania z tego samego jej obszaru. W innych przypadkach rodzi to zagrożenie integralności programów, a nawet samego SO. System musi zatem udostępniać porcje pamięci w różny sposób różnym użytkownikom.
5. **Magazynowanie długoterminowe.** Wiele programów użytkowych wymaga środków długotrwałego przechowywania informacji po wyłączeniu zasilania komputera.

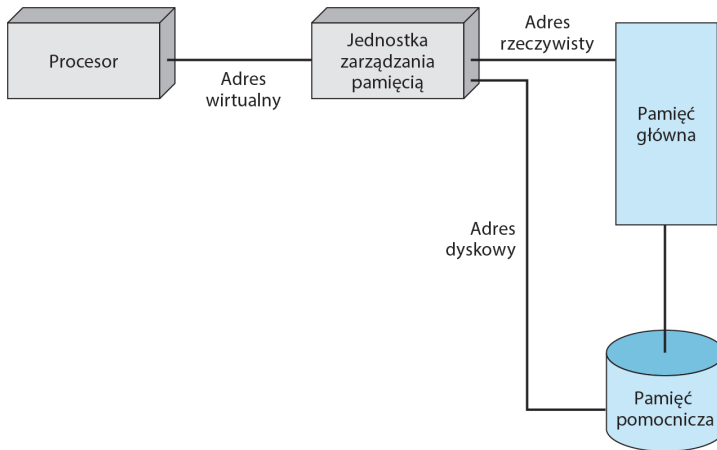
Systemy operacyjne na ogół spełniają te wymagania za pomocą pamięci wirtualnej i udogodnień zawartych w systemie plików. System plików urzeczywistnia pamięć długoterminową, w której informacje są przechowywane w ponazywanych obiektach, zwanych plikami. Plik stanowi wygodną koncepcję dla programisty i jest użyteczną jednostką kontrolowania dostępu i ochrony w SO.

Pamięć wirtualna (ang. *virtual memory*) jest udogodnieniem, które umożliwia programom adresowanie pamięci w aspekcie logicznym, bez zważania na ilość pamięci głównej dostępnej fizycznie. Pamięć wirtualną wymyślono po to, aby spełnić wymaganie jednoczesnej obecności w pamięci głównej zadań wielu użytkowników, co umożliwia unikanie przerw między wykonaniami kolejnych procesów, gdy jeden proces jest zapisywany w pamięci drugorzędnej, a następny jest wczytywany. Ponieważ procesy różnią się rozmiarami, ich ściśle upakowywanie w pamięci głównej sprawia trudności, gdy procesor jest przełączany od jednego procesu do drugiego. Wprowadzono więc systemy stronicowania, które umożliwiają rozmieszczenie procesów w pewnej liczbie stałej długości bloków, zwanych **stronami**. Program odwołuje się do słowa za pomocą **adresu wirtualnego** składającego się z numeru strony i odległości na stronie. Każda strona procesu może być ulokowana gdziekolwiek w pamięci głównej⁸. System stronicowania realizuje dynamiczne odwzorowanie między adresem wirtualnym używanym w programie a **adresem rzeczywistym**, czyli fizycznym adresem w pamięci głównej.

W obecności sprzętu realizującego dynamiczne odwzorowywanie następnym logicznym krokiem było wyeliminowanie wymagania, aby wszystkie strony procesu jednocześnie rezydowały w pamięci. Wszystkie strony procesu są utrzymywane na dysku. Podczas działania procesu w pamięci głównej znajdują się niektóre z jego stron. Jeśli wystąpi odwołanie do strony nieobecnej w pamięci głównej, sprzęt zarządzający pamięcią wykrywa to i w uzgodnieniu z SO przechodzi do załadowania brakującej strony. Taki schemat jest określany mianem **pamięci wirtualnej**. Przedstawiono go na rysunku 2.9.

Sprzęt procesora wraz z systemem operacyjnym dostarcza użytkownikowi „procesora wirtualnego”, który ma dostęp do pamięci wirtualnej. Ta pamięć może być liniową przestrzenią adresów lub kolekcją segmentów będących różnej długości blokami kolejnych adresów. W obu przypadkach instrukcje języka programowania mogą się odnosić do miejsc w programie i do danych w obszarze pamięci wirtualnej. Izolowanie procesów można osiągnąć przez danie każdemu procesowi jednoznacznie określonej, niezachodzącej na inną pamięci wirtualnej. Dzielenie pamięci można uzyskać przez zachodzące na siebie porcje dwu pamięci wirtualnych. Pliki są utrzymywane w pamięci

⁸ A dokładniej: musi trafić do którejś z tak zwanych ramek, na które pamięć stronicowana jest sprzętowo dzielona — *przyjp. tłum.*



Rysunek 2.10. Adresowanie pamięci wirtualnej

Ochrona informacji i bezpieczeństwo

Wzrost zastosowań systemów z podziałem czasu oraz — w czasach bliższych współczesności — sieci komputerowych spowodował zwiększenie troski o ochronę informacji. Natura zagrożeń, którym poświęcają uwagę instytucje i przedsiębiorstwa, zależy od okoliczności. Istnieją jednak pewne ogólne narzędzia, nadające się do wbudowania w komputery i systemy operacyjne, które umożliwiają różnorodne mechanizmy ochrony i bezpieczeństwa. Ogólnie biorąc, interesuje nas problem kontrolowania dostępu do systemów komputerowych i przechowywanych w nich informacji.

Większość prac dotyczących ochrony i bezpieczeństwa w odniesieniu do systemów operacyjnych można z grubsza zaliczyć do następujących kategorii:

1. **Dostępność** (ang. *availability*). Dotyczy ochrony systemu przed przerwami w pracy.
2. **Poufność** (ang. *confidentiality*). Zapewnienie, że użytkownicy nie mogą czytać danych, jeśli nie zostali upoważnieni, aby po nie sięgać.
3. **Nienaruszalność danych** (ang. *data integrity*). Ochrona danych przed modyfikacją bez upoważnienia.
4. **Uwierzytelnianie** (ang. *authenticity*). Dotyczy odpowiedniego sprawdzania tożsamości użytkowników i ważności (wiarygodności) komunikatów lub danych.

Planowanie operowania zasobami

Podstawowym obowiązkiem systemu operacyjnego jest zarządzanie różnymi zasobami, którymi rozporządza (obszarem pamięci głównej, urządzeniami zewnętrznymi, procesorami), oraz planowanie ich użycia przez różne aktywne procesy. Przydział każdego zasobu i polityka planowania muszą uwzględniać trzy czynniki:

1. **Uczciwość**. Zwykle dążymy do tego, aby wszystkie procesy rywalizujące o możliwość użycia pewnego zasobu otrzymywały do niego dostęp w przybliżeniu na równych prawach i sprawiedliwie. Jest tak szczególnie w przypadku zadań tej samej klasy, to znaczy zadań o podobnych wymaganiach.

- 2. Zróżnicowana reakcja.** Z drugiej strony, może istnieć konieczność rozróżniania przez SO klas zadań o różnych wymaganiach dotyczących obsługi. System operacyjny powinien starać się podejmować decyzje o przydziale i planowaniu w taki sposób, aby spełniać cały zbiór wymagań. Powinien również podejmować te decyzje dynamicznie. Jeśli na przykład proces czeka na użycie urządzenia wejścia-wyjścia, to SO może dążyć do zaplanowania tego procesu do wykonywania najszybciej jak to możliwe. Proces może wtedy natychmiast skorzystać z urządzenia i zwolnić je, oddając do dyspozycji innych procesów, którym będzie ono potrzebne później.
- 3. Sprawność.** SO powinien starać się maksymalizować przepustowość, minimalizować czas odpowiedzi i — w przypadku podziału czasu — przyjąć tylu użytkowników, ilu się da. Te kryteria są wzajemnie sprzeczne; znalezienie właściwej równowagi w konkretnej sytuacji jest nieustającym problemem w badaniach systemów operacyjnych.

Planowanie zasobów i zarządzanie nimi są w istocie problemami badań operacyjnych, jednak matematyczne wyniki osiągnięte w tej dyscyplinie nadają się do zastosowania również w systemach operacyjnych. Ważne jest ponadto mierzenie aktywności systemu, aby można było monitorować i regulować jego działanie.

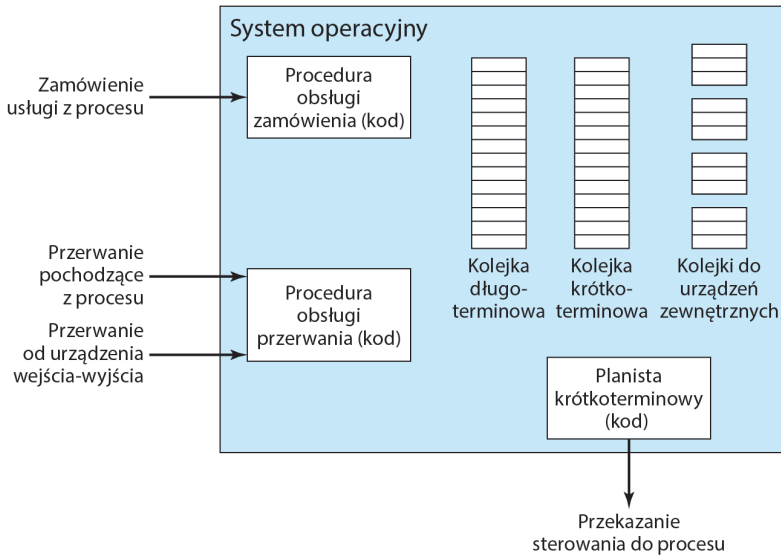
Na rysunku 2.11 naszkicowano główne elementy SO zaangażowane w procesy planowania i przydziału zasobów w środowisku wieloprogramowym. System operacyjny utrzymuje pewną liczbę kolejek będących po prostu listami procesów oczekujących na jakiś zasób. Kolejka krótkoterminowa składa się z procesów, które znajdują się w pamięci głównej (lub przynajmniej znajduje się tam istotna porcja każdego z nich) i są gotowe do działania, gdy tylko procesor stanie się dostępny¹⁰. Każdy z tych procesów mógłby natychmiast użyć procesora. Od planisty krótkoterminowego, czyli ekspedytora (dyspozytora, ang. *dispatcher*), zależy, który zostanie wybrany¹¹. Typowa strategia polega na dawaniu każdemu procesowi w kolejce pewnej porcji czasu po kolei; określa się to mianem techniki **rotacyjnej** (ang. *round-robin*). Skutkuje ona traktowaniem kolejki jako cyklicznej. Inna strategia polega na przydzieleniu poszczególnym procesom poziomów priorytetów. Planista wybiera wtedy procesy w porządku priorytetów.

Kolejka długoterminowa jest listą nowych zadań czekających na użycie procesora. System operacyjny dodaje nowe zadania do wykonania przez przeniesienie procesu z kolejki długoterminowej do kolejki krótkoterminowej. W tym czasie nowo przybyłemu procesowi musi być przydzielona porcja pamięci głównej. SO musi zatem mieć pewność, że nie przekroczy pamięci lub możliwości przetwarzania przez przyjęcie zbyt wielu procesów. Z każdym urządzeniem zewnętrznym jest związana kolejka wejścia-wyjścia. To samo urządzenie może być zamówione przez więcej niż jeden proces. Wszystkie procesy czekające na użycie dowolnego z urządzeń są ustawiane jeden za drugim w kolejce do odpowiedniego urządzenia. Tu również SO musi rozstrzygać, któremu procesowi przydzielić dostępne urządzenie.

Jeśli wystąpi przerwanie, SO uzyskuje kontrolę nad procesorem w procedurze obsługi przerwania. Proces może określać, o którą usługę mu chodzi (np. o pobudzenie do działania procedury obsługi urządzenia wejścia-wyjścia), za pomocą wywoływania usług. W tej sytuacji procedura obsługi danego zamówienia jest punktem wejścia do SO. W każdym przypadku po obsłużeniu przerwania lub wywołanej usługi uruchamiany jest planista krótkoterminowy, aby wybrać proces do wykonania.

¹⁰ Opisana sytuacja sprawia, że inni autorzy używają nazwy kolejka gotowych (kolejka procesów gotowych do działania, ang. *ready queue*) — *przyp. tłum.*

¹¹ Inni autorzy nie utożsamiają tych funkcji: planista krótkoterminowy wybiera proces do wykonywania, po czym ekspedytor ładuje kontekst tego procesu do rejestrów sprzętowych — *przyp. tłum.*



Rysunek 2.11. Podstawowe elementy systemu operacyjnego służące do organizowania wieloprogramowości

To, co tu powiedziano, jest opisem funkcjonalnym; szczegóły i modularny projekt tej części SO będą się różniły w różnych systemach. Znaczna część badań i prac rozwojowych w obszarze systemów operacyjnych była ukierunkowana na znalezienie algorytmów i struktur danych umożliwiających planowanie operowania zasobami, które zapewniałyby uczciwość, różnicowanie reakcji i efektywność działania.

2.4. DROGA DO WSPÓŁCZESNYCH SYSTEMÓW OPERACYJNYCH

Z biegiem czasu następowała stopniowa ewolucja struktury i zdolności SO. Jednakże w ostatnich latach zarówno do projektów nowych systemów operacyjnych, jak i do nowych wydań systemów istniejących wprowadzono liczne nowe elementy, które spowodowały poważne zmiany w ich naturze. Nowoczesne systemy operacyjne są odpowiedzią na postęp w rozwoju sprzętu, nowe zastosowania i nowe zagrożenia bezpieczeństwa. Do kluczowych osiągnięć w sprzęcie należą systemy wieloprocessorowe istotnie zwiększające szybkość procesora, osprzęt sieci dużych prędkości i rosnąca pojemność różnego rodzaju urządzeń magazynowania informacji. Na polu zastosowań wpływ na projekt SO wywarły aplikacje multimedialne, dostęp do Internetu i usługi WWW oraz model klient-serwer. Jeśli chodzi o bezpieczeństwo, dostęp z Internetu do komputerów spowodował znaczny wzrost potencjalnych zagrożeń, a rosnąca liczba wymyślnych ataków (jak wirusy, robaki i techniki hakierskie) odcisnęła poważne piętno na projektowaniu SO.

Tempo zmian w oczekiwaniach stawianych przed systemami operacyjnymi wymaga już nie tylko modyfikacji i ulepszeń w istniejących architekturach, lecz nowych sposobów organizacji SO. Wypróbowano szeroki zakres różnych podejść i elementów projektowych, tak w eksperymentalnych, jak i komercyjnych systemach operacyjnych, jednak znaczna część tych działań mieści się w następujących kategoriach:

- architektura mikrojądra,
- wielowątkowość,

- przetwarzanie symetryczne,
- rozproszone systemy operacyjne,
- projektowanie obiektowe.

Do niedawna większość systemów operacyjnych przybierała postać wielkiego **monolitycznego jądra** (ang. *monolithic kernel*). Większość tego, co pojmujemy jako funkcjonalność SO, mieści się w tych wielkich jądrach: planowanie, system plików, praca sieciowa, moduły sterujące urządzeniami, zarządzanie pamięcią — i to jeszcze nie wszystko. Jądro monolityczne jest zazwyczaj implementowane w postaci pojedynczego procesu, którego wszystkie elementy dzielą tę samą przestrzeń adresową. W architekturze **mikrojądra** (ang. *microkernel*) jądro przypada w udziale tylko kilka istotnych funkcji, w tym zarządzanie przestrzenią adresową, komunikacja międzyprocesowa (ang. *interprocess communication* — IPC) i podstawowe planowanie. Inne usługi SO są świadczone przez procesy, niekiedy zwane serwerami, pracujące w trybie użytkownika i traktowane przez mikrojądro na równi z innymi aplikacjami. Koncepcja mikrojądra upraszcza implementację, zapewnia elastyczność i dobrze pasuje do środowiska rozproszonego. Reasumując, mikrojądro współdziała z lokalnymi i zdalnymi procesami serwerów w jednakowy sposób, ułatwiając konstruowanie systemów rozproszonych.

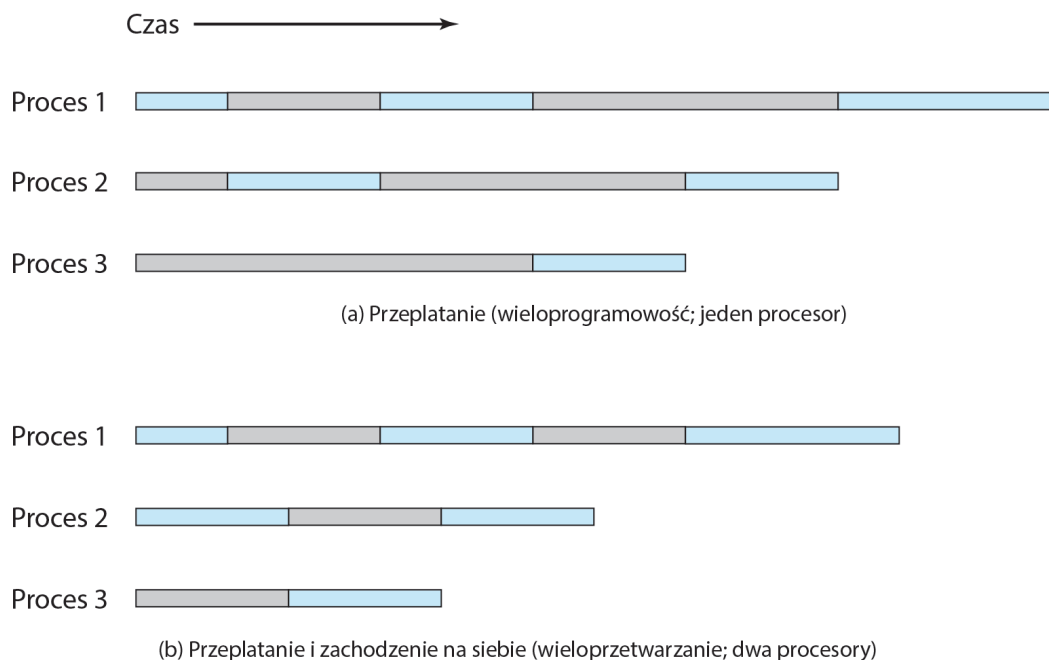
Wielowątkowość (ang. *multithreading*) jest techniką polegającą na tym, że proces wykonujący aplikację jest dzielony na wątki, które mogą działać współbieżnie. Możemy dokonać następującego rozróżnienia:

- **Wątek** (ang. *thread*) jest ekspediowalną jednostką pracy. Zawiera kontekst procesora (m.in. licznik programu i wskaźnik stosu) i własny obszar danych na stos (aby umożliwić wywołanie podprogramów). Wątek działa sekwencyjnie i jest przerywalny, więc procesor może przejść do innego wątku.
- **Proces** to zbiór złożony z jednego lub więcej wątków i skojarzonych z nimi zasobów systemowych (takich jak pamięć zawierająca kod i dane, otwarte pliki i urządzenia). Jest to bliskie koncepcji programu w trakcie wykonywania. Rozbijając aplikację na wiele wątków, programista zyskuje dużą kontrolę nad modularnością aplikacji i przebiegiem związanych z nią zdarzeń.

Wielowątkowość jest pożyteczna w aplikacjach wykonujących wiele w gruncie rzeczy niezależnych **prac** (ang. *tasks*), których nie trzeba szeregować. Przykładem jest będący na nasłuchu serwer bazy danych przetwarzający zamówienia licznych klientów. W procesie z wieloma wątkami przełączanie między jednym wątkiem a innym (i z powrotem) zużywa mniej czasu procesora niż poważniejsze w realizacji przełączanie między różnymi procesami. Wątki są również przydatne do strukturalizacji procesów będących częścią jądra SO, co zostanie opisane w następnych rozdziałach.

Wieloprzetwarzanie symetryczne (ang. *symmetric multiprocessing* — SMP) jest pojęciem odnoszącym się do architektury sprzętowej komputera (opisanej w rozdziale 1.), a także do zachowania SO, który korzysta z tej architektury. System operacyjny realizujący SMP planuje procesy lub wątki na wszystkich procesorach. SMP ma wiele potencjalnych przewag nad architekturą jednoprocessorową, w tym:

- **Wydajność.** Jeśli pracę do wykonania na komputerze można zorganizować tak, aby pewne jej fragmenty można było wykonać równolegle, to system z wieloma procesorami będzie działał efektywniej niż z jednym procesorem tego samego typu. Zilustrowano to na rysunku 2.12.



■ Zablokowany ■ Wykonywany

Rysunek 2.12. Wieloprogramowość i wieloprzetwarzanie

W wieloprogramowości tylko jeden proces może pracować w danej chwili, wszystkie inne procesy czekają na procesor. W wieloprzetwarzaniu może działać jednocześnie więcej procesów — każdy na innym procesorze.

- **Dostępność.** Ponieważ wszystkie procesory wieloprocesora symetrycznego mogą wykonywać te same funkcje, awaria jednego procesora nie zatrzymuje systemu; system może kontynuować działanie ze zmniejszoną sprawnością.
- **Stopniowy wzrost.** Użytkownik może poprawić działanie systemu przez dodanie kolejnego procesora.
- **Skalowanie.** Dostawcy mogą oferować szerszy asortyment wyrobów o różnych cenach i charakterystykach działania, bazując na liczbie procesorów skonfigurowanych w systemie.

Należy zauważyć, że są to korzyści raczej potencjalne niż gwarantowane. System operacyjny musi udostępniać narzędzia i funkcje do wykorzystania równoległości w systemie SMP.

Wielowątkowość i przetwarzanie SMP są często omawiane razem, lecz są to dwa niezależne udogodnienia. Nawet w systemie jednoprocessorowym wielowątkowość jest przydatna do strukturalizacji procesów użytkowych i jądra. System SMP jest pożyteczny nawet w wypadku procesów bez wątkowości, gdyż kilka procesów może działać równolegle. Wszakże oba rozwiązania uzupełniają się wzajemnie i można z nich skutecznie korzystać razem.

Atrakcyjną cechą SMP jest to, że istnienie wielu procesorów jest dla użytkownika przezroczyste. Planowaniem wątków lub procesów na poszczególnych procesorach i synchronizacją między nimi zajmuje się system operacyjny. W tej książce omawiamy mechanizmy planowania i synchronizacji

stosowane do wywierania na użytkownika wrażenia jednego systemu. Odmiennym problemem jest tworzenie wrażenia jednego systemu w przypadku grona osobnych komputerów — systemu wielokomputerowego. W tym wypadku mamy do czynienia ze zbiorem komputerów: każdy z nich ma własną pamięć główną, pamięć drugorzędną i inne moduły wejścia-wyjścia. **Rozproszony system operacyjny** (ang. *distributed operating system*) dostarcza iluzji jednej przestrzeni pamięci głównej i jednej przestrzeni pamięci pomocniczej oraz innych ujednoliconych udogodnień dostępu, takich jak rozproszony system plików. Chociaż grona (klastry) coraz bardziej zyskują na popularności i na rynku jest wiele produktów tego typu, stan sztuki, jeśli chodzi o rozproszone systemy operacyjne, zostaje w tyle za systemami operacyjnymi jednoprocessorowymi i SMP. W części VIII książki przyjrzymy się takim systemom.

Innowacją w projektowaniu SO jest również zastosowanie technologii obiektowych. **Projektowanie obiektowe** (ang. *object-oriented design*) pomaga dyscyplinować dodawanie modularnych rozszerzeń do małego jądra. Na poziomie SO struktura oparta na obiektach umożliwia osobom programującym dokonywanie przeróbek SO bez naruszania całości systemu. Obiektowość ułatwia również opracowywanie narzędzi rozproszonych i rozwiniętych rozproszonych systemów operacyjnych.

2.5. TOLEROWANIE AWARII

Tolerowanie awarii oznacza zdolność systemu lub komponentu do kontynuowania normalnego działania mimo występowania wad w sprzęcie lub oprogramowaniu. Zazwyczaj wymaga to zastosowania pewnego rodzaju nadmiarowości. Tolerowanie awarii ma służyć zwiększaniu niezawodności systemu. Polepszeniu zdolności do tolerowania awarii (a więc i zwiększeniu niezawodności) na ogół towarzyszy wzrost kosztów rozumianych w kategoriach finansowych albo wydajnościowych (albo w obu). Dlatego jest konieczne ustalenie miar zakresu stosowania tolerowania awarii wedle tego, do jakiego stopnia dany zasób ma krytyczne znaczenie w systemie.

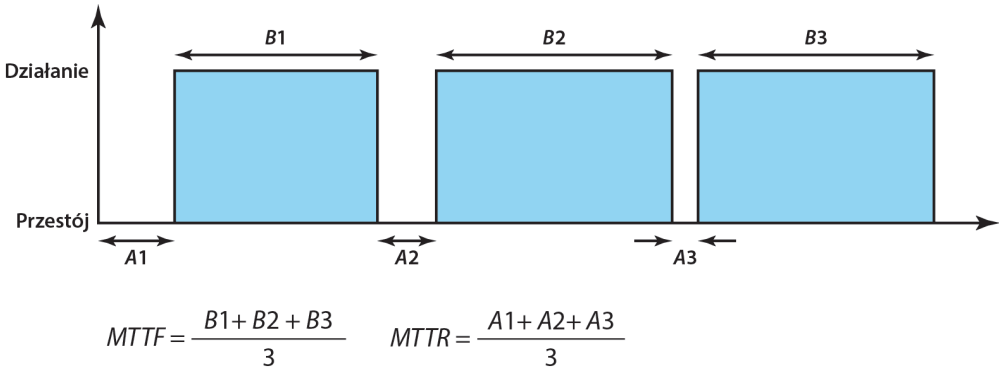
Podstawowe pojęcia

Z tolerowaniem awarii wiązą się trzy podstawowe miary jakości działania systemu: niezawodność, średni czas do wystąpienia awarii (MTTF) oraz dostępność. Te pojęcia powstały ze szczególnym odniesieniem do wad sprzętowych, lecz stosują się ogólniej zarówno do wad sprzętowych, jak i programowych.

Niezawodność (ang. *reliability*) $R(t)$ systemu jest definiowana jako prawdopodobieństwo jego poprawnego działania do chwili t przy założeniu, że system działał poprawnie w chwili $t = 0$. W odniesieniu do systemów komputerowych i systemów operacyjnych termin *poprawne działanie* oznacza poprawne wykonywanie zbioru programów i ochronę danych przed niepożądanymi modyfikacjami. **Średni czas do wystąpienia awarii** (ang. *mean time to failure* — MTTF) jest definiowany jako:

$$\text{MTTF} = \int_0^{\infty} R(t) dt$$

Średni czas naprawy (ang. *mean time to repair* — MTTR) jest średnim czasem zużywanym na zreperowanie lub wymianę wadliwego elementu. Na rysunku 2.13 pokazano związek między MTTF i MTTR.



Rysunek 2.13. Stany operacyjne systemu

Dostępność (ang. *availability*) systemu lub usługi jest ułamkiem czasu, w którym system jest zdolny obsługiwać żądania użytkowników. Równoważnie dostępność oznacza prawdopodobieństwo, że jednostka działa poprawnie, spełniając zadane warunki przez określony czas. Czas, w którym system nie jest dostępny, jest nazywany **przestojem** (ang. *downtime*). Czas, w którym system jest dostępny, określany jest jako **czas pracy** (ang. *uptime*). Dostępność systemu można wyrazić wzorem:

```
$JOB
$FTN
.
.
.
} Instrukcje w Fortranie
$LOAD
$RUN
.
.
.
} Dane
$END
```

W tabeli 2.4 pokazano niektóre powszechnie rozróżniane poziomy dostępności i odpowiednie okresy przestoju w skali roku.

Tabela 2.4. Klasy dostępności

Klasa	Dostępność	Przestój w skali roku
Ciągła	1,0	0
Tolerująca awarie	0,99999	5 minut
Znosząca awarie	0,9999	53 minuty
Wysoka dostępność	0,999	8,3 godziny
Zwykła dostępność	0,99 – 0,995	44 – 87 godzin

Często lepszym niż dostępność wskaźnikiem jest MTTF, czyli średni czas pracy. Połączenie krótkotrwałych przestojów z krótkimi czasami pracy mogłoby dać wysoką miarę dostępności, lecz użytkownicy nie byłoby w stanie skorzystać z żadnej usługi, gdyby czas pracy był mniejszy niż czas potrzebny do jej wykonania.

Wady

Słownik standardów IEEE definiuje **wadę** (ang. *fault*) jako błędny (zafałszowany, ang. *erroneous*) stan sprzętu lub oprogramowania wynikający z wady komponentu, pomyłki operatora, fizycznego wpływu środowiska, błędu projektowego, błędu programu lub błędu w strukturze danych. Standard głosi również, że wada objawia się 1) defektem urządzenia lub komponentu sprzętowego, na przykład krótkim spięciem lub zerwanym przewodem, lub 2) niepoprawnym krokiem, procesem lub definicją danych w programie komputerowym.

Wady możemy podzielić na następujące kategorie:

- **Trwała** (ang. *permanent*). Wada, która pojawia się i nie ustępuje. Wada trwa dopóty, dopóki wadliwa składowa nie zostanie wymieniona lub naprawiona. Przykładem może być awaria głowicy dyskowej, błędy w oprogramowaniu lub przepalony komponent komunikacyjny.
- **Chwilowa** (ang. *temporary*). Wada, która nie jest obecna nieustannie, w każdych warunkach działania. Wady chwilowe można dalej sklasyfikować następująco:
 - **Przejęciowa** (ang. *transient*). Wada występująca tylko jeden raz. Do przykładów należą błędy transmisji bitu spowodowane impulsowym szumem, zakłóceniami zasilania lub promieniowaniem, które zmienia bit pamięci.
 - **Nieciągła** (ang. *intermittent*). Wada pojawiająca się wielokrotnie w nieokreślonych momentach. Przykładem wady przejściowej jest luźny styk w łączówce.

Ogólnie biorąc, tolerowanie awarii jest wbudowywane do systemu przez dodanie nadmiarowości. Do rodzajów nadmiarowości należą:

- **Nadmiarowość przestrzenna (fizyczna)**. Fizyczna nadmiarowość polega na zastosowaniu wielu komponentów, które albo wykonują tę samą funkcję jednocześnie, albo są tak skonfigurowane, że jeden komponent jest dostępny jako zapasowy w wypadku awarii drugiego. Przykładem pierwszego rozwiązania jest zastosowanie wielu równoległych obwodów, w których za wynik wyjściowy uznaje się rezultat większościowy. Przykładem drugiej możliwości jest zapasowy serwer nazw w Internecie.
- **Nadmiarowość chwilowa**. Nadmiarowość chwilowa polega na powtórzeniu funkcji lub operacji w sytuacji wykrycia błędu. To podejście jest skuteczne w wypadku wad chwilowych, lecz nieprzydatne przy wadach trwałych. Przykładem jest powtórzenie transmisji bloku danych po wykryciu błędu, jak to się robi w protokołach łącza danych.
- **Nadmiarowość informacji**. Nadmiarowość informacji umożliwia tolerowanie awarii przez zwielokrotnienie lub zakodowanie danych w taki sposób, że błędy bitów można zarówno wykryć, jak i skorygować. Przykładami są obwód kodujący z kontrolą błędów, używany w systemach pamięci, i techniki korygowania błędów w dyskach RAID, o czym powiemy w następnych rozdziałach.

Mechanizmy systemu operacyjnego

Oprogramowanie SO może zawierać wiele mechanizmów wspierających tolerowanie awarii. W książce napotkamy sporo tego rodzaju przykładów. Tu podajemy niektóre:

- **Izolowanie procesu.** Jak już wspomniano w tym rozdziale, procesy są z zasady izolowane od siebie w rozumieniu miejsca ich pobytu w pamięci głównej, dostępu do plików i przebiegu sterowania. Struktura dostarczana przez SO na użytek administrowania procesami umożliwia pewien stopień ochrony innych procesów przed procesem powodującym błąd.
- **Sterowanie współbieżnością.** W rozdziałach 5. i 6. będą omówione niektóre trudności i wady mogące powstawać podczas komunikacji i współpracy procesów. W tych rozdziałach zostaną również omówione sposoby stosowane do zapewniania poprawnego działania i wychodzenia z sytuacji wadliwych, w rodzaju zakleszczenia.
- **Maszyny wirtualne.** Maszyny wirtualne, o czym powiemy w rozdziale 14., umożliwiają większy stopień izolowania aplikacji, więc również tolerowania awarii. Można je także zastosować do tworzenia nadmiarowości na tej zasadzie, że jedna maszyna wirtualna służy jako zapasowa dla innej.
- **Punkty kontrolne i wycofania.** Punkt kontrolny (ang. *checkpoint*) jest kopią stanu aplikacji zmagazynowaną w jakiejś pamięci odpornej na przewidywane awarie. Wycofanie polega na wznowieniu wykonywania od wcześniej przechowanego punktu kontrolnego. Gdy dojdzie do awarii, aplikację cofa się do stanu z poprzedniego punktu kontrolnego i od tego miejsca podejmuje się jej wykonywanie. Tej techniki można użyć do wychodzenia z przejściowych, a także trwałych awarii oraz z pewnych rodzajów awarii oprogramowania. Rozwiązania takie są zazwyczaj wbudowane w systemy baz danych i systemy transakcyjne.

Na uwagę zasługiwałoby znacznie więcej technik, lecz pełne omówienie tolerowania awarii w systemach operacyjnych wykracza poza ramy, które tutaj sobie nakreśliliśmy.

2.6. PROBLEMY PROJEKTOWANIA SYSTEMÓW OPERACYJNYCH WIELOPROCESORÓW I KOMPUTERÓW WIELORDZENIOWYCH

Rozważania dotyczące wieloprocesorowego symetrycznego SO

W systemie SMP (wieloprocesorowego przetwarzania symetrycznego) jądro może działać na każdym procesorze i na ogół każdy procesor sam planuje procesy lub wątki z dostępnej puli. Jądro może być skonstruowane z wielu procesów lub wątków, co umożliwia równoległe wykonywanie jego części. Podejście SMP komplikuje system operacyjny. Osoba projektująca SO musi poradzić sobie ze złożonością wynikającą z dzielenia zasobów (np. struktur danych) i działań koordynujących (np. związanych z dostępem do urządzeń) pochodzących z wielu części SO wykonywanych w tym samym czasie. Trzeba stosować odpowiednie techniki zamawiania i synchronizowania zasobów.

System operacyjny SMP zarządza procesorem i innymi zasobami komputera tak, aby użytkownik postrzegał go jako wieloprogramowy system jednoprocessorowy. Użytkownik może budować aplikacje wieloprocesowe lub z wieloma wątkami w procesach, nie troszcząc się o to, czy ma do dyspozycji jeden, czy wiele procesorów. Tak więc wieloprocesorowy SO musi zapewnić całą funkcjonalność systemu wieloprogramowego, a ponadto dodatkowe właściwości umożliwiające korzystanie z wielu procesorów. Do podstawowych zagadnień projektowych należą tu:

- **Jednoczesne współbieżne procesy lub wątki.** Procedury jądra muszą być wznowialne, aby umożliwiać jednoczesne wykonywanie na kilku procesorach tego samego kodu jądra. Jeśli wiele procesorów wykonuje te same lub różne części jądra, to zarządzanie tablicami i strukturami jądrowymi musi odbywać się tak, aby uniknąć uszkodzenia danych lub niedozwolonych operacji.
- **Planowanie.** Dowolny procesor może realizować planowanie, co komplikuje zadanie wymuszania polityki planowania i zapewnianie, że nie dojdzie do naruszenia struktur danych planisty. Jeśli na poziomie jądra stosuje się wielowątkowość, to istnieje możliwość planowania wielu wątków tego samego procesu jednocześnie na wielu procesorach. Planowanie wieloprocessorów będzie omówione w rozdziale 10.
- **Synchronizacja.** Jeśli wiele aktywnych procesów ma potencjalny dostęp do wspólnych przestrzeni adresowych lub dzielonych zasobów wejścia-wyjścia, należy zadbać o skuteczną synchronizację. Synchronizacja wymusza wzajemne wykluczanie i porządkowanie zdarzeń. Typowym mechanizmem synchronizacji, używanym w wieloprocessorowych systemach operacyjnych, są blokady, które opisemy w rozdziale 5.
- **Zarządzanie pamięcią.** Zarządzanie pamięcią wieloprocessora musi uwzględniać wszystkie problemy występujące na komputerach jednoprocessorowych i zostanie przedstawione w części III książki. System operacyjny musi ponadto wykorzystywać istniejącą w sprzęcie równoległość, aby osiągać najlepsze działanie. Mechanizmy stronicowania na poszczególnych procesorach muszą być koordynowane, aby wymuszać spójność w sytuacjach, gdy ileś procesorów dzieli stronę lub segment, i rozstrzygać o zastępowaniu stron. Ponowne użycie stron fizycznych stanowi największy problem: trzeba zagwarantować, że strona fizyczna ze starą zawartością nie będzie dostępna, zanim nie zostanie jej wyznaczony zastosowanie.
- **Niezawodność i tolerowanie awarii.** W wypadku awarii procesora system operacyjny powinien zapewniać łagodną degradację. Planista i inne części SO muszą rozpoznawać utratę procesora i odpowiednio restrukturyzować tablice zarządzania.

Ponieważ projektowanie wieloprocessorowego SO w zasadzie polega na rozszerzeniu rozwiązań problemów projektowania wieloprogramowego jednoprocessora, nie traktujemy wieloprocessorowych systemów operacyjnych oddzielnie. Zamiast tego specyficzne zagadnienia dotyczące wieloprocessorów są omawiane w książce w odpowiednich miejscach.

Rozważania dotyczące wielordzeniowych SO

Wszystkie zagadnienia projektowe dotyczące systemów SMP, omówione dotychczas w tym podrozdziale, odnoszą się również do systemów wielordzeniowych. Niemniej występują tu również inne problemy. Jednym z nich jest skala potencjalnej równoległości. Obecni dostawcy komputerów mają do zaoferowania dziesięć lub więcej rdzeni na jednym chipie. W każdej następnej technologii kolejnych generacji procesorów liczba rdzeni i ilość wspólnej i wydzielonej pamięci pod ręcznej rośnie, wchodzimy więc w erę systemów „licznych rdzeni”.

Niemalym wyzwaniem w projektowaniu systemów wielordzeniowych o wielkiej liczbie rdzeni jest efektywne zaprzęgnięcie do roboty mocy obliczeniowej licznych rdzeni i inteligentne zarządzanie sporymi zasobami pomieszczonymi w jednym układzie. W centrum zainteresowań znajduje się pytanie, jak dopasować wewnętrzny paralelizm systemu mnogich rdzeni do zapotrzebowań wydajnościowych aplikacji. Potencjalna równoległość istnieje w istocie na trzech poziomach współcze-

snego systemu wielordzeniowego. Po pierwsze, wewnątrz każdego procesora rdzeniowego istnieje równoległość sprzętowa, znana jako równoległość na poziomie rozkazów, z której mogą, lub nie, skorzystać programiści aplikacji i kompilatory. Po drugie, możliwość wykonywania wieloprogramowego i wielowątkowego można realizować w każdym procesorze. I na koniec potencjalnie każda aplikacja z osobna może być wykonana we współbieżnych procesach lub wątkach na wielu rdzeniach. Bez mocnego i skutecznego wsparcia ze strony SO obu ostatnich rodzajów równoległości zasoby sprzętowe nie będą wykorzystane efektywnie.

Z nastaniem technologii wielordzeniowej projektanci SO zmagają się w istocie z problemem możliwie najlepszego wydobycia równoległości z prac, którymi są obciążane komputery. Sprawdzają się rozmaite podejścia do budowy systemów operacyjnych nowej generacji. W tym podrozdziale przedstawimy dwa ogólne sposoby postępowania, o innych szczegółach opowiemy w dalszych rozdziałach.

RÓWNOLEGŁOŚĆ W OBRĘBIE APLIKACJI

Większość aplikacji można w zasadzie podzielić na wiele prac, które mogą być wykonywane równolegle przez realizowanie tych prac w postaci wielu procesów, być może wielowątkowych. Trudność polega na tym, że o podziale pracy aplikacji na niezależnie wykonywalne zadania musi decydować jej twórca. To znaczy musi on zdecydować, które fragmenty mogą lub powinny być wykonane asynchronicznie lub równolegle. W procesie projektowania programowania równoległego przede wszystkim pomagają właściwości kompilatora i języka programowania. Jednak SO może wspomagać ten proces projektowania przynajmniej przez skuteczne przydzielanie zasobów równoległym zadaniom zdefiniowanym przez twórcę aplikacji.

Jedną z najskuteczniejszych inicjatyw mających wesprzeć budowniczych jest Grand Central Dispatch (GCD, z ang. Wielka Centralna Ekspedycja) zrealizowana w ostatnich wydaniach opartych na UNIX-ie systemów Mac OS X oraz iOS. GCD jest środkiem wspomagania wielordzeniowości. Nie pomaga twórcom oprogramowania w rozstrzyganiu, jak podzielić zadanie lub aplikację na osobne współbieżne części. Jednak gdy budowniczy zidentyfikuje coś, z czego można wydzielić osobne zadanie, GCD umożliwia urzeczywistnienie tego zamiaru w łatwy i nieinwazyjny sposób.

W gruncie rzeczy GCD jest mechanizmem puli wątków, w którym SO odwzorowuje prace na wątki reprezentujące osiągalny stopień współbieżności (plus wątki blokujące się na wejściu-wyjściu). System Windows ma również mechanizm puli wątków (od 2000 roku), a w aplikacjach serwerowych pule wątków są używane od lat. Nowością w GCD jest to, że stanowi rozszerzenie wobec języków programowania, umożliwiające stosowanie funkcji anonimowych (zwanymi blokami) jako sposób specyfikowania prac. Wielka Centralna Ekspedycja nie jest zatem poważnym krokiem ewolucyjnym, niemniej stanowi nowe i wartościowe narzędzie wykorzystywania możliwej równoległości w systemie wielordzeniowym.

Jeden ze sloganów rodem z Apple'a głosi, że GCD stanowi „wyspy serializacji na morzu współbieżności”. Ujmuje to praktyczny aspekt coraz większego uwspółbieżniania przeciętnych aplikacji komputerów biurowych. Owe wyspy są tym, co izoluje twórców od drażliwych problemów jednoczesnego dostępu do danych, zakleszczeń i innych pułapek wielowątkowości. Budowniczowie aplikacji są zachęceni do identyfikowania w swoich wyrobach funkcji, które mogłyby być lepiej wykonywane poza głównym wątkiem — nawet jeśli składają się z pewnej liczby sekwencyjnych lub w inny sposób częściowo współzależnych prac. GCD ułatwia podzielenie całej jednostki pracy z utrzymaniem istniejącego porządku i zależności między podzadaniami. W dalszych rozdziałach przyjrzymy się niektórym detalom GCD.

PODEJŚCIE OPARTE NA MASZYNIE WIRTUALNEJ

Rozwiązanie alternatywne polega na uświadomieniu sobie, że wraz z nieustannie rosnącą liczbą rdzeni w układzie dążenie do wieloprogramowania poszczególnych rdzeni w celu obsługi wielu aplikacji może być mylnym użyciem zasobów [JACK10]. Jeśli zamiast tego pozwolimy na przypisanie jednego lub więcej rdzeni do konkretnego procesu i zostawimy procesor do jego wyłącznej dyspozycji, to unikniemy większości nakładów związanych z przełączaniem prac i decyzjami planistycznymi. Wielordzeniowy SO mógłby się wtedy przekształcić w hipernadzorcę podejmującego wysokopoziomowe decyzje o przydziale rdzeni do aplikacji, czyniąc niewiele ponad to w kwestii sposobu przydzielania zasobów.

Za takim rozwiązaniem przemawiają następujące argumenty. W zaraniu komputeryzacji jeden program działał na jednym procesorze. Z nadejściem wieloprogramowości każdej aplikacji stworzono iluzję działania na osobnym procesorze. Wieloprogramowość opiera się na koncepcji procesu będącego abstrakcją środowiska wykonywania. Aby zarządzać procesami, SO wymaga chronionej przestrzeni, wolnej od wpływu użytkownika i programu. W tym celu wprowadzono rozróżnienie między trybem jądra a trybem użytkownika. W rezultacie tryby jądra i użytkownika wyabstrahowały procesor w dwa procesory. W obecności tych procesorów przyszło jednak zmagać się z problemem, któremu z nich powierzyć prawdziwy procesor. Nakłady na przełączanie między tymi procesorami zaczęły rosnąć, aż osiągnęły punkt, w którym zaczęło się to negatywnie odbijać na gotowości do reagowania, szczególnie gdy wprowadzono wiele rdzeni. Wszakże w systemach z licznymi rdzeniami możemy dać sobie spokój z rozróżnianiem między trybami jądra i użytkownika. System operacyjny działa w tym podejściu jak swoisty hipernadzorca. Wiele obowiązków dotyczących zarządzania zasobami przechodzi na same programy. SO przydziela aplikację, procesor i jakieś zasoby, a program, korzystając z metadanych wygenerowanych przez kompilator, będzie sam najlepiej wiedział, jak z tych zasobów korzystać.

2.7. PRZEGLĄD SYSTEMU MICROSOFT WINDOWS

Rodowód

Microsoft początkowo (w 1985 roku) użył nazwy Windows na określenie rozszerzenia środowiska operacyjnego nader elementarnego systemu MS-DOS, który jako system operacyjny wcześniejszych komputerów osobistych okazał się bardzo udany. Kombinacja Windows/MS-DOS została w końcu zastąpiona nową wersją systemu Windows, znaną jako Windows NT, której pierwsze wydanie ukazało się w 1993 roku z przeznaczeniem do laptopów i systemów biurkowych. Choć podstawowa architektura wewnętrzna pozostawała z grubsza ta sama od czasów Windows NT, system był rozwijany dalej, obrastając w funkcje i właściwości. Ostatnie wydanie, pochodzące z czasu pisania tej książki, to Windows 10. Windows 10 łączy cechy poprzednich wydań biurkowo-laptopowych, wersji Windows 8.1, a także wersji systemu Windows przeznaczonych dla urządzeń mobilnych w Internecie rzeczy (ang. *Internet of things* — IoT). W Windows 10 wcielono również oprogramowanie z systemu Xbox One. Wynikowy, ujednolicony Windows 10 obsługuje komputery biurkowe, laptopy, smartfony, tablety i Xbox One.

Architektura

Na rysunku 2.14 przedstawiono ogólną budowę systemu Windows. Jak niemal wszystkie systemy, Windows oddziela oprogramowanie użytkowe od oprogramowania stanowiącego rdzeń SO. Ten ostatni, zawierający egzekutora, jądro, moduły sterujące i warstwę abstrakcji sprzętu, działa w trybie jądra. Oprogramowanie z trybu jądra ma dostęp do danych systemowych i do sprzętu. Pozostałe oprogramowanie, działające w trybie użytkownika, ma ograniczony dostęp do danych systemu.

ORGANIZACJA SYSTEMU OPERACYJNEGO

Windows ma wysoce modularną architekturę. Każda funkcja systemowa jest zarządzana przez tylko jedną składową SO. Reszta SO i wszystkie aplikacje uzyskują dostęp do tej funkcji za pośrednictwem odpowiedzialnej za nią składowej, korzystając ze standardowych interfejsów. Kluczowe dane systemu są dostępne tylko za pośrednictwem odpowiedniej funkcji. W zasadzie każdy moduł może być uaktualniony lub zastąpiony bez przepisywania całego systemu lub jego standardowych interfejsów programów użytkowych (ang. *application program interfaces* — API¹²). Oto składowe systemu Windows działające w trybie jądra:

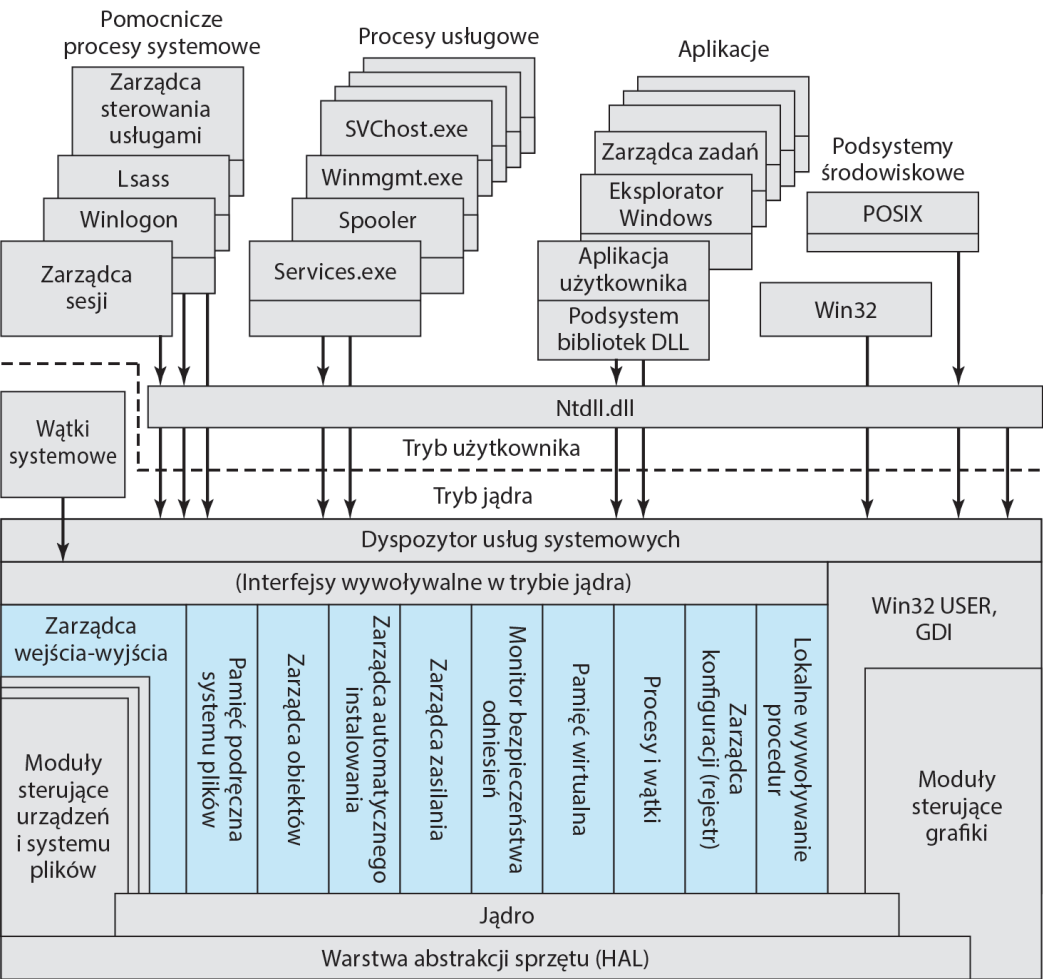
- **Egzekutor** (ang. *Executive*¹³). Zawiera podstawowe usługi SO, takie jak zarządzanie pamięcią, zarządzanie procesami, bezpieczeństwo, wejście-wyjście i komunikacja międzyprocesowa.
- **Jądro** (ang. *Kernel*). Steruje działaniem procesorów. Jądro zarządza planowaniem wątków, przełączaniem procesów, obsługą wyjątków i przerwania i synchronizacją wieloprocessora. W odróżnieniu od reszty¹⁴ egzekutora i poziomów użytkowych kod jądra nie działa w wątkach.
- **Warstwa abstrakcji sprzętu** (ang. *hardware abstraction layer* — HAL). Dokonuje odwzorowań między ogólnymi poleceniami i odpowiedziami sprzętowymi a specyficznymi dla danej platformy. Izoluje SO od różnic sprzętowych zależnych od platformy. HAL sprawia, że szyna systemowa każdego komputera, sterownik bezpośredniego dostępu do pamięci (DMA), sterownik przerwania, czasomierze systemowe i sterownik pamięci od strony egzekutora i składowych jądrowych wyglądają tak samo. Dostarcza również zaplecza potrzebnego do SMP (komunikacji międzyprocesowej), co zostanie wyjaśnione później.
- **Moduły sterujące urządzeń** (ang. *device drivers*)¹⁵. Są to biblioteki dynamiczne, które rozszerzają funkcjonalność egzekutora. Należą do nich moduły sterujące urządzeniami sprzętowymi, tłumaczące wywołania wejścia-wyjścia użytkownika na żądania wejścia-wyjścia specyficzne dla danego sprzętu, oraz komponenty programowe do implementowania systemu plików, protokoły sieciowe i wszelkie inne rozszerzenia systemu, które muszą działać w trybie jądra.
- **System okien i grafiki** (ang. *windowing and graphics system*). Realizuje funkcje GUI, takie jak operowanie oknami, kontrolkami interfejsu użytkownika i rysowaniem.

¹² Jak widać, również w języku angielskim akronim API jest rozwijany różnie, por. 2.1 — *przyp. tłum.*

¹³ Zachowujemy za oryginałem rozpoczynanie nazw niektórych składowych Windows dużą literą, choć nie jest to powszechnie przyjęte w innych opisach tego systemu — *przyp. tłum.*

¹⁴ Reszty w tym sensie, że również wywołania egzekutora są realizowane w trybie jądra (por. rysunek 2.14) — *przyp. tłum.*

¹⁵ Potocznie nazywane „drajwerami” lub sterownikami. Należy zaznaczyć, że sterownikami (ang. *controller*) nazywamy sprzętowe elementy systemu komputerowego; moduły sterujące są elementami oprogramowania; jeszcze inna spotykana nazwa: programy obsługi — *przyp. tłum.*



Lsass = serwer uwierzytelnień bezpieczeństwa lokalnego (ang. *local security authentication server*)
POSIX = interfejs przenośnego systemu operacyjnego (ang. *portable operating system interface*¹⁶)
GDI = interfejs urządzeń graficznych (ang. *graphics device interface*)
DLL = biblioteka konsolidowana dynamicznie (ang. *dynamic link library*)

Obszarami zabarwionymi zaznaczono egzekutora

Rysunek 2.14. Wewnętrzna architektura systemu Windows [RUS511]

Egzekutor systemu Windows zawiera składowe dotyczące specyficznych funkcji systemowych i dostarcza API oprogramowaniu działającemu w trybie użytkownika. Poniżej podajemy krótki opis poszczególnych modułów egzekutora:

- **Zarządca wejścia-wyjścia.** Tworzy ramę udostępniania aplikacjom urządzeń wejścia-wyjścia i odpowiada za ekspediowanie do odpowiednich modułów sterujących urządzeń w celu dalszego przetwarzania. Zarządca wejścia-wyjścia realizuje wszystkie API wejścia-wyjścia w systemie Windows i wymusza bezpieczeństwo oraz nazewnictwo urządzeń, protokołów sieciowych i systemów plików (używając zarządcy obiektów). Wejście-wyjście w systemie Windows będzie omówione w rozdziale 11.

- **Zarządca pamięci podręcznej.** Poprawia działanie wejścia-wyjścia opartego na plikach, umieszczając w pamięci głównej w celu szybkiego dostępu dane plików, do których następowały niedawno odniesienia, oraz opóźniając zapisy na dysk przez krótkotrwałe utrzymywanie uaktualnień w pamięci, nim zostaną wysłane na dysk w wydajniejszych partiach.
- **Zarządca obiektów.** Tworzy obiekty egzekutora systemu Windows, administruje nimi oraz je usuwa; obiekty te są używane do reprezentowania zasobów takich jak procesy, wątki i obiekty synchronizacji. Wymusza jednolite reguły zachowywania, nazywania i określania bezpieczeństwa obiektów. Zarządca obiektów tworzy również wpisy w tablicy uchwytów procesowych, zawierające informacje służące do kontroli dostępu i wskaźnik do obiektu. Obiekty systemu Windows będą omówione w dalszej części tego podrozdziału.
- **Zarządca automatycznego instalowania** (ang. *plug-and-play manager*). Określa, które moduły sterujące są potrzebne do obsługi danego urządzenia, i ładuje te moduły.
- **Zarządca zasilania** (ang. *power manager*). Koordynuje zarządzanie zasilaniem różnych urządzeń i może być konfigurowany w celu zmniejszenia zużycia energii przez wyłączanie bezczynnych urządzeń, usypianie procesora, a nawet przepisywanie całej pamięci operacyjnej na dysk i wyłączanie zasilania całego systemu.
- **Monitor bezpieczeństwa odniesień** (ang. *security reference monitor*). Wymusza reguły kontrolowania prawomocności dostępu i generowania audytów. Model obiektowy systemu Windows umożliwia spójne i ujednolicone ujęcie bezpieczeństwa aż po elementarne jednostki tworzące egzekutora. Tak więc Windows stosuje te same procedury sprawdzania legalności dostępu i kontroli w odniesieniu do wszystkich chronionych obiektów, włączając w to pliki, procesy, przestrzenie adresowe i urządzenia zewnętrzne. Bezpieczeństwo w systemie Windows będzie omówione w rozdziale 15.
- **Zarządca pamięci wirtualnej** (ang. *virtual memory manager*). Administruje adresami wirtualnymi, pamięcią fizyczną i plikami stronicowania na dysku. Nadzoruje sprzęt zarządzający pamięcią i struktury danych odwzorowujące wirtualne adresy z przestrzeni adresowej procesu na strony fizyczne w pamięci komputera. Zarządzanie pamięcią wirtualną w systemie Windows będzie omówione w rozdziale 8.
- **Zarządca procesów lub wątków** (ang. *process/thread manager*). Tworzy i usuwa obiekty procesów i wątków oraz administruje nimi. Zarządzanie procesami i wątkami będzie opisane w rozdziale 4.
- **Zarządca konfiguracji** (ang. *configuration manager*). Odpowiada za zarządzanie rejestrem systemu, będącym magazynem ustawień różnych parametrów, zarówno ogólnosystemowych, jak i poszczególnych użytkowników.
- **Udogodnienie ulepszanego wywołania procedury lokalnej** (ang. *advanced local procedure call facility* — ALPC). Realizuje sprawny mechanizm międzyprocesowych wywołań procedur między lokalnymi procesami implementującymi usługi i podsystemy. Jest podobne do zdalnego wywołania procedury (ang. *remote procedure call* — RPC) używanego w przetwarzaniu rozproszonym.

PROCESY TRYBU UŻYTKOWNIKA

Windows tworzy zaplecze czterech podstawowych rodzajów procesów działających w trybie użytkownika:

1. **Specjalne procesy systemowe.** Usługi trybu użytkownika potrzebne do zarządzania systemem, takie jak zarządca sesji, podsystem uwierzytelniania, zarządca usług i proces rozpoczęcia sesji (logowania, *logon*).
2. **Procesy usługowe.** Spooler drukarki, rejestrator zdarzeń, składowe trybu użytkownika współpracujące z modułami obsługi urządzeń, różnorodne usługi sieciowe i wiele innych. Z usług korzystają zarówno twórcy oprogramowania w Microsoftzie, jak i konstruktorzy zewnętrzni, aby rozszerzać funkcjonalność systemu, jako że są one w systemie Windows jedynym sposobem wykonywania drugoplanowych działań w trybie użytkownika.
3. **Podsystemy środowiskowe.** Udostępniają różne osobiste prezencje (środowiska) SO. Dostępne są podsystemy Win32 i POSIX. Każdy podsystem środowiskowy zawiera proces podsystemu dzielony między wszystkie aplikacje korzystające z danego podsystemu i dynamicznie konsolidowane biblioteki (DLL), które zamieniają wywołania aplikacji użytkownika na wywołania ALPC procesu podsystemu i (lub) rodzime wywołania Windows.
4. **Aplikacje użytkownika.** Pliki wykonywalne (EXE-ki) i DLL-e reprezentujące działania wykonywane przez użytkowników w celu zastosowania systemu. Pliki EXE i DLL są na ogół nakierowane na konkretny podsystem środowiskowy, aczkolwiek niektóre programy dostarczane jako część SO używają rodzimych interfejsów systemu (NT API). Istnieje też zaplecze wykonywania programów 32-bitowych w systemach 64-bitowych.

Strukturalizacja systemu Windows ma umożliwiać użytkowanie aplikacji pisanych dla wielu prezencji („osobowości”, ang. *personalities*) SO. Windows tworzy te możliwości za pomocą wspólnego zbioru komponentów trybu jądra, które leżą u podłoża podsystemów środowiskowych SO. Implementacja każdego podsystemu środowiskowego zawiera osobny proces, który mieści dzielone struktury danych, przywileje i uchwytty do obiektów egzekutora potrzebne do realizacji danej prezencji. Proces ów jest rozpoczynany przez zarządcę sesji Windows, kiedy rusza pierwsza aplikacja danego typu. Proces podsystemu działa jako użytkownik systemowy, więc egzekutor będzie chronił jego przestrzeń adresową przed procesami wykonywanymi przez zwykłych użytkowników.

Podsystem środowiskowy realizuje graficzny interfejs użytkownika lub interfejs poleceń pisanych, dający użytkownikowi poczucie określonego SO. Ponadto każdy podsystem dostarcza API konkretnego środowiska. Oznacza to, że aplikacje utworzone pod kątem konkretnego środowiska operacyjnego wystarczy ponownie skompilować, aby mogły działać pod systemem Windows. Ponieważ interfejs systemu operacyjnego oglądany przez aplikacje jest taki sam jak ten, dla którego zostały napisane, nie trzeba modyfikować kodu źródłowego.

Model klient-serwer

Usługi SO Windows, podsystemy środowiskowe i aplikacje tworzą strukturę z wykorzystaniem modelu obliczeniowego klient-serwer, popularnego w obliczeniach rozproszonych — będzie on omówiony w części VI książki. Tę samą architekturę można zaadaptować na wewnętrzny użytek jednego systemu, jak w przypadku SO Windows.

Rodzimy interfejs NT API jest zbiorem usług opartych na jądrze, które realizują podstawowe abstrakcje używane przez system: procesy, wątki, pamięć wirtualną, wejście-wyjście i komunikację. Dzięki zastosowaniu modelu klient-serwer do realizacji poszczególnych funkcji w procesach trybu użytkownika Windows dostarcza znacznie bogatszego zbioru usług. Zarówno podsystemy środowiskowe, jak i usługi Windows w trybie użytkownika są implementowane w postaci procesów komunikujących się z klientami za pośrednictwem RPC (zdalnych wywołań procedur). Każdy proces serwera czeka na zamówienie przez klienta któregoś z jego usług (np. usług dotyczących pamięci, tworzenia procesów lub sieci). Klient, który może być programem użytkowym lub programem innego serwera, zamawia usługę, wysyłając komunikat. Komunikat jest kierowany przez egzekutora do odpowiedniego serwera. Serwer wykonuje żadaną operację i zwraca wyniki lub informację o stanie za pomocą innego komunikatu, który — kierowany przez egzekutora — dociera do klienta.

Do zalet architektury klient-serwer należą:

- **Uproszczenie egzekutora.** Jest możliwe skonstruowanie wielu interfejsów API zrealizowanych w serwerach trybu użytkownika bez konfliktów lub podwojeń w egzekutorze. Nowe API można łatwo dodawać.
- **Poprawianie niezawodności.** Każdy nowy serwer działa na zewnątrz jądra, z własną częścią pamięci, chroniony przed innymi serwerami. Pojedynczy serwer może ulec awarii bez załamывania lub naruszania reszty SO.
- **Dostarczanie jednolitych środków komunikowania się aplikacji z usługami za pośrednictwem RPC bez ograniczania elastyczności.** Proces przekazywania komunikatów jest ukryty przed aplikacjami klienta za pomocą namiastek funkcji, będących niewielkimi fragmentami kodu opakowującego wywołanie RPC. Gdy aplikacja wywołuje za pomocą API podsystem środowiskowy lub usługę, namiastka w aplikacji klienta pakuje parametry wywołania i przesyła je w komunikacie do procesu serwera, który realizuje wywołanie.
- **Tworzenie bazy odpowiedniej do obliczeń rozproszonych.** W obliczeniach rozproszonych z reguły używa się modelu klient-serwer ze zdalnymi wywołaniami procedur implementowanymi za pomocą rozproszonych modułów klientów i serwerów i wymiany komunikatów między klientami i serwerami. W systemie Windows lokalny serwer może przekazywać zdalnemu serwerowi komunikaty z zamówieniami przetwarzania pochodzącymi od lokalnych aplikacji klientów. Klienci nie muszą wiedzieć, czy zamówienie jest obsługiwane lokalnie, czy zdalnie. W rzeczywistości to, czy zamówienie jest obsługiwane lokalnie, czy zdalnie, może się zmieniać dynamicznie, wedle bieżących obciążeń i stosownie do dynamicznych zmian konfiguracji.

Wątki i SMP

Ważną cechą systemu Windows jest tworzenie zaplecza wątkowości i wieloprzetwarzania symetrycznego (SMP); obie techniki przedstawiliśmy w podrozdziale 2.4. W pracy [RUSS11] wymieniono następujące cechy systemu Windows umożliwiające realizację wątków i SMP:

- Podprogramy SO mogą działać na dowolnym z dostępnych procesorów, a różne podprogramy mogą być wykonywane jednocześnie na różnych procesorach.
- Windows umożliwia użycie wielu wątków wykonania w jednym procesie, przy czym poszczególne wątki mogą działać jednocześnie na różnych procesorach.

- Procesy serwera mogą stosować wiele wątków do jednoczesnego przetwarzania zamówień wielu klientów.
- Windows udostępnia mechanizmy dzielenia danych i zasobów między procesami oraz elastyczne możliwości komunikacji międzyprocesowej.

Obiekty systemu Windows

Chociaż rdzeń Windows jest napisany w języku C, przyjęte zasady projektowe są mocno osadzone w paradygmacie projektowania obiektowego. Takie podejście ułatwia współużytkowanie zasobów i danych przez procesy oraz ochronę zasobów przed nieuprawnionym dostępem. Do najważniejszych koncepcji obiektowych zastosowanych w systemie Windows należą:

- **Obudowywanie** (ang. *encapsulation*). Obiekt składa się z jednej lub więcej jednostek danych, zwanych *atrybutami*, i jednej lub więcej procedur, które mogą być wykonywane na tych danych, zwanych *usługami*. Jedynym sposobem dostępu do danych obiektu jest wywołanie którejś z jego usług. Dzięki temu dane w obiekcie można łatwo chronić przed nieupoważnionym lub niewłaściwym użyciem (np. przed próbą wykonania niewykonywalnego fragmentu danych).
- **Klasa obiektów i konkretny obiekt** (ang. *object class and instance*). Klasa obiektów jest szablonem (ang. *template*)¹⁶, w którym wymieniono atrybuty i usługi obiektu i zdefiniowano pewne cechy charakterystyczne obiektu. System operacyjny może tworzyć konkretne obiekty („instancje”) danej klasy stosownie do potrzeb. Mamy na przykład jedną klasę obiektu procesu i po jednym jej obiekcie dla każdego aktualnie czynnego procesu. To podejście ułatwia tworzenie obiektów i zarządzanie nimi.
- **Dziedziczenie** (ang. *inheritance*). Choć implementacja jest kodowana ręcznie, egzekutor korzysta z dziedziczenia, aby rozszerzać klasy obiektów przez dodawanie nowych właściwości. Każda klasa egzekutora wywodzi się z klasy bazowej, która określa metody wirtualne tworzenia, nazywania, zabezpieczania i usuwania obiektów. Obiekty dyspozytora są obiektami egzekutora dziedziczącymi własności obiektu zdarzeń, mogą więc używać wspólnych metod synchronizacji. Inne swoiste typy obiektów, jak te z klasy urządzeń, umożliwiają klasom specyficznych urządzeń dziedziczenie z klasy bazowej i dokładanie dodatkowych danych i metod.
- **Wielopostaciowość** (ang. *polymorphism*). Windows utrzymuje typowy zbiór funkcji API do manipulowania obiektami dowolnego typu; jest to własność polimorfizmu, jak zdefiniowano w dodatku D. Jednak Windows nie jest w pełni polimorficzny, gdyż istnieje wiele interfejsów API specyficznych dla obiektu określonego typu.

Czytelnik niezaznajomiony z zasadami obiektowości powinien zajrzeć do dodatku D.

Nie wszystkie elementy Windows są obiektami. Obiekty są używane tam, gdzie przewiduje się korzystanie z danych w trybie użytkownika, lub wówczas, gdy dostęp do danych odbywa się wspólnie lub jest ograniczany. Wśród jednostek reprezentowanych w postaci obiektów znajdują się pliki, procesy, wątki, semaforey, czasomierze i okna graficzne. Windows tworzy i kontroluje wszystkie typy obiektów w sposób jednolity, przy udziale zarządcy obiektów. Zarządca obiektów odpowiada za tworzenie i likwidowanie obiektów na życzenie aplikacji oraz za udzielanie dostępu do usług i danych obiektu.

¹⁶ Użycie słowa szablon jest tu trochę niefortunne, zważywszy na jego znaczenie w rozumieniu klasy parametrycznej w języku C++; klasa jest matrycą obiektu, przepisem, według którego obiekt jest konkretyzowany — *przypr. tłum.*

Każdy obiekt w egzekutorze, niekiedy nazywany obiektem jądra (w celu odróżnienia od obiektów poziomu użytkownika, niemających związku z egzekutorem), istnieje w postaci bloku pamięci przydzielonego przez jądro i jest bezpośrednio osiągalny tylko przez komponenty trybu jądra. Niektóre elementy tej struktury danych są wspólne dla obiektów wszystkich typów (np.: nazwa obiektu, parametry bezpieczeństwa, konto użycia), inne są swoiste dla konkretnego typu obiektu (np. priorytet obiektu wątku). Ponieważ struktury danych tych obiektów znajdują się w części przestrzeni adresowej każdego procesu dostępnej tylko dla jądra, aplikacja nie może się do nich odnieść bezpośrednio, aby je czytać lub zapisywać. Zamiast tego aplikacje manipulują obiektami pośrednio, za pomocą zbioru funkcji działań na obiekcie udostępnianych przez egzekutora. Przy tworzeniu obiektu aplikacja zamawiająca tworzenie otrzymuje w odpowiedzi uchwyt do obiektu. Najkrócej mówiąc, **uchwyt** (ang. *handle*) określa pozycję w procesowej tablicy egzekutora zawierającą wskaźnik do potrzebnego obiektu. Ten uchwyt może być potem używany przez dowolny wątek tego samego procesu do wywoływania funkcji Win32 działających na obiektach lub jego kopię można przekazać do innych procesów.

Z obiektami mogą być kojarzone informacje dotyczące ich bezpieczeństwa w postaci **deskryptora bezpieczeństwa** (ang. *security descriptor* — SD). Z tych informacji można korzystać, aby ograniczać dostęp do obiektu oparty na **dowodzie** (ang. *token*) obiektu opisującym konkretnego użytkownika. Na przykład proces mógłby utworzyć nazwany obiekt semafora z zamiarem, aby tylko pewni użytkownicy mogli ten semafor otwierać i z niego korzystać. Deskryptor SD obiektu takiego semafora może wymieniać użytkowników, którym wolno (lub jest zabronione) sięgać po obiekt semafora, wraz z wykazem dozwolonych praw dostępu (czytanie, pisanie, zmiana itp.).

Obiekty w systemie Windows mogą być nazwane lub nienazwane. Jeśli proces tworzy obiekt nienazwany, zarządca obiektu zwraca uchwyt do tego obiektu i ten uchwyt jest jedynym sposobem odwołania się do niego. Uchwyty mogą być dziedziczone przez procesy potomne lub poddawane między procesami. Obiekty nazwane mają również nadawaną nazwę, za pomocą której inne, postronne procesy mogą uzyskać uchwyt do obiektu. Gdyby na przykład proces A życzył sobie synchronizacji z procesem B, mógłby utworzyć nazwany obiekt zdarzenia i przekazać nazwę zdarzenia do B. Proces B mógłby wówczas otworzyć ten obiekt zdarzenia i go używać. Gdyby jednak proces A chciał po prostu posłużyć się zdarzeniem do synchronizacji dwu swoich wątków, mógłby utworzyć nienazwany obiekt zdarzenia, ponieważ nie ma potrzeby, aby inne procesy mogły z tego zdarzenia robić użytek.

Istnieją dwie kategorie obiektów stosowanych przez Windows do synchronizacji użycia procesora:

- **Obiekty dyspozytora** (ang. *dispatcher objects*). Podzbiór obiektów egzekutora, których wątki mogą czekać, aby posterować ekspediowaniem i synchronizacją operacji opartych na systemie wątków. Zostanie to opisane w rozdziale 6.
- **Obiekty sterujące** (ang. *control objects*). Używane przez komponent jądra do zarządzania działaniem procesora w obszarach niepodlegających normalnemu planowaniu wątków. W tabeli 2.5 ujęto obiekty sterujące jądrem.

Windows nie jest w pełni obiekowym systemem operacyjnym. Nie jest zrealizowany w języku obiekowym. Struktury danych pozostające w całości w jednej składowej egzekutora nie są reprezentowane jako obiekty. Niemniej Windows stanowi ilustrację potencjału technologii obiektowej i reprezentuje wzmożony trend do stosowania tej technologii w projektowaniu SO.

Tabela 2.5. Obiekty sterujące jądra Windows

Asynchroniczne wywołanie procedury	Używane do wkraczania w wykonywanie określonego wątku i powodowania wywołania procedury w określonym trybie procesora
Odroczone wywołanie procedury	Używane do odwlekania przetworzenia przerwania w celu uniknięcia opóźniania przerwzeń sprzętowych. Stosowane również do realizacji czasomierzy i do komunikacji międzyprocesorowej
Przerwanie	Używane do połączenia źródła przerwania z procedurą obsługi przerwania za pomocą wpisu w tablicy rozdzielczej przerwzeń (ang. <i>interrupt dispatch table</i> — IDT). Każdy procesor ma IDT, która jest używana do ekspediowania występujących w nim przerwzeń
Proces	Reprezentuje przestrzeń adresów wirtualnych i informacje kontrolne niezbędne do wykonywania zbioru obiektów wątków. Proces zawiera wskaźnik do mapy adresów, listę wątków gotowych do działania zawierającą obiekty wątków, listę wątków należących do procesu, całkowity nagromadzony czas wszystkich wątków działających w procesie i priorytet bazowy
Wątek	Reprezentuje obiekty wątków, zawiera priorytet planowania i przyznany kwant czasu oraz wykaz procesorów, na których dany wątek może działać
Profil	Używany do mierzenia dystrybucji czasu wykonywania w obrębie bloku kodu. Profilować można zarówno kod użytkownika, jak i kod systemu

2.8. TRADYCYJNE SYSTEMY UNIKSOWE

Historia

UNIX powstał w Bell Labs i wszedł do eksploatacji na maszynie PDP-7 w 1970 roku. Prace nad UNIX-em w Bell Labs i później, w wielu różnych ośrodkach, doprowadziły do wytworzenia całej serii odmian UNIX-a. Pierwszym godnym odnotowania kamieniem milowym okazało się przeniesienie systemu UNIX z maszyny PDP-7 na PDP-11. To była pierwsza oznaka, że UNIX mógłby stać się systemem dla wszystkich komputerów¹⁷. Następnym ważnym wydarzeniem było przepisanie UNIX-a w języku programowania C. W owym czasie była to nieznana nikomu strategia. W powszechnym odczuciu rzecz tak złożona jak system operacyjny, mający do czynienia ze zdarzeniami, w których czas ma znaczenie przesądzające, powinna być pisana wyłącznie w assemblerze. Sąd ten uzasadniano następująco:

- Pamięć (zarówno RAM, jak i pomocnicza) była mała i kosztowna jak na dzisiejsze standardy, więc efektywne jej użycie miało pierwszorzędne znaczenie. Uwzględniano przy tym rozmaite techniki pamięci nakładkowej, zawierającej różny kod i różne segmenty danych, w tym również kod dokonujący samomodyfikacji.
- Mimo że kompilatory były w użyciu od lat 50 XX wieku, przemysł komputerowy odnosił się na ogół z rezerwą do jakości kodu generowanego automatycznie. Mała pojemność zasobów sprawiała, że wydajny kod, zarówno pod względem czasu, jak i przestrzeni, był nieodzowny.

¹⁷ Maszyna PDP-11 miała odpowiednik produkowany w ZSRR pod nazwą SM-4. Również na nią dotarła taśma z UNIX-em, co ciekawe, jeszcze w 1981 roku, oficjalną przesyłką — *przypr. tłum.*

- Szybkości procesora i szyny były stosunkowo małe, więc oszczędzanie cykli zegarowych mogło mieć istotny wpływ na czas działania.

Realizacja w C wykazała przewagę zastosowania języka wysokiego poziomu w odniesieniu do większości (jeśli nie do całego) kodu systemu. Dziś w zasadzie wszystkie implementacje UNIX-a są pisane w języku C.

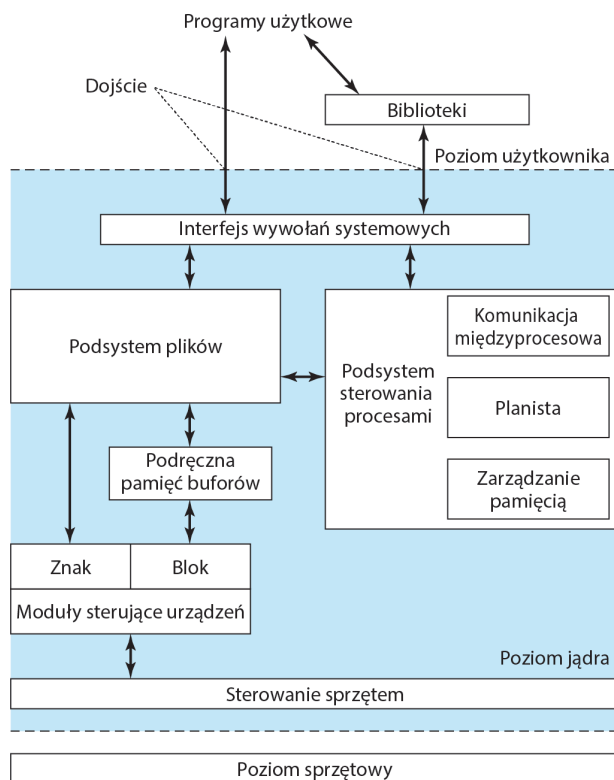
Owe wczesne wersje UNIX-a były popularne wewnątrz Bell Labs. W 1974 roku system UNIX został po raz pierwszy opisany w czasopiśmie technicznym [RITC74]. Wzbudziło to wielkie zainteresowanie tym systemem. Licencje na UNIX-a były udzielane instytucjom komercyjnym, a także uniwersytetom. Pierwszą wersją szeroko upowszechnioną poza Bell Labs była Version 6 z 1976 roku. Następna, Version 7, wydana w roku 1978, jest protoplastą większości współczesnych systemów uniksowych. Najważniejszy z systemów opracowanych poza AT&T¹⁸ był wykonany na Uniwersytecie Kalifornijskim w Berkeley. Opatrzony nazwą UNIX BSD (Berkeley Software Distribution), pracował najpierw na PDP, a potem na komputerach VAX. W AT&T dalej rozwijano i ulepszano system. W 1982 roku w Bell Labs połączono kilka wariantów UNIX-a AT&T w jeden system, któremu nadano handlowe oznaczenie UNIX System III. Później dodano do tego systemu kolejne właściwości i tak powstał UNIX System V.

Opis

Typową architekturę UNIX-a można zobrazować w trzech poziomach: sprzętu, jądra i użytkownika. SO jest często nazywany jądrem systemu, lub po prostu jądrem, aby zaakcentować jego odizolowanie od użytkownika i aplikacji. Jądro współpracuje bezpośrednio ze sprzętem. To właśnie na jądrze UNIX-a będziemy się koncentrować w naszych odwołaniach do UNIX-a jako do przykładowego systemu w tej książce. UNIX bywa również wyposażony w sporo usług i interfejsów użytkownika, uważanych za część systemu. Można je pogrupować w **powłokę** (ang. *shell*) dostarczającą wywołań systemowych z aplikacji, oprogramowanie innych interfejsów i składowe kompilatora C (kompilator, assembler, ładowacz). Poziom ponad tym stanowią aplikacje użytkownika i interfejs użytkownika do kompilatora C.

Zarys jądra przedstawiono na rysunku 2.15. Programy użytkownika mogą wywoływać usługi SO bezpośrednio lub za pomocą programów systemowych. Interfejs wywołań systemowych graniczy z użytkownikiem i umożliwia oprogramowaniu wyższego poziomu uzyskiwanie dostępu do poszczególnych funkcji jądra. Na drugim końcu SO zawiera elementarne procedury, które współpracują bezpośrednio ze sprzętem. Między tymi dwoma interfejsami system dzieli się na dwie główne części: jedna dotyczy sterowania procesami, druga zarządzania plikami i wejściem-wyjściem. Podsystem sterowania procesami odpowiada za zarządzanie pamięcią, planowanie i ekspediovanie procesów oraz za synchronizację i komunikację międzyprocesową. System plików wymienia dane między pamięcią a urządzeniami zewnętrznymi za pomocą strumieni znaków lub bloków. Aby to osiągnąć, stosuje się rozmaite moduły sterujące urządzeń. W przypadku przesyłania blokowego znajduje zastosowanie dyskowa pamięć podręczna: między przestrzenią adresową użytkownika a urządzeniem zewnętrznym lokowany jest w pamięci głównej bufor systemowy.

¹⁸ Zakłady Bell Labs podlegały wówczas konsorcjum AT&T — *przyp. tłum.*

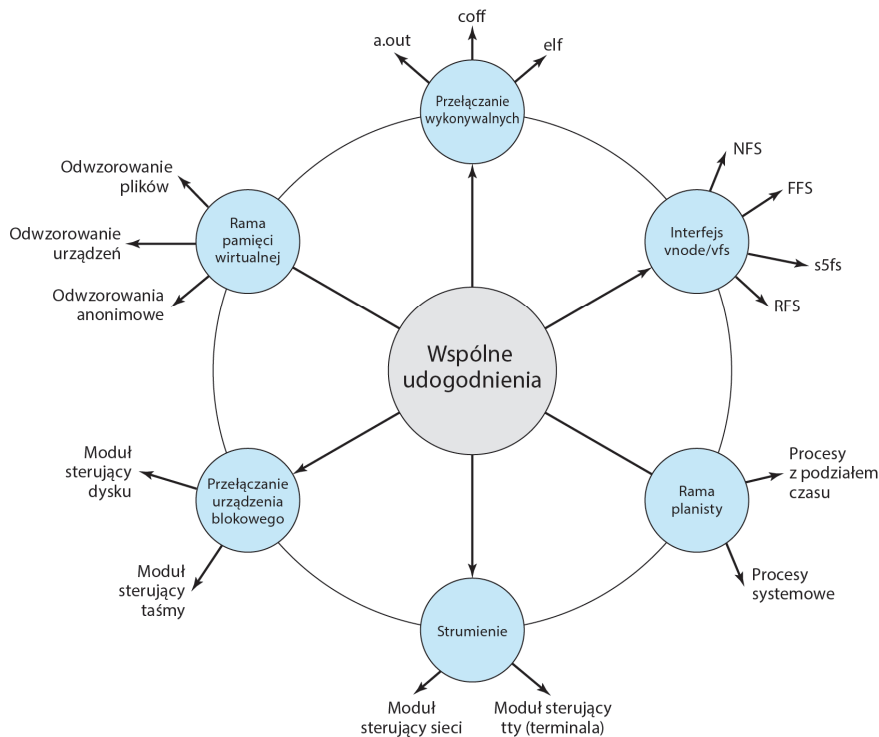


Rysunek 2.15. Architektura tradycyjnego systemu UNIX

To, co opisujemy w tym rozdziale, odnosi się do tak zwanych *tradycyjnych systemów unixowych*; w [VAHA96] użyto tego terminu w odniesieniu do wersji System V Release 3 (SVR3), 4.3BSD i wcześniejszych. O tradycyjnym (konwencjonalnym) systemie UNIX można wypowiedzieć następujące stwierdzenia. Jest zaprojektowany do działania na pojedynczym procesorze i brakuje mu zdolności chronienia swoich struktur danych przed współbieżnym dostępem wielu procesorów. Jego jądro nie jest zbyt uniwersalne, realizując jeden typ systemu plików, politykę planowania procesów i format pliku wykonywalnego. Tradycyjne jądro UNIX-a nie jest zaprojektowane z myślą o rozszerzaniu i ma niewiele udogodnień, jeśli chodzi o wtórne wykorzystanie kodu. W rezultacie w miarę dodawania nowych cech do różnych wersji UNIX-a trzeba było dodawać dużo nowego kodu, co doprowadziło do przerośniętego i niemodularnego jądra.

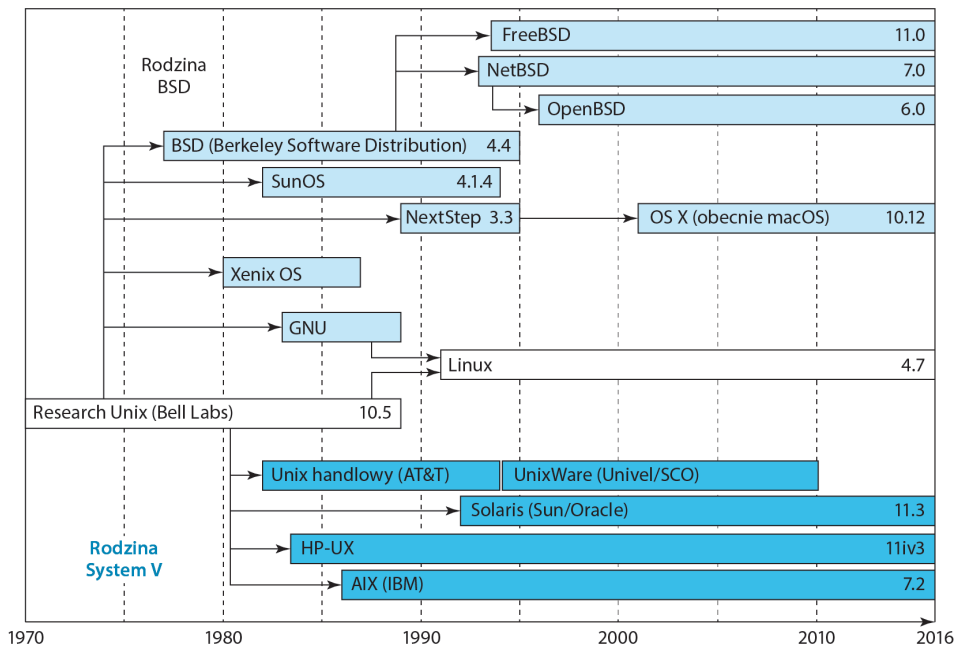
2.9. NOWOCZESNE SYSTEMY UNIKSOWE

W miarę ewolucji UNIX-a szybko mnożyły się liczne i różne jego implementacje, każda z pewnymi użytecznymi właściwościami. Pojawiła się potrzeba wytworzenia nowej implementacji, w której ujednolicono by wiele ważnych innowacji, dodano cechy projektowe innych nowoczesnych SO i utworzono bardziej modułarną architekturę. Typową architekturę współczesnego jądra UNIX-a przedstawia rysunek 2.16. Mamy tu niewiele rdzennych właściwości, zapisanych modułarnie, stanowiących funkcje i usługi wymagane przez pewną liczbę procesów SO. Każde z zewnętrznych kół reprezentuje funkcje i interfejs, który można zrealizować na różne sposoby.



Rysunek 2.16. Jądro nowoczesnego UNIX-a

Przejdziemy teraz do kilku przykładów nowoczesnych systemów uniksowych (rysunek 2.17).



Rysunek 2.17. Drzewo genealogiczne UNIX-a

System V Release 4 (SVR4)

SVR4 (z ang. System V, Wydanie 4), opracowany przez AT&T razem z Sun Microsystems, łączy cechy systemów SVR3, 4.3BSD, Microsoft Xenix System V i SunOS. Było to niemal całkowicie od nowa napisane jądro Systemu V, co doprowadziło do powstania czystej, mimo że złożonej, implementacji. Do nowych cech w tym wydaniu należą: przetwarzanie w czasie rzeczywistym, klasy planowania procesów, dynamiczny przydział struktur danych, zarządzanie pamięcią wirtualną, wirtualny system plików i wywłaszczalne jądro.

System SVR4 czerpie z wysiłków zarówno projektantów komercyjnych, jak i akademickich. Opracowano go w celu dostarczenia ujednoliconej platformy rozwoju handlowego UNIX-a. Pod tym względem odniósł sukces i jest być może najważniejszą odmianą UNIX-a. Łączy większość ważnych cech, jakie kiedykolwiek opracowano w którymkolwiek z systemów UNIX, przy czym robi to w sposób zintegrowany i opłacalny z handlowego punktu widzenia. SVR4 działa na szerokim asortymencie procesorów: od 32-bitowych aż po superkomputery.

System BSD

Ciąg wydań UNIX-a określanych mianem Berkeley Software Distribution (BSD) odegrał kluczową rolę w rozwoju teorii projektowania systemów operacyjnych. 4.xBSD jest szeroko stosowany w instalacjach akademickich i posłużył jako podstawa wielu handlowych produktów UNIX-a. Można zapewne bezpiecznie rzec, że UNIX swą popularność zawdzięcza po większej części systemowi BSD oraz że najwięcej ulepszeń w UNIX-ie ujrzało światło dzienne w wersjach BSD.

System 4.4BSD był ostateczną wersją BSD ukazującą się w Berkeley, której projekt i organizacja implementacji przeniknęły dalej. Jest on istotnym uaktualnieniem wersji 4.3BSD i zawiera nowy system pamięci wirtualnej, zmiany w strukturze jądra oraz długą listę innych ulepszeń.

Istnieje kilka szeroko rozpowszechnionych otwartych wersji BSD. FreeBSD jest popularny w serwerach internetowych i zaporach sieciowych, jest także używany w wielu systemach wbudowanych. NetBSD jest dostępny na wielu platformach, w tym w wielkoskalowych systemach serwerów, systemach biurkowych i urządzeniach ręcznych, znajduje też liczne zastosowania w systemach wbudowanych. OpenBSD jest otwartą wersją SO, w której specjalny nacisk położono na bezpieczeństwo.

Najnowsza wersja SO dla komputerów Macintosh, pierwotnie nazywana OS X, obecnie zaś macOS, opiera się na FreeBSD 5.0 i mikrojądrze systemu Mach 3.0.

Solaris 11

Solaris jest wydaniem UNIX-a pochodzącym z Oracle'a¹⁹ i opartym na SVR4, jego ostatnia wersja ma numer 11. Solaris udostępnia wszystkie własności systemu SVR4, ma także liczne bardziej zaawansowane właściwości: w pełni wywłaszczalne, wielowątkowe jądro, kompletne zaplecze przetwarzania symetrycznego (SMP) i obiektowy interfejs do systemu plików. Solaris jest jedną z najszerzej stosowanych i najbardziej udanych realizacji systemu UNIX.

¹⁹ Solaris jako taki pochodzi z Sun Microsystems, wchłoniętej później przez Oracle — *przyp. tłum.*

2.10. LINUX

Historia

Linux wystartował jako odmiana UNIX-a dla architektury IBM PC (Intel 80386). Pierwszą jego wersję napisał Linus Torvalds, fiński student informatyki. Torvalds zamieścił wczesną wersję Linuxa w Internecie w 1991 roku. Od tego czasu wiele osób, współpracując przez Internet, wniosło wkład w rozwój Linuxa, a wszystko to pod kontrolą Torvaldsa. Ponieważ Linux jest bezpłatny, a jego kod źródłowy dostępny, od początku istnienia system ten stał się alternatywną możliwością wobec innych uniksowych stacji roboczych, jak te oferowane przez Sun Microsystems i IBM. Dziś jest pełnowartościowym systemem uniksowym, działającym na niemal wszystkich platformach.

Kluczem do sukcesu Linuxa była dostępność bezpłatnych pakietów oprogramowania upowszechnianych pod auspicjami Free Software Foundation (FSF, z ang. Fundacja Wolnego Oprogramowania). Celem FSF jest stabilne oprogramowanie, niezależne od platformy, bezpłatne, wysokiej jakości i cenione przez społeczność użytkowników. Projekt GNU²⁰ dostarcza narzędzi twórcom oprogramowania, a licencja GNU GPL (ang. GNU Public Licence) jest udzielaną przez FSF gwarancją jakości. Torvalds, pracując nad jądrem, używał narzędzi GNU i opublikował je wtedy na licencji GPL. Tak więc dystrybucje Linuxa, które oglądacie dzisiaj, są rezultatem projektu GNU fundacji FSF, indywidualnych starań Torvaldsa i wysiłków wielu uczestników tego przedsięwzięcia na całym świecie.

Prócz wykorzystania przez wielu indywidualnych twórców Linux znacząco przeniknął obecnie do świata korporacji. Jest to efektem nie tylko bezpłatności oprogramowania, lecz również jakości jądra Linuxa. Wielu utalentowanych budowniczych dołożyło się do bieżącej wersji, tworząc produkt robiący wrażenie swoją fachowością. Co więcej, Linux jest wysoce modularny i łatwy do konfigurowania. To ułatwia „wyciskanie” optymalnego działania z rozmaitych platform sprzętowych. Zważywszy jeszcze na dostępność kodu źródłowego, dostawcy mogą podrasowywać aplikacje i narzędzia tak, aby spełniały specyficzne wymagania. Są też przedsiębiorstwa, jak Red Hat lub Canonical, które na długi czas zapewniają wysoce profesjonalne i niezawodne wsparcie swoich opartych na Linuxie dystrybucji. W różnych miejscach książki przedstawimy szczegóły wewnętrznych rozwiązań jądra Linuxa, opierając się na jądrze 4.7 upublicznonym w 2016 roku.

System operacyjny Linux dużą część sukcesu zawdzięcza swojemu modelowi rozwojowemu. Nowo wnoszony kod trafia na jedną główną listę pocztową, zwaną LKML (ang. *Linux Kernel Mailing List*). Prócz tego jest wiele innych list pocztowych przeznaczonych poszczególnym podsystemom jądra Linuxa (jak *netdev* — lista pocztowa dotycząca pracy sieciowej, *linux-pci* — podsystem PCI, *linux-acpi* — podsystem ACPI i wiele, wiele innych). Łaty wysyłane na te listy powinny spełniać ściśle reguły (przede wszystkim konwencje stylistyczne kodowania jądra Linuxa); są one przeglądane przez fachowców z całego świata, subskrybujących wiadomości z danych list. Pod adresem tych list łaty może wysłać każdy; statystyki (np. publikowane od czasu do czasu w witrynie *lwn.net*) wykazują, że wiele łat jest przesyłanych przez deweloperów ze znanych firm handlowych, między innymi: Intel, Red Hat, Google i Samsung. Wielu opiekunów (pielęgnatorów, ang. *maintener*) również pochodzi z firm komercyjnych (np. David Miller, konserwator sieci zatrudniony w Red Hat).

²⁰ GNU jest rekurencyjnym akronimem pochodzącym od *GNU's Not Unix* (z ang. GNU nie jest Unixem). Projekt GNU jest zbiorem będących w wolnym obiegu pakietów oprogramowania i narzędzi służących do budowy uniksopodobnego systemu operacyjnego; jest często używany w powiązaniu z jądrem Linuxa.

Niejednokrotnie łaty takie są poprawiane dzięki wzajemnym kontaktom i dyskusjom na listach pocztowych, po czym są wysyłane ponownie i po raz kolejny sprawdzane. Dany opiekun decyduje w końcu, czy przyjąć łaty, czy je odrzucić, a ponadto każdy opiekun podsystemu od czasu do czasu wysyła zamówienie ściągnięcia swojego drzewa do głównego drzewa jądra, utrzymywanego przez Linusa Torvaldsa. Sam Linus publikuje nową wersję jądra co około 7 – 10 tygodni, a każde takie wydanie ma około 5 – 8 wersji wydań kandydujących (ang. *Release Candidate* — RC)²¹.

Przy okazji warto się zastanowić, dlaczego inne systemy operacyjne o otwartym kodzie, jak różne odmiany BSD lub OpenSolaris, nie osiągnęły sukcesu i popularności, które stały się udziałem Linuxa. Powodów może być wiele, natomiast ponad wszelką wątpliwość przyczyniła się do tego otwartość modelu rozwojowego Linuxa. Jednak ten temat wybiega poza zakres naszej książki.

Struktura modularna

Większość jąder UNIX-a jest monolityczna. Przypomnijmy, co zostało już powiedziane w tym rozdziale, że jądro monolityczne zawiera w zasadzie całą funkcjonalność SO w jednym wielkim bloku kodu, działającym jako jeden proces w jednej przestrzeni adresowej. Wszystkie składowe funkcjonalne jądra mają dostęp do wszystkich jego wewnętrznych struktur danych i podprogramów. Jeśli zmiana następuje w jednym fragmencie typowego monolitycznego SO, to wszystkie moduły i podprogramy muszą być na nowo skonsolidowane, ponownie zainstalowane, po czym system musi zostać uruchomiony od początku. Dopiero wtedy skutki zmiany staną się widoczne. W rezultacie każda modyfikacja, jak dodanie modułu obsługi nowego urządzenia lub funkcji systemu plików, staje się trudna. Ten problem jest szczególnie dokuczliwy w Linuxie, którego rozwój ma skalę globalną i przebiega z udziałem luźno powiązanych grup niezależnych twórców.

Choć w Linuxie nie przyjęto koncepcji mikrojądra, osiągnięto w nim wiele potencjalnych zalet tego podejścia, dzięki jego drobiazgowo modularnej architekturze. Linux jest ustrukturyzowany w postaci kolekcji modułów, z których część może być automatycznie ładowana i rozładowywana na żądanie. Te względnie niezależne bloki są określane jako **moduły ładowalne** (ang. *loadable modules*) [GOYE99]. Zasadniczo moduł jest plikiem wynikowym, którego kod można skonsolidować z jądrem lub od niego odłączyć w trakcie działania systemu. Moduł zazwyczaj implementuje pewną specyficzną funkcję: system plików, oprogramowanie obsługi urządzenia lub jakąś inną własność górnej warstwy jądra. Moduł nie działa na zasadzie samoistnego procesu lub wątku, chociaż stosownie do potrzeb może tworzyć w różnych celach wątki jądrowe. Zamiast tego jest wykonywany w trybie jądra na zlecenie bieżącego procesu.

W ten sposób, choć Linux może być traktowany jako monolityczny, jego modularna struktura pokonuje pewne trudności związane z rozwijaniem i ewoluowaniem jądra. Ładowalne moduły jądra wyróżniają dwie ważne cechy:

1. **Dynamiczna konsolidacja** (łączenie, ang. *linking*). Moduł jądra można załadować i skonsolidować z jądrem, kiedy jądro jest już w pamięci i działa. Moduł można też od jądra odłączyć i pozbyć się go z pamięci w dowolnej chwili.
2. **Stosowość modułów** (ang. *stackable modules*). Moduły są zorganizowane w hierarchię. Poszczególne moduły służą jako biblioteki, jeśli odwołują się do nich moduły klienta położone wyżej w hierarchii, a gdy odwołują się do modułów położonych niżej, są traktowane jak klienci.

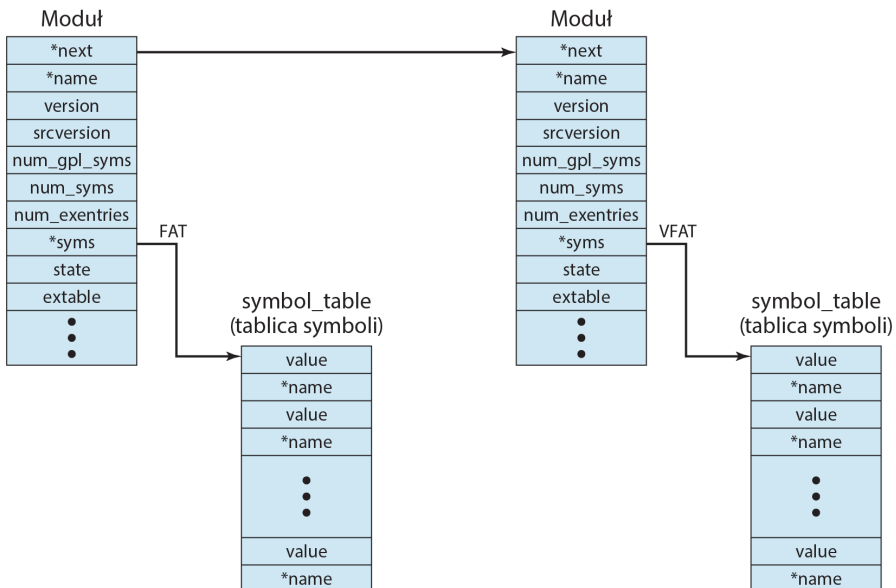
²¹ Dodajmy jeszcze, że wszystko to przebiega w rozproszonym systemie kontroli wersji Git, którego autorem jest również Linus Torvalds; por. J. Loeliger, M. McCullough, *Kontrola wersji z systemem Git. Narzędzia i techniki programistów*, Helion 2014 — przyp. tłum.

Dynamiczna konsolidacja ułatwia konfigurowanie i oszczędza pamięć jądra [FRAN97]. W Linuxie program użytkownika lub użytkownik może jawnie ładować i rozładowywać moduły jądra, korzystając z poleceń `insmod` lub `modprobe` i `rmmod`. Jądro samo monitoruje zapotrzebowanie na poszczególne funkcje i może wprowadzać i usuwać moduły zależnie od potrzeb. Dzięki stosowości modułów można definiować zależności między modułami. Ma to dwie zalety:

1. Kod wspólny dla zbioru podobnych modułów (np. obsługujących podobny sprzęt) można przenieść do odrębnego modułu, zmniejszając zwielokrotnienie.
2. Jądro może zapewnić obecność niezbędnych modułów, powstrzymując się od rozładowania modułu, od którego zależą inne moduły, i ładując wszelkie dodatkowo wymagane moduły podczas ładowania nowego modułu.

Na rysunku 2.18 podano przykład ilustrujący struktury używane przez Linuxa do zarządzania modułami. Rysunek ukazuje listę modułów jądra po załadowaniu tylko dwóch modułów: FAT i VFAT. Każdy moduł jest zdefiniowany przez dwie tablice: tablicę modułu i tablicę symboli²². Tablica modułu zawiera następujące elementy:

- ***name** — nazwa modułu.
- **refcnt** — licznik modułu; licznik jest zwiększany, gdy rozpoczyna się operacja zawierająca funkcje modułu, i zmniejszany, gdy ta operacja się kończy.
- **num_syms** — liczba eksportowanych symboli.
- ***syms** — wskaźnik do tablicy symboli modułu.



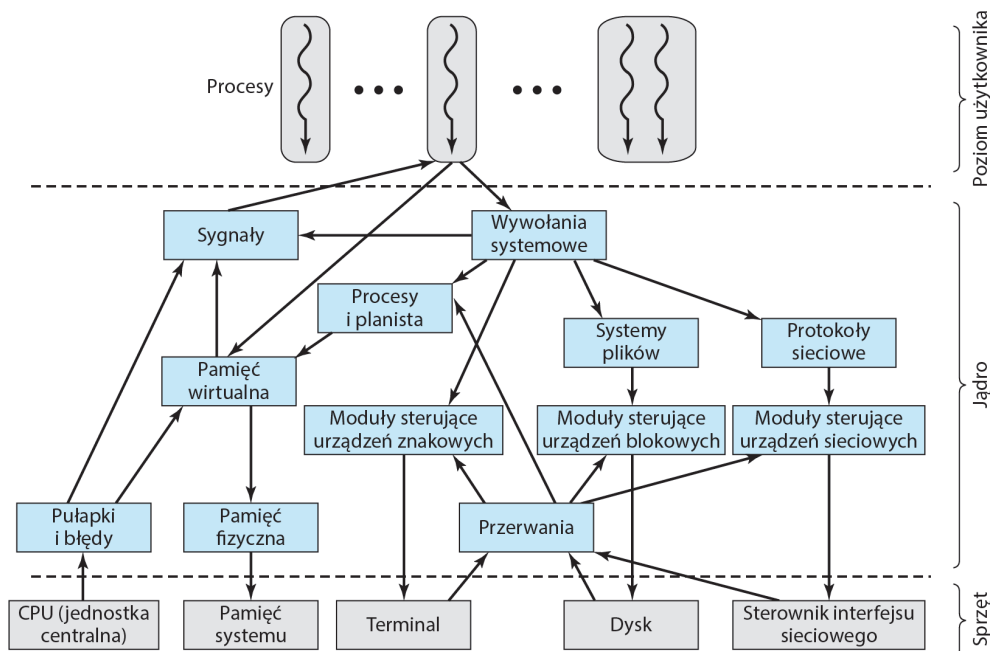
Rysunek 2.18. Przykładowa lista modułów jądra Linuxa

Tablica symboli zawiera wykaz symboli zdefiniowanych w danym module i używanych gdziekolwiek.

²² Do czego dochodzi jeszcze tablica symboli jądra (ang. *kernel symbol table*) rejestrująca wszystkie symbole eksportowane przez ładowany moduł — *przyp. tłum.*

Składowe jądra

Na rysunku 2.19, zaczerpniętym z [MOSB02], przedstawiono główne komponenty typowej implementacji jądra Linuxa. Pokazano na nim kilka procesów działających w szczytowej partii jądra. Każda ramka symbolizuje osobny proces, natomiast linia falista ze strzałką reprezentuje wykonanie wątku. Jądro składa się ze zbioru współpracujących składowych (strzałkami zaznaczono główne interakcje). Usytuowany poniżej sprzęt jest również wyobrażony w postaci zbioru komponentów ze strzałkami wskazującymi, które z nich używają innych lub nimi sterują. Wszystkie komponenty jądra działają oczywiście na procesorze. Tych związków dla prostoty nie pokazano.



Rysunek 2.19. Składowe jądra Linuxa

W skrócie — do podstawowych składowych jądra należą:

- **Sygnały.** Jądro używa sygnałów do wkraczania do procesu. Sygnały są w szczególności wykorzystywane do powiadamiania procesu o pewnych błędach, takich jak dzielenie przez zero. W tabeli 2.6 podano kilka przykładów sygnałów.
- **Wywołania systemowe.** Wywołanie systemowe służy procesowi do zamawiania w jądrze jakiejś usługi. Wywołań systemowych jest kilkaset, można je z grubsza pogrupować w sześć kategorii: dotyczące systemu plików, procesów, planowania, komunikacji międzyprocesowej, gniazd (praca w sieci) i inne. W tabeli 2.7 zdefiniowano po kilka przykładów z każdej kategorii.
- **Procesy i planista.** Tworzy procesy oraz zajmuje się ich planowaniem i zarządzaniem nimi.
- **Pamięć wirtualna.** Przydziela pamięć wirtualną na użytek procesów i zarządza nią.
- **Systemy plików.** Dostarczają globalnej hierarchicznej przestrzeni nazw plików, katalogów i innych obiektów związanych z plikami oraz udostępniają funkcje systemu plików.
- **Protokoły sieciowe.** Realizują interfejs gniazd dla użytkowników zestawu protokołów TCP/IP.

Tabela 2.6. Niektóre sygnały Linuxa

SIGHUP	Zawieszenie terminala	SIGCONT	Kontynuuj
SIGQUIT	Zaniechanie z klawiatury	SIGTSTP	Stop z klawiatury
SIGTRAP	Pułapka śladu	SIGTTOU	Pisanie na terminalu
SIGBUS	Błąd szyny	SIGXCPU	Przekroczony limit CPU
SIGKILL	Sygnał likwidacji	SIGVTALARM	Wirtualny alarm zegarowy
SIGSEGV	Naruszenie segmentacji	SIGWINCH	Rozmiar okna niezmieniony
SIGPIPT	Przerwany potok	SIGPWR	Awaria zasilania
SIGTERM	Zakończenie	SIGRTMIN	Pierwszy sygnał czasu rzeczywistego
SIGCHLD	Stan potomka niezmieniony	SIGRTMAX	Ostatni sygnał czasu rzeczywistego

Tabela 2.7. Niektóre wywołania systemowe Linuxa

Odnoszące się do systemu plików	
close	Zamknij deskryptor pliku
link	Utwórz nową nazwę pliku
open	Otwórz, a być może utwórz, plik lub urządzenie
read	Czytaj z deskryptora pliku
write	Pisz do deskryptora pliku
Odnoszące się do procesów	
execve	Wykonaj program
exit	Zakończ wywołujący proces
getpid	Pobierz identyfikator procesu
setuid	Ustaw tożsamość użytkownika bieżącego procesu
ptrace	Zapewnij możliwość obserwowania i kontrolowania przez proces rodzicielski wykonania innego procesu oraz sprawdzania i zmiany jego obrazu pamięci i rejestrów
Odnoszące się do planowania	
sched_getparam	Ustaw parametry planowania związane z zasadami planowania procesu identyfikowanego przez pid
sched_get_ ↳priority_max	Zwróć maksymalną wartość priorytetu, który może być użyty w algorytmie planowania identyfikowanym przez policy
sched_setscheduler	Ustaw zarówno politykę planowania (np. FIFO), jak i związane z nią parametry w odniesieniu do procesu pid
sched_rr_get_ ↳interval	Zapisz w strukturze timespec wskazanej przez parametr kwant czasu algorytmu rotacyjnego w odniesieniu do procesu pid
sched_yield	Za pośrednictwem tego wywołania systemowego proces może zwolnić procesor dobrowolnie, bez blokowania. Proces może potem być przeniesiony na koniec kolejki z uwagi na jego statyczny priorytet, a działanie rozpocznie nowy proces

Tabela 2.7. Niektóre wywołania systemowe Linuxa — ciąg dalszy

Odnoszące się do komunikacji międzyprocesowej (IPC)	
msgrev	Przydziel strukturę bufora komunikatu do odbioru komunikatu. Wywołanie systemowe czyta wówczas komunikat z kolejki komunikatów określonej przez msgid do nowo utworzonego bufora
semctl	Wykonaj operację sterującą określoną przez cmd na zbiorze semaforów semid
semop	Wykonaj operacje na wybranych elementach zbioru semaforów semid
shmat	Przydziel segment pamięci dzielonej identyfikowany przez semid do segmentu danych procesu wywołującego
shmctl	Zezwól użytkownikowi na odebranie informacji z segmentu pamięci dzielonej, ustaw właściciela, grupę i pozwolenia dotyczące segmentu pamięci dzielonej albo zlikwiduj segment
Odnoszące się do gniazd (pracy sieciowej)	
bind	Przydziel gniazdu lokalny adres IP i port. Zwróć 0 w przypadku powodzenia lub 1, gdy wystąpił błąd
connect	Nawiąż połączenie między danym gniazdem a gniazdem zdalnym skojarzonym z adresem gniazda (sockaddr)
gethostname	Zwróć lokalną nazwę komputera sieciowego (hosta)
send	Wyślij danym gniazdem bajty zawarte w buforze wskazywanym przez *msg
setsockopt	Ustaw opcje gniazda
Inne	
fsync	Skopiuj na dysk wszystkie części pliku przebywające w pamięci i czekaj, aż urządzenie zraportuje, że wszystkie części są w pamięci trwałej
time	Zwróć czas w sekundach odliczony od 1 stycznia 1970 roku
vhangup	Zasymuluj zahamowanie na aktualnym terminalu. To wywołanie pozwala innym użytkownikom mieć „czysty” terminal (tty) w czasie rozpoczynania sesji (logowania)

- **Moduły sterujące urządzeń znakowych.** Zarządzają urządzeniami, które wymagają od jądra przesyłania do lub z urządzenia po jednym bajcie jednorazowo, takimi jak terminale, modemy i drukarki.
- **Moduły sterujące urządzeń blokowych.** Zarządzają urządzeniami czytającymi i zapisującymi dane blokami; zaliczają się do nich różne odmiany pamięci pomocniczych (dyski magnetyczne, CD-ROM-y itd.).
- **Moduły sterujące urządzeń sieciowych.** Zarządzają kartami interfejsów sieciowych i portami komunikacyjnymi łączącymi z urządzeniami sieci, takimi jak mosty i rutery.
- **Pułapki i błędy.** Obsługują pułapki i błędy generowane przez procesor, takie jak błąd pamięci.
- **Pamięć fizyczna.** Zarządza pulą ramek stron w rzeczywistej pamięci i rozmieszcza strony pamięci wirtualnej.
- **Przerwania.** Ta składowa zajmuje się obsługą przerwania pochodzących od urządzeń zewnętrznych.

2.11. ANDROID

Oparty na Linuxie system operacyjny Android pierwotnie był zaprojektowany do telefonów przenośnych. Jest najpopularniejszym mobilnym SO, i to z dużą przewagą: sprzedaż telefonicznych zestawów Androida przewyższyła iPhone'y Apple'a w stosunku 4 do 1 [MORR16]. To tylko jeden aspekt rosnącej dominacji Androida. W coraz większym stopniu staje się on systemem operacyjnym prawie wszystkich urządzeń z układem komputera innych niż serwery i PC-ty. Android jest szeroko używanym SO w Internecie rzeczy.

Pierwsze prace nad SO Android wykonano w firmie Android Inc., kupionej w 2005 roku przez Google'a. Pierwsza komercyjna wersja, Android 1.0, ukazała się w 2008 roku. W czasie pisania tych słów najnowszą wersją jest Android 7.0 (Nougat). Android ma aktywną społeczność twórców i entuzjastów używających kodu źródłowego Android Open Source Project (AOSP) do opracowywania i rozpowszechniania własnych, zmodyfikowanych wersji systemu operacyjnego. Otwartość kodu źródłowego Androida stała się kluczem do jego popularności.

Architektura oprogramowania Androida

Android jest zdefiniowany w postaci stosu oprogramowania, które zawiera zmodyfikowaną wersję Linuxa, jądro, warstwę pośrednią i podstawowe aplikacje. Na rysunku 2.20 przedstawiono architekturę oprogramowania Androida nieco szczegółowiej. Na Androida należy zatem spoglądać jak na kompletny zestaw oprogramowania, nie tylko sam SO.

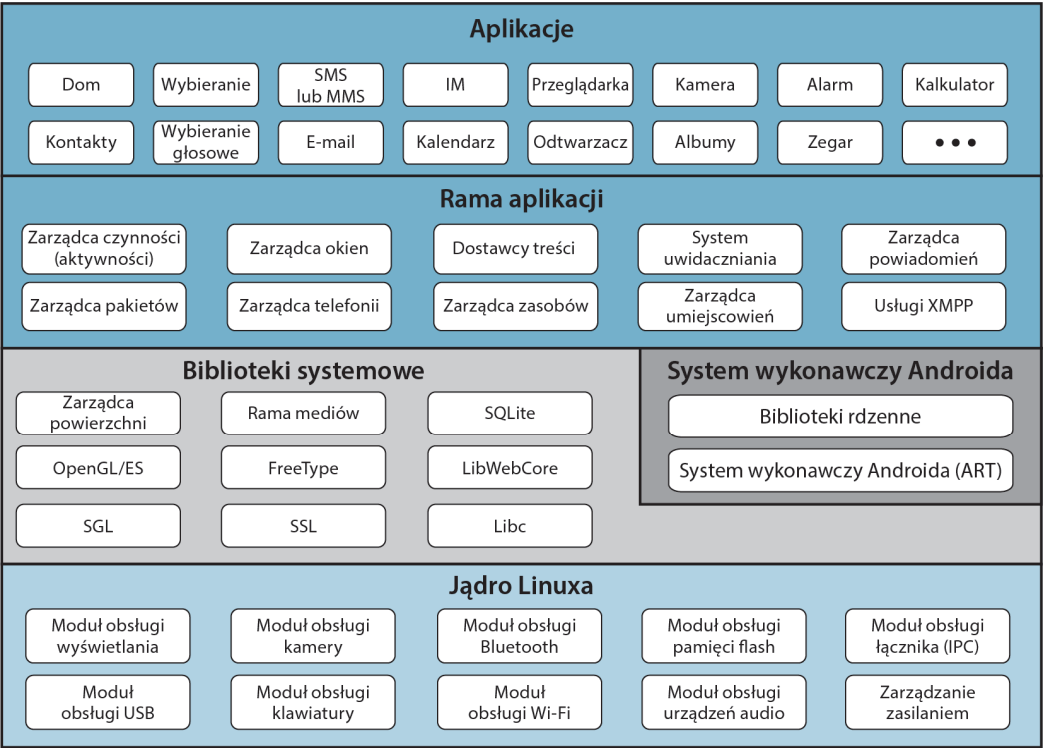
APLIKACJE

Wszystkie aplikacje, z którymi użytkownik ma bezpośredni kontakt, są częścią warstwy zastosowań. Mamy tutaj rdzenny zestaw aplikacji ogólnego przeznaczenia: klienta poczty elektronicznej, program SMS, kalendarz, mapy, przeglądarkę, kontakty i inne aplikacje będące standardowo na wyposażeniu każdego mobilnego telefonu. Aplikacje są z reguły zrealizowane w Javie. Zasadniczym celem otwartości architektury kodu źródłowego Androida jest ułatwienie budowniczym realizowania nowych aplikacji dla specyficznych urządzeń i specyficznych wymagań użytkownika docelowego. Zastosowanie Javy uwalnia budowniczych od zagadnień związanych ze sprzętem i jego osobliwościami oraz umożliwia szybkie robienie użytku z właściwościami zawartych w Javie, języku wyższego poziomu, takich jak predefiniowane klasy. Na rysunku 2.20 podano przykłady podstawowych rodzajów aplikacji występujących na platformie Android.

RAMA APLIKACJI

Warstwa **ramy aplikacji** (ang. *Application Framework*) dostarcza wysokopoziomowych bloków konstrukcyjnych dostępnych za pośrednictwem ustandaryzowanych API, z których osoby programujące tworzą nowe aplikacje. Architektura jest zaprojektowana tak, aby uprościć ponowne użycie komponentów. Oto niektóre podstawowe komponenty ramy aplikacji:

- **Zarządca czynności.** Zarządza cyklem eksploatacyjnym aplikacji. Odpowiada za uruchamianie, pauzowanie i wznowianie różnych aplikacji.
- **Zarządca okien.** Jest to abstrakcja Javy usytuowana niżej zarządcy powierzchni. **Zarządca powierzchni** (ang. *Surface Manager*) manipuluje interakcją bufora ramek i niskopoziomym rysowaniem, natomiast zarządca okien tworzy warstwę ponad nim, umożliwiając aplikacjom deklarowanie obszarów swoich klientów i używanie właściwości takich jak pasek stanu.



Implementacja:



Aplikacje, rama aplikacji: Java



Biblioteki systemowe, system wykonawczy Androida: C i C++



Jądro Linuxa: C

Rysunek 2.20. Architektura oprogramowania Androida

- **Zarządca pakietów.** Instaluje i usuwa aplikacje.
- **Zarządca telefonii.** Umożliwia interakcję z telefonem, usługi SMS-owe i MMS-owe.
- **Dostawcy treści.** Te funkcje obudowują dane aplikacji, które muszą być dzielone między aplikacjami, takie jak kontakty.
- **Zarządca zasobów.** Zarządza zasobami aplikacji, takimi jak zlokalizowane napisy (ang. *strings*) i bitmapy.
- **System uwidaczniania.** Udostępnia elementy interfejsu użytkownika (ang. *user interface* — UI) w rodzaju przycisków, wykazów, terminarzy i innych kontrolerek, a także zdarzenia UI (ang. *UI Events*), takie jak dotknięcia i gesty.
- **Zarządca umiejscowienia.** Umożliwia budowniczym wprowadzanie do komputera usług opartych na lokalizacji: czy to przez GPS, czy przez identyfikatory wież telefonów komórkowych, czy też za pomocą lokalnych baz danych Wi-Fi (rozpoznane hotspoty Wi-Fi i ich stan).

- **Zarządca powiadomień.** Zarządza zdarzeniami w rodzaju nadejścia komunikatów lub terminów.
- **XMPP.** Udostępnia standaryzowane funkcje przesyłania komunikatów (także czat) między aplikacjami.

BIBLIOTEKI SYSTEMOWE

Warstwa poniżej ramy aplikacji składa się z dwu części: bibliotek systemowych i **systemu wykonawczego Androida** (ang. *Android Runtime*). Składowa **biblioteki systemowe** (ang. *System Libraries*) jest zbiorem przydatnych funkcji systemowych napisanych w C lub C++ i używanych przez różne komponenty systemu Android. Są one wywoływane z ramy aplikacji i z samych aplikacji za pośrednictwem interfejsu Javy. Te właściwości są uwidaczniane przed budowniczymi w ramie aplikacji Androida. Do niektórych z podstawowych bibliotek systemowych należą:

- **Zarządca powierzchni** (ang. *Surface Manager*). Android używa kompozycyjnego zarządcy okien, podobnego do Visty lub Compiza, lecz znacznie prostszego. Zamiast rysowania bezpośrednio w buforze ekranu, polecenia rysowania trafiają do pozaekranowych bitmap, łączonych następnie z innymi bitmapami w celu utworzenia zawartości ekranu oglądanej przez użytkownika. Umożliwia to systemowi tworzenie najróżniejszych ciekawych efektów, jak widzenie poprzez okna i fantazyjne przejścia.
- **OpenGL.** OpenGL (*Open Graphics Library*, z ang. otwarta biblioteka graficzna) jest międzyjęzykowym, wieloplatformowym API do renderowania dwu- i trójwymiarowej grafiki komputerowej. OpenGL/ES (OpenGL do systemów wbudowanych) jest podzbiorem OpenGL zaprojektowanym z przeznaczeniem dla systemów wbudowanych.
- **Rama mediów.** Rama mediów (ang. *Media Framework*) wspomaga wideozapis i odtwarzanie w różnych formatach, w tym ACC, AVC (H.264), H.263, MP3 i MPEG-4.
- **Baza danych SQL.** Android zawiera lekki mechanizm²³ bazy danych SQLite do przechowywania trwałych danych. SQLite jest omówiona w następnym podrozdziale.
- **Mechanizm przeglądarki** (ang. *Browser Engine*). Do szybkiego wyświetlania treści w HTML-u Android używa biblioteki WebKit, będącej w istocie tą samą biblioteką co stosowana w przeglądarce Safari i iPhone'ach. Biblioteka ta była też używana w przeglądarce Google Chrome, dopóki Google nie przeszedł na Blink.
- **Bionic LibC.** Jest to zredukowana wersja standardowej biblioteki systemu C, dostosowana do urządzeń wbudowanych opartych na Linuxie. Interfejsem jest standardowy JNI (Java Native Interface, z ang. rodzimy interfejs Javy).

JĄDRO LINUXA

Jądro SO zastosowane w Androidzie jest podobne do dystrybucji standardowego jądra Linuxa. Godną odnotowania zmianą w jądrze Androida jest brak modułów sterujących urządzeniami nienadającymi się do środowisk mobilnych, dzięki czemu jądro jest mniejsze. Ponadto Android wzbogaca jądro Linuxa o cechy przykrojone na miarę środowiska mobilnego i na ogół niezbyt przydatne na platformach biurowych czy w laptopach.

²³ Ang. *engine*. Powszechnie określa się to mianem „silnik”, jednak z uwagi na ochronę środowiska proponujemy ogólniejsze określenie (by nie używać terminu „maszyneria”) — *przyp. tłum.*

Android bazuje na swoim linuksowym jądrze, jeśli chodzi o rdzenne usługi systemowe: bezpieczeństwo, zarządzanie pamięcią, zarządzanie procesami, stos sieciowy i model programów obsługi urządzeń. Jądro działa również jako warstwa abstrakcji pomiędzy sprzętem a pozostałym stosem oprogramowania i umożliwia Androidowi korzystanie z szerokiego asortymentu udostępnianych w Linuxie modułów sterujących sprzętem.

Środowisko wykonawcze Androida

Większość systemów operacyjnych używanych w urządzeniach mobilnych, jak iOS lub Windows, stosuje oprogramowanie kompilowane bezpośrednio na konkretną platformę sprzętową. W Androidzie większość oprogramowania jest natomiast odwzorowywana na format bajtokodu, następnie przekształcany na rodzime instrukcje danego urządzenia. We wcześniejszych wydaniach Androida stosowano schemat o nazwie Dalvik. Jednak Dalvik ma sporo ograniczeń, jeśli chodzi o skalowanie na większe pamięci i architektury wielordzeniowe, toteż nowsze wydania Androida bazują na schemacie zwanym Android Runtime (ART). ART jest w pełni zgodny z istniejącym formatem bajtokodowym Dalvika (dex, od Dalvik Executable), więc twórcy aplikacji nie muszą zmieniać swojego kodu, aby działał on pod ART-em. Przyjrzymy się najpierw Dalvikowi, a potem spojrzymy na ART.

MASZYNA WIRTUALNA DALVIK

Dalvik VM (DVM) wykonuje pliki w formacie *.dex*, optymalizowanym pod kątem wydajnego magazynowania i wykonywania odwzorowywalnego w pamięci. VM może wykonywać klasy skompilowane przez kompilator języka Java, przekształcone na jego rodzimy format z zastosowaniem dołączonego narzędzia „dx”. VM działa powyżej jądra Linuxa, wykorzystując jego funkcjonalność (np. wątkowość i niskopoziomowe zarządzanie pamięcią). Rdzenna biblioteka klas Dalvika ma w założeniu dostarczać znanej bazy rozwojowej tym, którzy przywykli do programowania w Java Standard Edition, lecz jest nastawiona specjalnie na wymagania małych urządzeń przenośnych.

Każda aplikacja Androida działa we własnym procesie, z własnym egzemplarzem maszyny wirtualnej Dalvik. Dalvika napisano w taki sposób, aby urządzenie mogło efektywnie wykonywać wiele maszyn wirtualnych.

FORMAT PLIKU DEX

Maszyna wirtualna DVM wykonuje aplikacje i kod napisany w Javie. Standardowy kompilator Javy przekształca kod źródłowy (zapisany w postaci pliku tekstowego) na bajtokod. Bajtokod jest następnie kompilowany na plik *.dex*, nadający się do czytania i użytku przez DVM. W istocie pliki klas są zamieniane na pliki *.dex* (na podobieństwo pliku *.jar*, gdyby używano standardowej maszyny wirtualnej Javy — JVM), po czym są czytane i wykonywane przez DVM. Powielone dane używane w plikach klas są włączane do pliku *.dex* tylko raz, co oszczędza pamięć i jest mniej kosztowne. Pliki wykonywalne można ponownie modyfikować podczas instalowania aplikacji, aby jeszcze bardziej optymalizować je pod kątem urządzenia mobilnego.

KONCEPCJE ŚRODOWISKA WYKONAWCZEGO ANDROIDA

ART jest aktualnym środowiskiem wykonawczym aplikacji używanym przez Androida, wprowadzonym w wersji Android 4.4 (KitKat). Pierwszy projekt Androida zakładał, że system ten będzie pracował na jednym rdzeniu (z minimalnym wsparciem sprzętowym wielowątkowości) i z urzą-

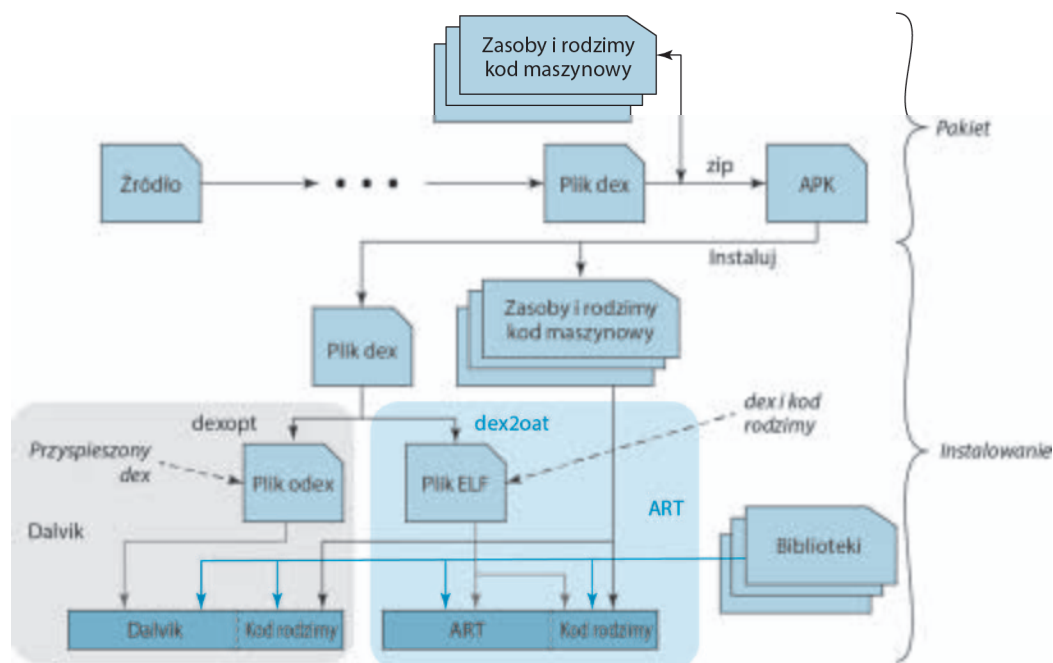
dzeniami o małych pamięciach, dla których Dalvik wydawał się odpowiednim środowiskiem wykonawczym. Jednak ostatnimi czasy urządzenia, na których działa Android, mają procesory wielordzeniowe i więcej pamięci (o względnie umiarkowanej cenie), co spowodowało, że w Google'u powtórnie przemyślano projekt środowiska wykonawczego, aby dostarczyć budowniczym aplikacji i użytkownikom więcej wrażeń przez umożliwienie wykorzystania własności będącego w obiegu sprzętu wysokiej klasy.

Zarówno na użytek Dalvika, jak i ART-a aplikacje pisane w Javie są kompilowane na bajtkod dex. Jako że Dalvik używa formatu bajtkodu dex ze względu na przenośność, format ten musi być zamieniany (kompilowany) na kod maszynowy, aby mógł być naprawdę wykonany przez procesor. Środowisko wykonawcze Dalvik dokonuje konwersji z bajtkodu dex na rodzimy kod maszynowy podczas biegu aplikacji, a proces ten nazwano JIT (ang. *just-in time* — kompilacja terminowa)²⁴. Ponieważ JIT kompiluje tylko część kodu, w mniejszym stopniu odciska się w pamięci i używa mniej fizycznej przestrzeni w urządzeniu. (W pamięci trwale są przechowywane tylko pliki dex, w przeciwieństwie do rzeczywistego kodu maszynowego). Dalvik identyfikuje często wykonywaną sekcję kodu i przechowuje ją po jednorazowym skompilowaniu w pamięci podręcznej, dzięki czemu kolejne wykonania takiej sekcji są szybsze. Strony pamięci fizycznej przechowujące podręcznie kod nie są wymienne ani nie podlegają stronicowaniu, co również odciska się nieco na pamięci, jeśli system jest już w takim stanie. Nawet z tymi optymalizacjami Dalvik musi wykonywać kompilację terminową przy każdym przebiegu aplikacji, co zabiera sporo zasobów procesora. Zauważmy, że procesor jest używany nie tylko do faktycznego wykonywania aplikacji, lecz także do konwersji bajtkodu dex na kod rodzimy, co jest przyczyną „drenowania” mocy. Takie użycie procesora było również powodem słabego odbioru interfejsu przez użytkowników w niektórych aplikacjach mocno obciążanych w fazie początkowej.

Aby obejść część tych problemów i zwiększyć efektywność zastosowania dostępnego sprzętu wysokiej klasy, w Androidzie wprowadzono ART. Środowisko ART również wykonuje bajtkod dex, lecz zamiast kompilowania bajtkodu w fazie wykonywania ART kompiluje bajtkod na kod maszynowy podczas instalowania aplikacji. Nazywa się to kompilacją zawczasu (ang. *ahead of time* — AOT). Do wykonywania tej kompilacji w fazie instalowania ART używa narzędzia `dex2oat`. Narzędzie to daje na wyjściu plik, który jest wykonywany podczas działania aplikacji.

Na rysunku 2.21 pokazano realizowany cykl życia APK, pakietu aplikacji przychodzącego od jej wykonawcy do użytkownika. Cykl rozpoczyna się od skompilowania kodu źródłowego na format `.dex` i połączenia go z odpowiednim kodem uzupełniającym w celu utworzenia APK. Po stronie użytkownika otrzymany APK jest rozpakowywany. Zasoby i rodzimy kod są na ogół instalowane wprost w katalogu aplikacji. Jednak kod `.dex` wymaga dalszego przetwarzania, zarówno w przypadku Dalvika, jak i ART-a. W Dalviku do wyprodukowania z pliku `.dex` zoptymalizowanej wersji formatu dex (odex), określanego jako przyspieszony dex, jest używana funkcja o nazwie `dexopt`. Chodzi o to, aby uczynić kod dex szybszy w interpretacji przez interpreter dexa. W ART funkcja `dex2oat` wykonuje optymalizację tego samego typu: również kompiluje kod dex na kod maszynowy docelowego urządzenia. Wynikiem funkcji `dex2oat` jest plik typu *Executable and Linkable Format* (ELF, z ang. format wykonywalny i konsolidowalny), który działa bezpośrednio, bez interpretera.

²⁴ Czyli w samą porę, na czas; przed blisko dwiema dekadami tak samo nazwano ten sposób interpretacji kodu w Javie — *przyp. tłum.*



Rysunek 2.21. Cykl życia APK (pakietu aplikacji)

ZALETY I WADY

Korzyści z zastosowania ART-a są następujące:

- Zmniejsza się czas rozruchu aplikacji, ponieważ jest wykonywany bezpośrednio kod rodzimy (maszynowy).
- Poprawia się żywotność baterii, gdyż unika się używania procesora do kompilacji terminowej (JIT).
- Ślad wymagany w pamięci do działania aplikacji jest mniejszy (bo nie ma zapotrzebowania na pamięć podręczną JIT). Co więcej, skoro nie ma niestroniczowalnego kodu JIT do podręcznego przechowywania, uelastycznia to wykorzystanie pamięci RAM w sytuacjach, gdy zaczyna jej brakować.
- Wraz z ART-em przybyło parę optymalizacji przy łączeniu nieużytków i ulepszeń dotyczących uruchamiania kodu (oczyszczania go z błędów).

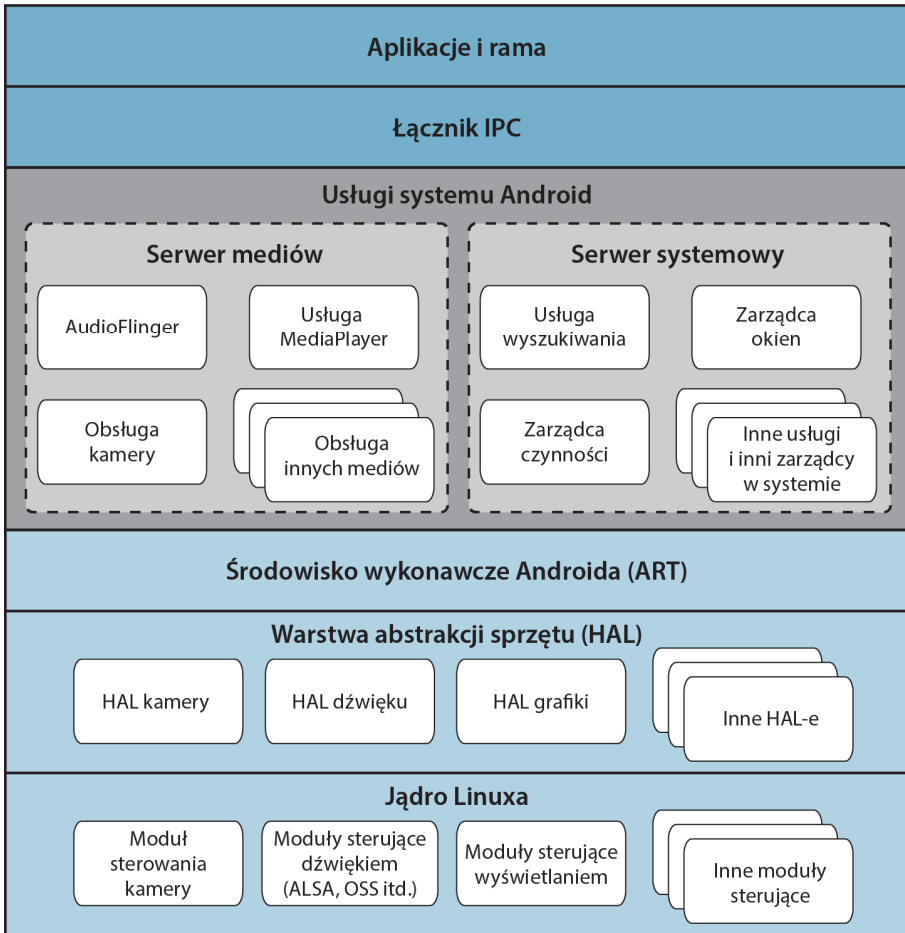
ART ma również kilka potencjalnych wad:

- Ponieważ zamiana bajtokodu na kod maszynowy jest wykonywana w czasie instalowania, instalowanie aplikacji zajmuje więcej czasu. Dla twórców aplikacji Androida, którzy ładują aplikację wiele razy podczas testowania, ten czas może być zauważalny.
- Podczas pierwszego, świeżego rozruchu po przywróceniu ustawień fabrycznych wszystkie aplikacje zainstalowane na urządzeniu są kompilowane na kod rodzimy za pomocą funkcji dex2opt. Dlatego w porównaniu z Dalvikiem pierwszy rozruch może trwać znacznie dłużej (3 – 5 sekund), nim pojawi się ekran powitalny (ang. *Home Screen*).

- Wygenerowany w ten sposób kod rodzimy zostaje przechowany w pamięci wewnętrznej, co wymaga znacznej dodatkowej jej ilości.

Architektura systemu Android

Pożyteczne jest przedstawienie Androida z perspektywy twórcy aplikacji, co ukazano na rysunku 2.22. Architektura tego systemu jest uproszczoną abstrakcją architektury oprogramowania z rysunku 2.20. Android oglądany w tym ujęciu składa się z następujących warstw:



Rysunek 2.22. Architektura systemu Android

- Aplikacje i rama.** Twórcy aplikacji są skoncentrowani przede wszystkim na tej warstwie i na interfejsach API, które umożliwiają dostęp do usług niższej warstwy.
- Łącznik IPC.** Mechanizm łącznika o nazwie Binder Inter-Process Communication umożliwia ramie aplikacji przekraczanie granic procesu i wywoływanie kodu usług systemu Android. To zasadniczo umożliwia ramowym interfejsom API wysokiego poziomu interakcję z usługami systemu Android.

- **Usługi systemu Android.** Większość funkcji uzewnętrznianych za pośrednictwem ramy interfejsów API wywołuje usługi systemu, które z kolei zwracają się do usytuowanych niżej funkcji sprzętowych i jądrowych. Usługi można postrzegać jako zorganizowane w dwie grupy: usługi medialne odnoszą się do odtwarzania i nagrywania mediów, a usługi systemowe realizują funkcje na poziomie systemu, takie jak zarządzanie energią, położeniem (lokalizacją) i powiadomianiem.
- **Warstwa abstrakcji sprzętu (HAL).** HAL (ang. *hardware abstraction layer*) realizuje standardowy interfejs do modułów sterujących urządzeń na poziomie jądra, więc kod górnej warstwy nie musi uwzględniać szczegółów implementacji specyficznych modułów sterujących i sprzętu. Warstwa HAL w zasadzie nie różni się od występującej w standardowej dystrybucji Linuxa. Służy do abstrahowania możliwości poszczególnych urządzeń (udostępnianych przez sprzęt i eksponowanych przez jądro) z przestrzeni użytkownika. Przestrzeń użytkownika może zawierać usługi Androida lub aplikacje. Celem warstwy HAL jest utrzymywanie przestrzeni użytkownika spójnej w odniesieniu do różnych urządzeń. Dostawcy mogą również dokonywać własnych ulepszeń i zamontowywać je w warstwie HAL bez oddziaływania na przestrzeń użytkownika. Przykładem takiego postępowania jest HwC (ang. *Hardware Composer* — kompozytor sprzętu) będący zależną od dostawcy implementacją HAL, która rozumie zdolności renderowania sprzętu znajdującego się poniżej. Zarządca powierzchni gładko współpracuje z różnymi implementacjami HwC, pochodzącymi od różnych dostawców.
- **Jądro Linuxa.** Jądro Linuxa jest przykrojone na miarę wymagań środowiska mobilnego.

Czynności

Czynność (aktywność, ang. *activity*) jest pojedynczym wizualnym komponentem interfejsu użytkownika, zawierającym obiekty takie jak wybory z menu, ikony i pola kontrolne. Każdy ekran aplikacji jest rozszerzeniem klasy *Activity*. Czynności korzystają z widoków (*Views*), aby tworzyć graficzne interfejsy użytkownika, wyświetlające informacje i reagujące na działania użytkownika. Czynności omówimy w rozdziale 4.

Zarządzanie zasilaniem

Android uzupełnia jądro Linuxa o dwie cechy mające usprawnić zdolność zarządzania zasilaniem. Są to alarmy i blokowanie czuwania (ang. *wakelocks*).

Możliwość alarmów jest zrealizowana w jądrze Linuxa i widoczna dla budowniczego aplikacji za pośrednictwem zarządcy alarmów (*AlarmManager*) w rodzimych bibliotekach środowiska wykonawczego (*RunTime*). Poprzez zarządcę alarmów aplikacja może zamówić usługę terminowego budzenia. Możliwość alarmów jest implementowana w jądrze, zatem alarm może zadziałać nawet wówczas, gdy system jest w trybie uśpienia. Umożliwia to systemowi przechodzenie w stan uśpienia i oszczędzanie energii także wtedy, gdy jest jakiś proces, który wymaga budzenia.

Blokowanie czuwania chroni system Android przed wejściem w stan uśpienia. Aplikacja może utrzymywać jedną z następujących blokad czuwania:

- `Full_Wake_Lock` (z ang. pełna blokada czuwania) — procesor jest włączony, pełna jasność ekranu, klawiatura czynna.
- `Partial_Wake_Lock` (z ang. częściowa blokada czuwania) — procesor włączony, ekran wyłączony, klawiatura wyłączona.
- `Screen_Dim_Wake_Lock` — procesor włączony, ekran przyciemniony, klawiatura wyłączona.
- `Screen_Bright_Wake_Lock` — procesor włączony, ekran jasny, klawiatura wyłączona.

Te blokady są zamawiane poprzez API, ilekroć aplikacja zażąda, aby któreś z urządzeń zewnętrznych będących pod kontrolą pozostało włączone. Jeśli nie istnieje żadna blokada czuwania blokująca urządzenie, zostaje ono wyłączone ze względu na dbałość o żywotność baterii.

Te obiekty jądra są widoczne dla aplikacji w przestrzeni użytkownika za pomocą plików `/sys/power/wake_lock`. Plików `wake_lock` i `wake_unlock` można użyć do zdefiniowania i przełączania blokady przez odpowiedni zapis w stosownym pliku.

2.12. PODSTAWOWE POJĘCIA, PYTANIA SPRAWDZAJĄCE I ZADANIA

Podstawowe pojęcia

Adres fizyczny	Planowanie	System z podziałem czasu
Adres rzeczywisty	Podział czasu	Średni czas do wystąpienia awarii (MTTF)
Adres wirtualny	Praca	Średni czas naprawy (MTTR)
Czas pracy	Proces	Tryb jądra
Jądro (ang. <i>kernel</i> , <i>nucleus</i>)	Programowanie seryjne	Tryb użytkownika
Jądro monolityczne	Projektowanie obiektowe	Wada
Jednoprogramowość	Przerwanie	Wątek
Język sterowania zadaniami (JCL)	Przestój	Wieloprogramowość
Kontekst wykonywania	Przetwarzanie wsadowe	Wieloprogramowy system operacyjny
Kwantowanie czasu	Rotacyjny (algorytm)	Wieloprzetwarzanie symetryczne (SMP)
Maszyna wirtualna	Rozkaz uprzywilejowany	Wielowątkowość
Mikrojądro	Rozproszony system operacyjny	Wielozadaniowość
Moduły ładowalne	Stan	Zadanie
Monitor	Stan procesu	Zarządzanie pamięcią
Monitor rezydentny	System operacyjny	
Niezawodność	System wsadowy	
Pamięć wirtualna		

Pytania sprawdzające

- 2.1. Jakie są trzy cele projektowania SO?
- 2.2. Co to jest jądro SO?
- 2.3. Co to jest wieloprogramowość?
- 2.4. Co to jest proces?
- 2.5. Jak wygląda kontekst wykonywania procesu używany przez SO?
- 2.6. Wymień i krótko omów pięć powinności typowego SO dotyczących zarządzania pamięcią.
- 2.7. Wyjaśnij różnicę między adresem rzeczywistym a adresem wirtualnym.
- 2.8. Opisz technikę planowania rotacyjnego.
- 2.9. Wyjaśnij różnicę między jądrem monolitycznym a mikrojądrem.
- 2.10. Co to jest wielowątkowość?
- 2.11. Wymień podstawowe zagadnienia projektowania systemu operacyjnego typu SMP.

Zadania

- 2.1. Przypuśćmy, że mamy wieloprogramowany komputer, w którym każde zadanie ma identyczną charakterystykę. W jednym przydzielonym zadaniu okresie obliczeń T połowa czasu jest spędzana na operacjach wejścia-wyjścia, a połowa na zatrudnianiu procesora. Każde zadanie działa przez N takich okresów. Zakładając proste planowanie rotacyjne oraz możliwość zachodzenia w tym samym czasie działań na wejściu-wyjściu i w procesorze, określ następujące wielkości:

- czas cyklu przetwarzania = faktyczny czas do zakończenia zadania;
- przepustowość = średnia liczba zadań kończonych w okresie T ;
- wykorzystanie procesora = procent czasu, w którym procesor jest aktywny (nie czeka).

Oblicz te wielkości dla jednego, dwóch i czterech zadań wykonywanych jednocześnie, zakładając, że okres T podlega następującej dystrybucji:

- a. Pierwsza połowa idzie na wejście-wyjście, druga na procesor.
- b. Wejście-wyjście w pierwszej i czwartej ćwierci, procesor w drugiej i trzeciej ćwierci.

- 2.2. **Program ograniczony przez wejście-wyjście** (ang. *I/O-bound program*) to taki, który działając samotnie, spędzałby więcej czasu na oczekiwaniu na wejście-wyjście niż na używaniu procesora. **Program ograniczony przez procesor** (ang. *CPU-bound program*) zachowuje się odwrotnie. Załóżmy, że algorytm planowania krótkoterminowego faworyzuje programy, które ostatnio zużywały mniej czasu procesora. Wyjaśnij, dlaczego ten algorytm faworyzuje programy ograniczone przez wejście-wyjście, a jednak nie odmawia stale czasu procesora programom ograniczonym przez procesor.

- 2.3. Porównaj zasady planowania nadające się do użycia, gdy chodzi o optymalizację systemu z podziałem czasu, z tymi, z których należałoby skorzystać do optymalizowania **wieloprogramowego systemu wsadowego**.
- 2.4. Co jest celem wywołań systemowych? Jak wywołania systemowe mają się do SO i do koncepcji dualnego trybu działania (trybu jądra i trybu użytkownika)?
- 2.5. W systemie operacyjnym OS/360 komputerów głównych IBM jednym z ważniejszych modułów w jądrze jest zarządca zasobów systemu (ang. *System Resource Manager* — SRM). Moduł ten odpowiada za rozdzielnictwo zasobów między przestrzenie adresowe (procesy). SRM sprawia, że stopień wyrafinowania OS/360 jest niespotykany w systemach operacyjnych. Żaden inny SO komputera głównego (ang. *mainframe*) i z pewnością żaden SO innego typu nie może się równać z funkcjami wykonywanymi przez SRM. Pojęcie zasobu obejmuje procesor, pamięć rzeczywistą i kanały wejścia-wyjścia. SRM gromadzi statystyki dotyczące wykorzystania procesora, kanału i rozmaitych kluczowych struktur danych. Jego celem jest osiągnięcie optymalnego działania na podstawie monitorowania i analizy tych danych. W danej instalacji z góry ustala się różne cele efektywności służące jako wskazówki dla SRM, który dynamicznie modyfikuje instalację i charakterystykę wydajnościową zadań na podstawie wykorzystania systemu. Z drugiej strony, SRM dostarcza raporty, które doświadczonemu operatorowi umożliwiają doskonalenie konfiguracji i ustawień parametrów w celu poprawiania obsługi użytkowników.

Zadanie dotyczy jednego przykładu aktywności SRM. Rzeczywista pamięć jest podzielona na bloki równej długości, zwane ramkami, przy czym mogą ich być tysiące. W każdej ramce może znajdować się blok pamięci wirtualnej określany jako strona. SRM otrzymuje sterowanie w przybliżeniu 20 razy na sekundę i sprawdza każdą ramkę bez wyjątku. Jeśli do strony nie było odniesień ani nie została ona zmieniona, licznik jest zwiększany o 1. W miarę upływu czasu SRM uśrednia te wartości, aby określić średnią liczbę sekund, przez które ramka strony w systemie pozostawała nietknięta. Jaki może być tego cel i jakie działanie mógłby podjąć SRM?

- 2.6. Wieloprocesor o ośmiu procesorach ma przydzielonych 20 przewijaków taśmy. W systemie zlecono do wykonania dużą liczbę zadań, przy czym każde wymaga do zakończenia działania najwyżej czterech przewijaków. Załóżmy, że każde zadanie rozpoczyna działanie tylko z trzema przewijakami na długo przed zapotrzebowaniem czwartego krótko przed końcem. Załóżmy również nieustanne dostarczanie takich zadań.
- a. Zakładamy, że planista w SO nie rozpocznie zadania dopóty, dopóki nie będą dostępne cztery przewijaki taśmy. Po rozpoczęciu zadania cztery przewijaki są przydzielane natychmiast i nie są zwalniane aż do skończenia zadania. Ile zadań maksymalnie będzie w trakcie jednoczesnego wykonywania? Ile wynosi maksymalna i minimalna liczba przewijaków taśmy, które mogą pozostawać bezczynne wskutek zastosowania takiej polityki przydziału?
- b. Zaproponuj alternatywną politykę, aby poprawić wykorzystanie przewijaków taśmy, a jednocześnie uniknąć zakleszczenia w systemie. Ile wyniesie maksymalna liczba zadań, które mogą być wykonywane jednocześnie? Jakie są ograniczenia liczby próżnujących przewijaków taśmy?

SKOROWIDZ

A

- ABAC, 727
- ABI, 81
- AC, 41, 170
- access control, *Patrz:* dostęp kontrolowane
- ACE, 745
- activity, *Patrz:* czynność
- activity working set, *Patrz:* zbiór roboczy aktywności
- adresowanie, 367, 383
 - bezpośrednie, 285, 286
 - pośrednie, 286
- AF, 169
- age variable, *Patrz:* zmienna wieku
- ahead of time, *Patrz:* AOT
- aktywność, *Patrz:* czynność
- akumulator, 41, 171
- algorytm
 - bankiera, 324
 - C-SCAN, 563
 - częstości braków stron, 432, 434
 - Dekkera, 246, 248, 249, 250, 275
 - FCFS, 470, 471, 472, 479, 500, 501
 - FIFO, 418, 420, 421, 425, 470, 471, 472, 562, 563
 - dostęp do dysku, 559
 - FSCAN, 563
 - HRRN, 477, 478
 - LFU, 574
 - LIFO, 562
 - LLF, 538
 - LOOK, 562
 - LRU, 574, 576
 - MUF, 539
 - najdawniej używane, *Patrz:* algorytm LRU
 - najkrótszego pozostałego czasu, *Patrz:* algorytm SRT
 - najpierw najkrótszy czas obsługi, *Patrz:* algorytm SSTF
 - najpierw najkrótszy proces, *Patrz:* algorytm SPN
 - najpierw najmniejsza liczba wątków, 503
 - najpierw najwyższy wskaźnik reakcji, *Patrz:* algorytm HRRN
 - najrzadziej używane, *Patrz:* algorytm LFU
 - N-krokowy-SCAN, 563
 - omiatanie cykliczne, *Patrz:* algorytm C-SCAN
 - oparty na priorytetach, 559
 - optymalny, 418, 420
 - Petersona, 250, 275
 - planowania czasu rzeczywistego, 513
 - planowania dysku, 545
 - podzielonego LRU, 443
 - rotacyjny, 471, 472, 479, 513
 - samolubny, 492
 - SCAN, 562, 563
 - SPN, 473, 476, 477, 478
 - wyłączający, 476
 - SRT, 476, 477, 478
 - SSTF, 562, 563
 - windy, 562
 - wykrywania zakleszczenia, 328
 - wypatrywania, 562
 - zastępowania, 60, 420, 426, 435
 - najdawniej używanych, 60, 418, 420, 425
 - oparty na częstości, 575, 576
 - stron, 443
 - zbioru roboczego, 431, 434
 - zegara dwuwskazówkowego, 438
 - zegarowy, 418, 420, 421, 422, 424, 425, 435, 443, 471
- alignment check, *Patrz:* AC

alokator pamięci jądra, 439
 ALPC, 113
 ALU, 64
 Amdahla prawo, *Patrz:* prawo Amdahla
 Android, 129, 229, 665
 alarm, 136
 aplikacja, *Patrz:* aplikacja w systemie Android
 architektura, 129, 130, 131, 132, 135, 136
 biblioteka systemowa, 131
 blokowanie czuwania, 136
 Dalvik, *Patrz:* Dalvik
 katalog
 danych, *Patrz:* katalog danych Android
 systemowy, *Patrz:* katalog systemowy Android
 pamięć, 447
 system
 plików, *Patrz:* system plików Android
 wykonawczy, 131
 środowisko wykonawcze, 132
 warstwa
 ramy aplikacji, 129, 130
 zastosowań, 129
 wbudowany, 665
 współbieżność, *Patrz:* współbieżność Android
 anticipatory scheduler, *Patrz:* planista przewidujący
 AOT, 133
 API, 81, 111
 APK, 133
 aplikacja, 79
 Android, 229
 czynność, 229, 230, 231
 dostawca treści, 229
 komponent, 229
 likwidowanie, 232
 odbiorca ogłoszeń, 230
 proces, 232
 usługa, 229
 wątek, 232
 Javy, 209
 programista, *Patrz:* programista aplikacji
 wielogzemplarzowa, 209
 wieloprotokowa, 209
 wielowątkowa, 209
 Windows, 212
 cykl istnienia, 213
 drugoplanowa, 214
 pierwszoplanowa, 213
 wirtualna, 688
 bezpieczeństwa, *Patrz:* SVA
 application binary interface, *Patrz:* ABI

application processor, *Patrz:* procesor aplikacji
 application-specific integrated circuit, *Patrz:* ASIC
 architektura
 Android, *Patrz:* Android architektura
 klient-serwer, *Patrz:* model klient-serwer
 Linux, *Patrz:* Linux architektura
 nadrzędny-podrzędny, 499
 partnerska, 499
 UNIX, *Patrz:* UNIX architektura
 Windows, *Patrz:* Windows architektura
 zbioru rozkazów, *Patrz:* ISA
 architektura chmury obliczeniowej, *Patrz:* chmura
 obliczeniowa architektura
 ART, 132, 133, 134
 ASIC, 649
 atak
 z przepełnieniem bufora, 717, 723, 724
 z przepełnieniem stosu, 724
 application programming interface, *Patrz:* API
 atomic operation, *Patrz:* operacja niepodzielna
 audio, 40
 authenticity, *Patrz:* bezpieczeństwo uwierzytelnianie
 auxiliary carry flag, *Patrz:* AF
 auxiliary memory, *Patrz:* pamięć pomocnicza
 availability, *Patrz:* SO dostępność, bezpieczeństwo
 dostępność
 awaria
 sprzętu, 43
 tolerowanie, 104, 108

B

banker's algorithm, *Patrz:* algorytm bankiera
 Barbican, 779
 bariera pamięciowa, 344
 base register, *Patrz:* rejestr bazowy
 basic file system, *Patrz:* system plików podstawowy
 basic I/O supervisor, *Patrz:* nadzorca wejścia-
 wyjścia podstawowy
 baza danych, 595, 596
 SQLite, 641
 B-drzewo, 606, 607
 bezpieczeństwo, 99, 587, 712, 736, 737, 738, 740, 741
 chmura obliczeniowa, 778
 deskryptor, *Patrz:* deskryptor bezpieczeństwa
 dostępność, 99, 103, 105
 nienaruszalność danych, 99
 polityka, *Patrz:* polityka bezpieczeństwa
 poufność, 99
 rejestrwanie zdarzeń, 741

rozszerzenia językowe, 722
 testowanie, 740
 uwierzytelnianie, 99
 Windows, 742, 743, 747
 personifikacja, 743
 wybór
 bezpiecznych bibliotek, 722
 języka programowania, 721
 biblioteka
 bezpieczna, 722
 dołączana dynamicznie, *Patrz:* DLL
 P-wątków, 225
 RIOT, *Patrz:* RIOT biblioteka
 μ CLIBC, *Patrz:* μ CLIBC
 Binder, *Patrz:* łącznik
 bit
 modyfikacji, 402
 niewykonywania, 724
 odwołania, 438
 przylepności, 735
 użycia, 421
 zamknięcia, 420
 BitLocker, 587
 block-oriented device, *Patrz:* urządzenie wejścia-
 wyjścia blokowe
 blok, 598, 599, 603, 615, 770
 adres, 609
 definiowanie, 233
 długość, 615, 616
 dostęp, 605
 dwupośredni, 628
 jednopośredni, 628
 lista wolnych, 623
 porcja, 617
 przydział ciągły, 618
 przydział dynamiczny, 618, 628
 przydział indeksowy, 620
 przydział łańcuchowy, 620
 przydział z góry, 617, 618
 rozmiar, 617, 618
 trójpśredni, 628
 blokada
 czytelnicy – pisarze, 347, 348, 351
 wąska, 351
 wirująca, 339, 340, 341, 350
 wykluczająca, 350
 wzajemnego wykluczania, 345, 346
 blokowanie, *Patrz:* blok długość
 błąd, 41, 80, 94, 153
 pamięci, 80, 175

 programowy, 80
 sprzętowy, 80
 bomba logiczna, 713
 boot loader, *Patrz:* program rozruchowy
 bottom-half kernel, *Patrz:* jądro dolnej połowy
 bounds register, *Patrz:* rejestr graniczny
 buffer overflow, *Patrz:* bufor przepełnienie
 bufor
 przepełnienie, 717, 719
 obrona, 719, 720, 721, 722, 723, 724
 sforsowanie, *Patrz:* bufor przepełnienie
 TLB, *Patrz:* TLB
 translacji adresów stron, *Patrz:* TLB
 wejścia-wyjścia, 553, 554, 555
 UNIX, 578, 580
 zamienianie, *Patrz:* operacja wejścia-wyjścia
 buforowanie podwójne
 busy waiting, *Patrz:* czekanie aktywne

C

cache, *Patrz:* pamięć podręczna
 carry flag, *Patrz:* CF
 Celiometr, 779
 central processing unit, *Patrz:* CPU
 CF, 169
 c-grupa, 226, 228, 689
 checkpoint, *Patrz:* punkt kontrolny
 child process, *Patrz:* proces potomny
 chip multiprocessor, *Patrz:* komputer
 wielordzeniowy
 chmura obliczeniowa, 754, 755, 784
 architektura NIST, 760
 audytor, 761, 762
 bezpieczeństwo, 778
 dostawca, *Patrz:* CSP
 hybrydowa, 759
 model, 756, 757, 758, 759
 pośrednik, 761, 762
 prywatna, 759
 publiczna, 758
 społecznościowa, 759
 transporter, 761, 762
 Cinder, 778
 circular buffer, *Patrz:* bufor wejścia-wyjścia
 cykliczny
 circular SCAN, *Patrz:* algorytm C-SCAN
 cloud computing, *Patrz:* chmura obliczeniowa
 Cloud Native Computing Foundation, *Patrz:* CNCF
 cloud service consument, *Patrz:* CSC

CNCF, 692

coefficient of variation, *Patrz:* współczynnik wahań

commercial off-the-shelf, *Patrz:* COTS

Completely Fair Queuing I/O scheduler, *Patrz:* planista CFQ

Completely Fair Scheduler, *Patrz:* planista CFS

confidentiality, *Patrz:* bezpieczeństwo poufność

Congress, 780

constrained device, *Patrz:* urządzenie ograniczone

consumable resource, *Patrz:* zasoby zużywalne

container virtualization, *Patrz:* wirtualizacja kontenerowa

control group, *Patrz:* c-grupa

core, *Patrz:* rdzeń

coroutine, *Patrz:* współprocedura

coscheduling, *Patrz:* współplanowanie

COTS, 766

counting semaphore, *Patrz:* semafor zliczający

CPL, 173

CPU, 38

critical resource, *Patrz:* zasób krytyczny

critical section, *Patrz:* program sekcja krytyczna

CSC, 755, 758, 760, 764, 768

CSP, 760, 761, 764

current privilege level, *Patrz:* CPL

cykl

przetwarzania, 467, 470

rozkazowy, 41, 45

czas

cyklu przetwarzania, 467, 470

gotowości, 516

pobyty, 467

porcjowanie, 471

pracy, 105

przetwarzania, , 467, 470, 516

średni

naprawy, *Patrz:* MTTR

wystąpienia awarii, *Patrz:* MTTF

średni międzynaaprawczy, *Patrz:* MTTR

wirtualny, 429

wykonania wirtualny, 527

czasomierz, 86

czekanie

aktywne, 246, 262

wirujące, 246, 262, 264

czujnik bezprzewodowy, 665

sieć, 666, 668

czynność, 136

D

DAC, 727, 728

DACL, 744, 745

Dalvik, 132, 133

dane

archiwizowanie, 742

macierz, *Patrz:* macierz danych

przetwarzanie, 41

składowanie, 742

struktura task_struct, 223

DAS, 769

data integrity, *Patrz:* bezpieczeństwo nienaruszalność danych

DDR, *Patrz:* pamięć główna sterownik DDR

deadlock, *Patrz:* zakleszczenie

deadlock avoidance, *Patrz:* zakleszczenie unikanie

deadlock prevention, *Patrz:* zakleszczenie zapobieganie

debuger, 80

dedicated processor, *Patrz:* procesor dedykowany

deeply embedded system, *Patrz:* system głęboko wbudowany

Dekkers algorytm, *Patrz:* algorytm Dekkera

dentry, *Patrz:* k-wpis

descheduling, *Patrz:* planowanie wątków wypadanie z harmonogramu

Designate, 780

deskryptor bezpieczeństwa, 117, 744

DF, 169

diagram

kolejkowania, 151

połączonego postępu, 313, 316, 320

region fatalny, 315, 320

przebieg stanów, 158, 159, 459

digital signal processor, *Patrz:* DSP

direct addressing, *Patrz:* adresowanie bezpośrednie

direct attached storage, *Patrz:* DAS

direct memory access, *Patrz:* DMA

direction flag, *Patrz:* DF

disabled interrupt, *Patrz:* przerwanie zablokowane

disk allocation table, *Patrz:* dysk tablica przydziału

dispatcher, *Patrz:* dyspozytor, ekspedytor

dispatcher object, *Patrz:* obiekt dyspozytora

distributed operating system, *Patrz:* SO rozproszony

DLL, 396

DMA, 61, 546, 547, 548, 580

Docker, 694, 695

dostęp

kontrolowanie, 716, 725, 727, 747

lista, 745, 746

obligatoryjne, 727

UNIX, 733

uznaniowe, 727, 728

w systemie plików, 725

według atrybutów, 727

według ról, 727, 731, 732

korzeniowy, 712

double indirect block, *Patrz:* blok dwupośredni

dowiązanie, 626

downtime, *Patrz:* przestój

drukarka, 580, 581

DSP, 40

dumb terminal, *Patrz:* terminal niemy

DVM, 132, 133

dynamic linking, *Patrz:* program konsolidacja dynamiczna

dysk, 580

czas przesyłania, 556, 557

czas wyszukiwania, 556

dostęp losowy, 558

lista wolnych bloków, 624

logiczny, *Patrz:* tom

niezależny tablica nadmiarowa, *Patrz:* RAID

opóźnienie obrotowe, 557

organizacja sekwencyjna, 558

pamięć podręczna, *Patrz:* pamięć podręczna dysku

planowanie losowe, 545

losowe, 559

podwojenie, 587

szybkość, 556

tablica przydziału, 622, 623, 625

tom, 587

wirtualny, 684

wykrywanie rotacyjno-pozycyjne, *Patrz:* RPS

dyspozytor, 462, 512

dziedziczenie, 116

E

edge computing, *Patrz:* przetwarzanie na obrzeżu

edytor, 80

ekspedytor, 147, 149, 151

Elastic Compute Cloud, 775

elevator scheduler, *Patrz:* planista windy

elewator linuksowy, 581

embedded Linux, *Patrz:* Linux wbudowany

embedded operating system, *Patrz:* SO wbudowany

event flag, *Patrz:* znacznik zdarzeń

event object, *Patrz:* obiekt zdarzenia

execution context, *Patrz:* proces kontekst

wykonywania

exhaustive search, *Patrz:* plik przeszukiwanie pełne

F

fair-share scheduler, *Patrz:* planista uczciwych

udziałów

fair-share scheduling, *Patrz:* planowanie uczciwych

udziałów

FAT, 617, 625, 630

fatal region, *Patrz:* diagram połączonego postępu

region fatalny

fault, *Patrz:* wada

fiber, *Patrz:* włókno

field-programmable gate array, *Patrz:* FPGA

file allocation table, *Patrz:* FAT

file directory, *Patrz:* katalog plików

file management system, *Patrz:* system zarządzania plikami

filtr wiaderka żetonów, *Patrz:* TBF

firewall, *Patrz:* zaporą sieciową

fog computing, *Patrz:* mgła obliczeniowa

FPGA, 649

Free Software Foundation, *Patrz:* FSF

FSF, 123

FSS, *Patrz:* planista uczciwych udziałów

funkcja

cnotify, 281

csignal, 278, 279, 281

cwait, 278, 279

czekania, 348

DMA, *Patrz:* DMA

ExitThread, 212

fork, 186

odwzorowująca, 59

wyboru, 465, 466

G

gang scheduling, *Patrz:* planowanie wątków zespołowe

GCD, 109, 233

kolejka, 233

general semaphore, *Patrz:* semafor ogólny

Glance, 777

głodzenie, 246, 257, 258, 291, 477, 581

unikanie, 331

GPU, 40
 graf przydziału zasobów, 318, 319
 Grand Central Dispatch, *Patrz:* GCD
 graphical processing unit, *Patrz:* GPU
 grono, 62, 496
 grupa sterowania, *Patrz:* c-grupa
 jądrowa, *Patrz:* jądrowa grupa sterowania
 guard page, *Patrz:* pamięć strona strażnicza
 guest OS, *Patrz:* SO goszczony
 GUI, 111

H

haker, 713
 HAL, 789, 792
 hardware abstraction layer, *Patrz:* HAL
 Heat, 779
 highest response ratio first, *Patrz:* algorytm HRRN
 high-level language, *Patrz:* HLL
 hiperwizor, 681, 683, 684, 696, 697
 Hyper-V, 704
 typ
 1, 685, 686, 688
 2, 685, 686, 688, 704
 VMware ESXi, 701, 702, 703
 Xen, 703
 zadania, 685
 hit ratio, *Patrz:* współczynnik trafień H
 HLL, 81
 Hugepage, *Patrz:* pamięć strona wielka

I

I/O manager, *Patrz:* zarządca wejścia-wyjścia
 I/O privilege level, *Patrz:* IOPL
 IaaS, 756, 757, 759, 761, 763, 764, 770, 775
 ID, 170
 identification flag, *Patrz:* ID
 IDS, 714, 740
 IDT, 118
 IF, 170
 i-liczba, 629
 immediate preemption, *Patrz:* wywłaszczanie
 natychmiastowe
 index number, *Patrz:* i-liczba
 indirect addressing, *Patrz:* adresowanie pośrednie
 infrastruktura jako usługa, *Patrz:* IaaS
 instruction, *Patrz:* rozkaz
 instruction register, *Patrz:* rejestr rozkazów
 instruction set architecture, *Patrz:* ISA

instrukcja uprzywilejowana, 86
 interfejs
 binarny aplikacji, *Patrz:* ABI
 ISA, *Patrz:* ISA
 Metro, 213, 447
 NT API, 115
 programów użytkowych, *Patrz:* API
 semafora, 342
 interlocked operation, *Patrz:* operacja zazębiona
 internet rzeczy, *Patrz:* IR
 interrupt dispatch table, *Patrz:* IDT
 interrupt-driven I/O, *Patrz:* operacja wejścia-wyjścia
 sterowana przerwaniem
 interrupt enable flag, *Patrz:* IF
 interrupt handler, *Patrz:* procedura obsługi
 przerwania
 intruder, *Patrz:* intruz
 intrusion detection system, *Patrz:* IDS
 intruz, 712, 713
 i-number, *Patrz:* i-liczba
 IOPL, 170
 IPS, 740
 IR, 41, 754, 780, 781
 aktywator, 782
 czujnik, 782
 kontekst chmury, 782
 mikrokontroler, 782
 rozwój, 781
 system operacyjny, *Patrz:* SOIR
 Ironic, 780
 ISA, 80
 ISR, 511
 IT, 781
 i-węzeł, 626
 rozmiar, 628
 struktura, 626

J

Java virtual machine, *Patrz:* JVM
 jądro, 81, 119, 656
 Android, 131
 dolnej połowy, 530
 górnej połowy, 530
 Linux, 126, 128, 131
 monolityczne, 102
 nieprocesowe, 178
 Ra, 206
 Windows, 111
 wirtualizacji, 701

jądrowa grupa sterowania, 689
JCL, 85, 86
jednoprogramowość, 87
jednostka
 arytmetyczno-logiczna, *Patrz:* ALU
 centralna, *Patrz:* CPU
język
 sterowania zadaniami, *Patrz:* JCL
 wysokiego poziomu, *Patrz:* HLL
JIT, 133
job, *Patrz:* zadanie
job control language, *Patrz:* JCL
job object, *Patrz:* obiekt zadania
joint progress diagram, *Patrz:* diagram połączonego postępu
just-in time, *Patrz:* JIT
JVM, 705

K

kanarek, 723
katalog, 625
 danych Android, 640
 macierzysty, 612
 plików, 609, 610, 611, 612
 roboczy, 612
 systemowy Android, 640
 ścieżka, 612
 w systemie UNIX, 629
kernel, *Patrz:* jądro
kernel control group, *Patrz:* jądrowa grupa sterowania
kernel mode, *Patrz:* tryb jądra
kernel-level threads, *Patrz:* wątek KLT
Keystone, 778
klaser, *Patrz:* grono
kod
 pasożytniczy, 713
 programu, 146
 wynikowy, 85, 390
kolejka
 GCD, 233
 komunikatów, 287
 obsługa, 478
 procesów, *Patrz:* proces kolejka
 semaforów, 336
kompilacja
 terminowa, *Patrz:* JIT
 zawczasu, *Patrz:* AOT
kompilator, 85, 118

komputer
 jednoukładowy, *Patrz:* mikrokontroler
 wieloprocessorowy, *Patrz:* wieloprocessor
 wielordzeniowy, 62, 64, 108, 109
 rozwój, 64
komunikat, 259, 335, 337, 799
 format, 287
 kolejka, 287
 przekazywanie, 263, 283, 288
konsolidator, 394
konsument usług chmurowych, *Patrz:* CSC
kontener, 689
 Dockera, *Patrz:* Docker
 instalowanie, 690
 konfigurowanie, 690
 Linux, *Patrz:* LXC
 system plików, *Patrz:* system plików
 kontenerowy
 wady, 691
 wirtualizacji, 688
 wirtualny, 691
 zalety, 691
 zarządzanie, 690
kraker, 713
Kubernetes, 692
k-wpis, 629

L

laxity, *Patrz:* luz czasowy
least privilege principle, *Patrz:* zasada najmniejszych przywilejów
least-recently used, *Patrz:* algorytm zastępowania najdawniej używanych
Libsafe, 722
licencja GNU GPL, 123
liczba indeksowa, *Patrz:* i-liczba
licznik programu, *Patrz:* PC
lightweight process, *Patrz:* proces lekki
linia komunikacyjna, 580, 581
linkage editor, *Patrz:* program łączący
linker, *Patrz:* konsolidator
Linux, 123
 architektura, 124, 125, 126, 128
 c-grupa, *Patrz:* c-grupa
 do mikrokontrolerów, *Patrz:* µClinux
 dystrybucja, 123
 biurkowa, 655
 serwerowa, 655
 interfejs semafora, 342

Linux

- jądro, 123, 126, 128
- kontener, 228
- LKML, 123
- pamięć
 - podręczna, 584
 - wirtualna, *Patrz:* pamięć wirtualna Linux
 - zarządzanie, 440, 443, 444
- proces, 223, 226
- przestrzeń nazw, 226
 - ipc, 226, 227
 - mnt, 226, 227
 - net, 226, 228
 - pid, 226, 227
 - user, 226, 228
 - uts, 226, 227
- system plików, 630, 661
- urządzenie wejścia-wyjścia, *Patrz:* urządzenie wejścia-wyjścia Linux
- wątek, 225, 338
- wbudowany, 655, 658, 659, 661
 - Android, 665
 - rozmiar pamięci, 659
 - system plików, 661
 - współbieżność, *Patrz:* współbieżność Linux
- Linux Containers, *Patrz:* LXC
- Linux Elevator, *Patrz:* elewator linuksowy
- Linux VServer, *Patrz:* VM Linux VServer
- lista
 - aplikacji biała, 740
 - kontroli dostępu, 727
 - UNIX, 735
- livelock, *Patrz:* szamotanina
- locality of reference, *Patrz:* lokalność odwołań
- log file, *Patrz:* plik dziennika, plik rejestrowy
- logical I/O, *Patrz:* moduł logicznego wejścia-wyjścia
- lokalność
 - odwołań, 55, 59, 70, 72, 401
 - w czasie, 72
 - w przestrzeni, 72
- LRU, *Patrz:* algorytm zastępowania najdawniej używanych
- luz czasowy, 538
- LXC, 695

Ł

- ładowacz, 390, 391
- łącznik, 352

M

- MAC, 727
- MAC OS, 233
- macierz
 - danych, 40
 - dostępów, 725, 728, 732
- magistrala, *Patrz:* szyna
- Magnum, 780
- Manila, 780
- MANO, 772
- mapa bitowa klastrów, 637
- mapping function, *Patrz:* funkcja odwzorowująca
- MAR, 39
- master file table, *Patrz:*
- master-slave architecture, *Patrz:* architektura nadrzędny-podrzędny
- maszyna wirtualna, *Patrz:* VM
 - Javy, *Patrz:* JVM
 - monitor, *Patrz:* VMM
- matryca bramkowa programowalna według dziedziny, *Patrz:* FPGA
- MBR, 39
- mean time to failure, *Patrz:* MTTF
- mean time to repair, *Patrz:* MTTR
- mechanizm
 - ciągnięcia, 533
 - pchania, 533
 - puli wątków, 109
- memory address register, *Patrz:* MAR
- memory buffer register, *Patrz:* MBR
- memory hierarchy, *Patrz:* pamięć hierarchia
- memory management, *Patrz:* pamięć zarządzanie
- metoda
 - dynamiczna
 - czyń co tylko możliwe, 514, 515
 - oparta na planowaniu, 513, 515
 - statyczna, 513, 515
 - wyłączająca, 467, 513, 515
- MFT, 637
- mgła obliczeniowa, 783, 784
- microservice, *Patrz:* mikrousluga
- Microsoft Windows, *Patrz:* Windows
- mikrojądro, 101, 102
- mikrokontroler, 652, 653, 782
- mikroprocesor, 40, 650
- mikrousluga, 694
- model
 - klient-serwer, 115
 - piaskownicy, 230
 - usług, 756

modem, 40

moduł

DMA, *Patrz:* DMA

logicznego wejścia-wyjścia, 599

ładowny, 124

ładowania, 390

sterujący urządzeń, 598

wejścia-wyjścia, 38

ogólność, 550

projektowanie, 549, 550

wydajność, 549

monitor, 84, 86, 263, 264, 277, 279, 333

Hoare'a, 277, 281, 283

Lampsona-Redella, *Patrz:* monitor Mesy

maszyny wirtualnej, *Patrz:* VMM

Mesy, 281, 282, 283

rezydentny, 84

struktura, 278

zmienna warunkowa, 278, 279

MTTF, 104, 105

MTTR, 104, 573

multicore, *Patrz:* komputer wielordzeniowy

multiprogramming, *Patrz:* wieloprogramowość

multiprogramming level, *Patrz:* wieloprogramowość
stopień

multitasking, *Patrz:* wielozadaniowość

multithreading, *Patrz:* wielowątkowość

Murano, 780

muteks, 264, 267, 345, 346

mutex object, *Patrz:* obiekt muteksu

mutual exclusion, *Patrz:* wykluczenie wzajemne

N

nadmiarowość, 106

nadzorca wejścia-wyjścia podstawowy, 599

named pipe, *Patrz:* potok nazwany

napęd dysków, *Patrz:* dysk

NAS, 769, 771

nested task flag, *Patrz:* NT

network attached storage, *Patrz:* NAS

Neutron, 777

Nova, 775, 776

NT, 170

NTFS, 634, 638

grono, 635, 636

klaster, 635, 636

odtwarzalność, 638

sektor, 635

wolumin, 636, 637

nucleus, *Patrz:* jądro

NUMA, 419

O

obiekt, 116, 146, 770

czasomierza do czekania, 350

dyspozytora, 117, 348, 349, 350

jądra, 117

klasa, 116

muteksu, 350

nazwany, 117

nienazwany, 117

procesu, 215, 216

semaforów, 350

uchwyt, 117

VFS, 631, 632, 633

wątku, 215, 216

wielopostaciowość, 116

zadania, 212

zdarzenia, 350

OF, 169

OpenStack, 772, 773, 775, 777, 779

punkt końcowy, 778

operacja

cwait, 278

mutex-enter, 346

niepodzielna, 245, 338, 339

semSignal, 264, 265, 275, 277, 336

semSignalB, 266

semWait, 264, 265, 275, 277, 336

semWaitB, 265

wejścia-wyjścia, 60, 544, 574, 615

bezpośredni dostęp do pamięci, *Patrz:* DMA

blokowa, 553, 554, 580, 581, 599

buforowanie, 552, 555, 554, 578, 580

czytania, 582, 583

ewolucja, 547

Linux, 581

niebuforowana, 578, 580

ogólność, 550

organizacja fizyczna, 552

pisania, 582

planowanie, 551, 552

programowanego, 60, 61, 546

projektowanie, 549, 550

RAID, *Patrz:* RAID operacja wejścia-wyjścia

sterowana przerwaniem, 60, 61, 546

strumieniowa, 553, 554

UNIX, 578

operacja

wejścia-wyjścia

Windows, 585, *Patrz też:* zarządca wejścia-wyjścia

wydajność, 549, 576

znakowa, 581

zazębiona, 351

oprogramowanie

antywirusowe, 740

jako usługa, *Patrz:* SaaS

złośliwe, 712, 713, 740

OT, 781

overflow flag, *Patrz:* OF

P

PaaS, 756, 757, 761, 763

page fault frequency, *Patrz:* algorytm częstości braków stronpage table entry, *Patrz:* PTE

pamięć, 39, 41, 53

adres

fizyczny, 367, 369, 380, 383, 399, 414

logiczny, 367, 369, 380, 383, 399, 414

rzeczywisty, 97

wiązanie, 392

wirtualny, 97, 118, 398

względny, 380

balonowanie, 699

bloków, 778

bloków sieciowych, 773

cena, 53, 55, 56

czas

cyklu, 56

dostępu, 53

dostęp niejednolity, *Patrz:* NUMA

drugorzędna, 55, 164, 616

dwupoziomowa, 53

działanie, 72

wydajność, 54, 69, 73

dyskowa, *Patrz:* dysk

dzielona, 335, 337

fragmentacja

wewnętrzna, 372, 401, 410

zewnętrzna, 375, 376, 377, 378, 401

główna, 38, 81, 164, 369, 401

ochrona, 716

sterownik DDR, 64

hierarchia, 53, 54, 55, 56

kontrolowanie dostępu, 97

nakładkowanie, 369

obiektów, 771, 773, 777

obrazu maszyny wirtualnej, 773

ochrona, 86, 97, 366, 368, 415

oparta na plikach, 771

operacyjna, *Patrz:* pamięć główna

organizacja, 367, 369

podłączona do sieci, *Patrz:* NAS

podręczna, 40, 56, 113, 409

dysku, 574

dzielenie, 508

Linux, 584

pobieranie z wyprzedzeniem, 64, 554

projektowanie, 58

SMP, 63

stron, 584

struktura, 57

VFS, *Patrz:* VFS pamięć podręczna

zasady działania, 56

podstawowa, *Patrz:* pamięć główna

podział, 370

dynamiczny, 370, 374, 375, 376, 377, 378

stały, 370, 371, 372, 373, 383

pomocnicza, 55, 164

przeciążanie, 699

przemieszczanie, 366, 367, 380

przestrzeń adresowa, 398

przydział, 96

płytowy, 444

stały, 426

zmienny, 427

ramka, 366, 381, 382, 401

zamknięta, 420

rzeczywista, *Patrz:* pamięć główna

segment, 366, 385, 413

tablica, 385

segmentacja, 371, 378, 385, 398, 401, 413, 414, 415, 417

strona, 97, 366, 381, 401

brak, 407, 411, 420, 421, 432

buforowanie, 425, 433

dzielenie, 698

fizyczna, 440

odwzorowanie asocjacyjne, 409

rozmiar, 410, 412

strażnicza, 724

tablica, 382, 402

tablica odwrócona, 405, 406

wielka, 440

zastępowanie, 438

- stronicowanie, 370, 378, 381, 383, 398, 401, 402,
 - 414, 415, 417, 436
 - na żądanie, 418
 - wstępne, 418
- szybkość, 53, 55, 56
- tablica, 164
- upakowanie, 375, 376, 377
 - system kumpłowski, 377, 378
 - system kumpłowski leniwy, 439, 440
- wielkość, 53, 59
- wirtualna, 97, 98, 113, 164, 166, 370, 398, 400,
 - 402, 418
 - Linux, 441, 442, 443
 - ochrona, 716
 - przestrzeń, 398
 - segmentacja, 371, 401
 - stronicowanie, 371, 381, 401
 - Windows, 445
- współużytkowanie, 367, 368, 369, 415
- wydajność, 59
- zarządzanie, 96, 108, 113, 164, 366, 370, 371, 372,
 - 373, 374, 375, 376, 377, 378, 417
 - Android, 447
 - Linux, 440, 443, 444
 - polityka zastępowania, 419, 426, 427, 428
 - Solaris, 436
 - UNIX, 436, 438, 439
 - Windows, 445, 447
- parawirtualizacja, 686
- parent process, *Patrz:* proces macierzysty
- parity flag, *Patrz:* PF
- partycja, 624
 - sektor rozruchowy, 637
- PC, 41, 49, 95
- PCB, *Patrz:* płytką drukowaną
- PCI Express, *Patrz:* szyna PCI Express
- pCPU, 696, 697
- peer architecture, *Patrz:* architektura partnerska
- Petersona algorytm, *Patrz:* algorytm Petersona
- PF, 169
- PFF, *Patrz:* algorytm częstości braków stron
- physical CPU, *Patrz:* pCPU
- pierścień ochrony, 697
- pipe, *Patrz:* potok
- planista
 - bezimpulsowy, 791
 - całkowicie uczciwego kolejkowania wejścia-wyjścia, *Patrz:* planista CFQ
 - CFQ, 584
 - CFS, 527, 528
 - długoterminowy, 461, 462
 - elewatorowy, *Patrz:* planista windy
 - krótkoterminowy, 462, 512
 - niewywłaszczający, 516
 - priorytetowy, 479
 - przewidujący, 583
 - średnioterminowy, 462
 - terminów ostatecznych wejścia-wyjścia, 582
 - TinyOS, 670, 671
 - uczciwych udziałów, 484
 - windy, 581, 582
- planowanie, 458, 459, 465, 478, 480
 - CFS grupowe, 528
 - czasu rzeczywistego, 509, 510, 513, 515, 516, 525
 - długoterminowe, 458, 459, 461
 - dysku, 545, 559
 - dziedziczenie priorytetów, 523
 - efektywność, 478
 - grupowe, 499
 - koncentracja prac, 504
 - krótkoterminowe, 458, 459, 462
 - kryteria, 462, 463, 464
 - Linux, 525, 526, 527, 528
 - modelowanie analityczne, 482
 - monotonicznego tempa, *Patrz:* RMS
 - odwrócenie priorytetów, 522
 - procesów, 500
 - przewidywalność, 463
 - pułap priorytetów, 525
 - średnioterminowe, 458, 459, 461
 - terminów nieprzekraczalnych, 515, 516
 - uczciwych udziałów, 482
 - UNIX, 486
 - UNIX FreeBSD, 530, 531
 - UNIX SVR4, 528
 - w trybie użytkownika, *Patrz:* UMS
- wątków, 502
 - dynamiczne, 502, 507
 - dzielenie obciążeń, 502, 503
 - przydział wydzielonych procesorów, 502, 505, 507
 - wielordzeniowe, 508
 - wypadanie z harmonogramu, 507
 - zespołowe, 502, 504, 505, 507
- wejścia-wyjścia, 458
- wieloprocessorowe, 498, 536
- Windows, 534
- ze sprzężeniem zwrotnym, 477, 478
- zespołowe, *Patrz:* planowanie grupowe

platforma

- docelowa, 655
- gospzcząca, 707
- jako usługa, *Patrz:* PaaS
- wirtualna, 707

plik, 594, 595, 596, 770

- .dex, 132, 133
- alokacja, *Patrz:* plik lokowanie
- bezpośredniego dostępu, *Patrz:* plik haszowany
- czytanie, 595
- DLL, 114
- dostęp, 599, 600
 - bezpośredni, *Patrz:* plik haszowany
 - jednoczesny, 614
 - prawa, 612, 613, 614
- dzielenie, 612
- dziennika, 637
- EXE, 114
- haszowany, 601, 605
- indeks, 603
 - częściowy, 605
 - drzewiasty, 605
 - pełny, 605
- indeksowy, 601, 604
- indeksu, 604
- katalog, *Patrz:* katalog plików
- katalogowy, 625
- lokowanie, 616
- MFT2, 637
- nadmiarowy, 603
- nazwa symboliczna, 612
- otwieranie, 595
- pisanie, 595
- potoku, 626
- przeszukiwanie pełne, 601
- rejestrowy, 603
- sekwencyjny, 601, 603
 - indeksowany, 601, 603, 604, 605
- struktura, 599, 600, 601, 603
 - sterta, 601, 603
- system zarządzania, *Patrz:* system zarządzania plikami, system plików
- ścieżka, 612
- tablica główna, *Patrz:* MFT
- tablica przydziału, *Patrz:* FAT
- transakcji, *Patrz:* plik rejestrowy
- tworzenie, 595, 610
- umiejscawianie, *Patrz:* plik lokowanie
- usuwanie, 595, 610
- w systemie Android, 639

w systemie Linux, 630

w systemie UNIX, 625

lokalizacja, 628, 632

w systemie Windows, 634

właściciel, 614

współużytkowanie, *Patrz:* plik dzielenie

wyszukiwanie, 610

zamykanie, 595

płyta

główna, 651

rozszerzeń, 651

płytką drukowaną, 651

pole, 595, 599

polecenie

chroot, 706

clone, 226

fork, 226

receive, 284, 285

send, 284, 285

polityka

bezpieczeństwa, 716

czyszczenia, 433

LOOK, 562

przydziału, 426, 427

VSWS, 432

zapisywania, 60

potok, 335, 337

nazwany, 335, 626

nienazwany, 335

powłoka, 119

praca, 147, 194

praca krokowa, *Patrz:* TF

prawo

Amdahla, 207

Moore'a, 681, 696

preemption point, *Patrz:* punkt wywłaszczania

primary memory, *Patrz:* pamięć główna

primary token, *Patrz:* żeton podstawowy

principle of locality, *Patrz:* zasada lokalności

printed circuit board, *Patrz:* płytką drukowaną

priority ceiling, *Patrz:* planowanie pułap

priorytetów

priority inheritance, *Patrz:* planowanie

dziedziczenie priorytetów

priority inversion, *Patrz:* planowanie odwrócenie

priorytetów

problem

czytelników i pisarzy, 290, 291

obiadujących filozofów, 331, 332, 333

ograniczonego buforowania, 803

- Pathfindera, 522, 525
- producenta-konsumenta, 269, 274, 290
- zakładu fryzjerskiego, 805
- procedura
 - lokalna wywołanie, *Patrz:* ALPC
 - obsługi przerwania, 45, 47, 49
 - wywołanie zdalne, *Patrz:* RPC
- proces, 93, 102, 113, 118, 144, 145, 146, 223, 232
 - atrybuty, 166, 168
 - blok kontrolny, 146, 150, 172, 173
 - czasu rzeczywistego, *Patrz:* zadanie
 - drugoplanowy, 233
 - identyfikator numeryczny, 154, 165, 166, 168, 174
 - informacje kontrolne, 170
 - interakcja, 255, 256, 258
 - izolowanie, 96, 107
 - jednowątkowy, 196
 - kolejka, 100, 150, 153
 - kontekst wykonywania, 95
 - kończenie, 152, 153, 154
 - kooperacja
 - oparta na dzieleniu, 258
 - oparta na komunikacji, 259
 - lekki, 194, 219
 - LWP, 219
 - macierzysty, 152
 - mieszanka, 461
 - mnożenie, 152
 - model dwustanowy, 150
 - obraz, 165
 - pierwszoplanowy, 233
 - planowanie, *Patrz:* planowanie procesów
 - położenie, 165
 - potomny, 152
 - zawieszenie, 162
 - przełączanie, 174, 177
 - przynależność, 194
 - przywileje, 712
 - pusty, 233
 - rodzenie, 152
 - rywalizacja, *Patrz:* rywalizacja
 - stan, 95
 - gotowy, 153, 155, 156, 157, 158, 160, 161
 - nowy, 154, 155, 160
 - wyjściowy, 154, 155, 156, 161
 - wykonywany, 153, 155, 157
 - zablokowany, 154, 155, 156, 157, 158, 160, 198
 - zawieszony, 158, 160, 161, 162
 - sterowanie, 150, 163, 165, 166, 168, 172, 173
 - UNIX SVR, 181, 183, 184, 185, 186
 - synchronizacja, 283, 284
 - ścieżka, 194
 - śląd, 147
 - tablica
 - odwrócona stron, 405, 406
 - segmentów, 413
 - stron, 382, 402
 - tryb, 176, 177
 - tworzenie, 151, 154, 174, 196, 218, 226
 - usługowy, 233
 - uwidoczniony, 233
 - w Linuxie, *Patrz:* Linux proces
 - wielowątkowy, 196, 197
 - Windows, *Patrz:* Windows proces
 - własność, *Patrz:* proces przynależność
 - wprowadzenie ponowne, 461
 - wymiana, 158
 - wymiany, 182
 - zawieszanie, 435
- procesor, 38, 41, 81, 84
 - aplikacji, 650
 - dedykowany, 650
 - fizyczny, *Patrz:* pCPU
 - fragmentacja, 507
 - graficzny, *Patrz:* GPU
 - logiczny, 40
 - powinowactwo, 531, 533, 536
 - przepustowość, 463
 - przydział czasu, 462
 - specjalizowany, 496
 - sygnałowy, *Patrz:* DSP
 - szybkość, 61
 - tryb, 173
 - wielordzeniowy, 652
 - wirtualizacja, *Patrz:* wirtualizacja procesora
 - wirtualny, *Patrz:* vCPU
 - wyspecjalizowany wydzielony, *Patrz:* procesor dedykowany
- process image, *Patrz:* proces obraz
- process mix, *Patrz:* proces mieszanka
- process spawning, *Patrz:* proces mnożenie
- process state, *Patrz:* proces stan
- processor affinity, *Patrz:* procesor powinowactwo
- processor fragmentation, *Patrz:* procesor fragmentacja
- processor status register, *Patrz:* PSR
- profil, 118
- program, 40
 - biblioteczny, 79
 - konsolidacja dynamiczna, 395

program

licznik, *Patrz:* PC

łączący, 394

narzędziowy, 79, 80, 81

obsługi przerwania, *Patrz:* ISR

rozruchowy, 655

sekcja krytyczna, 257, 264, 339, 350

słowo stanu, *Patrz:* PSW

sterowanie, 41

systemowy, 79

uruchomieniowy, 80

użytkowy, 81

program counter, *Patrz:* PCprogram status word, *Patrz:* PSW

programista aplikacji, 79

programmed I/O, *Patrz:* operacja wejścia-wyjścia

programowanego

programowanie

modularne, 97

współbieżne, 332

projektowanie obiektowe, 102, 104, 116

protection ring, *Patrz:* pierścień ochrony

protokół

dziedziczenia priorytetów, 523

HTTP, 771

NFS, 771

pułapu priorytetów, 523, 525

sieciowy, 126

SMB, 771

przerwanie, 43, 45, 86, 93, 118, 175, 176

awaria sprzętu, 43

blokowanie, 260, 261, 263

czasomierza, 43

efekt zagnieżdżenia, 511

jako wątek, 223

nieobsłużone, 45, 51

obsługa, 45, 47

priorytet, 52

program obsługi, *Patrz:* ISR

programowe, 43

systemowe, 175

tablica rozdzielcza, *Patrz:* IDT

wejścia-wyjścia, 43, 175, 462

wielokrotne, 50, 51, 52

zablokowane, 51

zegarowe, 175, 462, 471

żądanie, 45

przestój, 105

przestrzeń nazw Linux, *Patrz:* Linux przestrzeń nazw

przetwarzanie

asynchroniczne, 197

krawędziowe, *Patrz:* przetwarzanie na obrzeżu na obrzeżu, 783

rozproszone, 244

seryjne, 83

w chmurze, *Patrz:* chmura obliczeniowa

wsadowe, 90, 461

z kwantowaniem czasu, 91

z podziałem czasu, 90, 92, 93

uniwersalnym, 93

przewijak taśmy, 580

PSR, 173

PSW, 49, 168

PTE, 402

pułapka, 176

punkt

kontrolny, 107

wywłaszczania, 529

R

RAID, 564

0, 565, 569

1, 570, 587

2, 571

3, 571, 572

4, 571, 572

5, 571, 572, 573, 587

6, 571, 572, 573

dostęp równoległy, 571, 572

koszt, 570

nadmiarowość, 570, 571

paskowanie danych, 572, 573

poziom, 564, 565, 568

realizacja, 587

rekonstrukcja danych, 572

szybkość, 569, 571, 572

rate monotonic scheduling, *Patrz:* RMS

RBAC, 727, 731, 732

RCU, 344

rdzeń, 40, 64, 650

styczny, 508

read-copy update, *Patrz:* RCUreal memory, *Patrz:* pamięć głównareal-time user, *Patrz:* użytkownik czasu

rzeczywistego

redundant array of independent disks, *Patrz:* RAID

rejestr, 64
 adresowy
 pamięci, *Patrz:* MAR
 we-wy, 39
 bazowy, 380
 bufora we-wy, 39
 EFLAGS, 168, 169
 graniczny, 380
 pamięci
 adresowy, *Patrz:* MAR
 buforowy, *Patrz:* MBR
 rozkazów, *Patrz:* IR
 stanu
 wejścia-wyjścia, 60
 procesora, *Patrz:* PSR
 rekord, 595, 596, 599, 615
 blokowanie, *Patrz:* blok długość
 pole kluczowe, 603
 reliability, *Patrz:* SO niezawodność
 replacement algorithm, *Patrz:* algorytm
 zastępowania
 replacement policy, *Patrz:* zasada zastępowania
 residence time, *Patrz:* czas pobytu
 resident set, *Patrz:* zbiór rezydujący
 resource allocation graph, *Patrz:* graf przydziału
 zasobów
 resume flag, *Patrz:* RF
 reusable resource, *Patrz:* zasoby nieużywalne
 RF, 170
 RIOT, 790
 biblioteka, 792
 HAL, 792
 jądro, 791
 planista, 791
 RMS, 519, 520
 root access, *Patrz:* dostęp korzeniowy
 root file system, *Patrz:* system plików korzeniowy
 rotacja, 471, 472
 rotational positional sensing, *Patrz:* RPS
 round robin, *Patrz:* rotacja
 Rozdzielnia, 779
 rozkaz, 40
 cykl, 41, 45
 faza
 pobierania, 41
 przerwania, 45
 wykonywania, 41
 przetwarzanie, 40
 RPC, 113, 115, 199, 352
 RPS, 557, 557

rywalizacja, 256, 257
 szkodliwa, 246, 254

S

SaaS, 756, 761
 SACL, 744, 745
 Sahara, 780
 SAN, 769, 770
 sandboxing model, *Patrz:* model piaskownicy
 scheduling, *Patrz:* planowanie
 secondary memory, *Patrz:* pamięć drugorzędna
 seek time, *Patrz:* dysk czas wyszukiwania
 sekcja krytyczna, 245
 program, *Patrz:* program sekcja krytyczna
 selection function, *Patrz:* funkcja wyboru
 semafor, 263, 264, 269, 275, 277, 279, 291, 332, 345,
 462, 800
 binarny, 264, 265, 267, 271, 342
 czytelniczy – pisarze, 342, 343
 kolejka, 336
 mocny, 267
 ogólny, 267, 274
 słaby, 267
 wartość, 336
 zliczający, 267, 274, 342, 343, 347
 semaphore object, *Patrz:* obiekt semaforów
 serwer, 102
 plików, 197
 skomercjalizowany, 766
 wirtualny, 707
 SF, 169
 shared memory, *Patrz:* pamięć dzielona
 shell, *Patrz:* powłoka
 shortest process next, *Patrz:* algorytm SPN
 shortest remaining time, *Patrz:* algorytm SRT
 shortest-service-time-first, *Patrz:* algorytm SSTF
 sieć
 bezprzewodowa osobistego zasięgu, *Patrz:*
 WPAN
 rdzenia, 784
 szkieletowa, 784
 sign flag, *Patrz:* SF
 SIMD, 40
 single indirect block, *Patrz:* blok jednopiętrowy
 skrzynka pocztowa, 264, 286, 335
 slack time, *Patrz:* luz czasowy
 słowo stanu programu, *Patrz:* PSW
 SMP, 62, 102, 103, 107, 115
 dostępność, 62
 organizacja, 63

SMP

pamięć podręczna, *Patrz:* pamięć podręczna SMP
skalowanie, 62
wydajność, 62

SO, 38, 45, 78, 80, 81

Android, *Patrz:* Android

architektura, 101, 102

bezpieczne działanie w warunkach awarii, 512

BSD, 124

chmurowy, 763

architektura, 766

komponent zarządzania i instrumentacji, *Patrz:*
MANO

o otwartym źródle, 772, 790

pamięć, 768, 769, 770

wirtualizacja, 766, 768, 769

wymagania, 765, 766

Clouds, 206

CTSS, 91

czas reakcji, 511

czasu rzeczywistego, 510, 511, 512, 513

deterministyczny, 511

dostępność, 105

Emerald, 206

ewolucja, 82, 101

guszczony, 681, 687

liczba, 682

hartowanie, 736, 737, 738, 740

IBSYS, 84

instalowanie, 737

IR, *Patrz:* SOIR

Linux, *Patrz:* Linux

MAC OS, *Patrz:* MAC OS

Mach, 504

MS-DOS, 110

niezawodność, 104, 511

oparty na procesach, 180

OpenSolaris, 124

OpenStack, *Patrz:* OpenStack

oprogramowanie, 106

podstawowy system plików, *Patrz:* system plików
podstawowy

projektowanie, 92, 93, 94, 95, 106, 109

reaktywność, 511

RIOT, *Patrz:* RIOT

rozproszony, 102, 104

rozwój, 83

Solaris, 205, 219

bezpieczeństwo, 723

proces, 219, 220

przerwanie, 223

wątek, 219, 220, 221, 222

wątek jądrowy, 219

współbieżność, *Patrz:* współbieżność Solaris

zarządzanie pamięcią, 436

sprawność, 100

stabilność, 512

testowanie bezpieczeństwa, 740

TRIX, 205

UNIX, *Patrz:* UNIX

w chmurze, 763

wbudowany, 648, 654

adaptacja, 657

jądro, 656

konstruowanie, 657

organizacja, 649

program rozruchowy, 655

system plików, 656

Windows, *Patrz:* Windows

wsadowy, 84, 86

wykonywanie, 178, 179

Xbox One, 110

zarządzanie procesami, *Patrz:* proces sterowanie

SoC, 40

socket, *Patrz:* układ gniazdowy

SOir, 790

SOIR, 785, 786

architektura, 789

bezpieczeństwo, 788

wymagania, 787, 788

spatial locality, *Patrz:* lokalność w przestrzeni

spin waiting, *Patrz:* czekanie wirujące

sprzężenie zwrotne wielopoziomowe, 477

SQLite, 641

stack overflow, *Patrz:* stos przepełnienie

stackguard, *Patrz:* stos strażnik

stan procesu, *Patrz:* proces stan

starvation, *Patrz:* głodzenie

sterta niewykonywalna, 724

storage area network, *Patrz:* SAN

stos

niewykonywalny, 724

przepełnienie, 717

obrona, 724

strażnik, 723

strategia zbioru roboczego, 429, 430, 431, 434

stream-oriented device, *Patrz:* urządzenie wejścia-

wyjścia strumieniowe

strong semaphore, *Patrz:* semafor mocny

stronicowanie, 82

struktura danych task_struct, 223
 superużytkownik, 735
 supervisor call, *Patrz:* wywołanie nadzorcy
 Surge, 671
 SVA, 688
 swapper process, *Patrz:* proces wymiany
 swapping, *Patrz:* proces wymiana
 Swift, 777
 sygnał, 336, 337, 462
 czasu rzeczywistego, *Patrz:* sygnał RT
 RT, 337
 żądania przerwania, 45
 symmetric multiprocessor, *Patrz:* SMP
 synchronization granularity, *Patrz:* wieloprocessor
 synchronizacja ziarnistość
 system
 antywłamaniowy, *Patrz:* IPS
 głęboko wbudowany, 653
 jednoukładowy, *Patrz:* SoC
 operacyjny, *Patrz:* SO
 plików, 551, 552, 594, 595, 599
 Android, 639
 architektura, 598
 blok danych, *Patrz:* blok
 cramfs, 661
 jffs2, 661
 kontenerowy, 693
 kontrola dostępu, 725
 korzeniowy, 656
 Linux, 630
 NTFS, *Patrz:* NTFS
 odtwarzalność, 638
 podstawowy, 598
 squashfs, 661
 ubifs, 661
 UNIX, 625, 628, 629, 630, 632
 Windows, 634
 wirtualny, *Patrz:* VFS
 yaffs2, 661
 transakcji, 93
 wykrywania włamań, *Patrz:* IDS
 zarządzania
 bazą danych, 596
 plikami, 594, 597, 598, 599
 system bus, *Patrz:* szyna systemowa
 System on Chip, *Patrz:* SoC
 szamotanina, 245, 249, 400, 434
 szeregowanie, *Patrz:* planowanie
 szyfrowanie, 587

szyna
 PCI Express, 64
 systemowa, 39, 61

T

tablica
 definicji atrybutów, 637
 główna plików, *Patrz:* MFT
 plików, 164
 procesów, 164
 przydziałów dysku, *Patrz:* dysk tablica przydziału
 segmentów pamięci, *Patrz:* pamięć segment
 tablica
 stron, 382, 402, 405, 406
 translacji, 698
 wejścia-wyjścia, 164
 target platform, *Patrz:* platforma docelowa
 task, *Patrz:* praca
 TAT, 467, 467
 TBF, 707, 708
 technika czujnik-aktywator, 781
 temporal locality, *Patrz:* lokalność w czasie
 terminal, 580, 581
 graficzny, 82
 niemy, 554
 TF, 170
 thrashing, *Patrz:* szamotanina
 thread, *Patrz:* wątek
 thread pool, *Patrz:* wątek pula
 throughput, *Patrz:* procesor przepustowość
 time sharing, *Patrz:* przetwarzanie z podziałem
 czasu
 time slicing, *Patrz:* czas porcjowanie, przetwarzanie
 z kwantowaniem czasu
 timer, *Patrz:* czasomierz
 time-sharing user, *Patrz:* użytkownik z podziałem
 czasu
 TinyOS, 665, 785, 790
 cele, 666
 interfejs, 673
 komponent, 668
 planista, 670, 671
 polecenie, 669
 TimerM, 670
 zasoby, 673, 674
 zdarzenie, 669
 TLB, 406, 407, 409
 efektywność, 412
 token bucket filter, *Patrz:* TBF

tolerowanie awarii, 104, 108
 tom, 624, 630
 Torvalds Linus, 123, 124
 trace, *Patrz:* proces ślad
 transakcja w czasie rzeczywistym, 93
 transfer time, *Patrz:* dysk czas przesyłania
 translation lookaside buffer, *Patrz:* TLB
 trap flag, *Patrz:* TF
 triple indirect block, *Patrz:* blok trójpośredni
 Trove, 779
 tryb
 jądra, 86, 110, 172, 204
 sterowania, *Patrz:* tryb jądra
 systemu, *Patrz:* tryb jądra
 użytkownika, 86, 110, 172
 planowanie, *Patrz:* UMS
 proces, 114
 turnaround time, *Patrz:* TAT
 tylne wejście, 713

U

uaktualnienie czytaj-kopiuj, *Patrz:* RCU
 układ
 gniazdowy, 40
 scalony, 652
 dopasowany do aplikacji, *Patrz:* ASIC
 UMS, 212
 unbounded priority inversion, *Patrz:* planowanie
 odwrócenie priorytetów nieograniczone
 uniprogramming, *Patrz:* jednoprogramowość
 UNIX
 architektura, 119, 120, 124
 historia, 118
 pamięć, 436, 438, 439
 planowanie, 486, 584
 system plików, *Patrz:* system plików UNIX
 tom, 630
 urządzenie wejścia-wyjścia, *Patrz:* urządzenie
 wejścia-wyjścia UNIX
 wersja
 BSD, 122
 FreeBSD, 530, 531
 Solaris, 122
 SVR4, 122, , 181, 183, 184, 185, 337, 528, 584
 współbieżność, 334
 zarządzanie procesami, 181, 183, 184, 185, 186
 uptime, *Patrz:* czas pracy
 urządzenie
 głęboko wbudowane, 781

 mobilne, 40
 ograniczone, 785, 786
 wejścia-wyjścia, 80, 544, 545, 550, 626
 blokowe, 553, 554, 580, 581
 dyskowe, *Patrz:* dysk
 strumieniowe, 553, 554
 UNIX, 578, 580
 znakowe, 581
 user mode, *Patrz:* tryb użytkownika
 user-level threads, *Patrz:* wątek ULT
 user-mode scheduling, *Patrz:* UMS
 użytkownik
 czasu rzeczywistego, 530
 przywileje, 712, 734, 739, 743
 uwierzytelnianie, 715
 z podziałem czasu, 531

V

variable-interval sampled working set, *Patrz:* zbiór
 roboczy próbkowany w zmiennych odstępach
 vCPU, 696
 VFS, 228, 630, 631
 obiekt, 631, 632, 633
 pamięć podręczna, 634
 virtual appliance, *Patrz:* aplikacja wirtualna
 virtual file system, *Patrz:* VFS
 virtual machine, *Patrz:* VM
 Virtual Machine Extensions, *Patrz:* VMX
 virtual machine monitor, *Patrz:* VMM
 virtual memory, *Patrz:* pamięć wirtualna
 virtual round robin, *Patrz:* rotacja wirtualna
 virtual time, *Patrz:* czas wirtualny
 virtualization container, *Patrz:* kontener
 wirtualizacji
 virtualization service provider, *Patrz:* VSP
 VM, 681, 684, *Patrz też:* wirtualizacja
 Javy, *Patrz:* JVM
 konkretyzacja, 684
 Linux VServer, 706, 707
 monitor, *Patrz:* VMM
 obraz, 773
 organizacja wejścia-wyjścia, 699, 700, 701
 pamięć, 697, 698, 699
 plik konfiguracyjny, 684
 szablon, 684
 VMkernel, 701, 702
 VMM, *Patrz:* hiperwizor
 VMX, 687
 VRR, *Patrz:* rotacja wirtualna

VSP, 704

VSWs, *Patrz:* zbiór roboczy próbkowany
w zmiennych odstępach

W

wada, 106

waitable timer object, *Patrz:* obiekt czasomierza do
czekania

warstwa

abstrakcji, 680

abstrakcji sprzętu, *Patrz:* HAL

ramy aplikacji, 129, 130

tłumaczenia oprogramowania, *Patrz:* warstwa
abstrakcji

zastosowań, 129

wątek, 95, 102, 113, 115, 118, 194, 196, 212, 232

blokowanie, 203

czasu rzeczywistego, 534

interaktywny, 532

jądrowy, 219

KLT, 200, 203, 204

kombinowany, 204

planowanie, *Patrz:* planowanie wątków

połowy jądra, 530

POSIX-owy, 225

poziomu jądra, *Patrz:* wątek KLT

poziomu użytkownika, *Patrz:* wątek ULT

priorytet, 534

przełączanie, 203

pula, 109, 212, 213, 233

stan, 198, 217

czekania, 217

działania, 217

gotowy, 217

kończenie, 198

odblokowanie, 198

pogotowia, 217

przejścia, 217

rozmnażanie, 198

zablokowanie, 198

zakończony, 218

tworzenie, 196

ULT, 200, 203, 204, 219

w Linuxie, *Patrz:* Linux wątek

w wielu domenach, 205

wątkowość

drobnoziarnista, 210

gruboziarnista, 210

hybrydowa, 210

weak semaphore, *Patrz:* semafor słaby

wideo, 40

wieloprocessor, 197, 496, 536

jednoukładowy, 64

planowanie, 498, 536

powiązany

luźno, *Patrz:* wieloprocessor rozproszony

ściśle, 496

rozproszony, 496

równoległość

niezależna, 497

ziarnistość drobna, 498

ziarnistość gruba, 497

ziarnistość średnia, 498, 500

symetryczny, *Patrz:* SMP

synchronizacja, 497

z niejednocelowym dostępem do pamięci, 419

wieloprogramowość, 87, 89, 90, 92, 93, 103, 110,
199, 244

jednoprocessorowa, 251, 252

planowanie, *Patrz:* planowanie

stopień, 434, 435

wieloprocessorowa, 251, 252

wieloprzetwarzanie, 197, 198, 244

symetryczne, *Patrz:* SMP

wielordzeniowość, 207, 209, 536

wydajność, 207

wielowątkowość, 101, 102, 103, 195, 197, 198, 205,
207, 209, 217

wielozadaniowość, 87

Windows, 110

aplikacja, *Patrz:* aplikacja w systemie Windows

architektura, 111

egzekutor, 111

jądro, *Patrz:* jądro Windows

klaster mapa bitowa, 637

moduł sterujący urządzenia, 111

obiekt, *Patrz:* obiekt

pamięć, 445, 447

podsystem, 218

powiadomienia wypychane, 214

proces, 214, 215

tworzenie, 218

system plików, 634

warstwa abstrakcji sprzętu, *Patrz:* HAL

wersja 10, 110

współbieżność, *Patrz:* współbieżność Windows
zarządca, 113

Windows Live Tiles, 213

Windows Notification Service, *Patrz:* WNS

Windows Server, 587
 Windows Store, 447
 wireless personal area networks, *Patrz:* WPAN
 wireless sensor network, *Patrz:* czujnik
 bezprzewodowy sieć
 wirtualizacja, 680, 681, 682, 684, *Patrz też:* VM
 dostawca usług, *Patrz:* VSP
 kontenerowa, 688
 lekka, 226
 liczba procesorów, 696
 połowiczna, *Patrz:* parawirtualizacja
 procesora, 696
 sprzętu, 680
 z udziałem sprzętu, 687
 wirus, 713
 włókno, 212
 WNS, 214
 wolumin, *Patrz:* tom
 WPAN, 786
 wpis tablicy stron, *Patrz:* PTE
 write policy, *Patrz:* polityka zapisywania
 WSN, *Patrz:* czujnik bezprzewodowy sieć
 współbieżność, 107, 108, 217, 244, 245, 255, 332, 798
 Android, 352
 Linux, 337, 338, 339, 341, 342, 343, 344
 Solaris, 345
 UNIX, 334
 Windows, 348, 349, 351
 współczynnik
 konsolidacji, 682
 trafień H, 53, 59, 74
 wahań, 500
 współplanowanie, 504
 współprocedura, 248
 wycofanie, 107
 wykluczenie
 wzajemne, 245, 246, 250, 257, 258, 269, 288, 290
 blokada wirująca, 264
 podejście sprzętowe, 260, 261, 263, 275
 wyścigi, 798, 799
 wyścigi, *Patrz:* rywalizacja szkodliwa
 wywłaszczanie natychmiastowe, 513
 wywołanie nadzorcy, 176

Z

zadanie, 83, 84, 510
 czas przygotowania, 84
 czasu rzeczywistego rygorystycznego, 510
 język sterowania, *Patrz:* JCL

nieokresowe, 510
 termin nieprzekraczalny, 510
 zakleszczenie, 94, 245, 257, 258, 312, 313, 315, 552
 unikanie, 318, 321, 322, 323, 330, 331
 odmowa przydziału zasobu, *Patrz:* algorytm
 bankiera
 odmowa wszczynania procesu, 323
 stan bezpieczny, 324
 stan zagrożenia, 324
 warunki występowania, 320
 brak wywłaszczeń, 320, 322
 czekanie cykliczne, 320, 322
 przetrzykiwanie i oczekiwanie, 320, 321
 wzajemne wykluczanie, 320, 321
 wykrywanie, 318, 321, 328
 algorytm, *Patrz:* algorytm wykrywania
 zakleszczenia
 rekonstrukcja, 329
 zapobieganie, 318, 321, 330
 zasoby niezużywalne, 316, 317
 zasoby zużywalne, 317
 zamek wzajemnego wykluczania, *Patrz:* muteks
 zapor sieciowa, 716, 740
 Zaqar, 779
 zarządca
 pamięci, 638
 wejścia-wyjścia, 585, 586
 zasada
 FCFS, 470
 lokalności *Patrz:* lokalność odwołań
 najmniejszych przywilejów, 732
 pierwszy zgłoszony – pierwszy obsłużony, 470
 zastępowania, 419, 426, 427, 428
 zasobnik wirtualizacji, *Patrz:* kontener wirtualizacji
 zasoby, 258
 dzielenie, 508
 graf przydziału, 318, 319
 gromadzenie, 755
 krytyczne, 257
 niezużywalne, 316
 zużywalne, 317
 zbiór
 rezydujący, 399, 419, 426, 427
 roboczy, 429, 430, 431, 434
 aktywności, 507
 próbkowany w zmiennych odstępach, 432
 rozmiar, 430
 zero flag, *Patrz:* ZF
 ZF, 169

zmienna

globalna, 252, 253

warunkowa, 345, 348

wieku, 443

znacznik

dopuszczalności przerwania, *Patrz:* IF

identyfikacji, *Patrz:* ID

kierunku, *Patrz:* DF

kontrolny, 169

nadmiaru, *Patrz:* OF

parzystości, 169

pomocniczego przeniesienia, *Patrz:* AF

przeniesienia, *Patrz:* CF

przepełnienia, *Patrz:* OF

pułapki, *Patrz:* TF

systemowy, 170

wznowienia, *Patrz:* RF

zagnieżdżonego zadania, *Patrz:* NT

zdarzeń, 264

zera, *Patrz:* ZF

znaku, *Patrz:* SF

Ż

żądanie przerwania, 45

żeton, 778

dostępu, 743

podstawowy, 214

wiaderko, 707

μ

μClibc, 664

μCLIBC, 663

μClinux, 662, 664

PROGRAM PARTNERSKI

GRUPY WYDAWNICZEJ 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 WYDAWNICZA

 **Helion SA**

SYSTEMY OPERACYJNE CAŁY CZAS SĄ WZBOGACANE O INNOWACJE I ULEPSZENIA. STAJĄ SIĘ TEŻ CORAZ BARDZIEJ WYSPECJALIZOWANE, CO WYNIKA Z ROSNĄCEJ RÓŻNORODNOŚCI MASZYN, KTÓRE MAJĄ OBSŁUŻYĆ. WYSTAR-CZY TU WSKAZAĆ SYSTEMY WBUDOWANE, SMARTFONY, KOMPUTERY OSOBISTE I KOMPUTERY GŁÓWNE (ANG. MAINFRAME) ORAZ SUPERKOMPUTERY, A TAKŻE SPECJALNE SYSTEMY CZASU RZECZYWISTEGO. ZNAJOMOŚĆ WEWNĘTRZNYCH MECHANIZMÓW SYSTEMU I ARCHITEKTURY JĄDRA OKAZUJE SIĘ ZATEM NIEZWYKLE ISTOT-NA DLA PROGRAMISTÓW I INŻYNIERÓW. BEZ TEJ WIEDZY TRUDNO MÓWIĆ O NIEZAWODNOŚCI TWORZONEGO OPROGRAMOWANIA.

TA KSIĄŻKA JEST KOLEJNYM, GRUNTOWNIE PRZEJRZANYM I ZAKTUALIZOWANYM WYDANIEM KLASYCZNEGO PODRĘCZNIKA, W KTÓRYM JASNO I WYCZERPUJĄCO WYJAŚNIONO KONCEPCJE, STRUKTURĘ I MECHANIZMY RZĄDZĄCE FUNKCJONOWANIEM NOWOCZESNYCH SYSTEMÓW OPERACYJNYCH. WNIKLIVIE OMÓWIONO RÓWNIEŻ PODSTAWOWE ZASADY PROJEKTOWANIA SYSTEMÓW OPERACYJNYCH I POWIĄZANO JE ZE WSPÓŁ-CZESNYMI ZAGADNIENIAMI PROJEKTOWYMI ORAZ KIERUNKAMI ROZWOJU SYSTEMÓW OPERACYJNYCH. ABY ZILUSTROWAĆ PREZENTOWANE TREŚCI, ODNIESIONO SIĘ DO CZTERECH SYSTEMÓW: WINDOWS, ANDROID, UNIX I LINUX. W TEN SPOSÓB KONCEPCJE PROJEKTOWE OMAWIANE W DANYM ROZDZIALE SĄ NATYCHMIAST POPIERANE RZECZYWISTYMI PRZYKŁADAMI.

NAJWAŻNIEJSZE ZAGADNIENIA:

- PRZEGLĄD SYSTEMÓW OPERACYJNYCH
- WSPÓŁBIEŻNOŚĆ I ROZPROSZONE ZARZĄDZANIE PROCESAMI
- ZARZĄDZANIE PAMIĘCIĄ I PAMIĘĆ WIRTUALNA
- BEZPIECZEŃSTWO SYSTEMÓW OPERACYJNYCH
- OPERACJE WEJŚCIA-WYJŚCIA I ZARZĄDZANIE PRZESTRZENIĄ DYSKOWĄ
- ODPORNOŚĆ NA AWARIE

DR WILLIAM STALLINGS JEST AUTOREM WIELOKROTNIENAGRADZANYCH KSIĄŻEK POŚWIĘCONYCH BEZPIECZEŃSTWU SYSTEMÓW INFORMATYCZNYCH, SIECIOM I ARCHITEKTURZE KOMPUTERÓW. PROJEKTOWAŁ I IMPLEMENTOWAŁ ZESTAWY PROTOKOŁÓW SIECIOWYCH NA RÓŻNE KOMPUTERY I SYSTEMY OPERACYJNE. DORADZAŁ AGENCJOM RZĄDOWYM ORAZ DOSTAWCOM KOMPUTERÓW I OPROGRAMOWANIA W KWESTIACH PROJEKTOWANIA I UŻYTKOWANIA OPROGRAMOWANIA I PRODUKTÓW SIECIOWYCH. JEST CZŁONKIEM KOLEGIUM REDAKCYJNEGO CZASOPISMA „CRYPTOLOGIA”, NAUKOWEGO PERIODYKU POŚWIĘCONEGO ZAGADNIE-NIOM SZYFROWANIA.

SYSTEMY OPERACYJNE: POZNAJ I ZAPROJEKTUJ!

Helion 	Sprawdź nasze szkolenia!  AKADEMIA IT & BUSINESS WWW.SZKOLENIA.HELION.PL	KOD KORZYŚCI Sięgnij po więcej! 
 helion.pl		ISBN 978-83-283-3759-6
 0 801 339900		
 0 601 339900		9 788328 337596
INFORMATYKA W NAJLEPSZYM WYDANIU		Cena: 129,00 zł