

Itzik Ben-Gan

Podstawy języka T-SQL

Microsoft SQL Server 2016
i Azure SQL Database

Przekład: Leszek Biolik, Marek Włodarz

APN Promise, Warszawa 2016

Podstawy języka T-SQL: Microsoft SQL Server 2016 i Azure SQL Database

Authorized Polish translation of the English language edition entitled
T-SQL Fundamentals, Third Edition, by Itzik Ben-Gan
ISBN: 978-1-5093-0200-0

Copyright © 2016 by Itzik Ben-Gan

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 APN PROMISE SA Copyright © 2016

Autoryzowany przekład z wydania w języku angielskim, zatytułowanego:
T-SQL Fundamentals, Third Edition, by Itzik Ben-Gan
ISBN: 978-1-5093-0200-0

Wszystkie prawa zastrzeżone. Żadna część niniejszej książki nie może być powielana ani rozpowszechniana w jakiejkolwiek formie i w jakikolwiek sposób (elektroniczny, mechaniczny), włącznie z fotokopiowaniem, nagrywaniem na taśmy lub przy użyciu innych systemów bez pisemnej zgody wydawcy.

APN PROMISE SA, ul. Domaniewska 44a, 02-672 Warszawa
tel. +48 22 35 51 600, fax +48 22 35 51 699
e-mail: mspress@promise.pl

Książka ta przedstawia poglądy i opinie autora. Przykłady firm, produktów, osób i wydarzeń opisane w niniejszej książce są fikcyjne i nie odnoszą się do żadnych konkretnych firm, produktów, osób i wydarzeń chyba że zostanie jednoznacznie stwierdzone, że jest inaczej. Ewentualne podobieństwo do jakiejkolwiek rzeczywistej firmy, organizacji, produktu, nazwy domeny, adresu poczty elektronicznej, logo, osoby, miejsca lub zdarzenia jest przypadkowe i niezamierzone.

Microsoft oraz znaki towarowe wymienione na stronie <http://www.microsoft.com/about/legal/en/us/IntellectualProperty/Trademarks/EN-US.aspx> są zastrzeżonymi znakami towarowymi grupy Microsoft. Wszystkie inne znaki towarowe są własnością ich odnośnych właścicieli.

APN PROMISE SA dołożyła wszelkich starań aby zapewnić najwyższą jakość tej publikacji. Jednakże nikomu nie udziela się rękojmi ani gwarancji. APN PROMISE SA nie jest w żadnym wypadku odpowiedzialna za jakiegokolwiek szkody będące następstwem korzystania z informacji zawartych w niniejszej publikacji, nawet jeśli APN PROMISE została powiadomiona o możliwości wystąpienia szkód.

ISBN: 978-83-7541-305-2 (druk), 978-83-7541-3

Przekład: Leszek Biolik, Marek Włodarz
Redakcja: Marek Włodarz
Korekta: Ewa Swędrowska
Skład i łamanie: MAWart Marek Włodarz

Dla Dato

*Żć w sercach tych, których zostawiliśmy,
to nie umrzeć*

– Thomas Campbell

Spis treści

<i>Wprowadzenie</i>	xiii
<i>Podziękowania</i>	xvii
1. Podstawy zapytańi programowania T-SQL	1
Podstawy teoretyczne.....	1
SQL	3
Teoria zbiorów.....	4
Logika predykatów.....	5
Model relacyjny	6
Typy systemów bazodanowych	13
Architektura SQL Server.....	15
Odmiany ABC produktu SQL Server.....	15
Instancje produktu SQL Server.....	18
Bazy danych	19
Schematy i obiekty.....	23
Tworzenie tabel i definiowanie integralności danych.....	24
Tworzenie tabel	25
Definiowanie integralności danych.....	27
Podsumowanie	31
2. Zapytania do pojedynczej tabeli.....	33
Elementy instrukcji <i>SELECT</i>	33
Klauzula <i>FROM</i>	36
Klauzula <i>WHERE</i>	38
Klauzula <i>GROUP BY</i>	39
Klauzula <i>HAVING</i>	43
Klauzula <i>SELECT</i>	44
Klauzula <i>ORDER BY</i>	49
Filtry <i>TOP</i> i <i>OFFSET-FETCH</i>	52
Szybki przegląd funkcji okna	56
Predykaty i operatory.....	58
Wyrażenia <i>CASE</i>	61
Znacznik <i>NULL</i>	64

Operacje jednoczesne – „all-at-once”	69
Stosowanie danych znakowych	71
Typy danych	71
Opcje sortowania (<i>collation</i>)	72
Operatory i funkcje	75
Predykat <i>LIKE</i>	84
Posługiwanie się danymi typu daty i czasu	87
Typy danych dotyczące daty i czasu	87
Literały	88
Rozdzielne stosowanie daty i czasu	92
Filtrowanie zakresów danych	94
Funkcje daty i godziny	95
Zapytania dotyczące metadanych	106
Widoki katalogowe	106
Informacyjne widoki schematu	107
Systemowe procedury składowane i funkcje	108
Podsumowanie	109
Ćwiczenia	110
Rozwiązania	115
3. Złączenia	121
Złączenia krzyżowe	122
Składnia ISO/ANSI SQL-92	122
Składnia ISO/ANSI SQL-89	123
Samo-złączenie krzyżowe (Self Cross Join)	123
Tworzenie tabel liczb	124
Złączenia wewnętrzne	126
Składnia ISO/ANSI SQL-92	126
Składnia ISO/ANSI SQL-89	127
Bezpieczeństwo złączenia wewnętrznego	128
Dodatkowe rodzaje złączeń	129
Złączenia złożone	129
Złączenie nierównościowe (Non-Equi Join)	130
Złączenia wielokrotne (multi-join)	132
Złączenia zewnętrzne	133
Podstawy złączeń zewnętrznych	133
Złączenia zewnętrzne – zagadnienia zaawansowane	136
Podsumowanie	144

Ćwiczenia	144
Ćwiczenie 1-2 (zaawansowane ćwiczenie opcjonalne)	145
Ćwiczenie 7 (zaawansowane ćwiczenie opcjonalne)	147
Ćwiczenie 8 (zaawansowane ćwiczenie opcjonalne)	148
Ćwiczenie 9 (zaawansowane ćwiczenie opcjonalne)	148
Rozwiązania	149
4. Podzapytania	155
Podzapytania niezależne	155
Przykłady skalarnych podzapytań niezależnych	156
Podzapytania niezależne o wielu wartościach	158
Podzapytania skorelowane	162
Predykat <i>EXISTS</i>	165
Zaawansowane aspekty podzapytań	167
Zwracanie poprzednich lub kolejnych wartości	167
Agregacje bieżące	168
Postępowanie w przypadku nieprawidłowo działających podzapytań	169
Podsumowanie	174
Ćwiczenia	175
Ćwiczenie 2 (zaawansowane ćwiczenie opcjonalne)	175
Ćwiczenie 7 (zaawansowane ćwiczenie opcjonalne)	177
Ćwiczenie 8 (zaawansowane ćwiczenie opcjonalne)	178
Ćwiczenie 10 (zaawansowane ćwiczenie opcjonalne)	178
Rozwiązania	179
5. Wyrażenia tablicowe	185
Tabele pochodne	185
Przypisywanie aliasów kolumn	187
Stosowanie argumentów	189
Zagnieżdżanie	190
Wielokrotne odwołania	191
Współne wyrażenia tablicowe	192
Przypisywanie aliasów kolumn w wyrażeniach CTE	192
Stosowanie argumentów w wyrażeniach CTE	193
Definiowanie wielu wyrażeń CTE	193
Wielokrotne odwołania w wyrażeniach CTE	194
Rekurencyjne wyrażenia CTE	195
Widoki	198

Widoki i klauzula <i>ORDER BY</i>	199
Opcje widoku	202
Wbudowane funkcje zwracające tabele	206
Operator <i>APPLY</i>	207
Podsumowanie	211
Ćwiczenia	211
Ćwiczenie 4 (zaawansowane ćwiczenie opcjonalne)	213
Ćwiczenie 5-2 (zaawansowane ćwiczenie opcjonalne)	215
Rozwiązania	216
6. Operatory zbiorowe	221
Operator <i>UNION</i>	222
Operator wielozbioru <i>UNION ALL</i>	222
Operator zbiorowy <i>UNION</i> z niejawną opcją <i>Distinct</i>	223
Operator <i>INTERSECT</i>	224
Operator <i>INTERSECT</i> (z ukrytą opcją <i>Distinct</i>)	225
Operator wielozbioru <i>INTERSECT ALL</i>	225
Operator <i>EXCEPT</i>	228
Operator zbiorowy <i>EXCEPT</i> (z opcją <i>Distinct</i>)	228
Operator wielozbioru <i>EXCEPT ALL</i>	229
Pierwszeństwo	230
Omijanie nieobsługiwanych faz logicznych	232
Podsumowanie	234
Ćwiczenia	234
Ćwiczenie 6 (zaawansowane ćwiczenie opcjonalne)	236
Rozwiązania	237
7. Zaawansowane zagadnienia tworzenia zapytań	241
Funkcje okna	241
Rankingowe funkcje okna	244
Offsetowe funkcje okna	248
Agregujące funkcje okna	251
Przestawianie danych	254
Przestawianie danych przy użyciu zapytania grupującego	256
Przestawianie danych przy użyciu operatora <i>PIVOT</i>	257
Odwrotne przestawianie danych	260
Odwrotne przestawianie danych przy użyciu operatora <i>APPLY</i>	261
Odwrotne przestawianie danych za pomocą operatora <i>UNPIVOT</i>	264

Zbiory grupujące	265
Klauzula pomocnicza <i>GROUPING SETS</i>	267
Klauzula pomocnicza <i>CUBE</i>	267
Klauzula pomocnicza <i>ROLLUP</i>	268
Funkcje <i>GROUPING</i> i <i>GROUPING_ID</i>	269
Podsumowanie	272
Ćwiczenia	272
Rozwiązania	277
8. Modyfikowanie danych	281
Wstawianie danych	281
Wyrażenie <i>INSERT VALUES</i>	281
Instrukcja <i>INSERT SELECT</i>	283
Instrukcja <i>INSERT EXEC</i>	284
Instrukcja <i>SELECT INTO</i>	285
Instrukcja <i>BULK INSERT</i>	286
Właściwość <i>Identity</i> i obiekt sekwencji	286
Usuwanie danych	296
Instrukcja <i>DELETE</i>	297
Instrukcja <i>TRUNCATE</i>	297
<i>DELETE</i> oparte na złączeniu	299
Aktualizowanie danych	300
Instrukcja <i>UPDATE</i>	301
<i>UPDATE</i> oparte na złączeniu	302
<i>UPDATE</i> z przypisaniem	305
Scalanie danych	306
Modyfikowanie danych przy użyciu wyrażeń tablicowych	311
Modyfikacje przy użyciu opcji <i>TOP</i> i <i>OFFSET-FETCH</i>	313
Klauzula <i>OUTPUT</i>	316
<i>INSERT</i> z klauzulą <i>OUTPUT</i>	316
<i>DELETE</i> z klauzulą <i>OUTPUT</i>	318
<i>UPDATE</i> z klauzulą <i>OUTPUT</i>	319
<i>MERGE</i> z klauzulą <i>OUTPUT</i>	320
Zagnieżdżone wyrażenia DML	321
Podsumowanie	323
Ćwiczenia	323
Rozwiązania	327

9.	Tabele temporalne	333
	Tworzenie tabel	334
	Modyfikowanie danych	338
	Odpytywanie danych	341
	Podsumowanie	348
	Ćwiczenia	348
	Rozwiązania	351
10.	Transakcje i współbieżność	357
	Transakcje	357
	Blokowanie	361
	Blokady	361
	Rozwiązywanie problemów związanych z blokadami	364
	Poziomy izolacji	372
	Poziom izolacji <i>READ UNCOMMITTED</i>	374
	Poziom izolacji <i>READ COMMITTED</i>	375
	Poziom izolacji <i>REPEATABLE READ</i>	377
	Poziom izolacji <i>SERIALIZABLE</i>	378
	Poziomy izolacji oparte na wersjonowaniu wierszy	380
	Podsumowanie poziomów izolacji	387
	Zakleszczenia	388
	Podsumowanie	391
	Ćwiczenia	391
11.	Obiekty programowalne	403
	Zmienne	403
	Wsady	406
	Wsad jako jednostka analizy	406
	Wsady i zmienne	407
	Instrukcje, których nie można łączyć w tym samym wsadzie	408
	Wsad jako jednostka rozpoznawania	408
	Opcja <i>GO n</i>	409
	Elementy kontroli przepływu wykonania	410
	Element kontroli przepływu <i>IF ... ELSE</i>	410
	Element kontroli przepływu <i>WHILE</i>	411
	Kursory	413
	Tabele tymczasowe	417

Lokalne tabele tymczasowe	417
Globalne tabele tymczasowe	419
Zmienne tablicowe	421
Typy tablicowe	422
Dynamiczny kod SQL	423
Polecenie <i>EXEC</i>	424
Procedura składowana <i>sp_executesql</i>	424
<i>PIVOT</i> w dynamicznym kodzie SQL	426
Procedury	427
Funkcje definiowane przez użytkownika	428
Procedury składowane	429
Wyzwalacze	432
Obsługa błędów	436
Podsumowania	440
A. Rozpoczynamy	441
Rozpoczynamy pracę w Azure SQL Database	442
Instalowanie produktu SQL Server w wersji dla siedziby	442
Ćwiczenie 1. Uzyskanie produktu SQL Server	442
Ćwiczenie 2. Instalowanie silnika bazy danych	443
Pobieranie i instalowanie SQL Server Management Studio	448
Pobieranie kodu źródłowego i instalowanie przykładowej bazy danych	448
Posługiwanie się programem SQL Server Management Studio	451
Korzystanie z SQL Server Books Online	457
<i>Informacje o autorze</i>	460
<i>Indeks</i>	461

Wprowadzenie

Książka ta pełni rolę przewodnika dla osób podejmujących pierwsze kroki w języku T-SQL (nazywanym także Transact-SQL), który jest opracowanym w firmie Microsoft dialektem języka SQL zdefiniowanego przez standardy ISO i ANSI. Poznamy teorię konstruowania zapytań i programowania w języku T-SQL oraz sposoby projektowania kodu T-SQL w celu uzyskiwania i modyfikowania danych, a także ogólny przegląd obiektów programowalnych.

Pomimo że książka pomyślana jest dla Czytelników początkujących, nie jest jedynie zbiorem procedur, według których mają postępować – wykracza poza elementy składni T-SQL i wyjaśnia logikę działającą w tle języka i jego elementów.

Od czasu do czasu w książce pojawiają się zagadnienia, które mogą być uważane za tematykę zaawansowaną – z tego też względu zapoznanie się z tymi fragmentami jest opcjonalne. Jeśli Czytelnik pewnie czuje się w omówionym do tej pory materiale, może przejść do tematów bardziej zaawansowanych; w przeciwnym razie spokojnie może opuścić te fragmenty i powrócić do nich, gdy już nabierze większego doświadczenia. Fragmenty uważane za bardziej zaawansowane są w tekście zaznaczone jako opcjonalne.

Wiele aspektów SQL jest unikatowych dla tego języka i znacznie odbiega od innych języków programowania. Książka ta ułatwi przyswojenie sobie właściwego sposobu myślenia i pozwoli dobrze poznać elementy języka. Czytelnik będzie mógł nauczyć się myśleć w kategoriach relacyjnych i postępować zgodnie z najlepszymi zaleceniami praktycznymi programowania w języku SQL.

Książka nie jest związana z konkretną wersją oprogramowania SQL Server; obejmuje jednak elementy języka, które zostały wprowadzone w ostatnich wersjach SQL Server, w tym SQL Server 2016. Przy omawianiu ostatnio wprowadzonych elementów języka wskazuję wersję produktu, w której dany element został dodany.

SQL Server, oprócz „klasycznego” rozwiązania instalowanego na lokalnym komputerze (serwerze), jest także dostępny jako usługa chmurowa o nazwie Windows Azure SQL Database (w skrócie SQL Database). Przykłady kodu przytaczane w książce były testowane zarówno w lokalnych instalacjach SQL Server, jak i w Azure SQL Database. Powiązana z książką witryna sieci Web (<http://aka.ms/T-SQLFund3e/downloads>) udostępnia informacje dotyczące problemów zgodności pomiędzy tymi rozwiązaniami.

W celu usprawnienia procesu nauczania książka zawiera ćwiczenia, które pozwalają poznaną informację utrwalić w praktyce. Od czasu do czasu pojawiają się ćwiczenia opcjonalne, które są bardziej zaawansowane. Ćwiczenia te przeznaczone są dla

Czytelników, którzy dobrze poznali omawiany materiał i chcą sami sprawdzić swoje umiejętności, rozwiązując trudniejsze problemy. Ćwiczenia opcjonalne dla Czytelników zaawansowanych są odpowiednio oznaczone.

Dla kogo przeznaczona jest ta książka

Niniejsza książka skierowana jest dla programistów korzystających z języka T-SQL, administratorów baz danych, osób zajmujących się rozwiązaniami BI, autorów raportów, analityków, architektów baz danych i zaawansowanych użytkowników, którzy dopiero rozpoczynają pracę z SQL Server i muszą tworzyć kwerendy albo kod przy użyciu języka Transact-SQL.

Założenia

Największe korzyści książka ta przyniesie osobom, które mają już doświadczenie w pracy z systemami Windows i aplikacjami opartymi na tych systemach. Ponadto osoby te powinny znać podstawowe pojęcia dotyczące systemów zarządzania relacyjnymi bazami danych. Przydatne będzie też podstawowe doświadczenie w tworzeniu oprogramowania.

Kto nie powinien czytać tej książki

Nie każda książka nadaje się dla każdego czytelnika. W książce omówiono podstawy języka i głównie skierowana jest ona do osób w praktyce korzystających z języka T-SQL, które nie mają w tym wielkiego doświadczenia. Nie będzie zapewne zbyt interesująca dla doświadczonych praktyków, od lat posługujących się tym językiem. Tym niemniej, wielu czytelników poprzedniego wydania tej książki uważa, że pomimo doświadczeń zdobytych w kolejnych latach pracy książka ta nadal jest przydatna i uzupełnia brakującą wiedzę.

Organizacja książki

Książka rozpoczyna się od przedstawienia teoretycznych podstaw konstruowania zapytań i programowania w języku T-SQL (rozdział 1), co stanowi fundament dla pozostałej części książki, a także dla procesów tworzenia tabel i definiowania integralności danych. W rozdziałach 2 do 9 poruszane są różnorodne aspekty uzyskiwania i modyfikowania danych. Rozdział 10 zawiera omówienie współbieżności i transakcji. Na koniec rozdział 11 stanowi przegląd obiektów programowalnych. Poniżej przedstawiono listę rozdziałów wraz z krótkim ich opisem:

- Rozdział 1 „Podstawy zapytań i programowania T-SQL” – teoretyczne podstawy SQL, teoria zbiorów i logika predykatów; analizy modelu relacyjnego; opisy

architektury SQL Server; wyjaśnienie sposobów tworzenia tabel i definiowania integralności danych.

- Rozdział 2 „Zapytania do pojedynczej tabeli” – różnorodne aspekty konstruowania zapytań dotyczących pojedynczej tabeli przy użyciu polecenia *SELECT*.
- Rozdział 3 „Złączenia” – opis zapytań dotyczących wielu tabel przy użyciu złączeń (*join*), w tym złączenia krzyżowe (*cross join* – iloczyn kartezjański), złączenia wewnętrzne i zewnętrzne.
- Rozdział 4 „Podzapytania” – omówienie zapytań zawartych wewnątrz innych zapytań, czyli *podzapytań*.
- Rozdział 5 „Wyrażenia tablicowe” – omówienie tabel pochodnych, wyrażeń CTE (Common Table Expression), widoków, wbudowanych funkcji zwracających tabele i operatora *APPLY*.
- Rozdział 6 „Operatory zbiorowe” – omówienie operatorów *UNION*, *INTERSECT* i *EXCEPT*.
- Rozdział 7 „Zaawansowane zagadnienia tworzenia zapytań” – omówienie funkcji okien, operatorów *PIVOT* i *UNPIVOT* oraz praca z operatorami *GROUPING SETS*.
- Rozdział 8 „Modyfikowanie danych” – wstawianie, aktualizowanie, usuwanie i scalanie danych.
- Rozdział 9 „Tabele temporalne” – omówienie wersjonowanych przez system tabel temporalnych (czasowych).
- Rozdział 10 „Transakcje i współbieżność” – omówienie kwestii współdziałania połączeń użytkowników, którzy jednocześnie korzystają z tych danych; rozdział opisuje takie pojęcia, jak transakcje, blokady, poziomy izolacji czy zakleszczenia.
- Rozdział 11 „Obiekty programowalne” – omówienie możliwości programowania przy użyciu T-SQL w SQL Server.

W książce zamieszczono także dodatek „Rozpoczynamy”, który ułatwia skonfigurowanie środowiska, pobranie kodów źródłowych książki, zainstalowanie przykładowej bazy danych *TSQV4*, rozpoczęcie pisania kodu dla SQL Server oraz poznanie sposobów uzyskania pomocy dzięki dokumentacji SQL Server Books Online.

Wymagania systemowe

Dodatek „Rozpoczynamy” zawiera informacje, których wersji produktu SQL Server 2016 można użyć do pracy z przykładowym kodem zamieszczonym w tej książce. Poszczególne wersje SQL Server mogą mieć różne wymagania systemowe i programowe; te wymagania są dokładnie opisane w dokumentacji SQL Server Books Online w sekcji „Hardware and Software Requirements for Installing SQL Server 2016” pod adresem <https://msdn.microsoft.com/en-us/library/ms143506.aspx>. W Dodatku wyjaśniono również, jak korzystać z dokumentacji SQL Server Books Online.

Jeśli korzystamy z usługi Azure SQL Database, sprzęt i oprogramowanie jest utrzymywane przez firmę Microsoft, zatem wymagania te nie są istotne.

Do uruchamiania przykładów kodu, zarówno w przypadku lokalnej instancji SQL Server 2016, jak i Azure SQL Database, konieczne jest zainstalowanie oprogramowania SQL Server Management Studio (SSMS). Oprogramowanie to można pobrać z witryny firmy Microsoft pod adresem <https://msdn.microsoft.com/en-us/library/mt238290.aspx>.

Instalowanie i korzystanie z kodu źródłowego

Większość rozdziałów książki zawiera ćwiczenia, które pozwalają interaktywnie wypróbować nowo poznany materiał zawarty w książce. Wszystkie przykłady kodu używane w książce, w tym ćwiczenia i rozwiązania dostępne są na poniższej stronie w zakładce Dodatkowe Informacje:

<http://www.książki.promise.pl/asp/produkt.aspx?pid=112055>

Dodatek „Rozpoczynamy” zawiera szczegółowe informacje na temat instalowania kodów źródłowych.

Podziękowania

Wiele osób przyczyniło się do powstania tej książki, czy to bezpośrednio, czy pośrednio i wszystkim im należą się serdeczne podziękowania i uznanie.

Dla Lilach, za zrozumienie i wsparcie wszystkich moich poczynañ oraz za wyrozumiałość dla niekończących się godzin spędzonych nad SQL.

Dla mojej mamy Mila oraz rodzeństwa Mickey i Ina, za stałe wsparcie i akceptowanie nieobecności, co teraz jest trudniejsze niż kiedykolwiek. Mamo, wszyscy wierzymy, że będzie dobrze, dzięki Twojej sile i determinacji. Tato, dzięki, że jesteś tak pomocny.

Dla recenzenta technicznego książki, Boba Beauchemina: jesteś częścią społeczności SQL Server od tak wielu lat, ale nadal jestem pod wrażeniem twojej wiedzy; byłem szczęśliwy, gdy zgodziłeś się współpracować przy tej książce.

Dla Steve'a Kassa, Dejana Sarka, Gianluca Hotza, i Herberta Alberta: Dzięki za wasze cenne rady podczas planowania i pisania tej książki. Musiałem podjąć kilka trudnych decyzji, co uwzględnić, a czego nie dołączać do książki i wasze wskazówki były bardzo pomocne.

Dla SolidQ, mojej firmy od przeszło dziesięciu lat: to wielka satysfakcja pracować w tak wspaniale prowadzonej firmie. Pracownicy firmy to nie tylko koledzy – to partnerzy, przyjaciele i rodzina. Fernando G. Guerrero, Douglas McDowell, Herbert Albert, Dejan Sarka, Gianluca Hotz, Jeanne Reeves, Glenn McCoin, Fritz Lechnitz, Eric Van Soldt, Joelle Budd, Jan Taylor, Marilyn Templeton, Berry Walker, Alberto Martin, Lorena Jimenez, Ron Talmage, Andy Kelly, Rushabh Mehta, Eladio Rincón, Erik Veerman, Jay Hackney, Richard Waymire, Carl Rabeler, Chris Randall, Johan Åhlén, Raoul Illyés, Peter Larsson, Peter Myers, Paul Turley – dziękuję Wam i wielu innym osobom.

Dla zespołu redakcyjnego produktu SQL Server Pro, czyli Megan Keller, Lavon Peters, Michele Crockett, Mike Otey i z pewnością jeszcze wiele innych osób; pisałem dla tego magazynu od ponad dekady i jestem wdzięczny, że mogłem podzielić się moją wiedzą z czytelnikami magazynu.

Dla osób z grupy MVP produktu SQL Server – Alejandro Mesa, Erland Sommarskog, Aaron Bertrand, Tibor Karaszi, Paul White i wielu innych oraz dla prowadzącego program MVP, Simona Tien; jest to wspaniały program i jestem wdzięczny, że mogłem w nim uczestniczyć. Niespotykany jest poziom umiejętności w tej grupie i zawsze cieszyłem się na nasze spotkania, zarówno by podzielić się pomysłami, jak i po prostu pogadać przy piwie. Jestem przekonany, że w dużej mierze dodanie przez Microsoft funkcjonalności T-SQL do produktu SQL Server zostało zainspirowane przez członków

MVP i całą społeczność produktu SQL Server. To wspaniale przekonać się, że ta synergia przekształca się w tak znaczące i ważne rezultaty.

Na koniec, podziękowania kieruję do moich studentów: nauczanie SQL jest moją siłą napędową, moją pasją. Dziękuję, że mogę spełniać moje powołanie i dziękuję za wszystkie interesujące pytania, które pozwalają pogłębiać wiedzę.

Itzik Ben-Gan

ROZDZIAŁ 1

Podstawy zapytań i programowania T-SQL

Zaczynamy podróż do krainy innej niż wszystkie – miejsca, które rządzi się swoim własnym zbiorem praw. Jeśli ta książka stanowi pierwszy krok na drodze do poznania języka T-SQL (Transact-SQL), powinniście się czuć tak jak Alicja przed rozpoczęciem jej przygód w Krainie Czarów. Dla mnie podróż ta nie ma końca, a po drodze stale odkrywam coś nowego. Zazdroszczę moim czytelnikom – niektóre z najbardziej ekscytujących odkryć wciąż są przed Wami!

Z językiem T-SQL związany jestem od wielu lat: nauczanie, seminaria, pisanie i konsultacje związane z tą tematyką. Dla mnie T-SQL to coś więcej niż tylko język programowania – to sposób myślenia. Często wykladałem i pisałem o (bardzo) zaawansowanych kwestiach, ale opis podstaw języka wciąż odkładałem na później. Nie dlatego, że fundamenty T-SQL są proste czy łatwe – wręcz przeciwnie: pozorna prostota języka jest bardzo myląca. Mogłbym skrótkowo przedstawić elementy składni języka – to wystarczy, by w kilka minut zacząć pisanie zapytań. W praktyce jednak takie podejście utrudnia zrozumienie istoty języka i wydłuża proces jego poznawania.

Rola przewodnika osób podejmujących pierwsze kroki to duża odpowiedzialność. Chciałem mieć pewność, że poświęcę wystarczająco dużo czasu i wysiłku na analizę i poznanie języka, zanim zabrałem się za opisywanie jego podstaw. Język T-SQL nie jest prosty; właściwe opanowanie podstaw to znacznie więcej, niż zrozumienie elementów składni i kodowania zapytań zwracających właściwe wyniki. W istocie trzeba wyrzucić z pamięci wszystko, co się wie o innych językach programowania i zacząć myśleć w kategoriach języka T-SQL.

Podstawy teoretyczne

SQL to akronim nazwy *Structured Query Language* (strukturalny język zapytań). Jest to standaryzowany język, opracowany do tworzenia zapytań zarządzania danymi w systemach zarządzania relacyjnymi bazami danych (RDBMS). Akronim RDBMS oznacza system zarządzania bazą danych oparty na modelu relacyjnym (model semantyczny przedstawiania danych), który z kolei bazuje na dwóch działach matematyki: teorii zbiorów i logice predykatów. Wiele innych języków programowania i różnych

aspektów przetwarzania komputerowego powstało i rozwijało się w dużej mierze w oparciu o intuicję. Inaczej jest w przypadku SQL. W takim stopniu, w jakim SQL bazuje na modelu relacyjnym, wznosi się na solidnym fundamencie – matematyce stosowanej. Tak więc T-SQL opiera się na pewnych podstawach. Firma Microsoft udostępnia język T-SQL jako dialekt (lub rozszerzenie) języka SQL, wykorzystywany w opracowanym przez nią produkcie zarządzania danymi – swoim systemie RDBMS, czyli Microsoft SQL Server.

Podrozdział ten zawiera krótki opis teoretycznych podstaw języka SQL, teorii zbiorów i logiki predykatów, modelu relacyjnego i typów systemów bazodanowych. Ponieważ książka ta nie jest ani podręcznikiem matematycznym, ani książką o modelowaniu danych i projektowaniu, przedstawione w niej informacje teoretyczne przedstawiam w swobodnej formie, a tym samym nie są one pełne. Celem jest opisanie kontekstu języka T-SQL i przedstawienie kluczowych punktów, które są niezbędne dla prawidłowego pojmowania działania języka T-SQL w dalszej części książki.

Niezależność języka

Model relacyjny jest niezależny od języka. Oznacza to, że można zaimplementować model relacyjny przy użyciu innego języka niż SQL, na przykład przy użyciu C# i modelu obiektowego. Obecnie typowym rozwiązaniem są systemy RDBMS obsługujące także inne języki programowania, niż pewien dialekt SQL, czego przykładem jest integracja CLR w SQL Server, pozwalająca obsłużyć zadania historycznie realizowane w języku T-SQL przy użyciu innych języków programowania.

Dodatkowo już od początku trzeba zdawać sobie sprawę, że SQL pod wieloma względami odbiega od modelu relacyjnego. Niektórzy nawet twierdzą, że SQL powinien zostać zastąpiony nowym językiem (takim, który byłby bardziej zgodny z modelem relacyjnym). Na razie jednak to SQL jest branżowym standardem, językiem używanym przez wszystkie wiodące systemy RDBMS.



ZOBACZ TAKŻE Informacje szczegółowe na temat niezgodności SQL z modelem relacyjnym, a także na temat metod używania SQL zgodnie z tym modelem, znaleźć można w następującej książce: *SQL and Relational Theory: How to Write Accurate SQL Code*, Third Edition, C. J. Date (O'Reilly Media, 2015).

SQL

SQL to definiowany przez standardy ANSI i ISO, bazujący na modelu relacyjnym język programowania, opracowany pod kątem tworzenia zapytań i zarządzania danymi w systemach RDBMS.

Na początku lat siedemdziesiątych firma IBM opracowała język SEQUEL (Structured English QUery Language) na potrzeby swojego systemu RDBMS nazwanego System R. Nazwa języka została później zmieniona na SQL ze względu na problemy dotyczące zastrzeżonego znaku towarowego. Język SQL stał się najpierw standardem ANSI (w roku 1986), a następnie standardem ISO (w roku 1987). Od roku 1986 instytucje ANSI (American National Standards Institute) i ISO (International Organization for Standardization) co kilka lat publikują kolejne wersje standardu SQL. Do tej pory opublikowane zostały następujące standardy: SQL-86 (1986), SQL-89 (1989), SQL-92 (1992), SQL:1999 (1999), SQL:2003 (2003), SQL:2006 (2006), SQL:2008 (2008) i SQL:2011 (2011). Standard SQL składa się z kilku rozdziałów. Część I (*Framework* – platforma) oraz Część 2 (*Foundation* – podstawy) dotyczą języka SQL, podczas gdy pozostałe części definiują rozszerzenia standardu, takie jak *SQL for XML* lub integracja SQL i Java.

Co interesujące, składnia języka SQL przypomina naturalny język angielski i jest również bardzo logiczna. Inaczej niż w przypadku wielu języków programowania, które stosują imperatywne wzorce programowania, SQL używa wzorców deklaratywnych. Oznacza to, że SQL wymaga określenia, *co* chcemy uzyskać, a nie *jak* ma to być wykonane, pozostawiając systemowi RDBMS wybór mechanizmów fizycznych niezbędnych do przetworzenia żądania użytkownika.

Język SQL obejmuje kilka kategorii instrukcji, w tym DDL (Data Definition Language – język definicji danych), DML (Data Manipulation Language – język manipulacji danymi) i DCL (Data Control Language – język sterowania danymi). Kategoria DDL dotyczy definicji obiektów i obejmuje takie polecenia, jak *CREATE*, *ALTER* czy *DROP*. Kategoria DML umożliwia tworzenie zapytań oraz modyfikowanie danych i obejmuje takie instrukcje, jak *SELECT*, *INSERT*, *UPDATE*, *DELETE*, *TRUNCATE* i *MERGE*. Typowym nieporozumieniem jest opinia, że DML zawiera tylko polecenia służące do modyfikowania danych; jak już wspomniałem, zawiera także instrukcję *SELECT* (wybierz). Innym częstym nieporozumieniem jest traktowanie instrukcji *TRUNCATE* (obetnij) jako polecenia DDL, kiedy w rzeczywistości jest poleceniem z kategorii DML. Kategoria DCL dotyczy uprawnień i obejmuje takie polecenia, jak *GRANT* czy *REVOKE*. Niniejsza książka skupia się na poleceniach DML.

T-SQL opiera się na standardowym języku SQL, ale wprowadza także pewne własne, niestandardowe rozszerzenia. Przy opisie elementu języka zazwyczaj zaznaczam, czy jest to element standardu, czy też nie.

Teoria zbiorów

Teoria zbiorów, znana też pod nazwą teorii mnogości, której twórcą był matematyk Georg Cantor, to jeden z działów matematyki, na której bazuje model relacyjny. Definicja zbioru sformułowana przez Cantora jest następująca:

Zbiorem jest spójenie w całość określonych rozróżnialnych podmiotów naszej pojędowości czy myśli, które nazywamy elementami danego zbioru.
– Wikipedia (https://pl.wikipedia.org/wiki/Georg_Cantor)

Każde słowo tej definicji ma głębokie i kluczowe znaczenie. Definicje zbioru i członkostwa w zbiorze są aksjomatami, które nie podlegają dowodzeniu (są to *pojęcia pierwotne* teorii mnogości). Każdy element jest częścią wszechświata i należy lub nie należy do zbioru.

Rozpocznijmy od słowa *całość* w definicji Cantora. Zbiór powinien być traktowany jako pojedyncza jednostka, a nie kolekcja elementów, które go tworzą. Powinniśmy się skupić na zestawie obiektów, a nie na poszczególnych obiektach, tworzących zestaw. Później, kiedy będziemy pisać zapytania T-SQL odwołujące się do tabeli w bazie danych (jak na przykład tabela pracowników), powinniśmy myśleć o zbiorze pracowników jako o całości, a nie o poszczególnych pracownikach. Może się to wydawać oczywiste i bardzo proste, jednak wielu programistów ma trudności z przyswojeniem sobie takiego sposobu myślenia.

Słowo *rozróżnialny* oznacza, że każdy element zbioru musi być niepowtarzalny. Wybiegając w przód do pojęcia tabeli w bazie danych, możemy wymusić unikatowość wierszy w tabeli, definiując ograniczenia klucza. Bez klucza nie będziemy mogli w sposób jednoznaczny identyfikować wierszy i dlatego tabela nie będzie mogła być kwalifikowana jako zbiór, a będzie raczej *wielozbiorem* czy *pojemnikiem*.

Wyrażenie *naszej myśli czy postrzegania* sugeruje subiektywność definicji zbioru. Weźmy pod uwagę salę lekcyjną: jedna osoba może dostrzegać zbiór osób, a inna zbiór uczniów oraz zbiór nauczycieli. Z tego względu w definiowaniu zbiorów mamy pokaźną dawkę swobody. Gdy projektujemy model dla naszej bazy danych, proces ten powinien uważnie uwzględniać subiektywne potrzeby aplikacji, by określić właściwe definicje występujących tam jednostek.

Jeśli chodzi o słowo *obiekt*, definicja zbioru nie jest ograniczona tylko do bytów fizycznych, takich jak samochody czy pracownicy, ale odnosi się także do tworów abstrakcyjnych, takich jak linie czy liczby pierwsze.

To, czego brakuje w definicji Cantora, jest zapewne równie ważne jak to, co zawiera. Zwróćmy uwagę, że definicja ta nie wspomina o jakimkolwiek uporządkowaniu elementów zbioru. Kolejność elementów zbioru nie jest istotna. Używany w matematyce formalny zapis wymieniający elementy zbioru korzysta z nawiasów klamrowych: {a, b, c}. Ponieważ kolejnoścnie ma znaczenia, ten sam zbiór można przedstawić jako

{b, a, c} lub {b, c, a}. Wyprzedzając opis, w zestawie atrybutów (nazywanych w SQL *kolumnami*) tworzących nagłówek relacji (w SQL nazywany *tabelą*) zakłada się, że element jest identyfikowany przez nazwę, a nie przez numer porządkowy.

Podobnie rozważmy zbiór krotek (w SQL nazywanych *wierszami*), który stanowi treść relacji; element jest identyfikowany przez jego wartości klucza, a nie na podstawie położenia. Wielu programistów ma trudności z przyswojeniem sobie tego sposobu myślenia, że z punktu widzenia zapytań dotyczących tabel nie istnieje żadna określona czy preferowana kolejność wierszy. Inaczej mówiąc, zapytanie odwołujące się do tabeli może zwrócić jej wiersze w dowolnej kolejności, chyba że jawnie zażądamy określonego posortowania wyników danych, na przykład w celu ich określonego zaprezentowania.

Logika predykatów

Logika predykatów, korzeniami sięgająca starożytnej Grecji, jest innym działem matematyki, na którym opiera się model relacyjny. Dr Edgar F. Codd tworząc model relacyjny dostrzegł możliwość połączenia logiki predykatów zarówno z zarządzaniem danymi, jak i tworzeniem zapytań. Mówiąc niezbyt ściśle, *predykat* jest pewną właściwością lub wyrażeniem, które albo jest spełnione, albo nie – mówiąc inaczej, albo jest prawdą, albo fałszem. Model relacyjny polega na predykatkach w celu utrzymywania logicznej integralności danych i definiowania struktury danych. Przykładem predykatu użytego do wymuszenia integralności jest ograniczenie zdefiniowane w tabeli nazwanej *Employees* (pracownicy), które zezwala na umieszczenie w tabeli jedynie takich danych pracowników, których pensja jest większa od zera. Predykatem jest tu wyrażenie „pensja większa niż 0” (wyrażenie T-SQL: *pensja > 0*).

Predykaty używane są również podczas filtrowania danych w celu zdefiniowania podzbiorów. Jeśli na przykład zachodzi potrzeba odpytania tabeli *Employees* i zwrócenia jedynie tych wierszy, które dotyczą pracowników działu sprzedaży, w naszym filtrze możemy użyć predykatu „działem jest (równa się) dział sprzedaży” (wyrażenie T-SQL: *dział = 'sprzedaż'*).

W teorii mnogości predykatów można używać do definiowania zbiorów. Metoda taka jest przydatna, ponieważ nie zawsze możemy zdefiniować zbiór poprzez wymienienie wszystkich jego elementów (przykładem mogą być zbiory nieskończone), a często wygodniej jest zdefiniować zbiór w oparciu o właściwość. Przykładem zbioru nieskończonego zdefiniowanego za pomocą predykatu może być zbiór wszystkich liczb pierwszych, który można zdefiniować przy użyciu następującego wyrażenia: „x jest dodatnią liczbą całkowitą większą niż 1, która podzielna jest tylko przez 1 i samą siebie”. Dla dowolnej wyspecyfikowanej wartości predykat jest albo prawdą, albo fałszem. Zbiorem wszystkich liczb pierwszych jest zbiór wszystkich elementów, dla których predykat jest prawdziwy. Przykładem skończonego zbioru zdefiniowanego za pomocą predykatu może być zbiór {0, 1, 2, 3, 4, 5, 6, 7, 8, 9}, definiowany jako zbiór wszystkich

elementów, dla których następujący predykat ma wartość prawda: „x jest liczbą całkowitą równą lub większą niż 0 i równą lub mniejszą niż 9”.

Model relacyjny

Model relacyjny jest modelem semantycznym zarządzania i modyfikowania danych, opartym na teorii zbiorów i logice predykatów. Jak już wspomniałem wcześniej, model ten został stworzony przez Edgara F. Codd, a następnie rozbudowany i opracowany przez Chrisa Date, Hugh'a Darwena i innych. Pierwsza wersja modelu relacyjnego została zaproponowana przez Codd w roku 1969 w postaci raportu badawczego firmy IBM pod tytułem „Derivability, Redundancy, and Consistency of Relations Stored in Large Data Banks”. Poprawiona wersja przedstawiona została przez Codd w roku 1970 w dokumencie zatytułowanym „A Relational Model of Data for Large Shared Data Banks”, opublikowanym w czasopiśmie *Communications of the ACM*.

Celem modelu relacyjnego jest udostępnienie spójnej reprezentacji danych przy minimalnej redundancji lub jej braku i bez utraty kompletności oraz zdefiniowanie integralności danych (wymuszanie spójności danych) jako części modelu. System RDBMS powinien implementować model relacyjny i dostarczać środków do realizowania zapytań oraz do przechowywania, zarządzania i wymuszania integralności danych. Model relacyjny opiera się na silnych podstawach matematycznych, a nie na intuicji, co oznacza, że mając pewną instancję modelu danych (na podstawie których zostanie później wygenerowana fizyczna baza danych) możemy dokładnie określić, czy projekt ten jest podatny na wady, nie polegając wyłącznie na intuicji.

Model relacyjny korzysta z takich pojęć jak tezy, predykaty, relacje, krotki, atrybuty i wiele innych. Dla osób nie zajmujących się matematyką pojęcia te mogą wydawać się początkowo przytłaczające, ale w kolejnych podrozdziałach przedstawię najważniejsze aspekty modelu w sposób przystępny i nie-matematyczny oraz wyjaśnię, jak pojęcia te odnoszą się do baz danych.

Tezy, predykaty i relacje

Dość powszechne przekonanie, że pojęcie *relacyjny* wywodzi się z zależności pomiędzy tabelami, nie jest poprawne. „Relacyjny” w rzeczywistości odnosi się do matematycznego pojęcia *relacji*. W teorii mnogości relacja jest przedstawieniem pewnego zbioru. W modelu relacyjnym relacją jest zbiór powiązanych informacji, którego odpowiednikiem w SQL jest tabela – aczkolwiek nie jest to dokładny odpowiednik. Kluczowy punkt modelu relacyjnego to fakt, że pojedyncza relacja powinna reprezentować pojedynczy zbiór (na przykład *Klienci*). Warto zwrócić uwagę na to, że operacje na relacjach (oparte na algebrze relacji) dają w wyniku relacje (na przykład złączenie dwóch relacji).

UWAGA Model relacyjny rozróżnia pojęcie *relacji* i *zmiennej relacji*, ale dla uproszczenia nie będę ich rozróżniać – w obu przypadkach będę używać pojęcia *relacji*. Ponadto relacja zbudowana jest z nagłówka i treści. Nagłówek składa się ze zbioru atrybutów (nazywanych w SQL *kolumnami*), gdzie każdy element jest identyfikowany przez nazwę atrybutu i nazwę typu. Treść składa się ze zbioru krotek (w SQL nazywanych *wierszami*), gdzie każdy element jest identyfikowany przez klucz. Dla uproszczenia będę odnosił się do tabeli jak do zbioru wierszy.



Podczas projektowania modelu danych dla bazy danych przedstawiamy wszystkie dane za pomocą relacji (tabel). Rozpoczynamy od zidentyfikowania tez (ang. *proposition*), które powinny być reprezentowane w bazie danych. Teza jest twierdzeniem, które może być prawdziwe lub fałszywe. Na przykład tezą jest stwierdzenie „Pracownik Itzik Ben-Gan urodził się 12 lutego 1971 roku i pracuje w dziale IT”. Jeśli ta teza jest prawdą, sama uwidoczni się w postaci wiersza tabeli *Employees* (pracownicy). Fałszywa teza po prostu się nie uwidoczni. Takie wstępne założenie nazywane jest *założeniem o świecie zamkniętym* – CWA (*Close World Assumption*).

Następnym krokiem jest sformalizowanie tez. W tym celu bierzemy rzeczywiste dane (treść relacji) i definiujemy strukturę (nagłówek relacji) – na przykład poprzez tworzenie predykatów. Możemy traktować predykaty jako sparametryzowane tezy. Nagłówek relacji zawiera zbiór atrybutów. Zwróćmy uwagę na użycie pojęcia „zbiór”; w modelu relacyjnym atrybuty są nieuporządkowane i unikatowe. Atrybut jest identyfikowany przez nazwę atrybutu i nazwę typu. Na przykład nagłówek relacji *Employees* może składać się z następujących atrybutów (wyrażonych jako pary – nazwa atrybutu i nazwa typu): *employeeid integer* (identyfikator pracownika, liczba całkowita), *firstname character string* (imię, ciąg znaków), *lastname character string* (nazwisko, ciąg znaków), *birthdate date* (data urodzin, data), *departmentid integer* (identyfikator działu, liczba całkowita).

Typ jest jednym z najbardziej fundamentalnych bloków konstrukcyjnych relacji. Typ ogranicza atrybut do pewnego zestawu możliwych lub poprawnych wartości. Na przykład typ *int* jest zbiorem wszystkich liczb całkowitych z zakresu od $-2\,147\,483\,648$ do $2\,147\,483\,647$. Typ jest w bazie danych jedną z najprostszych form predykatu, ponieważ ogranicza dopuszczane wartości atrybutu. Na przykład baza danych nie będzie akceptować tezy, w której data urodzin pracownika to 31 luty 1971 (nie wspominając o dacie podanej jako coś w rodzaju „abc!”). Zwróćmy uwagę, że nie jesteśmy ograniczeni do typów podstawowych, takich jak liczby całkowite czy ciąg znaków; typ może być również wyliczeniem możliwych wartości, jak na przykład lista możliwych stanowisk pracy. Typ może być bardzo złożoną konstrukcją. Dla osób, które znają inne języki programowania, prawdopodobnie najlepszym sposobem myślenia o typie jest przyrównanie go do klasy – kapsułkującej zarówno dane, jak i sposób ich obsługi. Przykładem typu złożonego jest typ *geometry*, który obsługuje wielokąty.

Brakujące wartości

Istnieje pewien aspekt modelu relacyjnego, który jest źródłem wielu emocjonujących dyskusji – czy predykaty powinny być ograniczone do logiki dwuwartościowej. W przypadku dwuwartościowej logiki predykat może mieć tylko wartość prawda albo fałsz. Jeśli predykat nie jest prawdą, musi być fałszem. Używanie dwuwartościowej logiki predykatu jest zgodne z prawem matematycznym zwanym „zasadą wyłączonego środka”. Niektórzy jednak twierdzą, że istnieje przestrzeń dla stosowania logiki trójwartościowej (czy nawet logiki o czterech wartościach), co pozwoliłoby uwzględniać przypadki, w których brakuje wartości. Predykat zastosowany do brakującej wartości nie generuje prawdy lub fałszu, a wartość *nieznaną*.

Dla przykładu weźmy atrybut telefonu komórkowego relacji *Employees*. Załóżmy, że brak jest informacji o numerach telefonów niektórych pracowników. W jaki sposób przedstawić ten fakt w bazie danych? W przypadku zaimplementowania trójwartościowej logiki atrybut telefonu komórkowego powinien pozwalać na stosowanie specjalnego znaku dla brakującej wartości. Następnie predykat porównując atrybut telefonu z pewnym określonym numerem wygeneruje wartość *nieznany* dla przypadków, kiedy brak tej wartości. Trójwartościowa logika predykatu generuje jedną z trzech możliwych wartości logicznych – *prawda*, *fałsz* i *nieznany*.

Niektóre osoby są przekonane, że trójwartościowa logika predykatu jest nie-relacyjna, podczas gdy inni są zupełnie odmiennego zdania. W rzeczywistości Codd był zwolennikiem czterowartościowej logiki predykatu, twierdząc, że istnieją dwa różne przypadki wartości brakujących: brakująca, ale mająca zastosowanie (*A-Mark*, od *applicable*) oraz brakująca, ale i bez zastosowania (*I-Mark*, od *inapplicable*). Przykładem wartości *A-Mark* jest sytuacja, kiedy pracownik ma telefon komórkowy, ale nie znamy numeru tego telefonu. Z wartością *I-Mark* będziemy mieli do czynienia w sytuacji, kiedy pracownik w ogóle nie posiada telefonu. Według Codd'a do obsługi tych dwóch przypadków wartości brakujących powinny być używane dwa znaczniki specjalne.

Język SQL implementuje trójwartościową logikę predykatu, obsługując znacznik *NULL* dla oznaczania ogólnego pojęcia wartości brakującej. Obsługa znaczników *NULL* i trójwartościowej logiki predykatu w SQL jest źródłem różnych pomyłek i złożoności, chociaż można się upierać, że brakujące wartości to nieusuwalna cecha świata rzeczywistego. Ponadto podejście alternatywne, czyli stosowanie wyłącznie dwuwartościowej logiki predykatu, wcale nie jest mniej problematyczne.



UWAGA Jak wspomniałem, *NULL* nie jest wartością, ale znacznikiem wskazującym brak wartości. Tym samym, choć nieszczęśliwie często spotykane, sformułowanie w rodzaju „wartość *NULL*” jest błędem logicznym. Właściwa terminologia to „znacznik *NULL*” lub po prostu samo „*NULL*”. W tej książce posługuję się tym drugim wariantem, gdyż jest powszechnie używany w społeczności SQL.

Ograniczenia

Jedną z największych zalet modelu relacyjnego jest możliwość definiowania integralności danych jako części modelu. Integralność danych jest osiągana poprzez reguły nazywane *ograniczeniami* (*constraint*), które są definiowane w modelu danych i wymuszane przez system RDBMS. Najprostsza metoda wymuszania integralności danych to przypisanie wymaganego typu atrybutu oraz określenie możliwości obsługi znaczników *NULL* lub brakiem tej obsługi. Ograniczenia są wymuszane również poprzez sam model; na przykład relacja *Orders*(*orderid*, *orderdate*, *duedate*, *shipdate*) [Zamówienia(*id*, *data*, *data* płatności, *data* wysyłki)] dla jednego zamówienia dopuszcza trzy różne daty, podczas gdy relacje *Employees*(*empid*) [Pracownicy(*id*)] i *EmployeeChildren*(*empid*, *childname*) [DzieciPracownika(*id* pracownika, imię dziecka)] zezwalają na liczbę dzieci z zakresu od 0 do (przeliczalnej) nieskończoności.

Innymi przykładami ograniczeń są *klucze kandydujące* (*candidate keys*), zwane też *potencjalnymi*, które zapewniają integralność krotki* (wiersza w terminologii SQL), oraz *klucze obce* (*foreign keys*), które zapewniają integralność referencyjną. Klucz kandydujący to klucz zdefiniowany w oparciu o jeden lub kilka atrybutów, uniemożliwiający pojawianie się w relacji więcej niż jednej instancji tej samej krotki (wiersza). Predykat oparty na kluczu kandydującym jednoznacznie identyfikuje wiersz (na przykład pracownika). W relacji można definiować wiele kluczy kandydujących. Na przykład w relacji *Employees* możemy zdefiniować klucze kandydujące w oparciu o atrybut *employeeid* (*id* pracownika), *SSN* (numer ubezpieczenia) i inne. Zazwyczaj arbitralnie wybieramy jeden z kluczy kandydujących jako *klucz główny* (*primary key*), na przykład *employeeid* w relacji *Employees* i używamy tego klucza jako preferowanej metody identyfikowania wiersza. Wszystkie pozostałe klucze kandydujące nazywane są *kluczami alternatywnymi*.

Klucze obce są używane do wymuszania integralności referencyjnej. Klucz obcy definiowany jest w oparciu o jeden lub kilka atrybutów relacji (nazywanej *relacją referencyjną*, czyli odwołującą się) i odnosi się do klucza kandydującego w innej (lub niekiedy w tej samej) relacji. Ta więc ogranicza wartości mogące wystąpić w atrybutach klucza obcego relacji odwołującej się do wartości, które już występują w atrybutach klucza kandydującego relacji, do której następuje odwołanie. Załóżmy na przykład, że relacja *Employees* ma klucz obcy zdefiniowany w oparciu o atrybut *departmentid* (*id* działu), który odwołuje się do atrybutu *departmentid* klucza głównego w relacji *Departments*. Oznacza to, że wartości *Employees.departmentid* zostaną ograniczone do wartości, które występują w *Departments.departmentid*.

* Krotka (ang. *tuple*) – struktura danych będąca odzwierciedleniem matematycznej *n*-ki, tj. uporządkowanego ciągu wartości. Rekordy relacyjnych baz danych są krotkami, gdyż poszczególne pola rekordów mogą zawierać dane różnych typów. W bazach SQL rekordy są nazywane wierszami (ang. *row*), zaś zbiory krotek definiowane przez relacje – tabelami.

Normalizacja

Model relacyjny definiuje również *reguły normalizacji*, nazywane również *postaciami normalnymi* (*normal form* – NF). Normalizacja to formalny proces matematyczny, który gwarantuje, że każda jednostka będzie reprezentowana przez pojedynczą relację. W znormalizowanej bazie danych zapobiegamy powstawaniu nieprawidłowości i utrzymujemy redundancję na minimalnym poziomie bez utraty danych. Jeśli postępujemy zgodnie z modelowaniem ERM (Entity Relationship Modeling) i przedstawiamy każdą jednostkę i jej atrybuty, normalizacja prawdopodobnie nie będzie potrzebna; normalizację będziemy wtedy stosować jedynie do wzmocnienia poprawności działania modelu. W kolejnych podpunktach skrótowo przedstawię trzy postaci normalne (1NF, 2NF i 3NF) wprowadzone przez Codda.

1NF Pierwsza postać normalna określa, że krotka (wiersz) w relacji (tabeli) musi być unikatowa, a atrybuty powinny być niepodzielne na mniejsze wartości (atomowość danych). Jest to nadmiarowa definicja relacji; inaczej mówiąc, jeśli tabela rzetelnie odzwierciedla relację, spełnia już pierwszą postać normalną.

Niepowtarzalność wierszy osiągamy definiując dla tabeli unikatowy klucz.

Na atrybutach można wykonywać tylko takie operacje, które są zdefiniowane dla typu atrybutu. Atomowość atrybutów jest cechą subiektywną, podobnie jak subiektywna jest definicja zbioru. Powstaje na przykład pytanie, czy nazwa pracownika w relacji *Employees* powinna być wyrażana za pomocą jednego atrybutu (nazwisko), dwóch atrybutów (imię i nazwisko) czy trzech atrybutów (imię, drugie imię i nazwisko)? Odpowiedź zależy od zastosowań. Jeśli aplikacja musi oddzielnie operować na częściach nazwy użytkownika (na przykład w celu wyszukiwania), sensownie jest je rozdzielać; w przeciwnym razie – nie.

Podobnie jak atrybut może nie być wystarczająco „atomowy” z punktu widzenia potrzeb aplikacji, może również być „podatomowy”. Jeśli na przykład dla danej aplikacji adresu jest traktowany jak niepodzielny, niedołączenie miasta jako części adresu będzie niespełnieniem warunku pierwszej postaci normalnej.

Ta postać normalna jest często rozumiana nieprawidłowo. Niektórzy uważają, że próba naśladowania tablic arkuszy kalkulacyjnych narusza pierwszą postać normalną. Przykładem może być zdefiniowanie relacji *YearlySales* (sprzedaż roczna) za pomocą następujących atrybutów: *salesperson* (sprzedawca), *qty2013* (ilość w 2013), *qty2014* i *qty2015*. W tym przykładzie jednak tak naprawdę nie naruszamy pierwszej postaci normalnej; po prostu nakładamy ograniczenie – ograniczamy dane do trzech określonych lat: 2013, 2014 i 2015.

2NF Druga postać normalna zakłada dwie reguły. Pierwsza reguła określa, że dane muszą spełniać wymagania pierwszej postaci normalnej. Druga reguła dotyczy zależności pomiędzy atrybutami klucza kandydującego a atrybutami, które nie są częścią klucza. Dla każdego klucza kandydującego, każdy atrybut niebędący częścią klucza musi

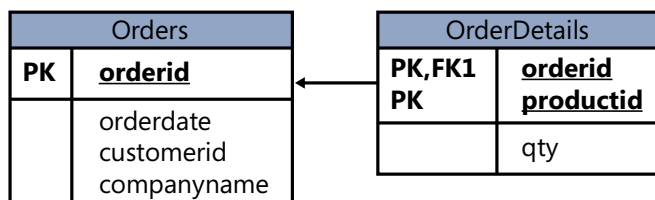
być w pełni funkcjonalnie zależny od całego klucza kandydującego. Inaczej mówiąc, atrybut niebędący częścią klucza nie może być w pełni funkcjonalnie zależny od części klucza kandydującego. Aby to trochę uprościć, jeśli potrzebujemy uzyskać wartość pewnego atrybutu niebędącego częścią klucza, trzeba dostarczyć wartości wszystkich atrybutów klucza kandydującego tej samej krotki. Możemy wyszukać dowolną wartość dowolnego atrybutu dowolnej krotki, jeśli znamy wszystkie wartości atrybutów klucza kandydującego.

Dla przykładu naruszenia drugiej postaci normalnej założmy, że zdefiniowaliśmy relację nazwaną *Orders* (zamówienia), która reprezentuje informacje o zamówieniach i jego pozycjach (rysunek 1-1). Relacja *Orders* zawiera następujące atrybuty: *orderid*, *productid*, *orderdate*, *qty*, *customerid* i *companyname* (identyfikator zamówienia, identyfikator produktu, data zamówienia, ilość, identyfikator klienta, nazwa firmy). Klucz główny zdefiniowany jest w oparciu o atrybuty *orderid* i *productid*.

Orders	
PK	<u>orderid</u>
PK	<u>productid</u>
	orderdate qty customerid companyname

RYСУNEK 1-1 Model danych przed zastosowaniem postaci 2NF

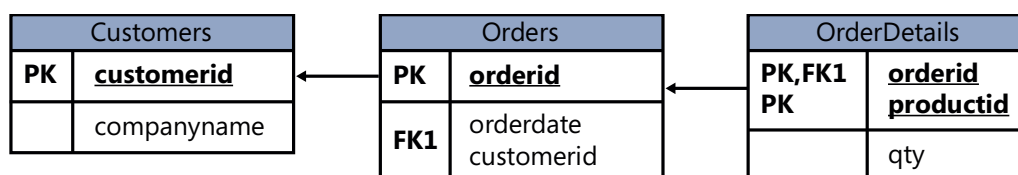
Relacja pokazana na rysunku 1-1 nie spełnia drugiej postaci normalnej, ponieważ istnieją atrybuty niebędące częścią klucza, które są zależne tylko od części klucza kandydującego (w tym przykładzie od klucza głównego). Na przykład, możemy znaleźć atrybut *orderdate* zamówienia, a także atrybuty *customerid* i *companyname*, bazując jedynie na samym atrybucie *orderid*. Aby spełniona była druga postać normalna, trzeba rozdzielić pierwotną relację na dwie relacje: *Orders* oraz *OrderDetails* (rysunek 1-2). Relacja *Orders* będzie zawierać atrybuty *orderid*, *orderdate*, *customerid* oraz *companyname*, z kluczem głównym zdefiniowanym za pomocą atrybutu *orderid*. Relacja *OrderDetails* będzie zawierać atrybuty *orderid*, *productid* oraz *qty*, wraz z kluczem głównym zdefiniowanym za pomocą atrybutów *orderid* i *productid*.



RYСУNEK 1-2 Model danych po zastosowaniu postaci 2NF, a przed zastosowaniem postaci 3NF

3NF Trzecia postać normalna również kieruje się dwoma regułami. Dane muszą spełniać drugą postać normalną. Ponadto wszystkie atrybuty niebędące częścią klucza muszą być zależne od kluczy kandydujących w sposób nieprzechodni. Upraszczając, reguła ta oznacza, że wszystkie atrybuty niebędące częścią klucza muszą być wzajemnie niezależne. Inaczej mówiąc, jeden atrybut niebędący częścią klucza nie może być zależny od innego takiego atrybutu.

Opisane poprzednio relacje *Orders* i *OrderDetails* obecnie spełniają drugą postać normalną. Pamiętamy, że w tym momencie relacja *Orders* zawiera atrybuty *orderid*, *orderdate*, *customerid* i *companyname*, wraz z kluczem głównym zdefiniowanym w oparciu o atrybut *orderid*. Oba atrybuty *customerid* i *companyname* zależne są od całego klucza głównego – *orderid*. Na przykład, aby wyszukać atrybut *customerid* reprezentujący klienta, który złożył zamówienie, potrzebny jest cały klucz główny. Podobnie potrzebny jest cały klucz główny, by wyszukać nazwę firmy klienta, który złożył zamówienie. Atrybuty *customerid* i *companyname* są jednak również zależne od siebie samych. Aby spełnione były wymagania trzeciej postaci normalnej, trzeba dodać relację *Customers* (rysunek 1-3) wraz z atrybutami *customerid* (klucz główny) i *companyname*. Następnie możemy usunąć atrybut *companyname* z relacji *Orders*.



RYСУNEK 1-3 Model danych po zastosowaniu postaci 3NF

Nieformalnie postaci 2NF i 3NF są często podsumowywane następującym stwierdzeniem: „Każdy atrybut niebędący częścią klucza jest zależny od klucza, od całego klucza i od niczego poza kluczem – tak mi dopomóż Codd”.

Istnieją wyższe postaci normalne, poza pierwszymi trzema opracowanymi przez Codda, które obejmują złożone klucze główne i tymczasowe bazy danych, ale tematyka ta wykracza poza zakres książki.

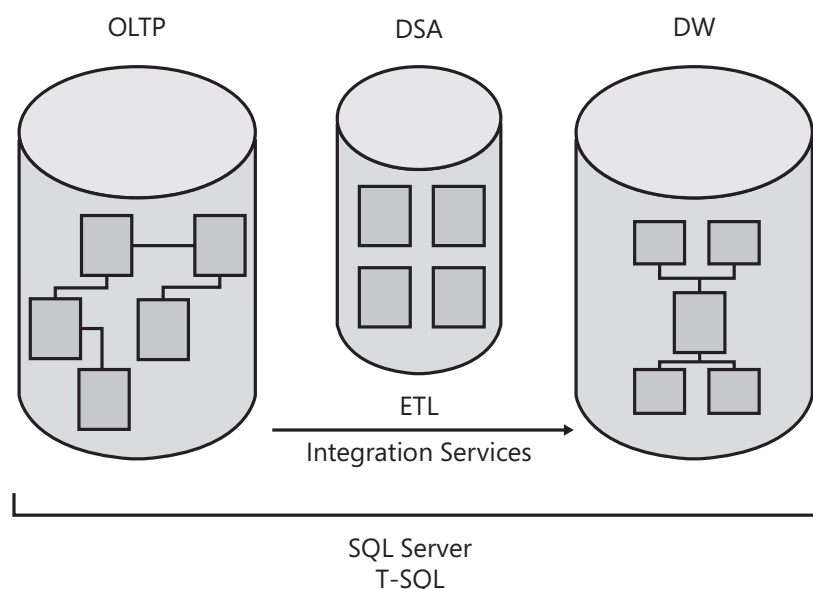


UWAGA Warto zauważyć, że SQL, a także T-SQL, pozwalają na naruszanie wszystkich form normalnych w rzeczywistych tabelach. To na twórcy modelu danych spoczywa odpowiedzialność za zaprojektowanie modelu znormalizowanego.

Typy systemów bazodanowych

Dwa główne typy systemów, czy też obciążeń które wykorzystują SQL Server jako mechanizm bazodanowy oraz T-SQL do zarządzania i manipulowania danymi, to przetwarzanie transakcyjne w trybie online oraz hurtownie danych. Rysunek 1-4 ilustruje te systemy i procesy transformacji, które zwykle zachodzą pomiędzy nimi. Użyte akronimy mają następujące znaczenie.

- OLTP: Online Transactional Processing (przetwarzanie transakcyjne w trybie online)
- DSA: Data Staging Area (obszar przygotowywania danych)
- DW: Data Warehouse (hurtownia danych)
- ETL: Extract, Transform, Load (ekstrakcja, transformacja i ładowanie)



RYSUNEK 1-4 Klasy systemów bazodanowych

Online Transactional Processing

Dane są wprowadzane początkowo do systemu przetwarzania transakcji w trybie online. System OLTP skupia się na wprowadzaniu danych, a nie na raportowaniu – transakcje dotyczą głównie wstawiania, aktualizowania i usuwania danych. Model relacyjny jest ukierunkowany głównie na systemy OLTP, gdzie znormalizowany model zapewnia zarówno dobrą wydajność wprowadzania danych, jak i ich spójność. W znormalizowanym środowisku każda tabela reprezentuje pojedynczą jednostkę i minimalizuje redundancję. Jeśli zachodzi potrzeba zmodyfikowania faktu, trzeba go zmienić tylko w jednym miejscu. Dzięki temu optymalizowana jest wydajność modyfikowania danych i zmniejszane prawdopodobieństwo wystąpienia błędu.

Środowisko OLTP nie jest jednak odpowiednie do raportowania, ponieważ znormalizowany model zazwyczaj korzysta z wielu tabel (jedna dla każdej jednostki) o złożonych powiązaniach. Nawet prosty raport wymaga złączenia wielu tabel, co generuje złożone zapytania o niskiej wydajności.

Bazę danych OLTP można zaimplementować w systemie SQL Server i zarówno zarządzać nimi, jak i wykonywać zapytania przy użyciu języka T-SQL.

Data Warehouse

Hurtownie danych (Data Warehouse) to środowisko opracowane do odczytywania danych i raportowania. Jeśli środowisko to służy całej organizacji, nazywane jest hurtownią danych; jeśli służy jedynie części organizacji (na przykład określonej działowi) lub tematycznemu obszarowi w organizacji, jest nazywane *data mart* (*mini-hurtownia*, *magazyn danych*). Model danych hurtowni danych został tak opracowany i zoptymalizowany, by przede wszystkim zapewnić obsługę funkcji odczytywania danych. Model zawiera zamierzoną redundancję, mniejszą liczbę tabel i uproszczone powiązania, co powoduje, że w porównaniu do środowiska OLTP zapytania są prostsze i działają efektywniej.

Najprostszy projekt hurtowni danych nazywany jest *schematem gwiazdy* (*star schema*). Schemat gwiazdy obejmuje tabelę faktów i kilka tabel wymiarów. Każda tabela wymiarów reprezentuje podmiot, który ma być podstawą analizy danych. Na przykład w systemie obsługi zamówień sprzedaży będziemy prawdopodobnie chcieli analizować dane według klientów, produktów, pracowników, czasu i temu podobnych wymiarów. W schemacie gwiazdy każdy wymiar jest implementowany w postaci pojedynczej tabeli z nadmiarowymi danymi. Na przykład wymiar produktu mógłby zostać zaimplementowany w postaci pojedynczej tabeli *ProductDim*, zamiast trzech znormalizowanych tabel: *Products*, *ProductSubCategories* i *ProductCategories* (*Produkty*, *Podkategorie_produktu* i *Kategorie_produktu*). Jeśli normalizujemy tabelę wymiarów, co spowoduje powstanie wielu tabel reprezentujących ten wymiar, otrzymamy strukturę, która nazywana jest wymiarem *płatka śniegu* (*snowflake dimension*). Schemat zawierający wymiary płatka śniegu nazywany jest *schematem płatka śniegu* (w odróżnieniu do schematu gwiazdy).

W tabeli faktów przechowywane są fakty i miary, takie jak ilość czy wartość każdego istotnego połączenia kluczy wymiarów. Na przykład dla każdego istotnego połączenia klienta, produktu, pracownika i dnia tabela faktów może zawierać wiersz zawierający ilość i wartość. Zwróćmy uwagę, że dane w hurtowni danych są zazwyczaj wstępnie przetworzone (posumowane) do pewnego poziomu rozdrobnienia (na przykład z podziałem na dni), inaczej niż w przypadku danych w środowisku OLTP, które zwykle są rejestrowane na poziomie transakcji.

Historycznie rzecz biorąc, wczesne wersje systemu SQL Server były głównie ukierunkowane na środowiska OLTP, jednak ostatecznie system SQL Server zaczął również

obsługiwać systemy hurtowni danych i zapewniać odpowiednią analizę. Możemy implementować hurtownię danych jako bazę danych systemu SQL Server oraz zarządzać danymi i wykonywać zapytania za pomocą T-SQL.

Proces wyciągania danych z systemów źródłowych (OLTP i inne), opracowywania danych i ładowania ich do hurtowni danych nazywany jest procesem ETL (Extract, Transform, Load). Do obsługi procesu ETL system SQL Server udostępnia narzędzie Microsoft SQL Server Integration Services (SSIS).

Proces ETL będzie często korzystał z warstwy danych DSA (Data Staging Area – DSA), znajdującej się pomiędzy OLTP a DW. Warstwa DSA umieszczona jest zazwyczaj w relacyjnej bazie danych, takiej jak SQL Server i jest używana jako obszar czyszczenia danych. Warstwa DSA nie jest dostępna dla użytkowników końcowych.

Architektura SQL Server

Niniejszy podrozdział jest wprowadzeniem do architektury systemu SQL Server, jego odmian i używanych jednostek – instancji SQL Server, baz danych, schematów i obiektów danych – oraz opisuje przeznaczenie wszystkich tych elementów.

Odmiany ABC produktu SQL Server

Przez wiele lat system SQL Server był dostępny tylko w jednej odmianie – w siedzibie (*on-premises*), czyli w wersji pudełkowej (*box*). Niedawno firma Microsoft zdecydowała się na udostępnienie kilku odmian, by klienci mogli wybrać taką, która im najbardziej odpowiada. Obecnie firma Microsoft oferuje trzy odmiany produktu SQL Server, które wewnętrznie są nazywane ABC: A dla Appliance, B dla Box oraz C dla Cloud.

Wersja Appliance

Pomysł związany z wyróżnieniem odmiany A wywodzi się z chęci udostępnienia kompletnego rozwiązania „pod klucz”, które obejmuje sprzęt, oprogramowanie i usługi. Szybkość działania jest osiągana dzięki spójnemu rozmieszczeniu komponentów z magazynem danych umieszczonym blisko jednostki przetwarzającej. Wersja ta jest utrzymywana lokalnie po stronie klienta. Firma Microsoft współpracuje z dostawcami sprzętu, takimi jak firmy Dell czy HP w celu oferowania takiego produktu. Eksperti z firmy Microsoft i dostawcy sprzętu zapewniają klientowi wydajność, bezpieczeństwo i dostępność jego systemu.

Obecnie dostępnych jest kilka rozwiązań klasy Appliance, z których wyróżnić można Microsoft Analytics Platform System (APS), skoncentrowany na hurtowniach danych i przetwarzaniu wielkich ilości danych (*big data*). Rozwiązanie to obejmuje silnik hurtowni danych o nazwie Parallel Data Warehouse (PDW), wykorzystujące technologię masowego przetwarzania równoległego (*massively parallel processing* – MPP).

Zawiera ono również mechanizm HDInsight, stanowiący opracowaną przez firmę Microsoft dystrybucję silnika bazodanowego Hadoop. APS zawiera również technologię zapytań o nazwie PolyBase, która pozwala na wykorzystanie zapytań T-SQL wobec danych relacyjnych zawartych w PDW i nierelacyjnych danych z HDInsight.

Wersja Box

Odmiana Box produktu SQL Server, formalnie nazywana systemem SQL Server w siedzibie (*on-premises*), jest tradycyjnym sposobem używania produktu, zazwyczaj instalowanego w siedzibie klienta. Klient jest odpowiedzialny za wszystko – uzyskanie sprzętu, zainstalowanie oprogramowania i obsługę aktualizacji, zapewnienie wysokiego poziomu dostępności i przywracania po awariach (HADR), bezpieczeństwo i wszystko inne.

Klient może instalować wiele instancji produktu na tym samym serwerze (więcej na ten temat w kolejnym podrozdziale) i może pisać zapytania, które współdziałają z wieloma bazami danych. Istnieje również możliwość przełączania się pomiędzy bazami danych, chyba że jedna z nich jest typu *contained* (zawarta) – definicję tego terminu podam później.

Podstawowym językiem zapytań jest T-SQL. Wszystkie przykłady kodu i ćwiczenia przytoczone w tej książce możemy uruchamiać na implementacji systemu SQL Server w siedzibie. W Dodatku znaleźć można informacje szczegółowe na temat uzyskania i instalowania wersji testowej SQL Server, a także na temat tworzenia przykładowej bazy danych.

Wersja Cloud

Przetwarzanie w chmurze polega na zapewnieniu zasobów obliczeniowych na żądanie z współużytkowanej puli zasobów. Technologie RDBMS firmy Microsoft mogą być udostępniane zarówno jako usługi chmury prywatnej, jak i publicznej. *Chmura prywatna* to infrastruktura chmurowa obsługująca pojedynczą organizację i zazwyczaj wykorzystuje technologię wirtualizacji. Zazwyczaj jest hostowana w siedzibie klienta i utrzymywana przez personel IT organizacji. Można patrzeć na to jak na samoobsługową elastyczność, pozwalającą użytkownikom na wdrażanie zasobów w miarę potrzeb. Zapewnia też standaryzację i pomiary wykorzystania. Silnikiem bazodanowym jest zazwyczaj instalacja pudełkowa (zatem ten sam język T-SQL jest używany do zarządzania i manipulowania danymi).

W przypadku *chmury publicznej* usługi udostępniane są poprzez sieć i dostępne dla każdego (gotowego zapłacić). Firma Microsoft udostępniła dwie formy publicznych usług chmurowych RDBMS: infrastruktura jako usługa (IaaS) oraz platforma jako usługa (PaaS). W przypadku IaaS wynajmujemy maszynę wirtualną (VM) zlokalizowaną w infrastrukturze chmurowej firmy Microsoft. Początkowo możemy wybrać spośród wielu wstępnie skonfigurowanych maszyn wirtualnych, które już zawierają

zainstalowaną określoną wersję i wydanie SQL Server (silnik typu *box*). Sprzęt jest utrzymywany przez Microsoft, ale to użytkownik jest odpowiedzialny za obsługę oprogramowania i instalowanie poprawek. Zasadniczo nie różni się to od utrzymywania własnej instalacji SQL Server – tyle, że ta zlokalizowana jest na sprzęcie należącym do firmy Microsoft.

W przypadku PaaS firma Microsoft udostępnia platformę bazodanową jako usługę. Jest ona hostowana w centrach danych firmy Microsoft. Sprzęt, instalowanie i utrzymanie oprogramowania, zapewnienie wysokiej dostępności i odzyskiwania po awariach, a także instalowanie poprawek – wszystko to stanowi odpowiedzialność firmy Microsoft. Jednak to klient nadal odpowiada za indeksy i dostrajanie zapytań

Microsoft udostępnia kilka wariantów baz danych w wariantcie PaaS. Na potrzeby systemów OLTP oferowana jest usługa Azure SQL Database. Często jest ona skrótowo określana terminem SQL Database. Klient może mieć wiele baz danych na serwerze w chmurze (jest to oczywiście serwer „konceptualistyczny” – w istocie nie wiemy, czy baza jest utrzymywana przez konkretną maszynę, a przede wszystkim nie jest to w ogóle istotne), ale nie może przełączać się pomiędzy tymi bazami w ramach pojedynczej sesji.

Co interesujące, firma Microsoft wykorzystuje ten sam bazowy kod dla SQL Database i SQL Server. Tak więc większość funkcjonalności języka T-SQL jest dostępna (ostatecznie) w obu środowiskach w taki sam sposób. Dlatego większość wiedzy o języku T-SQL, którą zawiera ta książka, ma zastosowanie w obydwu środowiskach. Występujące różnice zostały opisane w dokumencie dostępnym pod adresem <https://azure.microsoft.com/en-us/documentation/articles/sql-database-transact-sql-information>. Warto też zauważyć, że tempo aktualizacji i wdrożeń nowych wersji SQL Database jest większe, niż w przypadku SQL Server w wersji pudełkowej. Oznacza to, że niektóre nowe funkcje T-SQL mogą być dostępne w SQL Database, zanim pojawią się w wersji SQL Server dla siedziby.

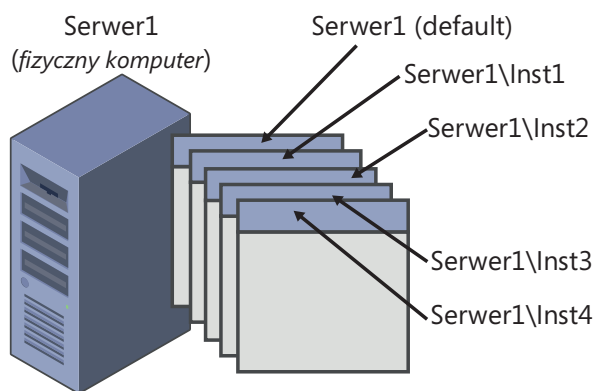
Microsoft udostępnia również ofertę typu PaaS dla systemów hurtowni danych, nazywaną Microsoft Azure SQL Data Warehouse (w skrócie SQL Data Warehouse). Usługa ta jest zasadniczo rozwiązaniem PDW/APS w chmurze. Firma Microsoft używa tego samego kodu bazowego dla wersji Appliance i usługi chmurowej. Do zarządzania i manipulowania danymi w APS i SQL Data Warehouse wykorzystywany jest język T-SQL, ale trzeba zauważyć, że nie jest to ta sama powłoka T-SQL, co w przypadku SQL Server i SQL Database.

Firma Microsoft udostępnia również inne chmurowe usługi danych, takie jak Data Lake dla rozwiązań klasy *big data*, Azure DocumentDB dla usług dokumentów NoSQL i inne.

Skomplikowane? Jeśli to kogoś pocieszy, nie jest osamotniony. Jak powiedziałem, firma Microsoft udostępnia przytłaczającą różnorodność technologii związanych z bazami danych. Osobliwe jest to, że jedynym wątkiem wspólnym dla większości z nich jest język T-SQL.

Instancje produktu SQL Server

W przypadku produktu pudełkowego *instancja* (wystąpienie) systemu SQL Server, jak ilustruje to rysunek 1-5, jest pełną, odrębną instalacją silnika bazy danych lub usługi SQL Server. Na tym samym komputerze (w tym samym systemie operacyjnym) możemy zainstalować wiele instancji produktu SQL Server w wersji dla siedziby. Każda instancja jest całkowicie niezależna od innych pod względem kontekstu zabezpieczeń, danych, którymi zarządza i w każdym innym aspekcie. Na poziomie logicznym dwie różne instancje umieszczone na tym samym komputerze nie mają więcej wspólnych elementów, niż dwie instancje umieszczone na różnych komputerach. Rzecz jasna, instancje umieszczone na tym samym komputerze współużytkują fizyczne zasoby serwera, takie jak CPU, pamięć czy dysk.



RYСУNEK 1-5 Wiele instancji SQL Server na tym samym komputerze

Jedną z wielu instancji na komputerze możemy skonfigurować jako *instancję domyślną*, natomiast pozostałe muszą być *instancjami nazwanymi*. To, czy instancja jest instancją domyślną, czy nazwaną, określane jest podczas instalacji; decyzji tej nie można później zmienić. W celu połączenia się z instancją domyślną aplikacja kliencka musi wskazać nazwę komputera lub jego adres IP. W celu połączenia się z instancją nazwaną program kliencki specyfikuje nazwę komputera lub adres IP, po których następuje odwrotny ukośnik (\) i nazwa instancji (określona podczas instalacji). Załóżmy dla przykładu, że mamy dwie instancje SQL Server zainstalowane na komputerze *Serwer1*. Jedna z tych instancji została zainstalowana jako domyślna, a druga jako instancja nazwana – *Inst1*. W celu podłączenia się do instancji domyślnej trzeba podać jedynie *Serwer1* jako nazwę serwera, natomiast aby połączyć się z instancją nazwaną, trzeba użyć zarówno nazwy serwera, jak i nazwy instancji: *Serwer1\Inst1*.

Istnieją różne powody, dla których instalowanych jest wiele instancji systemu SQL Server na tym samym komputerze, jednak wspomnę tylko o niektórych z nich. Jedną z nich jest chęć zmniejszenia kosztów obsługi. Na przykład w celu przetestowania działania funkcji w odpowiedzi na odebrane zgłoszenia lub odtworzenia napotkanych przez użytkowników błędów w środowisku produkcyjnym dla działu obsługi technicznej potrzebne są lokalne instalacje systemu SQL Server, które imitują środowisko