

Arthur MATEOS

Jothy ROSENBERG

The background of the cover is a vibrant blue sky with white clouds. In the center, a man in a dark suit is seen from behind, looking up at a large, glowing blue cloud. Inside this cloud, there is a complex arrangement of smaller images: a hand holding a globe, a person's head in profile, a person climbing a rope, and various abstract shapes and arrows. The overall theme is cloud computing and business solutions.

Chmura obliczeniowa

ROZWIĄZANIA DLA BIZNESU

**ODKRYJ WIELKĄ MOC CHMUR
OBLICZENIOWYCH I POTENCJAŁ,
JAKI KRYJĄ W SOBIE!**

- Co to jest chmura obliczeniowa?
- Kiedy korzystać z chmur, a kiedy ich unikać?
- Jak oszacować koszty korzystania z chmury obliczeniowej?

Helion



Tytuł oryginału: The Cloud at Your Service

Tłumaczenie: Justyna Walkowska

Projekt okładki: Jan Paluch

ISBN: 978-83-246-3416-3

Original edition copyright © 2011 by Manning Publications Co.

All rights reserved

Polish edition copyright © 2011 by Helion S.A.

All rights reserved

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from the Publisher.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w 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.

Materiały graficzne na okładce zostały wykorzystane za zgodą Shutterstock Images LLC.

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/chmura>

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

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

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

Spis treści

Słowo wstępne	9
Przedmowa	11
Podziękowania	13
O książce	17
1. Czym jest chmura obliczeniowa?	25
1.1. Pięć podstawowych zasad definiujących przetwarzanie w chmurze	27
1.1.1. Pula zasobów	28
1.1.2. Wirtualizacja zasobów obliczeniowych	29
1.1.3. Elastyczność wobec zmieniającego się zapotrzebowania	30
1.1.4. Automatyczne wdrażanie nowych zasobów	30
1.1.5. Naliczanie opłat: płacisz tylko za to, co faktycznie wykorzystasz	31
1.2. Zyski z przejścia na chmurę	31
1.2.1. Zyski ekonomiczne związane z zamianą wydatków inwestycyjnych na operacyjne	31
1.2.2. Zyski związane z elastycznością i brakiem zapotrzebowania na serwery	32
1.2.3. Zyski wydajnościowe dające przewagę nad konkurencją	33
1.2.4. Większe bezpieczeństwo w chmurze	33
1.3. Ewolucja w informatyce prowadząca do chmury obliczeniowej	33
1.3.1. Dlaczego „chmura”?	34
1.3.2. Zmiany paradygmatów przetwarzania: od samodzielnych jednostek, przez architektury klient-serwer, aż do sieci	35
1.3.3. Przechowywanie fizycznych zasobów obliczeniowych: ewolucja centrów danych	37
1.3.4. Modularyzacja oprogramowania i zdalny dostęp: wirtualizacja, SOA i SaaS	37

1.4. Klasyfikacja warstw chmury: różne typy do różnych zastosowań	38
1.4.1. Infrastruktura jako usługa (IaaS)	39
1.4.2. Platforma jako usługa (PaaS)	41
1.4.3. Oprogramowanie jako usługa (SaaS) i framework jako usługa (FaaS)	41
1.4.4. Chmury prywatne jako prekursorzy chmur publicznych	42
1.5. Podsumowanie	42
2. Klasyfikacja chmur obliczeniowych	43
2.1. Podstawy technologiczne przetwarzania w chmurze	44
2.1.1. Duże korzyści skali dzięki centrom danych w chmurze	45
2.1.2. Efektywne wykorzystanie serwerów w chmurze dzięki wirtualizacji	49
2.1.3. Sterowanie zdalnymi serwerami za pośrednictwem API chmury	52
2.1.4. Przechowywanie trwałych danych w chmurze	54
2.1.5. Przechowywanie danych aplikacji w chmurowej bazie danych	56
2.1.6. Elastyczność: skalowanie aplikacji w miarę zwiększania się lub zmniejszania popytu	62
2.2. Zrozumienie różnych typów chmur	63
2.2.1. Amazon EC2: IaaS	64
2.2.2. Microsoft Azure: IaaS	65
2.2.3. Google App Engine: PaaS	68
2.2.4. Ruby on Rails w chmurze: PaaS	69
2.2.5. Salesforce.com i Force.com: PaaS	70
2.2.6. Chmury prywatne: DaaS (centrum danych jako usługa)	70
2.3. Wybór chmury najlepiej dopasowanej do Twoich potrzeb	72
2.3.1. Amazon Web Services — chmura IaaS	72
2.3.2. Microsoft Azure — chmura IaaS i PaaS	73
2.3.3. Google App Engine — chmura PaaS	74
2.3.4. Ruby on Rails — chmura PaaS	74
2.3.5. Force.com — chmura PaaS	75
2.4. Podsumowanie	75
3. Analiza biznesowa chmury	77
3.1. Ekonomika przetwarzania w chmurze	78
3.1.1. Tradycyjna infrastruktura wewnętrzna, kolokacja, usługi zarządzane, a może model chmury?	79
3.1.2. Szczegółowe porównanie kosztów wdrażania w różnych modelach	81
3.2. Kiedy wdrożenie w chmurze ma sens?	86
3.2.1. Ograniczony czas życia lub zapotrzebowanie krótkoterminowe	87
3.2.2. Wąhnięcia skali	88
3.2.3. Aplikacje niestrategiczne	89
3.3. Kiedy wdrożenie w chmurze nie ma sensu?	90
3.3.1. Historyczne aplikacje	90
3.3.2. Aplikacje z krytycznymi scenariuszami czasu rzeczywistego	91
3.3.3. Aplikacje z dostępem do poufnych danych	91
3.4. Przedsiębiorstwa typu start-up bez kapitału zakładowego	92
3.4.1. Wtedy i teraz: tworzenie niewielkiego sklepu internetowego w 2000 i 2010 roku	92
3.4.2. Czy zewnętrzny kapitał inwestycyjny jest niezbędny?	93
3.4.3. Przykład 1.: FlightCaster — przewidywanie opóźnień lotów	94
3.4.4. Przykład 2.: analiza biznesowa jako SaaS	94

3.5. Małe i średnie przedsiębiorstwa	95
3.5.1. Prosty przykład: strona firmowa	95
3.5.2. Średnio skomplikowany przykład: kopie zapasowe i przechowywanie plików ...	96
3.5.3. Przykład zaawansowany: rozwijanie nowych produktów	96
3.6. Chmura w korporacjach	97
3.6.1. Eli Lilly: duży zbiór danych, obliczenia wysokowydajne	97
3.6.2. „The Washington Post”: duże problemy obliczeniowe z nieprzekraczalnymi terminami	98
3.6.3. Virgin Atlantic: obecność w sieci i zgromadzenie społeczności	99
3.7. Podsumowanie	99
 4. Bezpieczeństwo i chmura prywatna	 101
4.1. Bezpieczeństwo informacji w chmurze publicznej	102
4.1.1. Obawy o bezpieczeństwo spowalniające ekspansję chmury	103
4.1.2. Bezpieczeństwo największych centrów danych w chmurze	104
4.1.3. Środki kontroli dostępu w chmurze publicznej	106
4.1.4. Bezpieczeństwo sieciowe i bezpieczeństwo danych w dużych chmurach	111
4.1.5. Rola i zakres odpowiedzialności właściciela aplikacji	114
4.2. Przyczyny powstania chmury prywatnej	115
4.2.1. Definicja chmury prywatnej	115
4.2.2. Kwestie bezpieczeństwa	117
4.2.3. Pewność dostępności zasobów	117
4.2.4. Duża społeczność	118
4.2.5. Efekty skali	118
4.2.6. Potencjalne problemy z chmurą prywatną	119
4.2.7. Sposoby wdrożenia chmury prywatnej	119
4.3. Wirtualna chmura prywatna	124
4.3.1. Jak to działa?	124
4.3.2. API wirtualnej chmury prywatnej	125
4.3.3. Konsekwencje	126
4.4. Chmury prywatne w praktyce	126
4.4.1. Sprint: chmura prywatna dla aplikacji wykrywającej oszustwa	127
4.4.2. Project Services Network (PSN) firmy Bechtel	127
4.4.3. Rządowe chmury prywatne	128
4.5. Długoterminowa prognoza dla chmury prywatnej	129
4.6. Podsumowanie	130
 5. Projektowanie i architektura aplikacji w chmurze	 131
5.1. Wzorce aplikacji najlepiej pasujące do chmury	132
5.1.1. Przeniesienie	132
5.1.2. Skala internetowa	133
5.1.3. Ekspansja obliczeń	133
5.1.4. Elastyczne składowanie danych	134
5.1.5. Podsumowanie wzorców aplikacji	134
5.2. Projektowanie i architektura w skali internetowej: shardowanie	134
5.2.1. Cechy aplikacji blokujące skalowalność	136
5.2.2. Shardowanie: zrównoleglona architektura bazy danych umożliwiająca skalowanie	137
5.2.3. Jak shardowanie zmienia aplikację	139
5.2.4. Porównanie shardowania z tradycyjnymi architekturami baz danych	140

5.2.5.	Shardowanie w praktyce: najpopularniejsze schematy partycjonowania baz danych	143
5.2.6.	Trudności i problemy związane ze shardowaniem	145
5.2.7.	Shardowanie w praktyce: jak robi to Flickr?	148
5.3.	Zwiększenie mocy na życzenie: cloudbursting	150
5.3.1.	Cloudbursting: definicja	150
5.3.2.	Dwie pieczenie na jednym ogniu: wewnętrzne centrum danych oraz chmura	151
5.3.3.	Cloudbursting: analiza biznesowa	152
5.3.4.	Cloudbursting: architektura	154
5.3.5.	Jak zaimplementować cloudbursting?	156
5.3.6.	Cloudbursting: potrzeba standaryzacji	157
5.3.7.	Cloudbursting: problem dostępu do danych	157
5.4.	Jak przygotować się na wykładniczy przyrost ilości składowanych danych?	160
5.4.1.	Magazyn danych w chmurze: definicja	160
5.4.2.	Amazon S3	161
5.4.3.	Przykładowy interfejs magazynu danych w chmurze (S3)	161
5.4.4.	Koszty	164
5.4.5.	Montowalne systemy plików w chmurze	164
5.4.6.	Jak sobie radzić z opóźnieniami?	165
5.5.	Podsumowanie	166
6.	Niezawodność w skali chmury	167
6.1.	SOA jako prekursor chmury	168
6.1.1.	Systemy rozproszone	168
6.1.2.	Luźne sprzężenie	170
6.1.3.	SOA	172
6.1.4.	SOA i luźne sprzężenie	173
6.1.5.	SOA i usługi sieciowe	174
6.1.6.	SOA i przetwarzanie w chmurze	175
6.1.7.	Komunikacja między procesami w chmurze	176
6.2.	Niezawodność wysokowydajnych, rozproszonych aplikacji w chmurze	176
6.2.1.	Nadmiarowość	177
6.2.2.	MapReduce	178
6.2.3.	Hadoop: MapReduce w wersji open source	183
6.3.	Podsumowanie	184
7.	Testy, wdrożenie i działanie w chmurze	185
7.1.	Typowe wdrożenia	186
7.1.1.	Tradycyjna architektura wdrożeniowa	187
7.1.2.	Środowisko testowe i środowisko etapu pośredniego	188
7.1.3.	Wyliczenie kosztów	189
7.2.	Chmura na ratunek!	189
7.2.1.	Poprawa jakości produkcyjnej dzięki chmurze	190
7.2.2.	Szybsze wytwarzanie aplikacji oraz testowanie	192
7.3.	Siła równoległości	195
7.3.1.	Testy jednostkowe	196
7.3.2.	Testy funkcjonalne	198
7.3.3.	Testy obciążeniowe	201
7.3.4.	Testy wizualne	204
7.3.5.	Testy ręczne	206
7.4.	Podsumowanie	207

8. Kwestie praktyczne	209
8.1. Wybór dostawcy chmury	210
8.1.1. Kwestie biznesowe	210
8.1.2. Kwestie techniczne	212
8.2. Chmura publiczna i SLA	218
8.2.1. SLA dla Amazon AWS	219
8.2.2. SLA dla Microsoft Azure	220
8.2.3. SLA dla chmury Rackspace	221
8.3. Pomiary jakości operacji w chmurze	222
8.3.1. Widoczność i monitorowanie u dostawcy jakości świadczonych przez niego usług	222
8.3.2. Widoczność i monitorowanie jakości usług, które świadczy dostawca, za pomocą rozwiązań zewnętrznych firm	226
8.4. Podsumowanie	228
9. Przyszłość chmury	229
9.1. Najważniejsza transformacja w dziejach informatyki	231
9.1.1. Internet konsumentów oraz chmura	231
9.1.2. Chmura w przedsiębiorstwach	235
9.2. Dziesięć prognoz na temat ewolucji chmury	239
9.2.1. Tańsza, bardziej niezawodna, bezpieczniejsza i prostsza w użyciu	240
9.2.2. Motor napędzający wzrost prekursorów	241
9.2.3. Koszty niższe niż w firmowych centrach danych	241
9.2.4. Do 2020 roku — 500 tysięcy serwerów wartych miliard dolarów	242
9.2.5. Administratorzy a serwery — 1:10 000 w 2020 roku	243
9.2.6. Dominacja open source	243
9.2.7. Pragmatyczne standardy i rola Amazon API	244
9.2.8. Ostateczny standard ISO dla chmury	245
9.2.9. Rząd jako prekursor w chmurze	247
9.2.10. SaaS i standardy	247
9.3. Dziesięć prognoz na temat ewolucji sposobu wytwarzania aplikacji	248
9.3.1. Rola szkieletów aplikacji	248
9.3.2. Druga i trzecia warstwa działające w chmurze	249
9.3.3. Gwałtowna ewolucja mechanizmów składowania danych	250
9.3.4. Lepsza ochrona wrażliwych danych	251
9.3.5. Usługi wyższego poziomu o własnych API	252
9.3.6. Wzrost znaczenia aplikacji typu mashup	252
9.3.7. Dominacja PaaS i FaaS	254
9.3.8. Narzędzia do tworzenia aplikacji typu mashup	254
9.3.9. Sukces programistów spoza świata zachodniego	255
9.3.10. Koszty wytworzenia aplikacji nie są przeszkodą	256
9.4. Podsumowanie	256
9.4.1. Pięć podstawowych zasad przetwarzania w chmurze	256
9.4.2. Główne zyski z przejścia na chmurę	257
9.4.3. Chmura powstała na drodze ewolucji	257
9.4.4. Klasyfikacja chmur: od IaaS do SaaS	257
9.4.5. Podstawy technologiczne	258
9.4.6. Opłaty tylko za rzeczywiste zużycie	258
9.4.7. Przesadne obawy o bezpieczeństwo	259
9.4.8. Chmury prywatne jako zjawisko tymczasowe	259
9.4.9. Projektowanie z myślą o skali i shardowanie	260

9.4.10. Projektowanie z myślą o niezawodności i MapReduce	260
9.4.11. Lepsze testy, wdrożenia i działanie w chmurze	261
9.4.12. Wybór dostawcy	261
9.4.13. Monitorowanie chmur publicznych i SLA	261
9.4.14. Przyszłość chmury obliczeniowej	261
A. Powtórka z bezpieczeństwa	263
Sekretna komunikacja	264
Klucze	265
Kryptografia klucza współdzielonego	265
Kryptografia klucza publicznego	266
XML Signature	268
XML Encryption	268
Skorowidz	271

Klasyfikacja chmur obliczeniowych



W tym rozdziale:

- ♦ podstawy technologiczne wspólne dla wszystkich typów chmur,
- ♦ klasyfikacja typów chmur i ich możliwości,
- ♦ zasady wyboru chmury odpowiedniego typu i najlepszego dostawcy.

Poprzedni rozdział był wprowadzeniem do świata przetwarzania w chmurze. W tym zajmujemy się bardziej szczegółową analizą różnych typów chmur i charakterystycznych dla nich sposobów działania. Za przykład niech nam posłuży samochód — jeśli chmura jest samochodem, to nowoczesne centrum danych odgrywa rolę silnika, a virtualizacja odpowiada resorom, chroniącym nas na wyboistej drodze. API chmury przypomina deskę rozdzielczą, pozwalającą kontrolować chmurę. Magazyn danych działa jak bagażnik, umożliwiając transportowanie różnych rzeczy. Bazy danych w chmurze to system nawigacyjny — konkretne informacje niezbędne w podróży. W zależności od potrzeb samochód elastycznie przedstawia się na inny bieg — można porównać to do elastyczności aplikacji, która może obsługiwać jednego użytkownika, a chwilę później rozszerza się tak, że możliwa jest obsługa miliona zgłoszeń. Istnieje wiele modeli pojazdów i wiele typów chmur. Przyjrzymy się z bliska dominującym dzisiaj typom. Ocenimy, czy potrzebujesz wyścigówki, gdyż zależy Ci na prędkości, czy może wolisz osiemnastokółowego olbrzyma, który zapewni Ci mnóstwo przestrzeni.

Na początek przeanalizujemy najbardziej niezbędne aspekty technologiczne chmury, by dobrze zrozumieć, z jakich elementów jest ona zbudowana. W rozdziale 1. omówiliśmy wstępnie różne typy chmur i porównaliśmy je ze sobą. Teraz rozwinieśmy zagadnienie, by ułatwić Ci podjęcie decyzji na temat wyboru typu chmury oraz wspomóc najbardziej efektywne jej wykorzystanie.

2.1. Podstawy technologiczne przetwarzania w chmurze

Większość kierowców zna podstawy działania samochodu — niektórzy poznali je z czystej ciekawości, innym zależało na byciu lepszymi kierowcami i właścicielami auta. W przypadku chmury również możliwe jest wyodrębnienie ogólnych zasad działania niezależnych od typu. Omówimy je dla lepszego zrozumienia późniejszych kwestii:

- ◆ *Chmura potrzebuje serwerów sieciowych, te zaś potrzebują domu.* Serwery zgromadzone w pewnej fizycznej lokalizacji będziemy określali mianem centrum danych.
- ◆ *Serwery w chmurze należy zwirtualizować.* Celem tego procesu jest zwiększenie efektywności banku serwerów. W przeciwnym razie koszty utrzymania serwerów obniżą efektywność finansową chmury.
- ◆ *Chmura potrzebuje API.* Bez interfejsu dostępowego zwirtualizowane serwery w chmurze tkwiłyby w ciszy i samotności. Użytkownicy muszą mieć dostęp do chmury. Chcą zamawiać nowe serwery, wysyłać i pobierać dane, uruchamiać i zatrzymywać aplikacje, a także pozbywać się serwerów, które przestały już być potrzebne. Wszystko to ma się odbywać zdalnie, gdyż stopa użytkowników nigdy nie postanie w fizycznej serwerowni.

- ◆ *Chmura musi gdzieś przechowywać dane.* Musi składować obrazy maszyn wirtualnych, aplikacje użytkowników i wykorzystywane przez nie trwałe dane.
- ◆ *Chmura wymaga bazy danych.* Większość aplikacji do działania potrzebuje ustrukturyzowanych danych — w konsekwencji potrzebuje ich także chmura.
- ◆ *Chmura musi być elastyczna i skalować się dynamicznie zgodnie z potrzebami aplikacji i ich użytkowników.* Dla użytkowników chmury możliwość dostosowania ilości wykorzystywanych zasobów (oraz płatności za nie) do aktualnego obciążenia aplikacji to jedna z największych korzyści z przejścia na chmurę.

Omówimy teraz bardziej szczegółowo każdy z wymienionych sześciu aspektów technologicznych i infrastrukturalnych.

2.1.1. Duże korzyści skali dzięki centrom danych w chmurze

Wróćmy do analogii samochodowej — centrum danych to silnik samochodu. Centrum danych, posiadane obecnie przez każdą większą instytucję, to obiekt (nierzadko pilnie strzeżony), w którym znajduje się wiele komputerów i innego sprzętu sieciowo-komunikacyjnego. Duże korporacje internetowe, takie jak Amazon, Yahoo, Google, Intuit czy Apple, w ciągu lat zgromadziły *megacentra* danych z tysiącami serwerów. Stanowią one punkt wyjścia infrastruktury przetwarzania w chmurze.

Warto dobrze zrozumieć strukturę i ekonomikę tych ogromnych centrów danych. To na nich oparte są mechanizmy skalowania, niezawodność przetwarzania, bezpieczeństwo danych i przyszłość rozwoju chmur publicznych. Rozważ te kwestie, zwłaszcza jeśli myślisz o stworzeniu własnej chmury prywatnej. Jeszcze w tym rozdziale opowiemy co nieco o chmurach prywatnych. Dodatkowo cały rozdział 4. jest poświęcony chmurom prywatnym i bezpieczeństwu.

Struktura centrum danych

Centrum danych może mieścić się w jednym pokoju, zajmować jedno albo więcej pięter, a nawet objąć we władanie cały budynek. Większość jego wyposażenia z reguły stanowią serwery upakowane w dziewiętnastocalowych szafach, najczęściej ustawianych w rzędach, pomiędzy którymi tworzą się korytarze umożliwiające dostęp do serwerów z dwóch stron. Serwery różnią się rozmiarami. Serwer wielkości 1U (ang. *rack unit*) zajmuje jeden slot (z czterdziestu dwóch) w szafie montażowej, ale istnieją także samostojące serwery, których nie montuje się w szafach. Komputery typu *mainframe* oraz niektóre magazyny danych mogą być tak duże jak szafy i są ustawiane w szeregach razem z nimi. Duże centra danych czasami korzystają z kontenerów po tysiąc lub więcej serwerów każdy. Jeśli potrzebna jest naprawa lub aktualizacja, wymienia się cały kontener, a nie poszczególne serwery.

Niezbędny jest czysty, stabilny dopływ energii. Komputery w centrach danych działają bez przerwy. Muszą być odporne na spadki napięcia, a nawet przerwy w dostawach prądu. Serwerownia musi być zaopatrzona w stabilizator napięcia, zapasowe baterie oraz generatory prądu, zapewniające nieograniczony, nieprzerwany dopływ prądu. Ważne jest także chłodzenie sprzętu. Najczęściej stosuje się do tego schłodzone powietrze, jednak pojawiają się również systemy wykorzystujące wodę, jeśli jest ona łatwo dostępna. Wodą chłodzone są na przykład nowe centra danych położone wzdłuż rzeki Kolumbia w stanie Waszyngton. Klimatyzacja ma za zadanie nie tylko obniżenie temperatury, lecz także kontrolę wilgotności, by nie doszło do skraplania się lub wyładowań elektrostatycznych.

Centrum danych wymaga szerokopasmowego łącza umożliwiającego odbiór i wysyłanie danych. Istnienie serwerów i magazynów danych nie ma sensu, jeśli nikt nie może się z nimi połączyć.

Ważne jest także bezpieczeństwo, zarówno fizyczne, jak i logiczne. Duże centra danych są pod ciągłym ostrzałem hakerów z całego świata. Procedury bezpieczeństwa w niektórych centrach danych zaczynają się od tego, że ukrywane jest samo istnienie centrum w danej lokalizacji. Strażnicy, pułapki i najnowocześniejsze mechanizmy autentykacji mają za zadanie fizycznie powstrzymać nieproszonych gości. Firewalle, bramki VPN, oprogramowanie wykrywające włamania i inne tego typu rozwiązania chronią centrum przed włamaniem przez sieć (więcej na temat bezpieczeństwa w chmurze w rozdziale 4.).

Ponadto centra danych muszą być przygotowane na najgorsze oraz posiadać strategię wznowienia i utrzymania struktury po awarii, włamaniu lub innym zdarzeniu, by uniknąć utraty danych i zminimalizować okres nieaktywności.

Centra danych: skalowanie na potrzeby chmury

Tradycyjne duże centra danych przeznaczone do obsługi jednej dużej korporacji kosztują między 100 a 200 milionami dolarów¹. Dla porównania: całkowity koszt budowy największych megacentrów danych, świadczących usługi przetwarzania w chmurze, wynosi ponad 500 milionów dolarów^{2,3}. Co powoduje taki przyrost kosztów? Co mają największe centra przeznaczone dla chmury, czego nie mają tradycyjne centra dedykowane jednej firmie?

Najwięksi operatorzy, tacy jak Google, Amazon czy Microsoft, lokują swoje centra w niedużej odległości od rejonów, których mieszkańcy intensywnie z nich korzystają, co zmniejsza opóźnienia i ułatwia przełączanie się pomiędzy maszynami w przypadku awarii. Szukają także miejsc, w których energia elektryczna jest tania. W Stanach Zjednoczonych takim obszarem jest Północny Zachód — elektrownie wodne produkują najtańszą energię w kraju, a koszty chłodzenia są bliskie zeru. Sama konsumpcja energii przez duże centrum danych może kosztować więcej niż 30 milionów dolarów rocznie. Obecnie zużycie energii przez takie centra stanowi 1,2 procenta całkowitego zużycia energii w kraju i wciąż rośnie.

¹ <http://perspectives.mvdirona.com/2008/11/28/CostOfPowerInLargeScaleDataCenters.aspx>.

² <http://www.datacenterknowledge.com/archives/2007/11/05/microsoft-plans-500m-illinois-data-center>.

³ http://www.theregister.co.uk/2009/09/25/microsoft_chillerless_data_center.

Światowe serwery emitują więcej dwutlenku węgla niż cała Holandia

Firma doradcza McKinsey & Co. donosi, że 44 miliony działających na świecie serwerów są odpowiedzialne za 0,5 procenta światowej konsumpcji energii oraz 0,2 procenta emisji dwutlenku węgla, co odpowiada 80 megatonom rocznie. Zbliżone ilości dwutlenku węgla emitują takie kraje, jak Argentyna czy Holandia.

Plusem dla operatorów jest to, że przy takim zapotrzebowaniu na energię mogą negocjować pętelne zniżki.

Dodatkowo megacentra mogą negocjować ogromne zniżki na zakup sprzętu, wykraczające daleko poza zniżki oferowane największym nawet firmom budującym prywatne, dedykowane serwerownie. Przykładowo w 2008 roku firma Amazon zapłaciła około 90 milionów dolarów za 50 tysięcy serwerów firmy Rackable/SGI⁴. W normalnym handlu kosztowałyby one 215 milionów.

Serwery stanowią największą część kosztów ponoszonych przez centra danych. Dlatego Google i inni próbują ciąć koszty, samodzielnie konstruując serwery z gotowych komponentów. Google polega na tanich komputerach ze zwykłymi procesorami wielordzeniowymi. Jedno centrum danych Google posiada dziesiątki tysięcy takich niedrogich procesorów i dysków pospianych rzepami, co umożliwia łatwą wymianę komponentów.

By ograniczyć apetyt maszyn na energię, firma Google wprowadziła mechanizmy optymalizujące dopływ prądu, wentylatory o zmiennej prędkości oraz płyty główne wypatroszone ze wszystkich zbędnych elementów, w tym kart graficznych. Eksperymentowano także z **dynamicznym skalowaniem napięcia i częstotliwości**. Napięcie lub częstotliwość procesora są obniżane w pewnych okresach (np. gdy wynik obliczeń nie jest wymagany natychmiast). Serwer działa wolniej, co zmniejsza konsumpcję energii. Inżynierowie Google twierdzą, że w niektórych testach oszczędność energii wyniosła 20 procent.

W roku 2006 firma Google wybudowała w miejscowości Dalles w Oregonie dwa centra danych, z których każde zajmuje teren wielkości boiska futbolowego, ma cztery piętra i czteropiętrową chłodnię (rysunek 2.1). Strategiczne znaczenie dla centrum danych ma tama w Dalles, zapewniająca dopływ energii oraz chłodzenie. (Niektóre nowe centra danych korzystają z chłodni kominowych, które odprowadzają wodę wykorzystaną do chłodzenia, co zużywa o wiele mniej energii niż tradycyjne chłodziarki).

Centrum danych w Dalles jest doskonale połączone istniejącymi już wcześniej światłowodami z różnymi lokalizacjami w USA, Azji i Europie. Kable te są spuszczone po bańce internetowej.

W 2007 roku firma Google stworzyła co najmniej cztery nowe centra danych warte średnio 600 milionów dolarów. Weszły one w skład Googleplexu: wielkiej światowej sieci komputerowej, szacowanej na 450 tysięcy serwerów w dwudziestu pięciu lokalizacjach. Firma Amazon również zdecydowała się umieścić swoje największe centrum danych w Dalles, kawałek dalej wzdłuż rzeki.

⁴ <http://www.datacenterknowledge.com/archives/2009/06/23/amazon-adds-cloud-data-center-in-virginia>.



Rysunek 2.1. Zdjęcie przedstawia ściśle tajne centrum danych Google w miejscowości Dalles w Oregonie, zbudowane w pobliżu tamy w celu uzyskania dostępu do taniej energii. Zwróć uwagę na duże chłodnie kominowe na końcach znajdujących się po lewej stronie zdjęcia budynków wielkości boiska do futbolu. Kominy chłodzą budynki przez parowanie, bez używania kosztownych energetycznie chłodziarek. Źródło: Melanie Conner, „New York Times”

Yahoo! i Microsoft wybrały miejscowość Quincy w Waszyngtonie. Budynek Microsoftu ma powierzchnię porównywalną niemal z obszarem dziesięciu boisk futbolowych. Przedstawiciele firmy nabrali wody w usta, jeśli chodzi o liczbę serwerów, jednak wiadomo, że wykorzystuje się tam prawie pięć kilometrów rur chłodzących, prawie tysiąc kilometrów przewodów elektrycznych, niemal 10 tysięcy metrów kwadratowych płyt kartonowo-gipsowych i 1,6 tony baterii zapewniających prąd w razie awarii. Centrum zużywa 48 megawatów energii, co wystarczyłoby dla 40 tysięcy domów.

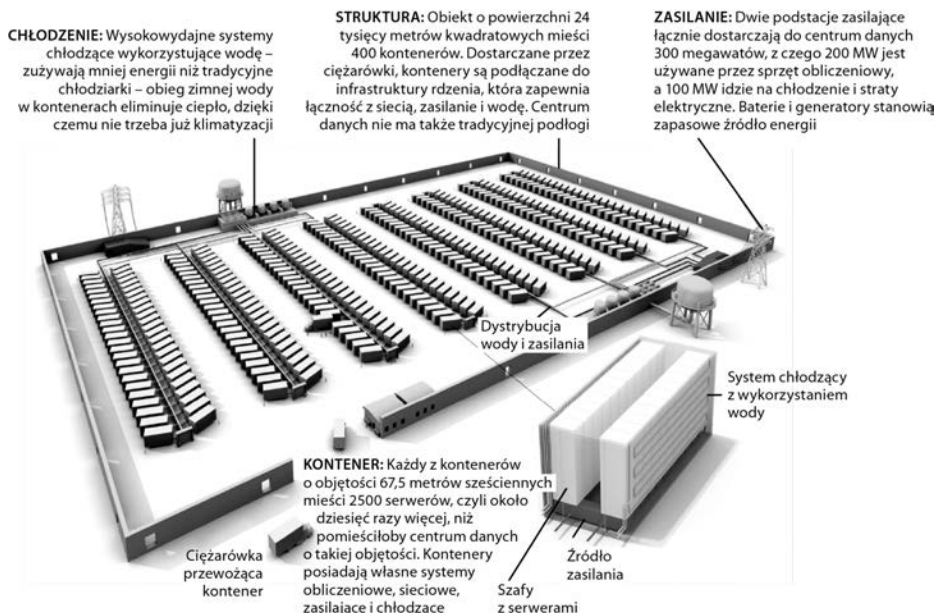
Chmurowe centra danych: większa efektywność i elastyczność dzięki modularyzacji

Już teraz, dzięki zakupom hurtowym, serwerom budowanym na zamówienie i starannie wybranym lokalizacjom, właściciele największych centrów danych ponoszą o wiele mniejsze koszty w przeliczeniu na operacje procesora niż zwykłe firmy. Robią, co mogą, by jeszcze zwiększyć tę różnicę. Centra danych w chmurze stają się coraz bardziej efektywne dzięki modularyzacji. Modułowe, skalowalne, efektywne centra danych są natychmiast dostępne za niewielką cenę z każdego miejsca na świecie.

Rysunek 2.2 przedstawia wizualizację modularnego centrum danych (fotografie takich obiektów są pilnie strzeżone). Firmowe centra danych zostały w tyle za wydajnymi megacentrami, a ich sytuacja będzie się jeszcze pogarszała.

Modularyzacja ma na celu ustandaryzowanie centrów danych i odstąpienie od własnych projektów, co ma ułatwić produkcję sprzętu dla takich centrów. Jedną z najbardziej uderzających postulowanych cech jest to, że takie centra znajdują się pod gołym niebem.

Podobnie jak Google Microsoft stara się ciąć koszty energii, ulega także naciskom ekologów, by ograniczać emisję i poprawić wydajność. Docelowy współczynnik wydajności energetycznej PUE (ang. *Power Usage Effectiveness*) to co najwyżej 1,125 we wszystkich centrach danych Microsoftu do 2012 roku.



Rysunek 2.2. Rozszerzalne, modułowe centrum danych dla chmury. Warto zwrócić uwagę na brak zadaszenia. Kontenery z serwerami, zasilaniem, sprzętem chłodzącym i monitorującym mogą być dodawane i usuwane w miarę potrzeb. Źródło: magazyn „IEEE Spectrum”

2.1.2. Efektywne wykorzystanie serwerów w chmurze dzięki wirtualizacji

Pozostawmy przy analogii samochodowej — wirtualizacja to zawieszenie. Zapewnia ona pożądane intensywne wykorzystanie serwerów. Łagodzi różnice pomiędzy aplikacjami, które prawie nie potrzebują mocy obliczeniowej (mogą one dzielić procesor z innymi aplikacjami), a tymi, które domagają się każdego istniejącego cyklu procesora. Z wszystkich rewolucyjnych technologii wykorzystywanych w chmurze to właśnie wirtualizacja i jej udane wdrożenie zadecydowały o przyjęciu się nowego trendu. Bez wirtualizacji i umożliwianego przez nią zużycia procesora wynoszącego ponad 60 procent chmura nie byłaby opłacalna.

Współczynnik wydajności energetycznej (PUE)

Współczynnik wydajności energetycznej (PUE, ang. *Power Usage Effectiveness*) określa wydajność centrum danych. PUE wyznacza się, dzieląc ilość energii dostarczanej przez centrum danych przez ilość niezbędną do działania infrastruktury obliczeniowej wewnątrz. Wydajność poprawia się, gdy PUE zmniejsza się w kierunku wartości 1.

Według organizacji Uptime Institute średnia wartość PUE typowego centrum danych to 2,5. Oznacza to, że z każdego 2,5 wata na liczniku tylko wat jest wykorzystywany do obliczeń. Z szacunków Uptime wynika, iż większość centrów mogłaby osiągnąć wynik 1,6 PUE, stosując najlepszy sprzęt i dobre praktyki. Google i Microsoft zbliżają się do wartości 1,125 — to wynik nieosiągalny dla firmowych, a nawet współdzielonych centrów danych.

WIRTUALIZACJA

W tej książce koncentrujemy się głównie na wirtualizacji *platformy*. Wirtualizacja platformy to technika abstrakcji zasobów komputera oddzielająca system operacyjny od sprzętu. System operacyjny nie odwołuje się bezpośrednio do sprzętu. Zamiast tego odwołuje się on do nowej warstwy programistycznej, nazywanej *monitorem maszyny wirtualnej*, która ma dostęp do sprzętu i przedstawia systemowi wirtualny zestaw zasobów sprzętowych. Oznacza to, że wiele obrazów maszyn wirtualnych lub instancji może działać na jednym fizycznym serwerze, a nowe instancje mogą być tworzone i uruchamiane na żądanie. W ten sposób tworzona jest podstawa elastycznych zasobów obliczeniowych.

Wspomnieliśmy wcześniej, że wirtualizacja wcale nie jest nowym pomysłem. Komputery typu *mainframe* produkowane przez IBM stosowały wirtualizację czasową już w latach sześćdziesiątych, umożliwiając wielu osobom korzystanie z jednej maszyny w różnym czasie bez żadnej interakcji i wchodzenia sobie w drogę. Wcześniej na użytkowników nakładany był obowiązek wykonania całości planowanej pracy w jednym, wyznaczonym okienku czasowym. Pojęcie pamięci wirtualnej, wprowadzone około 1962 roku, choć uznawane wówczas za radykalne, raz na zawsze uwolniło programistów od ciągłego zamartwiania się możliwością przekroczenia wielkości pamięci fizycznej. Dzisiaj wirtualizacja serwerów ma równie znaczący wpływ na wdrażanie i skalowanie aplikacji, będąc zarazem jednym z największych motorów rozwoju chmury. Jak do tego doszło?

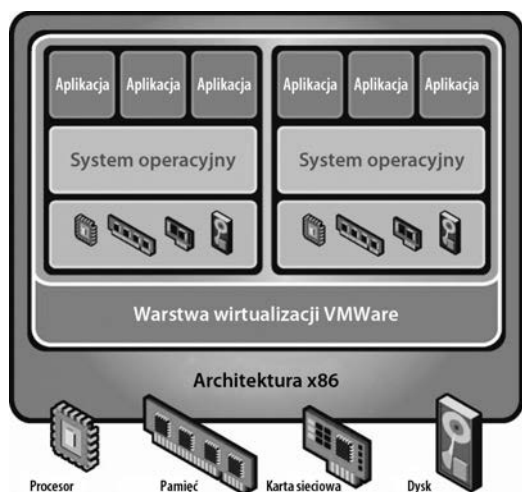
Przeciętny serwer w firmowym centrum danych jest wykorzystywany w około 6 procentach⁵. Nawet w chwilach maksymalnego obciążenia wykorzystanie nie przekracza 20 procent. W najlepiej zarządzanych centrach danych serwery wykorzystują średnio około 15 procent swoich możliwości. Jednak gdy takie centra zdecydują się na pełną wirtualizację, wykorzystanie mocy procesora wzrasta do 65 procent lub więcej. Z tego powodu w ciągu ostatnich paru lat w większości centrów danych wdrożono setki lub tysiące serwerów wirtualnych w miejsce poprzedniego modelu, w którym jeden serwer odpowiadał jednej fizycznej jednostce. Przyjrzyjmy się teraz temu, co sprawia, że wykorzystanie zmienia się aż tak radykalnie.

Jak to działa?

Wirtualizacja serwera przekształca (*wirtualizuje*) zasoby sprzętowe komputera — w tym czas procesora, pamięć RAM, dysk twardy i kontroler sieciowy — by stworzyć w pełni funkcjonalną maszynę wirtualną, na której system operacyjny i inne aplikacje mogą działać tak, jakby działały na fizycznym komputerze. Efekt ten osiąga się, dodając cienką warstwę oprogramowania działającą bezpośrednio nad sprzętem. Warstwa ta zawiera monitor maszyny wirtualnej (VMM, ang. *Virtual Machine Monitor*), który przydziela zasoby sprzętowe w dynamiczny i przejrzysty sposób. Dzięki temu na jednej fizycznej maszynie może działać wiele systemów operacyjnych, które współbieżnie korzystają z zasobów sprzętowych. Enkapsulacja całej maszyny, w tym procesora, pamięci i urządzeń sieciowych,

⁵ McKinsey & Company, 2008 *Data Center Efficiency* — raport na temat wydajności centrów danych.

sprawia, że maszyna wirtualna staje się w pełni kompatybilna ze standardowymi systemami operacyjnymi, aplikacjami i sterownikami urządzeń. Architektura maszyny wirtualnej VMware dla procesorów x86 została przedstawiona na rysunku 2.3.



Rysunek 2.3. Architektura maszyny wirtualnej na przykładzie VMware. Warstwa wirtualizacji odwołuje się bezpośrednio do komponentów sprzętowych, w tym do procesora. Warstwa ta prezentuje każdemu z goszczących na maszynie systemów operacyjnych jego własny zestaw zasobów sprzętowych. Taki system zachowuje się dokładnie tak samo jak system faktycznie mający dostęp do sprzętu. Dzięki wirtualizacji wiele systemów i działających w nich aplikacji może korzystać z jednej fizycznej maszyny, osiągając lepszą wydajność.
Źródło: VMware

Zastosowanie wirtualizacji w chmurze

Wirtualizacja szybko znalazła uznanie u architektów systemów i kierowników działów IT z bardzo prostego powodu — pozwala zaoszczędzić pieniądze. Radikalnie zwiększył się stopień wykorzystania zasobów sprzętowych. Bez większego wysiłku można było przejść od 5 czy 6 do 20 procent. Przy rozsądnym planowaniu możliwe stało się osiągnięcie wyniku 65 procent wykorzystania sprzętu.

Poza lepszym współczynnikiem wykorzystania sprzętu i związanymi z nim oszczędnościami wirtualizacja w firmowych centrach danych w dużej mierze przygotowała grunt dla przetwarzania w chmurze. Użytkownicy zostali oddzieleni od implementacji, poprawiła się szybkość, elastyczność i „zwinnosc”, a do tego zmienił się nienaruszalny dotąd model wyceny i licencjonowania oprogramowania. Tabela 2.1 zawiera objaśnienia poszczególnych elementów.

Tabela 2.1 przedstawia najważniejsze zyski z wprowadzenia chmury. Chmura zdobywa coraz większe uznanie i coraz więcej firm jest gotowych, by przejść na ten model obliczeniowy. Wiele firm już wcześniej korzystało z wirtualizacji, więc przejście na chmurę jest dla nich bezbolesne.

Rozpatrzmy scenariusz z wykorzystaniem tysięcy serwerów fizycznych. Każdy z nich został zwirtualizowany, w związku z czym może na nim działać dowolnie wiele systemów operacyjnych. Konfiguracja i wdrażanie trwają parę

Tabela 2.1. Wpływ wirtualizacji na firmowe centra danych

Zysk	Objaśnienie
Oddzielenie użytkowników od implementacji	Pojęcie serwera wirtualnego sprawia, że użytkownik nie musi, a nawet nie może martwić się o fizyczny serwer lub jego lokalizację. Zamiast tego koncentruje się na umowach i definicjach usług oraz wdrażanych aplikacjach.
Pozyskanie nowego serwera trwa parę minut, a nie kilka miesięcy	Zamówienie, instalacja, konfiguracja i uruchomienie fizycznego serwera w dużych firmach zajmuje 60, 90, a czasem nawet 120 dni. W modelu serwerów wirtualnych czas pomiędzy żądaniem a wdrożeniem gotowej aplikacji jest liczony w minutach lub godzinach, w zależności od stopnia automatyzacji.
Nowy model naliczania opłat i licencjonowania	Centrum danych nie może już naliczać opłat za cały serwer ani za wszystkie serwery, na których działa aplikacja. Zamiast tego nalicza opłaty za faktyczne zużycie — to rewolucja w informatyce.

minut, a opłaty są pobierane za każdą godzinę wykorzystania procesora. Zasoby megacentrów danych stały się dostępne dla zwykłych użytkowników dzięki połączeniu następujących elementów: wirtualizacji, automatycznego dodawania i usuwania sprzętu oraz nowych zasad naliczania opłat. Jako odpowiednik samochodowych rezerwacji wirtualizacja sprawia, że aplikacja może gwałtownie przyspieszyć, nie zabijając przy tym wszystkich siedzących w pojeździe po wjechaniu na pierwszy wybór.

Jednak potężny silnik (centrum danych) i dobre zawieszenie (wirtualizacja) to nie wszystko. Potrzebna jest jeszcze deska rozdzielcza — zestaw urządzeń pozwalających na ruszanie, zatrzymywanie oraz sterowanie autem. Niezbędny jest interfejs dający kontrolę nad chmurą.

2.1.3. *Sterowanie zdalnymi serwerami za pośrednictwem API chmury*

API (interfejs programowania aplikacji) jest dla chmury tym, czym deska rozdzielcza dla samochodu. Pod maską może czaić się prawdziwy potwór, jednak bez kontrolek i pokręteł nie dowiesz się, co właściwie dzieje się z pojazdem ani jak nim sterować. Potrzebujesz przynajmniej kierownicy, gazu i hamulca. Pamiętaj także o tym, że nie można jeździć szybko bez dobrych hamulców.

Do chmury trzeba się jakoś dostać. Chmury najwyższego poziomu — te oferujące aplikacje SaaS (oprogramowanie jako usługa) — z reguły mają interfejsy oparte na przeglądarce. Chmury niższego poziomu, IaaS (infrastruktura jako usługa), też muszą mieć dostęp do aplikacji. Chmura każdego typu musi oferować API, za pośrednictwem którego da się dodawać zasoby, zarządzać nimi, konfigurować je i zwalniać, gdy przestają być potrzebne.

API jest niezbędne do korzystania z usług dostawcy chmury. W ten sposób sprzedający prezentuje swoją ofertę, którą kupujący może ocenić według swoich potrzeb. Przykładowo API chmury EC2 firmy Amazon jest oparte na protokole SOAP i HTTP, za pomocą którego wysyła się właściwe tej firmie żądania tworzenia, składowania i dodawania obrazów maszyn Amazon (AMI) oraz zarzą-

dzania nimi. Z kolei API oferowanej przez firmę Sun⁶ chmury Kenai całkowicie opiera się na architekturze REST. To za pośrednictwem API tworzy się zasoby (moc obliczeniowa, składowanie, komponenty sieciowe) i zarządza nimi.

Architektura REST i zgodne z nią API REST (ang. *Representational State Transfer*, transfer stanów reprezentacyjnych) to styl architektoniczny dla rozproszonych systemów, takich jak World Wide Web. Opiera się on na protokole HTTP. Sieć WWW jest największą znaną implementacją zgodną z architekturą REST — mówi się nawet, że REST to *specyfikacja post hoc* tych właściwości sieci, które są odpowiedzialne za jej sukces. Architektura REST zakłada istnienie klientów i serwerów. Klienci wysyłają żądania do serwerów, a procesy serwerów przetwarzają te żądania i zwracają odpowiednie odpowiedzi. Żądania i odpowiedzi zawierają *reprezentacje zasobów*. Zasób to dowolny spójny i sensowny byt, który jest adresowalny. Reprezentacja zasobu to z reguły dokument opisujący aktualny lub zamierzony stan zasobu. Aplikacje i interfejsy zgodne z zasadami i ograniczeniami architektury REST określa się mianem *RESTful*.

Jako że aplikacja działająca w chmurze będzie głównym żywicielem Twojej firmy, musisz postarać się o jej ochronę przed nieuprawnionymi osobami. Gdyby działała ona w prywatnym, należącym do Twojej firmy centrum danych, chronionym przez wiele fizycznych i logicznych warstw zabezpieczających, miałbyś pewność, że nie dostanie się do niej nikt niepowołany. W chmurze wszystko jest z definicji dostępne przez internet. Amazon i inni rozwiązują ten problem, wydając klucze publiczne X.509 i żądając ich podania przy każdym wywołaniu API. W ten sposób serwer zyskuje pewność, że byt wywołujący API ma prawo dostępu do infrastruktury.

Aby zrozumieć API chmury — a nie istnieje jeszcze uznany standard — najlepiej przyjrzeć się API chmury Amazon, gdyż firma ta jest liderem i narzuca rozwiązania innym. Tabela 2.2 przedstawia najważniejsze definicje i operacje leżące u podstaw API chmury Amazon.

To oczywiście wierzchołek góry lodowej, jeśli chodzi o terminy i wywołania w API chmury Amazon. Dokumentacja jest dostępna pod adresem: <http://aws.amazon.com/documentation/>. API można podzielić na następujące obszary:

- ◆ adresowanie chmur,
- ◆ bezpieczeństwo sieciowe,
- ◆ regiony i strefy dostępności,
- ◆ usługa Amazon Elastic Block Store (EBS),
- ◆ automatyczne skalowanie, równoważenie obciążenia i Amazon Cloud Watch,
- ◆ publiczne zbiory danych,
- ◆ chmura prywatna Amazon.

⁶ Obecnie Oracle — *przyp. tłum.*

Tabela 2.2. Podstawowe pojęcia i operacje API chmury obliczeniowej Amazon EC2

Pojęcie	Opis
AMI	Obraz maszyny Amazon (AMI, ang. <i>Amazon Machine Image</i>) to zaszyfrowany i podpisany obraz maszyny, który można uruchomić w środowisku serwera wirtualnego. Przykładowo na obraz może składać się następujący zestaw: Linux, Apache, MySQL, PHP oraz aplikacja właściciela AMI. AMI mogą być publiczne (zapewniane przez Amazon), prywatne (zaprojektowane przez twórcę aplikacji), zakupione (od zewnętrznego producenta) lub współdzielone (stworzone za darmo przez pewną społeczność). AMI można składować w serwisie Amazon Simple Storage Service (S3).
Instancja	System działający w wyniku uruchomienia AMI to właśnie <i>instancja</i>. Gdy instancja kończy działanie, jej dane zostają utracone. Instancja pełni taką samą funkcję jak tradycyjny samodzielny komputer.
Standardowy przepływ	1. Dopasuj AMI do swoich potrzeb. 2. Przygotuj paczkę AMI i pozyskaj identyfikator. 3. Uruchom jedną lub więcej instancji AMI. 4. Zarządzaj działającymi instancjami i korzystaj z nich.
Podłączenie	W przeglądarce wpisz adres: <i>http://<nazwa-hosta></i>, w którym <i><nazwa-hosta></i> to publiczna nazwa Twojej instancji. Jeśli chcesz podłączyć się do dopiero co uruchomionego publicznego AMI, który nie został jeszcze zmodyfikowany, skorzystaj z polecenia <i>ec2-get-console-output</i>. W obu przypadkach masz możliwość zalogowania się jako administrator i sprawowania pełnej kontroli nad instancją — odpowiada to podejściu do komputera w centrum danych i zalogowaniu się do niego.

Do różnych aspektów API chmury będziemy jeszcze wracać na dalszych kartach tej książki. Na razie porzucimy ten temat i zajmiemy się kolejną ważną warstwą: składowaniem danych w chmurze.

2.1.4. Przechowywanie trwałych danych w chmurze

Bagażnik samochodu służy do przechowywania bagażu podróznego i całej masy innych rzeczy. Analogicznie, chmura zapewnia miejsce, w którym możesz przechować obrazy maszyn, aplikacje i wykorzystywane przez nie dane. Przechowywanie danych w chmurze zyskuje popularność z tych samych powodów, co przetwarzanie. Na żądanie otrzymujesz wirtualny magazyn danych w sieci, wymagasz także usługi określonej jakości. Inaczej niż w firmowym centrum danych nie musisz kupować dysków — czasami nie musisz nawet zamawiać miejsca przed wysłaniem danych. Z reguły płacisz za ilość przesyłanych danych oraz, co jakiś czas, za zajmowane przez nie miejsce.

Z przechowywania w chmurze można korzystać na różne sposoby. Przykładowo możesz tworzyć kopie bezpieczeństwa danych lokalnych, na przykład tych z prywatnego laptopa. Możesz zsynchronizować z chmurą dysk wirtualny i mieć dostęp do danych z różnych maszyn. Możesz również archiwizować dane zgodnie z pewną polityką, na przykład w celu nadzorczym.

Standard przechowywania danych w chmurze

Stowarzyszenie Storage Networking Industry Association powołało specjalną grupę roboczą, która miała zająć się wypracowaniem standardu składowania danych w chmurze. Nowy interfejs CDMI (ang. *Cloud Data Management Interface*, interfejs zarządzania danymi w chmurze) umożliwia uniwersalne, niezależne od środowiska zarządzanie takimi danymi. W CDMI magazyn danych jest widziany przez interfejsy jako abstrakcyjny kontener. Dodatkowo kontener ułatwia grupowanie danych i tworzenie agregatów.

Przechowywanie w chmurze sprawdzi się w przypadku aplikacji, które dostarczają dane bezpośrednio do klientów w sieci. Aplikacja przekierowuje klienta do odpowiedniej lokalizacji, gdzie pobiera on dane, na przykład pliki audio czy wideo. Wymagania związane z przesyłaniem strumieni danych poprzez sieć będą się skalowały niezależnie i nie zaburzają działania aplikacji.

Wykorzystuje się tu interfejs HTTP. Możesz pobrać plik za pomocą przeglądarki bez pisania dodatkowego kodu — automatycznie uruchomi się odpowiednia aplikacja. Tylko jak umieścić plik w chmurze i jak upewnić się, że magazyn jest odpowiedniego typu i zapewnia wymaganą przez Ciebie jakość? Znowu mamy do czynienia z szeroką ofertą. W tym wypadku nie dziwi, iż wielu dostawców korzysta z interfejsów zgodnych z zasadami architektury REST. Najczęściej mamy do czynienia z interfejsem pozwalającym na tworzenie, odczytywanie, aktualizację i usuwanie poszczególnych obiektów danych za pośrednictwem metod HTTP.

Kolejny raz potraktujemy API chmury Amazon jako przykład do analizy. Tabela 2.3 przedstawia najważniejsze elementy chmury S3.

Tabela 2.3. Podstawowe pojęcia i operacje chmury danych Amazon S3

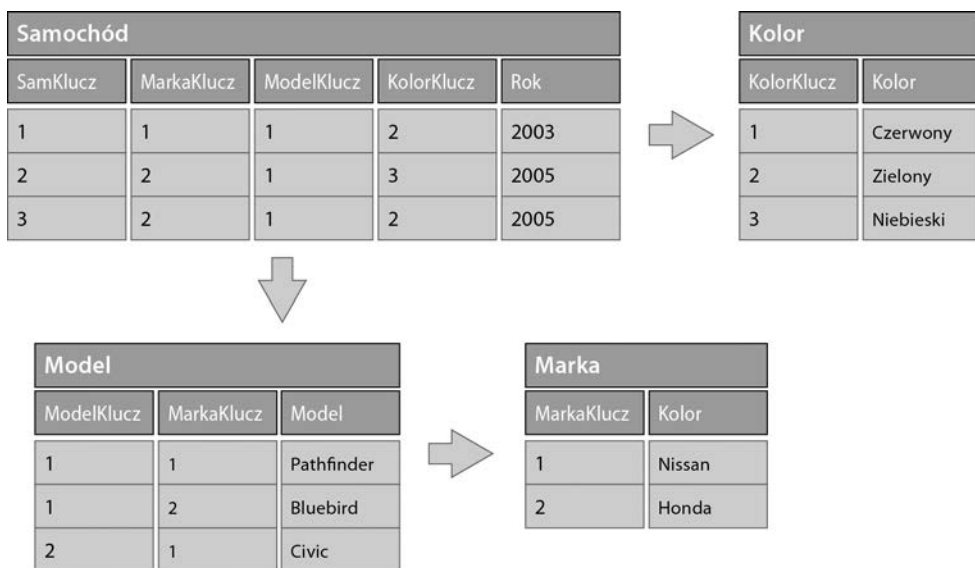
Pojęcie	Opis
Obiekt	Podstawowa jednostka składowana w chmurze danych S3. Obiekt może mieć rozmiar od 1 bajta do 5 gigabajtów. Każdy obiekt ma swoje metadane w postaci par klucz-wartość, opisujących obiekt.
Kubetek	Podstawowy kontener do przechowywania danych w chmurze S3. Obiekty wgrzywa się do kubetków. Kubetek (ang. <i>bucket</i>) to unikalna przestrzeń nazw, umożliwiająca zarządzanie całą zawartością. Nazwy kubetków są globalne. Jeden programista może korzystać jednocześnie z maksymalnie stu kubetków.
Klucz	Klucz to unikalny identyfikator obiektu wewnątrz kubetka. Nazwa kubetka połączona z kluczem jednoznacznie identyfikuje obiekt wewnątrz całej chmury S3.
Wykorzystanie	1. Stwórz kubetek, w którym będziesz przechowywać dane. 2. Załaduj (zapisz) dane (obiekty) do kubetka. 3. Pobierz (odczytaj) dane z kubetka. 4. Skasuj część danych przechowywanych w kubetku. 5. Wypisz obiekty znajdujące się w kubetku.

W wielu przypadkach „gruboziarnista” i nieuporządkowana natura usług przechowywania w chmurach takich jak S3 okazuje się niewystarczająca — potrzebna jest alternatywna metoda przechowywania o ustalonej strukturze. Spójrzmy zatem, jak działają (i nie działają) chmurowe bazy danych.

2.1.5. *Przechowywanie danych aplikacji w chmurowej bazie danych*

Samochodowy system nawigacji na bieżąco informuje Cię o Twoim aktualnym położeniu i odległości od celu podróży. Prowadzi Cię drogą, którą musisz przebyć. Tego typu dane, chociaż niezbędne w czasie podróży, po jej zakończeniu przestają być przydatne. System nawigacyjny jest dla samochodu tym, czym chmurowa baza danych dla aplikacji działającej w chmurze: dane transakcyjne są tworzone i wykorzystywane tylko w czasie działania aplikacji. Gdy mówimy o transakcyjnych danych przechowywanych w bazie, z reguły myślimy o relacyjnych bazach danych.

Czym jest RDBMS i dlaczego często słychać, że nie działa on w chmurze? RDBMS (ang. *Relational Database Management System*) to system zarządzania relacyjną bazą danych, w której przechowuje się dane w postaci tabel. Relacje pomiędzy danymi także reprezentowane są przez tabele. Rysunek 2.4 przedstawia prosty przykład.



Rysunek 2.4. Prosty przykład działania bazy relacyjnej. Cztery tabele odwzorowują zależności (relacje) pomiędzy danymi. Osobna tabela służy do prezentowania marek, inna pokazuje kolory — dzięki temu nie zachodzi konieczność oddzielnego reprezentowania na przykład czerwonego Nissana lub niebieskiego Forda. Jednak by w pełni zrozumieć, czym jest samochód o identyfikatorze SamKlucz 1, musisz połączyć tabele Samochód, Kolor, Model i Marka

RDBMS System zarządzania bazą danych (ang. *Database Management System*) opartą na modelu relacyjnym. *Relacyjne* bazy danych to bazy należące do dość rozległej klasy systemów bazodanowych, w których dane są prezentowane użytkownikowi w postaci relacji (zestaw tabel złożonych z wierszy i kolumn spełnia ten warunek), a on może modyfikować te dane za pomocą operatorów relacyjnych. Wszystkie współczesne relacyjne bazy danych wykorzystują język zapytań SQL, przez co bazy RDBMS często nazywa się *bazami SQL*.

Wyzwaniem dla systemów RDBMS działających w chmurze jest skalowalność. Aplikacje o znanej liczbie użytkowników i obciążeniu nie odnotują kłopotów. Większość dostawców chmur oferuje usługi RDBMS dla takich przypadków. Jednak gdy aplikacje działają w środowisku o dużym obciążeniu (jak w przypadku usług sieciowych *web services*), ich wymagania dotyczące skalowalności mogą zmieniać się (a zwłaszcza rosnać) bardzo szybko. Zmieniające się wymagania są problemem, gdy baza danych jest zainstalowana na pojedynczym serwerze — jeśli ilość danych codziennie się potraja, jak szybko będziesz nadążać z rozbudową sprzętu? Zbyt duża ilość danych może przekroczyć możliwości bazy relacyjnej — stanie się ona wąskim gardłem, uniemożliwiającym skalowanie aplikacji. Istniejącym rozwiązaniom tego problemu przyjrzymy się dokładniej w rozdziale 5.

Powiedzieliśmy już, że jedną z największych zalet chmury jest możliwość szybkiego (i automatycznego, co dopiero zaprezentujemy) dodawania nowych serwerów do aplikacji, gdy wzrasta obciążenie. Jednak skalowanie systemów RDBMS jest o wiele trudniejsze. Konieczne jest albo zreplikowanie danych na wszystkich nowych serwerach, albo podzielenie ich pomiędzy maszyny wirtualne. W obu przypadkach dodanie nowej maszyny wymaga kopiowania lub przenoszenia danych na nowy serwer. Jest to czasochłonny i drogi proces, dlatego bazy danych nie mogą być dynamicznie dostarczane na żądanie.

Duże wyzwanie podczas podziału lub replikacji RDBMS stanowi zachowanie zasady *integralności referencyjnej*. Baza spełnia tę zasadę, jeśli każda wartość atrybutu (kolumny) relacji (tabeli) istnieje jako wartość innego atrybutu w innej (lub tej samej) relacji (tabeli). Mniej formalnie: każde pole tabeli zadeklarowane jako klucz obcy może zawierać tylko wartości z klucza głównego lub potencjalnego tabeli rodzica. W praktyce oznacza to, że usunięcie rekordu zawierającego wartość, do której odwołuje się jakiś klucz obcy z innej tabeli jest złamaniem zasady integralności referencyjnej. Podczas dzielenia lub replikacji bazy danych prawie niemożliwe jest zachowanie integralności we wszystkich bazach. Bardzo przydatna cecha systemu RDBMS, jaką jest konstruowanie relacji z małymi, poindeksowanymi tabelami, do których odwołują się wartości rekordów, okazuje się nierealizowalna, gdy bazy mają się skalować podczas dużych obciążeń, mimo że możliwe jest skalowanie wszystkich innych elementów aplikacji w chmurze.

Ruch NoSQL

Od 1998 roku notuje się powolne, ale nabierające szybkości odchodzenie od baz SQL. Zamiast nich proponowana jest klasa nierelacyjnych magazynów danych, łamiących podstawowe zasady SQL, ale umożliwiającących osiągnięcie o wiele większej skali. Oczywiście jest to pożądana cecha w przypadku *niektórych* aplikacji w chmurze. Nierelacyjne bazy danych często nie wymagają sztywnych schematów danych i raczej unikają operacji łączenia (ang. *join*). Mówi się o nich, że skalują się **horyzontalnie** (poziomo). Używane jest także określenie **ustrukturyzowane przechowywanie** (ang. *structured storage*).

Bazy danych NoSQL, najczęściej w postaci klucz-wartość, faktycznie skalują się lepiej. Z tego powodu coraz chętniej są używane w chmurze. Podstawową jednostką takiej bazy są artykuły (ang. *items*) — wszystkie dane na temat pewnego artykułu są przechowywane wraz z nim. Tabela może zawierać bardzo różne artykuły, na przykład marki, modele i kolory samochodów. Oznacza to, że dane są często duplikowane w różnych artykułach w tabeli (przykładowo więcej niż jeden artykuł może mieć ten sam kolor). Przykład został przedstawiony na rysunku 2.5. W systemie RDBMS za coś takiego twórca zostałby wyklęty. W bazie NoSQL takie rozwiązanie jest akceptowane, gdyż przestrzeń dyskowa jest stosunkowo tania. W takim modelu pojedynczy artykuł zawiera wszystkie związane z nim dane, co poprawia skalowalność poprzez eliminację konieczności łączenia tabel. Chcąc przegrupować atrybuty w relacyjnej bazie danych, musielibyśmy połączyć tabele. To podstawowa zasada skalowalności — jeśli konieczne jest złączenie osobnych tabel, replikacja danych jest trudna i blokuje skalowalność.

Architektura NoSQL

Ograniczenia relacyjnych baz danych powodują, że nie najlepiej sprawdzają się one podczas przetwarzania dużych objętości danych w bezprecedensowej współczesnej skali. Przykłady ogromnej skali to nazwy najpopularniejszych stron: zielone plakietki Digg zajmują 3 TB (terabajty), wyszukiwanie w skrzynce odbiorczej na Facebooku to 50 TB, dane portalu eBay zajmują razem 2 PB (petabajty).

Systemy NoSQL z reguły dają słabe gwarancje spójności, na przykład na poziomie pojedynczych artykułów danych. W większości przypadków możliwe jest wymuszenie pełnych gwarancji ACID (ang. *Atomicity, Consistency, Isolation, Durability* — atomowość, spójność, izolacja, trwałość) przez dodanie pośredniczącej warstwy oprogramowania.

Niektóre systemy NoSQL charakteryzują się architekturą rozproszoną — dane są przechowywane w redundantny (nadmiarowy) sposób na kilku serwerach, często z użyciem rozproszonej tablicy mieszającej. Dzięki temu system można łatwo skalować, dodając nowe serwery, a awaria pojedynczego serwera nie jest problemem.

Niektórzy zwolennicy NoSQL postulują stosowanie prostych interfejsów, takich jak tablice asocjacyjne i pary klucz-wartość. Część systemów, jak bazy oparte na XML-u, wspiera się standardem XQuery.

Jak widać, chmura jest na wczesnym etapie ewolucji i wiele decyzji dopiero zostanie podjęte.

Samochód	
Klucz	Atrybuty
1	Marka: Nissan Model: Pathfinder Kolor: Zielony Rok: 2003
2	Marka: Nissan Model: Pathfinder Kolor: Niebieski Rok: 2005 Biegi: Automatyczne

Rysunek 2.5. Dane z rysunku 2.4, tym razem przedstawione w bazie danych typu klucz-wartość. Ponieważ wszystkie dane dla artykułu (wiersza) są zawarte wewnątrz niego, skalowanie staje się trywialne. Podział (przeniesienie części wierszy w inne miejsce) lub replikacja (skopiowanie całości danych) to proste czynności, niezaburzające integralności referencyjnej

Gdy firma decyduje się na stworzenie publicznej chmury obliczeniowej (jak Amazon) albo na budowanie silnie równoległych, redundantnych, a zarazem ekonomicznych aplikacji sterowanych danymi (Google), bazy relacyjne pozostają bezradne. Obie te firmy potrzebowały sposobu zarządzania danymi, który byłby niemalże nieskończenie skalowalny, godny zaufania i efektywny, także finansowo. W rezultacie obie zaproponowały nierelacyjne systemy bazodanowe oparte na parach klucz-wartość. Amazon nazywa swoją chmurową bazę danych SimpleDB, Google ma BigTable. Obie bazy powstały długo przed tym, zanim firmy uruchomiły chmurę. Stworzono je w celu rozwiązania wewnętrznych problemów. Po uruchomieniu chmury te same struktury stały się częścią oferty przetwarzania w chmurze.

Rozwiązanie BigTable to dość prosty system zarządzania danymi zapewniający szybki dostęp do petabajtów danych, które mogą być redundantnie rozproszone na tysiącach maszyn. Fizycznie BigTable przypomina tablicę z indeksem w formie B-drzewa, którego gałęzie i liście są przechowywane na różnych maszynach. Podobnie jak w B-drzewie wierzchołki w miarę przyrostu danych są dzielone. Dzięki temu, ponieważ węzły są rozproszone, baza skaluje się na wiele maszyn. Elementy danych w BigTable są identyfikowane za pomocą klucza głównego, nazwy kolumny i — opcjonalnie — znacznika czasowego. Wyszukiwanie przy użyciu klucza głównego jest przewidywalne i bardzo szybkie. Baza BigTable jest wykorzystywana przez Google App Engine. W dalszej części tego rozdziału dokładniej opiszemy to środowisko PaaS.

Google pobiera comiesięczną opłatę w wysokości 180 dolarów za każdy wykorzystany terabajt przestrzeni w BigTable. Poniżej przedstawiamy kilka przykładów kodu komunikującego się z tą bazą (w języku Python).

Deklaracja klasy przechowującej dane:

```
class Patient(db.Model):           //pacjent
    firstName = db.UserProperty() //imię
    lastName = db.UserProperty()  //nazwisko
    dateOfBirth = db.DateTimeProperty() //data urodzenia
    sex = db.UserProperty()       //pleć
```

Kod tworzący obiekt i zapisujący go w bazie:

```
patient = Patient()
patient.firstName = "Jan"
patient.lastName = "Kowalski"
patient.dateOfBirth = "2008-01-01"
sex = "M"
patient.put()
```

Zapytanie o obiekty danej klasy:

```
patients = Patient.all()

for patient in patients:
    self.response.out.write('Osoba: %s %s. ' % (patient.firstName,
patient.lastName))
```

Zapytanie o stu najmłodszych mężczyzn:

```
allPatients = Patient.all()
allPatients.filter('sex=', 'M')
allPatients.order('dateOfBirth')
patients = allPatients.fetch(100)
```

Baza SimpleDB, stanowiącą kluczową część środowiska obliczeń w chmurze AWS (Amazon Web Services), działa na bardzo podobnej zasadzie (analogiczną funkcjonalność oferuje także Microsoft przy użyciu SQL Server Data Services [SSDS] w chmurze Azure). SimpleDB również opiera się na parach klucz-wartość. Podstawową jednostką organizacyjną jest **domena** (ang. *domain*). Domeny to zbiory artykułów opisanych za pomocą par klucz-wartość. Tabela 2.4 przedstawia skrócony spis wywołań SimpleDB API wraz z ich opisem funkcjonalnym.

Przekształcenie istniejącej aplikacji tak, by wykorzystywała jedną z chmurowych baz danych, jest albo trudne, albo całkowicie niewarte wysiłku. Lepiej jest w przypadku aplikacji, które już teraz korzystają z frameworków ORM (ang. *Object-Relational Mapping*, mapowanie obiektowo-relacyjne) — w ich wypadku baza w chmurze może z łatwością realizować podstawowe funkcjonalności zarządzania danymi bez negatywnego wpływu na skalowalność (która jest jedną z podstaw istnienia chmury). Jednak, co ilustruje tabela 2.5, nowe typy chmurowych baz danych nie są pozbawione poważnych wad, które należy dokładnie przeanalizować, jeśli rozważa się przejście na chmurę.

Tabela 2.4. Streszczenie API bazy SimpleDB firmy Amazon

Wywołanie API	Opis funkcjonalny
CreateDomain	Tworzy domenę do przechowywania zbioru danych.
DeleteDomain	Usuwa domenę.
ListDomains	Wypisuje wszystkie domeny.
DomainMetadata	Pobiera informacje na temat czasu utworzenia domeny, liczbę nazw artykułów i atrybutów oraz całkowity rozmiar w bajtach.
PutAttribute(s)	Dodaje lub aktualizuje artykuł lub jego atrybuty, pozwala także dodawać pary klucz-wartość do już istniejących artykułów. Artykuły są automatycznie indeksowane po dodaniu do bazy.
BatchPutAttributes	Masowy zapis danych. Wykonuje do dwudziestu pięciu operacji PutAttribute w jednym wywołaniu.
DeleteAttribute(s)	Usuwa artykuł, atrybut lub wartość atrybutu.
GetAttribute(s)	Pobiera artykuł i wszystkie lub tylko niektóre atrybuty i wartości.
Select	Umożliwia odpytanie zbioru danych za pomocą składni <code>Select x from nazwa_domeny where warunki_zapytania</code> . Wartości można porównywać za pomocą operatorów <code>=, !=, <, >, <=, >=, like, not like, between, is null, isn't null</code> oraz <code>every()</code> . Przykład: <code>select * from mydomain where every(keyword) = "Książka"</code> Kolejność wyników można określić za pomocą operatora <code>SORT</code> . Operator <code>Count</code> zlicza wyniki pasujące do zapytania.

Tabela 2.5. Minusy chmurowych baz danych

Zastosowanie bazy	Wyzwania w przypadku bazy w chmurze
Transakcyjność i integralność referencyjna	Odpowiedzialność za integralność transakcji i relacji pomiędzy tabelami spoczywa w dużej mierze na aplikacji korzystającej z chmurowej bazy danych, a nie na samej bazie.
Dostęp do złożonych danych	Bazy w chmurze sprawdzają się w przypadku transakcji opartych na pojedynczych wierszach: pobierz wiersz, zapisz wiersz itd. Jednak nietrywialne aplikacje muszą wykonywać złączenia i inne operacje, które nie są możliwe w bazach danych tego typu.
Analityka biznesowa	Dane są potrzebne nie tylko do działania aplikacji — są także podstawą procesu analizy biznesowej. Nikt z własnej woli nie cofnie się do czasów sprzed powstania baz relacyjnych, gdy dane aplikacji były zakopane głęboko w jej trzewiach i nie była możliwa ich analiza i ocena.

Chmurowe bazy danych mogłyby zastąpić bazy relacyjne w sporym segmencie nowej generacji aplikacji korzystających z chmury. Jednak szefowie firm raczej nie podejść z entuzjazmem do architektury, która utrudnia lub uniemożliwia wykorzystanie danych w procesie analizy biznesowej i podejmowania decyzji — do tego potrzebna jest relacyjna baza danych. Rozwiązaniem byłaby architektura zapewniająca skalowalność i umożliwiającą wykorzystanie wszystkich zalet chmury, a jednocześnie dająca pełen dostęp do ustrukturyzowanych danych. W ciągu kilku kolejnych lat można się spodziewać dużych postępów i wielu innowacji.

Ostatnią z technologicznych podstaw chmury, której chcemy poświęcić więcej uwagi, jest elastyczność. W naszym przykładzie z samochodem elastyczność odpowiada skrzyni biegów.

2.1.6. Elastyczność: skalowanie aplikacji w miarę zwiększania się lub zmniejszania popytu

Automatyczna skrzynia biegów łagodnie dostosowuje prędkość kół pojazdu do prędkości silnika, gdy kierowca przyspiesza lub hamuje. Podobnie elastyczność w chmurze pozwala aplikacji na bezproblemowe radzenie sobie w chwilach różnego zapotrzebowania użytkowników na aplikację. Elastyczność polega na dodawaniu zasobów, gdy zwiększa się obciążenie, i usuwaniu ich, gdy nie są już potrzebne. Wiele dużych firm stanęło w obliczu porażki, kiedy nie były w stanie sprostać zainteresowaniu użytkowników.

Skalowalność w chmurze to zdolność platformy do obsłużenia zwiększonej liczby użytkowników korzystających z aplikacji. **Elastyczność** to zdolność skalowania w górę lub w dół bez przerywania normalnego działania aplikacji. Bez elastyczności przeniesienie aplikacji i działalności firmy w chmurę byłoby nieopłacalne.

Poniższy zestaw wywołań pokazuje, jak skonfigurować aplikację w chmurze EC2, by była automatycznie skalowana (czyli elastyczna) z wykorzystaniem co najmniej dwóch, a najwyżej dwudziestu instancji. Określamy, że aplikacja ma otrzymać dodatkową jedną instancję, gdy średnie wykorzystanie procesora przekroczy 80 procent. Jedna instancja ma zostać odjęta, gdy wynik ten spadnie poniżej 40 procent na dłużej niż dziesięć minut.

Elastyczność a śmierci gwiazd

W czerwcu 2009 roku tego samego dnia odeszły dwie gwiazdy. Najpierw umarła znana z *Aniołków Charliego* aktorka Farrah Fawcett — spowodowało to lekkie poruszenie w mediach. Trochę później, w godzinach popołudniowych, rozszalał się prawdziwy sieciowy sztorm, gdy w społecznej sieci pojawiły się doniesienia o śmierci Michaela Jacksona. Nieoczekiwanie okazało się, że serwis Twitter napotkał ogromne trudności ze skalowaniem, kiedy pojawiły się setki tysięcy wpisów o śmierci artysty. Twitter nie był jednak sam.

Na blogu TechCrunch podano, że: „Należąca do firmy AOL strona TMZ, która jako pierwsza opublikowała informację, co jakiś czas przestawała działać. Prawdopodobnie w wyniku tego popularny blog Pereza Hiltona również odnotował problemy, gdy ludzie rzucili się tam, by potwierdzić wiadomość. Następnie strona LA Times poinformowała, że Jackson nie umarł, tylko jest w śpiączce, więc wszyscy przenieśli się tam, ponownie doprowadzając do przeciążenia strony (w końcu reporterzy LA Timesa potwierdzili informację o śmierci)”.

Znane są liczne przypadki wiadomości, oświadczeń prasowych, nawet materiałów, takich jak niechlubna reklama Victoria's Secret podczas finałów Super Bowl, które spowodowały, że użytkownicy przypuścili szturm na stronę, a ta nie udźwignęła obciążenia. Zbyt duży ruch w połączeniu z niewystarczającymi zasobami oznacza katastrofę. Gdy użytkownik wchodzi na stronę, a ona nie działa, jest duże prawdopodobieństwo, że odpuści sobie dalsze wizyty. Takie problemy wyraźnie odciskają się na zyskach firmy. Wniosek jest taki: skalowanie w miarę wzrostu obciążenia ma kluczowe znaczenie dla powodzenia przedsięwzięcia.

Wywołaj `CreateLoadBalancer` z następującymi parametrami:

```
AvailabilityZones = eu-west-1a  
LoadBalancerName = MyLoadBalancer  
Listeners = lb-port=80,instance-port=8080,protocol=HTTP
```

Wywołaj `CreateLaunchConfiguration` z następującymi parametrami:

```
ImageId = myAMI  
LaunchConfigurationName = MyLaunchConfiguration  
InstanceType = m1.small
```

Wywołaj `CreateAutoScalingGroup` z następującymi parametrami:

```
AutoScalingGroupName = MyAutoScalingGroup  
AvailabilityZones = us-east-1a  
LaunchConfigurationName = MyLaunchConfiguration  
LoadBalancerNames = MyLoadBalancer  
MaxSize = 20  
MinSize = 2
```

Wywołaj `CreateOrUpdateScalingTrigger` z następującymi parametrami:

```
AutoScalingGroupName = MyAutoScalingGroup  
MeasureName = CPUUtilization  
Statistic = Average  
TriggerName = MyTrigger1a  
Namespace = AWS/EC2  
Period = 60  
LowerThreshold = 40  
LowerBreachScaleIncrement = -1  
UpperThreshold = 80  
UpperBreachScaleIncrement = 1  
BreachDuration = 600
```

W rozdziale 1. napisaliśmy, że istnieje więcej niż jeden typ przetwarzania w chmurze. Spróbujmy połączyć przedstawione wówczas informacje na temat różnych typów chmur z tym, co już wiesz o sześciu technologiach, bez których nie mogłaby istnieć chmura. W następnym podrozdziale postaramy się ułatwić Ci zrozumienie tego, czym różnią się typy, co mają do zaoferowania i jak działają. Umożliwi Ci to świadomy wybór najlepszego dla Ciebie rozwiązania.

2.2. Zrozumienie różnych typów chmur

Wiesz już, jakie są technologiczne podstawy chmury obliczeniowej. Rozumiesz, czym jest virtualizacja, elastyczność, jakie są problemy ze składowaniem danych w chmurze. Przyjrzymy się teraz rozwiązaniom zastosowanym w różnych typach chmur. Wróćmy do taksonomii z rozdziału 1.: IaaS, PaaS, DaaS i spróbujmy zaklasyfikować chmury największych przemysłowych dostawców.

2.2.1. Amazon EC2: IaaS

Chmura EC2 należy do klasy IaaS (infrastruktura jako usługa). Niektórzy mówią, że to HaaS (sprzęt jako usługa), jednak firma Amazon wprowadziła tyle dodatkowych usług, że nazywanie tej chmury HaaS to nieporozumienie. Jest to pierwsza i jak dotąd największa chmura tej kategorii. Została udostępniona w 2006 roku, gdy firma chciała jakoś spożytkować nadmiar sprzętu niewykorzystywanego do operacji handlowych. Pod koniec 2008 roku firma miała już ponad pół miliona użytkowników.

Wśród największych istniejących chmur EC2 jest chmurą najogólniejszego typu. Ma ona najmniejsze wsparcie automatyczne dla skalowania i korekty błędów — te wymagania musi obsłużyć sama aplikacja. Odróżnia to EC2 od automatycznie skalujących aplikacji chmur typu PaaS, takich jak Google App Engine, o której będzie mowa w punkcie 2.2.3. W chmurach typu IaaS, takich jak EC2, elastyczność musi zostać starannie zaprogramowana w oparciu o API chmury. Plusem jest to, że programista może wybrać dowolny język programowania i ma całkowitą kontrolę nad aplikacją. Konieczna jest ręczna obsługa, jednak w zamian dostaje się odpowiednik fizycznej maszyny, na której można zainstalować system i wszystkie inne wymagane elementy. Najpopularniejsza konfiguracja EC2 to tak zwany stos LAMP (tabela 2.6).

Tabela 2.6. Komponenty stosu LAMP w chmurze IaaS

L	Linux	System operacyjny
A	Apache	Serwer sieciowy
M	MySQL	Relacyjna baza danych
P	PHP	Obsługa strony internetowej po stronie serwera

Amazon zapewnia bardzo obszerne API dla wszystkich swoich usług, z których kilka przedstawiono w tabeli 2.7. API ma dwie wersje: jedna jest oparta o SOAP, druga o czysty HTTP (GET, POST). Żądanie i konfiguracja maszyny wirtualnej realizowane są w mniej niż dwudziestu wywołaniach.

Chmura EC2 i parawirtualizacja Xen

Chmura EC2 firmy Amazon korzysta ze zmodyfikowanej wersji otwartego monitora (VMM) Xen, stosując tak zwaną parawirtualizację. System operacyjny goszczący na maszynie parawirtualnej nie wykonuje samodzielnie operacji wymagających uprzywilejowanego dostępu, tylko korzysta z pośrednictwa monitora. W ten sposób zmniejsza się obciążenie procesora.

Parawirtualizacja jest wydajniejsza od zwykłej wirtualizacji, w której system operacyjny nie jest modyfikowany. Jednak system ten trzeba dostosować do parawirtualnego środowiska. Monitor powinien „oddać” niektóre wolniej przebiegające zadania systemowi operacyjnemu, by wykonał je bezpośrednio. Z tego powodu w chmurze Amazon nie można uruchomić dowolnego systemu operacyjnego, tylko te, które producent systemu dopasował i w pełni przetestował.

Tabela 2.7. Inne usługi Amazon związane z chmurą (zapewniające część funkcjonalności PaaS)

Zastosowanie bazy	Wyzwania w przypadku bazy w chmurze
Simple Storage Service (S3)	Magazyn danych w chmurze, w którym można przechowywać duże ilości danych. Dostęp jest możliwy z każdego miejsca w sieci za pośrednictwem prostego API. Dobrze zintegrowany z EC2: w S3 są przechowywane AML, przysyłanie danych pomiędzy S3 i EC2 jest zwolnione z opłat.
SimpleDB	Zapewnia podstawowe funkcje bazodanowe: indeksy (specjalne struktury organizacyjne przyspieszające wyszukiwanie) i zapytania. Pozwala uniknąć kosztów związanych z licencjami na relacyjne bazy danych i administracją, a także nie wymaga złożonej konfiguracji. Nie jest to relacyjna baza danych: nie ma schematu i nie można jej odpytywać za pomocą SQL.
CloudFront	Usługa sieciowa dostarczająca treści, konkurująca z Akamai. Pozwala na wygodną i szybką dystrybucję treści między użytkownikami w modelu opłat w miarę zużycia.
Simple Queue Service (SQS)	Zarządzana kolejka komunikatów przesyłanych pomiędzy komputerami. Ułatwia przysyłanie danych między rozproszonymi komponentami aplikacji wykonującymi różne zadania bez ryzyka utraty komunikatu i bez nakładania na komponenty obowiązku ciągłej dostępności.

Ceny EC2 zaczynają się od paru centów za godzinę używania procesora niewielkiej instancji linuksowej, najdroższe instancje to koszt około pół dolara za godzinę⁷. Ceny S3 zaczynają się od 15 centów za 1 GB na miesiąc. Im więcej przestrzeni wykorzysta użytkownik, tym mniejsza cena za 1 GB.

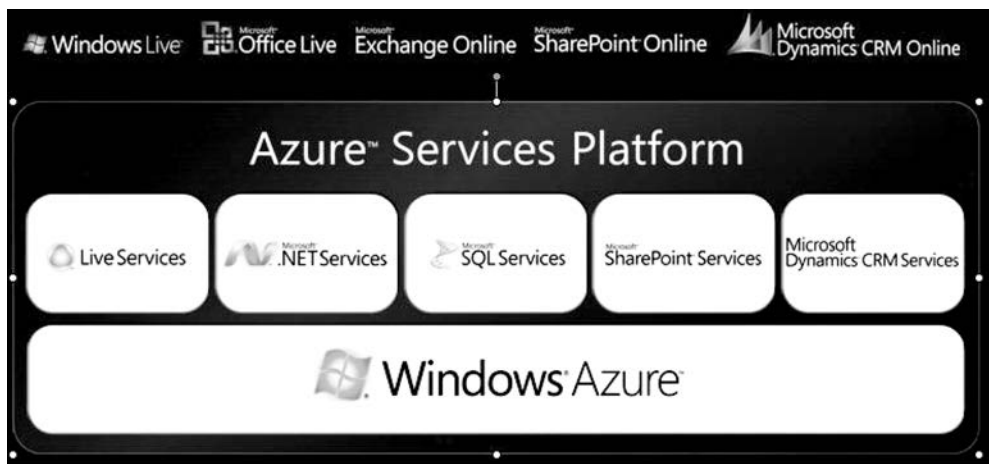
2.2.2. Microsoft Azure: IaaS

Azure, podobnie jak EC2, należy do klasy IaaS, ale oferuje także dodatkowe usługi, działające bardziej na poziomie PaaS. Wiele aplikacji dostarczanych przez Microsoft użytkownikom końcowym przenosi się obecnie do chmury. W rezultacie cała platforma zmierza w kierunku poziomu SaaS (oprogramowanie jako usługa), by zablokować zapędy Google do przejęcia rynku należącego do Microsoft Office za pomocą Dokumentów Google i Google Apps.

Element opisany jako Windows Azure na rysunku 2.6. to system Windows Server 2008 zmodyfikowany pod kątem pracy w chmurze. Jest to kolejny przykład parawirtualizacji, której celem jest efektywne działanie systemu w środowisku wirtualnym pod kontrolą narzędzia monitorującego Microsoft Hypervisor, uruchamianego na czystym sprzęcie w centrach danych Microsoftu.

Wewnętrznie warstwa systemowa — odziedziczona po Windows Server 2008 — składa się z czterech filarów. Są to: składowanie danych (jak system plików), kontroler odpowiadający za zarządzanie modelowaniem, wdrażaniem i realizacją zapotrzebowania na zasoby, maszyna wirtualna i środowisko programistyczne, które umożliwia programistom emulację Windows Azure na ich maszynach

⁷ <http://aws.amazon.com/ec2/pricing/>.



Rysunek 2.6. Architektura i framework Windows Azure. Na najniższym poziomie działa system operacyjny Azure. Jest on uruchamiany w wirtualnym środowisku stworzonym przez monitor Microsoft Hypervisor na czystym sprzęcie. W najwyższej warstwie działają aplikacje przeznaczone dla użytkowników końcowych, które Microsoft przekształca do postaci SaaS. Źródło: Microsoft

oraz wykorzystanie do pisania aplikacji takich narzędzi, jak Visual Studio czy Eclipse. Dzięki takiej architekturze wystarczy wdrożyć Azure na pojedynczej maszynie i zduplikować instancje na pozostałych serwerach z wykorzystaniem technologii wirtualizacyjnej.

Aplikacje działające w chmurze Azure pisze się w środowiskach programistycznych tej firmy, na przykład w Visual Studio, z wykorzystaniem bibliotek .NET i kompiluje się je do Common Language Runtime, czyli środowiska uruchomieniowego niezależnego od języka.

Windows Azure API

API Windows Azure jest oparte na architekturze REST i korzysta z certyfikatów X.509 do autentykacji użytkowników. Tabela 2.8 przedstawia fragment API, który daje pojęcie o tym, jak zarządza się aplikacjami działającymi w chmurze Azure. Są to wywołania związane z podstawowymi operacjami na wdrożonych usługach.

Podobnie jak w chmurze Amazon EC2 nad systemem Azure działa zestaw usług PaaS, które można wykorzystywać jako bloki do budowy aplikacji. Podstawowy zestaw takich usług to:

- ◆ Live Services,
- ◆ SQL Services,
- ◆ .Net Services,
- ◆ SharePoint Services,
- ◆ CRM Services.

Tabela 2.8. Fragment API Windows Azure. API jest zgodne z architekturą REST

Usługa	Opis
Wypisz dostępne usługi GET https://management.core.windows.net/<id-subskrypcji>/services/hostedservices	Wypisanie usług dostępnych w ramach aktualnej subskrypcji.
Wypisz właściwości usługi GET https://management.core.windows.net/<id-subskrypcji>/ ↳services/hostedservices/<nazwa-usługi>	Zwraca właściwości systemowe wskazanej usługi. Obejmują one nazwę i typ usługi oraz nazwę grupy, do której należy, lub lokalizację usługi, jeśli usługa nie jest częścią grupy. Opcjonalnie podawana jest informacja o wdrożeniach usługi.
Utwórz wdrożenie GET https://management.core.windows.net/<id-subskrypcji> ↳/services/hostedservices/<nazwa-usługi>/deploymentslots/ ↳<slot-wdrożenia>/ GET https://management.core.windows.net/<id-subskrypcji>/services/ ↳hostedservices/<nazwa-usługi>/deployments/<nazwa-wdrożenia>/	Wdrożenie specyfikuje się w pokazany niżej sposób. Można je usunąć, podając nazwę slotu wdrożenia (przygotowanie lub produkcja) albo unikalną nazwę wdrożenia.
Przerzuć wdrożenie POST https://management.core.windows.net/<id-subskrypcji>/hostedservices/ ↳<nazwa-usługi>	Rozpoczyna proces wymiany IP pomiędzy slotem przygotowawczym a slotem produkcyjnym usługi. Jeśli usługa obecnie działa w środowisku przygotowawczym, zostanie przeniesiona do środowiska produkcyjnego. Jeśli działa w środowisku produkcyjnym, trafi do przygotowawczego. Jest to operacja asynchroniczna, której status należy sprawdzać przy użyciu wywołania specjalnego wywołania.
Usuń wdrożenie DELETE https://management.core.windows.net/<id-subskrypcji>/ ↳services/hostedservices/<nazwa-usługi>/deploymentslots/ ↳<slot-wdrożenia> DELETE https://management.core.windows.net/<id-subskrypcji>/ ↳services/hostedservices/<nazwa-usługi>/deployments/ ↳<nazwa-wdrożenia>	Usuwa dane wdrożenie. Jest to operacja asynchroniczna.

Usługi te można traktować jako API (nie mają one interfejsów użytkownika) do budowy aplikacji w chmurze.

Ceny Azure są zbliżone do cen w chmurze Amazon. Czas obliczeniowy to koszt 12 centów za godzinę, przechowywanie danych to 15 centów za 1 GB, a przesyłanie danych zaczyna się od 1 centa za 10 KB. Baza danych to koszt 9,99 dolara za 1 GB miesięcznie w przypadku edycji sieciowej, a 99,99 dolara miesięcznie za 10 GB w edycji biznesowej. Wkrótce ma pojawić się warstwowy model typu „wszystko, co zdołasz zjeść” (w pewnych granicach).

2.2.3. **Google App Engine: PaaS**

App Engine to chmura całkowicie zgodna z definicją „platforma jako usługa”. Jest ona przeznaczona do uruchamiania tradycyjnych aplikacji internetowych, które są zmuszane do wyraźnego oddzielenia bezstanowej warstwy obliczeniowej od stanowej warstwy danych. Wirtualizacja i elastyczność, dobrze widoczne w modelu IaaS, tutaj pozostają prawie całkowicie niezauważalne, jednak za kulisami odgrywają naprawdę ważną rolę. Jedną z zalet tego modelu jest automatyczna elastyczność przy zmianie zapotrzebowania.

Języki programowania dozwolone w App Engine to Python i Java. App Engine nie jest chmurą ogólnego przeznaczenia. Najlepiej sprawdza się w przypadku aplikacji internetowych o strukturze żądanie-odpowiedź, w których pomiędzy wywołaniami występują długie okresy bez używania procesora (przykładowo użytkownik zastanawia się nad wyborem opcji). Z tego powodu Google surowo racjonuje czas korzystania z procesora przeznaczony do obsługi każdego z żądań.

Mechanizmy automatycznego skalowania i dużej dostępności oraz własny system składowania danych MegaStore (oparty na BigTable) zakładają zachowanie opisanych powyżej ograniczeń. Jednak jeśli Twoja aplikacja wpisuje się w te warunki, to najprawdopodobniej nie znajdziesz szybszego i tańszego sposobu budowania jej w taki sposób, by automatycznie się skalowała i działała w największej chmurze na planecie.

Środowisko programistyczne App Engine

- ◆ **Piaskownica** (ang. *sandbox*) — Aplikacje działają w bezpiecznym środowisku zapewniającym jedynie ograniczony dostęp do systemu operacyjnego. Te ograniczenia pozwalają App Engine na rozpraszanie żądań wysyłanych do aplikacji na wiele serwerów, uruchamianych i zatrzymywanych stosownie do zmieniającego się zapotrzebowania. Piaskownica izoluje aplikację we własnym, bezpiecznym środowisku niezależnym od sprzętu, systemu operacyjnego i fizycznej lokalizacji na serwerze.
- ◆ **Środowisko uruchomieniowe języka Java** — Możesz pisać aplikacje oparte o Javę 6 z użyciem popularnych narzędzi programistycznych i API do tworzenia aplikacji internetowych. Aplikacja komunikuje się ze środowiskiem uruchomieniowym przy użyciu standardu Java Servlet i może korzystać z popularnych technologii, takich jak Java Server Pages (JSP). Aplikacje odwołują się do większości usług App Engine za pośrednictwem standardowych interfejsów języka Java. Środowisko zawiera platformę i biblioteki Java SE Runtime Environment (JRE) 6. Ograniczenia związane z istnieniem piaskownicy zostały zaimplementowane wewnątrz maszyny wirtualnej Javy (JVM). Aplikacja może korzystać z dowolnej funkcjonalności kodu bajtowego lub biblioteki, o ile nie złamie w ten sposób narzuconych ograniczeń.

- ◆ **Środowisko uruchomieniowe języka Python** — Możesz pisać aplikacje w języku Python 2.5 i uruchamiać je przy użyciu interpretera. App Engine zawiera narzędzia i API do tworzenia aplikacji internetowych w Pythonie, w tym API do modelowania danych, framework aplikacji webowych i narzędzia umożliwiające zarządzanie danymi i dostęp do nich. Środowisko Pythona posiada standardowe biblioteki dostosowane do ograniczeń piaskownicy. Kod aplikacji uruchamianych w tym środowisku musi być napisany w czystym Pythonie. Środowisko Pythona daje dostęp do API przeznaczonego do magazynowania danych, pozyskiwania URK, Google Accounts i usług poczty elektronicznej.
- ◆ **Magazyn danych** — App Engine oferuje rozproszoną usługę składowania danych, obsługującą język zapytań i transakcje. Rozproszony magazyn skaluje się w zależności od zapotrzebowania. Zgodnie z tym, co powiedzieliśmy wcześniej na temat chmurowych baz danych, magazyn danych App Engine nie przypomina tradycyjnej relacyjnej bazy danych. Obiekty danych (**encje**) mają typ oraz zbiór właściwości. Zapytania pobierają encje danego typu, odfiltrowane i posortowane według wartości właściwości. Wartości właściwości mogą należeć do jednego ze wspieranych typów właściwości. Encje w magazynie danych są **pozbawione schematu**. To kod aplikacji odpowiada za określenie struktury encji. Może wykorzystać do tego interfejsy JDO/JPA albo interfejsy magazynu danych Pythona.

Chmura App Engine jest darmowa, jeśli spełnione są następujące warunki: 6,5 godziny czasu procesora i 1 GB przesyłanych w obie strony danych. Po przekroczeniu tych progów za dane wysyłane przez aplikację płaci się 12 centów za 1 GB, 10 centów za dane przychodzące. Czas procesora kosztuje 10 centów za godzinę, a składowanie danych 15 centów za 1 GB na miesiąc. Wysłanie wiadomości e-mail do jednej osoby kosztuje 0,0001 dolara.

2.2.4. **Ruby on Rails w chmurze: PaaS**

Ruby on Rails (RoR) to otwarty framework do tworzenia aplikacji internetowych w języku Ruby. Z założenia powinien być używany zgodnie z metodologią *agile* („zwinną”). Programiści często wybierają go ze względu na łatwość tworzenia niewielkich projektów skoncentrowanych wokół klienta. Podobnie jak w przypadku Google App Engine aplikacje muszą działać na zasadzie żądanie-odpowiedź.

Otwarte oprogramowanie (open source) Oprogramowanie *open source* wyróżnia się tym, że kod źródłowy oraz pewne prawa, które normalnie są chronione, udostępnia się użytkownikom na pewnej specjalnej licencji, umożliwiającej analizę, wprowadzanie zmian, a także poprawianie oprogramowania. Niektórzy nazywają *open source* filozofią, inni pragmatyczną metodologią. Zanim termin ten się przyjął, programiści i producenci używali różnych nazw na opisanie tego pomysłu. To określenie zyskało

popularność wraz z rozwojem internetu i związaną z nim potrzebą przekształcenia istniejących narzędzi i kodu. Otwarte oprogramowanie najczęściej jest tworzone publicznie przez grupę współpracujących programistów.

Język Ruby zaprojektowano tak, by łączył konceptualną elegancję Smalltalka, łatwość nauki i użycia Pythona oraz pragmatyzm Perla. Wiele zespołów odnotowało dziesięciokrotne przyspieszenie wytwarzania aplikacji internetowych z użyciem frameworku RoR. Jednak wielu programistów donosi o problemach ze skalowaniem, które mimo wszystko mają chyba źródło w architekturze i innych wyborach projektowych, a nie wynikają z charakterystyki samego języka.

Wiele małych firm wykorzystało sytuację, oferując stopy RoR działające w chmurze Amazon EC2. Przykłady to Heroku, Aptana i EngineYard.

2.2.5. *Salesforce.com i Force.com: PaaS*

Firma Salesforce.com stworzyła odnoszącą największe sukcesy aplikację SaaS, służącą do zarządzania relacjami z klientami (CRM). Działa ona jako typowa aplikacja w chmurze już od 1999 roku.

Force.com to oferta PaaS tej samej firmy. Programiści mogą tworzyć w języku Apex aplikacje — nakładki na podstawową aplikację Salesforce, działające także w chmurze. Firmy Google i Salesforce zintegrowały swoje chmury App Engine i Force.com w taki sposób, że możliwe jest stworzenie w oparciu o dowolną z nich aplikacji, która ma dostęp do repozytorium danych firmowych.

Force.com utrzymuje katalog aplikacji o nazwie AppExchange, umożliwiający zakup i dodanie do środowiska Salesforce aplikacji napisanych przez zewnętrzne podmioty. Dostępne jest ponad 800 aplikacji od ponad 450 niezależnych producentów. Cena za korzystanie z Force.com to 5 dolarów za logowanie, przy czym użytkownik ma prawo zalogować się maksymalnie pięć razy w miesiącu. Na stronie firmy napisano: „Chmura Force.com jest przeznaczona dla okazjonalnie korzystających z niej użytkowników i dużych aplikacji, jest dostępna jako platforma i nie służy do tworzenia aplikacji CRM”.

Opiszemy teraz jeszcze jedną, dość dziwną klasę chmur. Różni się ona od poprzednich typów nie środowiskiem działania aplikacji, ale strukturą własności. Nazwiemy tę klasę „centrum danych jako usługa” (DaaS, ang. *Datacenter as a Service*). Należą do niej prywatne chmury firm, przeznaczone do użytku wewnętrznego.

2.2.6. *Chmury prywatne: DaaS (centrum danych jako usługa)*

Chmura prywatna (nazywana także **chmurą wewnętrzną**, korporacyjną bądź firmową) to termin określający architekturę, w której usługi są dostępne dla ograniczonej liczby użytkowników znajdujących się za firewallem. W takiej chmurze stosuje się te same rozwiązania co w chmurach Amazon, Microsoft czy Google — wirtualizację, automatyzację, rozproszone przetwarzanie — tyle że jedynie na potrzeby klientów wewnątrz firmy.

Skorowidz

A

- Access Control List, ACL, 113
- AJAX, Asynchronous JavaScript and XML, 232
- Amazon AWS Start-Up Challenge, 94
- Amazon S3, 161
- Amazon Simple Queue Service (SQS), 176
- Amazon Virtual Private Cloud, VPC, 72, 121
- analiza kosztów, 84
- API chmury, 52
- API chmury Amazon, 53
 - AMI, Amazon Machine Image, 54
 - instancja, 54
 - podłączenie, 54
 - standardowy przepływ, 54
- API REST, Representational State Transfer, 53
- aplikacja
 - SaaS, 70
 - typu mashup, 252, 254
 - wewnętrzna, 151
 - wykrywająca oszustwa, 127
- aplikacje
 - działające w czasie rzeczywistym, 91
 - historyczne, 90
 - niestrategiczne, 89
 - o skali chmury, 132
 - z dostępem do poufnych danych, 91

architektura

- Flickr, 149
- maszyny wirtualnej, 51
- mechanizmu cloudbursting, 154
- NoSQL, 58
- ASP, Application Service Providers, 26
- atak ARP, 113
- atak DDoS, 112
- automatyzacja, 28, 30
- automatyzacja wdrożenia, 192

B

- bastion host, 112
- baza danych typu klucz-wartość, 59
- bezpieczeństwo
 - autentykacja wieloczynnikowa, 107
 - dane logowania, 108
 - filtry pakietów, 104
 - firewall, 104
 - Fort Knox, 251
 - klucz dostępowy, 108
 - kontrola dostępu, 106
 - mechanizm eskalacji przywilejów, 114
 - para kluczy, 110
 - przemieszczenie danych, 113
 - samodzielne generowanie par kluczy, 114

bezpieczeństwo
 szyfrowanie danych, 104
 szyfrowanie systemu plików, 113
 weryfikacja tożsamości przez telefon, 106
 wirtualne sieci lokalne (VLAN), 104
 bezpieczeństwo centrów danych, 104
 bezpieczeństwo danych w chmurach, 33, 111, 113
 bezpieczeństwo fizyczne, 104
 bezpieczeństwo informacji, 102
 bezpieczeństwo operacyjne, 216
 bezpieczeństwo sieciowe, 112
 certyfikat X.509, 109, 112
 IPtables, 112
 konfiguracja firewalla, 112
 bezpieczeństwo systemu, 111, 113
 BigTable, 59

C

CAPEX, 28, 79
 centrum danych, 45
 modularyzacja, 48
 skalowanie, 46
 struktura, 45
 centrum danych wewnętrzne, 151
 certyfikat SAS 70, 105
 certyfikat SAS 70 typu II, 211
 certyfikat X.509, 109, 112
 chmura, 37, 42, 115
 chmura App Engine, 68, 74
 środowisko programistyczne, 68
 chmura AWS, Amazon Web Services, 60, 72
 chmura Azure, 60, 65, 73
 API, 66
 architektura i framework, 66
 środowisko programistyczne, 66
 chmura EC2, 52, 64, 73
 chmura Force.com, 70, 75
 chmura Kenai, 53
 chmura obliczeniowa, 9, 26
 chmura prywatna, 42, 70, 102, 115
 bezpieczeństwo, 116, 117
 dostępność, 116
 dostępność zasobów, 117
 efekty skali, 116, 118
 oferta z hostingiem, 121
 open source, 120
 oprogramowanie komercyjne, 121
 potencjalne problemy, 119
 społeczność użytkowników, 116, 118
 sposoby wdrożenia, 119

chmura publiczna, 218
 chmura Ruby on Rails, 74
 chmura w przedsiębiorstwach, 235
 chmurowa baza danych, 56
 BigTable, 59
 minusy, 61
 NoSQL, 58
 RDBMS, 57
 integralność referencyjna, 57
 SimpleDB, 60
 chmury hybrydowe, 42
 chmury najwyższego poziomu, 52
 chmury niższego poziomu, 52
 cloudbursting, 150
 analiza biznesowa, 152
 architektura, 154
 implementacja, 156
 potrzeba standaryzacji, 157
 problem dostępu do danych, 157
 continuous integration, CI, 197
 cztery modele infrastruktury informatycznej, 79

D

dane CRM, 96
 dane transakcyjne, 56
 denormalizacja, 141
 domena, 60
 dostawca usług przetwarzania w chmurze, 34
 dostawca chmury, 28, 210
 dostawca internetu, 35
 dostawca SaaS, 97
 dostawca usługi, 174
 dostępność, 212
 dynamiczne skalowanie, 30

E

elastyczne łącza, 190
 elastyczne magazyny danych, 191
 elastyczność, 28, 30, 216
 elastyczność w chmurze, 62
 wywołania, 62
 enkapsulacja maszyny, 50
 etapy ewolucji, 36
 ewolucja centrów danych, 37
 ewolucja chmury, 239
 ewolucja sposobu wytwarzania aplikacji, 248

F

FaaS, 41, 254
firma
 Amazon, 161
 Bechtel, 127
 Eli Lilly, 97
 FlightCaster, 94
 GoodData, 94
 IDC, 103
 Savvis, 121
 Sprint, 127
 start-up, 92
 SunGuard, 121
 Virgin Atlantic, 99
Flickr, 148
FLOPS, FLoating point Operations Per
 Second, 37
framework jako usługa (FaaS), 41, 254
framework
 MapReduce, 178
 ORM, Object-Relational Mapping, 60
 RoR, 69, 70
 Selenium, 198

G

generowanie certyfikatów, 109
Googleplex, 47

H

hoteling, 79

I

IaaS, 39
implementacja chmury prywatnej, 123
infrastruktura fizyczna, 81
infrastruktura informatyczna, 79
 kolokacja, 79
 model bez outsourcingu, 79
 model chmury, 80
 usługa zarządzana, 80
infrastruktura jako usługa (IaaS), 39
infrastruktura wewnętrzna, 79
interoperacyjność, 217
inżynieria społeczna, 114
ISP, Internet Service Provider, 35

J

jakość operacji w chmurze, 222
jakość świadczonych usług, SLA, 10

K

klastrowanie, 177
klasyfikacja chmur, 41
klucz prywatny, 109
klucz publiczny, 109
klucz publiczny X.509, 53
kolokacja, 79
kompatybilność, 217
konfiguracja EC2, 64
konfiguracja małej aplikacji, 81
konfiguracja wirtualna, 83
konsument usługi, 175
koszt kontroli jakości, 96
koszty, 164
 inwestycyjne (CAPEX), 28, 79
 inwestycyjne dla usługi produkcyjnej, 190
 operacyjne (OPEX), 28, 79
 wdrażania, 81
 wdrożenia w chmurze publicznej, 86

L

LAMP, 64, 186

M

magazyn danych w chmurze, cloud storage, 160
MapReduce, 178
 Hadoop, 183
 krok map, 179
 krok reduce, 180
 zasada działania, 181
Mateos Arthur, 22
mechanizm równoważenia obciążenia, 152
mechanizm składowania danych, 250
metoda GetDatabaseFor(), 140
moc obliczeniowa, 37
model 95. percentyla, 85
model chmury, 80
model ciągłego wdrożenia, 208
model samodzielnego utrzymywania zasobów, 28
modele wdrożeniowe, 29
monitorowanie jakości usług, 226

N

nadmiarowość (redundancja), 177
nadmiarowy sprzęt, 178
naliczanie opłat, 28, 31
narzędzie

- Hypervisor, 65
- MapReduce, 168
- Pylot, 32

niedostępność regionalna, 219
normalizacja, 140
NoSQL, 58

O

obciążenie i skala, 88
OPEX, 28, 79
opóźnienie, 165
oprogramowanie jako usługa (SaaS), 31, 32, 38, 41
ORM, Object-Relational Mapping, 60
outsourcing, 9

P

PaaS, 41, 254
parawirtualizacja, 113
parawirtualizacja Xen, 64
pięć zasad przetwarzania, 27
platforma jako usługa (PaaS), 41, 254
przechowywanie danych w chmurze, 54

- interfejs CDMI, Cloud Data Management Interface, 55

przechowywanie danych w chmurze Amazon

- klucz, 55
- kubelek, 55
- obiekt, 55
- wykorzystanie, 55

przetwarzanie na żądanie, 230
przetwarzanie w chmurze, 27
przyszłość chmury, 236
PUE, Power Usage Effectiveness, 49
pula zasobów, 28

R

RDBMS, Relational Database Management System, 56
relacyjne bazy danych, 57
replikacja danych, 177
REST, Representational State Transfer, 160
RESTful, 53

Rosenberg Jothy, 21
równoważenie obciążeń, 177
Ruby on Rails (RoR), 69, 70
rządowe chmury prywatne, 128

S

S3, 54
SaaS, Software as a Service, 31, 32, 38, 41
Selenium, 198
serwer wirtualny, 29
shardowanie, 137

- denormalizowanie, 141
- integralność referencyjna, 147
- łączenie danych, 147
- ograniczone wsparcie, 147
- partycjonowanie oparte o funkcję skrótu, 144
- partycjonowanie oparte o klucze, 144
- partycjonowanie oparte o usługę przekierowującą, 145
- partycjonowanie oparte o zakres, 143
- partycjonowanie pionowe, 143
- równoważenie danych pomiędzy partycjami, 146
- schemat partycji bazy danych Flickr, 148
- szybkość zapisu, 138
- szybsze zapytania, 138
- trudności i problemy, 145
- wysoka dostępność, 138
- zrównoleglenie danych, 142

sieć VPN, 72
SimpleDB, 59

- API bazy, 61

skala internetowa, 134, 190
skalowalność, 62, 136
skalowanie bazy danych, 141
skalowanie za pomocą shardowania, 142
SLA, Service Level Agreement, 10, 218

- dla Amazon AWS, 219
- dla chmury Rackspace, 221
- dla Microsoft Azure, 220

SOA, Service Oriented Architecture, 38, 168, 172

- luźne sprzężenie, 173
- przetwarzanie w chmurze, 175
- usługi sieciowe, 174

SOA, Software Oriented Architecture, 26
sprzężenie luźne, 170
sprzężenie silne, 170
sprzężenie, coupling, 170
standaryzacja przeglądarek, 232
standaryzacja urządzeń, 233

stopa błędu, 219
stos LAMP, 64, 186
strona firmowa, 95
strona USA.gov, 157
system Eucalyptus, 122
system Salesforce.com, 96
systemy rozproszone, 168
sześć aspektów technologicznych i
infrastrukturalnych, 45
szkielety aplikacji, application frameworks, 249

Ś

środowisko wdrożeniowe, 186
etapu pośredniego (staging), 186
produkcyjne (production), 186
testowe (testing), 186
wytwórcze (development), 187

T

technologia i infrastruktura, 44
technologie przetwarzania w chmurze, 40
testy
ad hoc, 194
funkcjonalne, 194, 198
jednostkowe, 194, 196
obciążeniowe, 32, 194, 201
penetracyjne, 194
ręczne, 194, 206
Selenium, 199
użyteczności, 194
wizualne, 194, 204
The Washington Post, 98
transakcje, 177
tunel VPN, 85
tworzenie
chmury prywatnej, 116, 122
kopii zapasowej systemu plików, 96
wirtualnej chmury prywatnej, 125
zdalnych kopii zapasowych, 96
typy chmur, 63
DaaS, centrum danych jako usługa, 70
HaaS, sprzęt jako usługa, 64
IaaS, infrastruktura jako usługa, 64
PaaS, platforma jako usługa, 68
SaaS, oprogramowanie jako usługa, 65

U

usługa
Amazon Elastic Block Store (EBS), 164
SQS, 176
zarządzana, 80
usługi
biznesowe, 94
sieciowe, web services, 174
SOA, 174
w chmurze, 96
wyższego poziomu, 252
ustrukturyzowane przechowywanie,
structured storage, 58

V

VMM, Virtual Machine Monitor, 50
VPC, Virtual Private Cloud, 72

W

wahania skali, 89
walidacja adresu, 106
warstwa
danych, 68, 250
logiki aplikacji, 250
obliczeniowa, 68
programistyczna, 50
systemowa, 65
wirtualizacji, 51, 113
warstwy chmury, 38
wdrożenia chmury prywatnej, 120
wdrożenie w chmurze, 83
koszt gotowości, 85
koszt składowania danych, 85
opłata początkowa, 84
opłata rezerwacyjna, 84
opłata za aktywny tunel, 85
opłata za transfer danych, 84
opłata za żądania, 85
wdrożenie w modelu kolokacji, 82
wdrożenie w modelu usługi zarządzanej, 83
wdrożenie wewnętrzne, 81
węzeł nadrzędny, 179
wirtualizacja, 28, 29, 38, 50
zastosowanie w chmurze, 51
wirtualizacja platformy, 50
wirtualizacja zasobów obliczeniowych, 29

wirtualna chmura prywatna, VPC,
API, 124, 125
elastyczne skalowanie strony, 126
odzyskiwanie systemu, 126
przenoszenie aplikacji firmowych, 126
wirtualna chmura prywatna Amazon, 72, 121
właściciel chmury, 32
wydajność, 212
wynajmowanie mocy obliczeniowej, 34
wywołania API Amazon S3, 162
wzorce aplikacji w chmurze, 132, 135
wzorzec
elastyczne składowanie danych, 134
przeniesienie, 132
skala internetowa, 133
wybuch obliczeń, burst compute, 133

X

X jako usługa, 39

Z

zalety chmury, 189
oszczędności finansowe, 192
poprawa jakości produkcyjnej, 190
proste usuwanie awarii, 191
przyspieszenie testów, 194
skalowalność, 190
skalowanie danych, 191
wzrost przepustowości, 190
zapotrzebowanie krótkoterminowe, 87
zasoby w chmurze, 196
zasób, 53
zrównoległanie, 195
zysk czasowy, 32
zysk ekonomiczny, 31
zysk wydajnościowy, 33

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**

Jeszcze parę lat temu udostępnienie dużej, popularnej aplikacji wiązało się z ogromnymi wydatkami na infrastrukturę. Konieczne było posiadanie własnej serwerowni, wynajmowanie przestrzeni w centrum danych lub uciekanie się do innych kosztownych rozwiązań. W tej chwili na zwołanie można otrzymać dokładnie tyle mocy obliczeniowej i przestrzeni dyskowej, ile w danej chwili jest potrzebne. Zmartwienia związane z nagłymi i chwilowymi wzrostami obciążenia odeszły na zawsze, a dostępność Twoich aplikacji na poziomie bliskim 100% przez okrągły rok jest w zasięgu ręki. Jak to możliwe?

Książka ta wprowadzi Cię w świat, jakiego nie znałeś. Dowiesz się, czym są chmury obliczeniowe, kiedy z nich korzystać i co w nich umieszczać. Poznasz obecnych na rynku dostawców i ich platformy: Google App Engine, Amazon EC2, Windows Azure oraz Salesforce.com i Force.com. Każda z nich ma swoje mocne i słabe strony oraz sprawdza się najlepiej w innych rozwiązaniach. Po lekturze tej książki bezbłędnie wybierzesz najlepsze z nich — idealnie dopasowane do Twoich potrzeb. W kolejnych rozdziałach autorzy poruszają kwestie związane z bezpieczeństwem w chmurze, omawiają najlepsze wzorce dla działających w niej aplikacji oraz sposoby szacowania kosztów przechowywania danych. Znajdziesz tu też sposoby prowadzenia testów i wdrażania aplikacji w chmurach. Chmury obliczeniowe są przyszłością świata informatyki — głównie w sferze biznesu. Nie pozostawaj w tyle i już dziś sięgnij po kompendium, które otworzy przed Twoją firmą nowe możliwości!

- Zasady definiujące przetwarzanie danych w chmurze
- Historia chmur obliczeniowych i ich klasyfikacja
- Najnowsze technologie, najwięksi dostawcy, najlepsze praktyki
- Kiedy korzystać z chmury, a kiedy jej unikać
- Tworzenie skalowalnych i niezawodnych aplikacji w chmurze
- Przenoszenie do chmury istniejących aplikacji i infrastruktury
- Niezawodność i bezpieczeństwo chmur obliczeniowych
- Testy, wdrożenia i działanie w chmurze

**WYJDŹ NAPRZECIW NOWYM TECHNOLOGIOM
I PRZENIEŚ SWÓJ BIZNES DO CHMURY!**



Nr katalogowy: 7450



Księgarnia internetowa
<http://helion.pl>



Zamówienia telefoniczne:
0 801 339900



0 601 339900

Informatyka w najlepszym wydaniu



Helion

Sprawdź najnowsze promocje:

- <http://helion.pl/promocje>
- Książki najchętniej czytane:
- <http://helion.pl/bestsellery>
- Zamów informacje o nowościach:
- <http://helion.pl/nowosci>

Helion SA
ul. Kościuszki 1c, 44-100 Gliwice
tel.: 32 230 98 63
e-mail: helion@helion.pl
<http://helion.pl>

helion.pl
księgarnia
internetowa

Cena 49,00 zł

ISBN 978-83-246-3416-3



9 788324 634163