

Big Data a NoSQL databáze

EDICE
PROFESIONÁL

 **GRADA**

**Irena Holubová, Jiří Kosek
Karel Minařík, David Novák**

Tuto knihu bychom rádi věnovali:

Kryštofovi.

– Irena

*Rodině, která mne podpořila při práci na knize,
i když dobře věděla, co ji čeká.*

– Jirka

Mým učitelům z Ústavu filosofie a religionistiky FF UK.

– Karel

Sofince, která mi rostla před očima spolu s knihou.

– David

Upozornění pro čtenáře a uživatele této knihy

Všechna práva vyhrazena. Žádná část této tištěné či elektronické knihy nesmí být reprodukována a šířena v papírové, elektronické či jiné podobě bez předchozího písemného souhlasu nakladatele.

**Doc. RNDr. Irena Holubová, Ph.D., Ing. Jiří Kosek, Mgr. Karel Minařík
a RNDr. David Novák, Ph.D.**

Big Data a NoSQL databáze

Knih je monografie

Vydala Grada Publishing, a.s.

U Průhonu 22, 170 00 Praha 7

tel.: +420 234 264 401, fax: +420 234 264 400

www.grada.cz

jako svou 6041. publikaci

Odborní recenzenti:

Doc. Ing. Michal Krátký, Ph.D.

RNDr. Jiří Materna, Ph.D.

Vydání knihy schválila Vědecká redakce nakladatelství Grada Publishing, a.s.

Odpovědný redaktor Petr Somogyi

Sazba Jiří Kosek

Grafické zpracování obrázků Milan Vokál

Návrh a zpracování obálky Vojtěch Kočí

Počet stran 288

První vydání, Praha 2015

Knih byla připravena v XML formátu DocBook

a vysázena pomocí XSL-FO v programu XEP.

Vytiskla Tiskárna v Ráji, s.r.o., Pardubice

© 2015 Grada Publishing, a.s.

Cover Photo © fotobanka allphoto

ISBN 978-80-247-5939-5 (ePub)

ISBN 978-80-247-5938-8 (pdf)

ISBN 978-80-247-5466-6 (print)

Stručný obsah

0 autorech	13
Předmluva	15

I. Pojem Big Data a principy distribuovaného zpracování dat

1. Úvod	19
2. Datové formáty	29
3. Základní principy	47
4. Zpracování dat pomocí MapReduce	63

II. NoSQL databáze

5. Základní principy NoSQL databází	87
6. Databáze typu klíč-hodnota	95
7. Dokumentové databáze	109
8. Sloupcové databáze	127
9. Grafové databáze	143

III. Pokročilé aspekty zpracování Big Data

10. Další aspekty zpracování Big Data	171
11. Dotazování nad NoSQL databázemi	193
12. Transakce v distribuovaném prostředí	205
13. Pokročilé aspekty grafových databází	217
14. Další databáze pro Big Data	243
Závěr	261
Použitá literatura	263
Rejstřík	273

Obsah

0 autorech	13
Předmluva	15

I. Pojem Big Data a principy distribuovaného zpracování dat

1. Úvod	19
1.1 Jak velká jsou Big Data?	20
1.2 Historie a vznik NoSQL databází	22
1.2.1 Konec relačních databází?	25
1.3 O čem bude kniha	27
2. Datové formáty	29
2.1 JSON	30
2.1.1 JSON schéma	33
2.2 XML	35
2.2.1 XML schémata	37
2.3 YAML	39
2.4 Formáty Linked Data	41
2.4.1 RDF/XML	41
2.4.2 JSON-LD	42
2.5 CSV	43
2.6 Optimalizace ukládání a přenosu dat	44
2.6.1 Protocol Buffers	44
2.6.2 Apache Thrift	45
2.6.3 BSON	45
2.6.4 EXI a FastInfoset	45
2.6.5 ASN.1	46
2.7 Jaký formát vybrat	46
3. Základní principy	47
3.1 Škálovatelnost	48
3.2 Konzistence	49
3.2.1 Souběh transakcí	50
3.2.2 CAP teorém	52
3.2.3 Občasná konzistence	54

3.3	Distribuce	56
3.3.1	Rozdělení dat (sharding)	57
3.3.2	Master-slave replikace	58
3.3.3	Peer-to-peer replikace	59
3.3.4	Replikace + sharding	61
4.	Zpracování dat pomocí MapReduce	63
4.1	Funkce Map a Reduce	66
4.1.1	Další příklady	68
4.2	MapReduce framework	68
4.2.1	Další vlastnosti	70
4.3	Hadoop	72
4.3.1	HDFS	72
4.3.2	Hadoop MapReduce	74
4.3.3	Další nadstavby systému Hadoop	76
4.4	Kritika a ústup od MapReduce	82

II. NoSQL databáze

5.	Základní principy NoSQL databází	87
5.1	Společné principy NoSQL databází	88
5.2	Datové modely v NoSQL databázích	89
5.3	Typologie NoSQL databází	93
6.	Databáze typu klíč-hodnota	95
6.1	Principy	96
6.1.1	Základní operace a práce s klíči	96
6.1.2	Jmenné prostory klíčů	97
6.1.3	Druhy úložišť typu klíč-hodnota	98
6.2	Realizace a vlastnosti	99
6.2.1	Distribuce dat	99
6.2.2	Konzistence a dostupnost dat	102
6.2.3	Lokální organizace dat	105
6.3	Práce s daty	105
6.3.1	Sekundární indexy	106
6.3.2	Redis	106
7.	Dokumentové databáze	109
7.1	Datový model „dokument“	109
7.2	Dotazování a manipulace s daty	115
7.2.1	Dotazy	115
7.2.2	Modifikace databáze	116
7.2.3	Agregované dotazy a MapReduce	117
7.2.4	Shrnutí	117

7.3	Vlastnosti dokumentových databází	118
7.3.1	Indexy	118
7.3.2	Replikace dat a dostupnost systému	119
7.3.3	Rozdělení dat	122
7.3.4	ACID pro jednotlivé operace a transakce	123
7.4	Závěr	124
8.	Sloupcové databáze	127
8.1	Datový model	128
8.2	Cassandra: datový model sloupců v praxi	132
8.2.1	Data jako multidimenzionální pole	133
8.2.2	Data jako řídké tabulky	134
8.3	Struktura a vlastnosti systému	136
8.3.1	Distribuce a replikace dat	136
8.3.2	Lokální organizace dat	137
8.4	Dotazy, indexy a transakce	138
8.4.1	Dotazy	139
8.4.2	Indexy	140
8.4.3	Transakce	141
9.	Grafové databáze	143
9.1	Typy grafů a související pojmy	145
9.2	Databáze Neo4j	146
9.2.1	Datový model Neo4j	146
9.3	Přístup k databázi Neo4j	147
9.3.1	Java API	147
9.3.2	Gremlin	149
9.3.3	Cypher	152
9.4	Pokročilé rysy Neo4j	156
9.4.1	Neo4j HA	156
9.4.2	Transakce	157
9.4.3	Indexy	158
9.5	Další grafové databáze	162
9.5.1	Sparksee	163
9.5.2	InfiniteGraph	163
9.5.3	OrientDB	163
9.5.4	Titan	164
9.6	RDF databáze	164
9.7	Srovnání úložišť pro grafy	165
9.8	Závěr	167

III. Pokročilé aspekty zpracování Big Data

10. Další aspekty zpracování Big Data	171
10.1 Analytické zpracování Big Data	172
10.1.1 Schéma dat	172
10.1.2 Tvorba datových skladů	175
10.1.3 Analytické zpracování	176
10.2 Vizualizace Big Data	178
10.2.1 Vizualizace propojených dat	179
10.2.2 Nástroje pro vizualizaci	181
10.3 Invertovaný index jako databáze	182
10.3.1 Apache Lucene a jeho nástavby	183
10.3.2 Zpracování logů	186
10.4 Cloud computing	187
10.4.1 Cloud computing a Big Data	189
11. Dotazování nad NoSQL databázemi	193
11.1 Přímý přístup pomocí programového rozhraní	194
11.2 MapReduce	196
11.3 Specifické dotazovací jazyky	196
11.3.1 Elasticsearch Query DSL	197
11.4 Univerzální dotazovací jazyky	199
11.4.1 Deriváty SQL	199
11.4.2 Rozšíření SQL	199
11.4.3 XQuery	202
11.4.4 JSONiq	203
11.4.5 SPARQL	203
11.5 Závěr	204
12. Transakce v distribuovaném prostředí	205
12.1 Vlastnosti CAP podrobněji	205
12.2 Základní transakční modely	206
12.2.1 Ploché transakce	207
12.2.2 Zřetěžené transakce	207
12.2.3 Hnízděné transakce	207
12.3 Transakce v distribuovaném prostředí	208
12.3.1 2PC protokol	209
12.3.2 3PC protokol	210
12.4 Optimistické a pesimistické off-line zámky	210
12.4.1 Optimistický přístup	211
12.4.2 Pesimistický přístup	211
12.5 Uspořádání časových razítek	213
12.5.1 Pesimistické uspořádání	213
12.5.2 Optimistické uspořádání	214

®	
12.6	MVCC 215
12.7	Závěr 216
13.	Pokročilé aspekty grafových databází 217
13.1	Reprezentace grafů 217
13.1.1	Matice sousednosti 218
13.1.2	Seznam sousedů 218
13.1.3	Matice incidence 219
13.1.4	Laplaceova matice 219
13.2	Lokalita dat 220
13.3	Distribuce grafu 221
13.4	Dotazování nad grafy 223
13.4.1	Typy dotazů 224
13.4.2	Vyhodnocování dotazů a indexace grafových dat 225
13.4.3	Dotazovací jazyky pro grafy 234
13.5	Závěr 242
14.	Další databáze pro Big Data 243
14.1	Hybridní databáze 243
14.1.1	PostgreSQL 244
14.1.2	MarkLogic 249
14.2	Databáze ve webovém prohlížeči 250
14.2.1	Web Storage 250
14.2.2	Indexed Database 252
14.3	NewSQL databáze 254
14.3.1	VoltDB 255
14.4	Array databases 256
14.4.1	SciDB 257
	Závěr 261
	Použitá literatura 263
	Rejstřík 273

O autorech

Doc. RNDr. Irena Holubová, Ph.D. se habilitovala v roce 2014 v oboru informatika na MFF UK v Praze, kde v současné době působí jako docent na Katedře softwarového inženýrství. Současně externě působí na Katedře počítačů FEL ČVUT. Je autorkou více než 80 původních článků, které byly publikovány na mezinárodních konferencích a v impaktovaných časopisech, z oblasti analýz reálných dat a operací, odvozování XML schémat, XML benchmarkingu, generování testovacích dat a efektivní propagace změn v komplexních systémech týkajících se převážně semi-strukturovaných dat. Čtyři z nich získaly významná mezinárodní ocenění. Je spoluautorkou knihy „Technologie XML“ vydané v roce 2008 v nakladatelství Grada, která v tomtéž roce získala Cenu děkana MFF UK za nejlepší monografii. V rámci svého pedagogického působení (spolu)vytvářela na MFF UK a FEL ČVUT předmět „Technologie XML“ a „Pokročilé aspekty a nové trendy v XML“. V nedávné době vytvořila nový předmět „Big Data management a NoSQL databáze“. Na tyto oblasti také zaměřuje vedené bakalářské, diplomové a dizertační práce. Více informací je možné nalézt na stránce <http://www.ksi.mff.cuni.cz/~holubova/>.

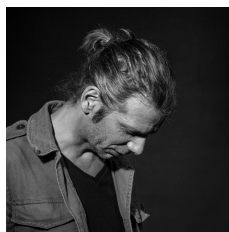


Ing. Jiří Kosek již více než 15 let poskytuje školení a konzultace v oblasti webových a XML technologií. Celá generace tvůrců webu vyrostla na jeho knížkách o HTML a PHP a je i autorem řady článků vydaných jak v Česku, tak v zahraničí. Na půdě Vysoké školy ekonomické v Praze vytvořil a učí předměty zaměřené na webové technologie a XML. Ve svém volném čase spolupořádá a programově zajišťuje konferenci XML Prague.¹ Jirka se podílí na tvorbě a údržbě důležitých standardů v několika organizacích – zejména W3C, OASIS a ISO a přispívá do několika open source projektů. Více se o jeho aktivitách můžete dozvědět na jeho stránkách <http://www.kosek.cz> a <http://xmlguru.cz>.



¹ <http://xmlprague.cz>

Mgr. Karel Minařík je webový designér a vývojář. Vystudoval filosofii na FF UK. Věnuje se programovacímu jazyku Ruby, využití nerelačních databází a vizualizaci dat. V současnosti pracuje pro společnost Elastic.² Žije v Praze. Více informací naleznete na webových stránkách <http://karmi.cz>.



RNDr. David Novák, Ph.D. získal doktorát z informatiky v roce 2008 na Fakultě informatiky MU v Brně, kde nyní pracuje jako vědecký pracovník. Ve svém výzkumu se věnuje zejména technikám pro podobnostní vyhledávání, vyhledávání v multimediálních datech a distribuovaným datovým strukturám. Pracoval na více než deseti národních a evropských výzkumných projektech a je autorem třiceti publikací na mezinárodních odborných fórech. V roce 2014 zavedl na FI MU předmět o NoSQL databázích. V roce 2015 získal Fulbrightovo stipendium na semestrální pobyt na University of Massachusetts Amherst v USA. Více informací naleznete na jeho stránce <http://disa.fi.muni.cz/david-novak/>.



² <http://elastic.co>

Předmluva

Kniha *Big Data a NoSQL databáze* má tři hlavní cíle: vysvětlit pojem Big Data, představit svět NoSQL databází a objasnit jeho souvislost s Big Data. Rozhodli jsme se ji napsat, protože zatím žádná ucelená publikace na dané téma v češtině neexistovala.

Knihu jsme rozdělili do tří částí. V první části vysvětlujeme fenomén Big Data a principy distribuovaného zpracování dat. Ve druhé představujeme několik typů databázových systémů označovaných jako NoSQL databáze. Poslední část knihy přináší přehled dalších nových typů databázových systémů, popisuje pokročilejší aspekty distribuovaného zpracování dat a nabízí přehled dalších souvisejících technologií.

Na první pohled různorodý autorský kolektiv spojuje právě dlouhodobý zájem o oblast zpracování dat. Irena Holubová a David Novák působí v akademickém prostředí. Big Data a související technologie vysvětlují v širších souvislostech. Přinášejí srovnání s tradičními technologiemi a ukazují, jak fenomén Big Data ovlivnil přístup ke zpracování dat a vývoj databázových systémů. Jiří Kosek se ve své praxi i v této knize věnuje zpracování strukturovaných dat, formátům pro jejich ukládání a možnostem dotazování. Karel Minařík sleduje vývoj NoSQL databází od jejich počátku a aktivně se ho účastní. Knihu obohatil zejména o cenné postřehy z praxe. Konkrétně se na obsahu knihy autoři podíleli následovně: Irena Holubová – kapitoly 1, 3, 4, 9, 12, 13 a sekce 10.2, 10.4, 14.3, 14.4; Jiří Kosek – kapitoly 2, 11 a sekce 14.1.2, 14.2; Karel Minařík – sekce 6.3.2, 10.3, 14.1.1 a většina šedých doplňujících rámečků; David Novák – kapitoly 5, 6, 7, 8 a sekce 10.1.

Na tomto místě bychom rádi poděkovali prof. RNDr. Jaroslavu Pokornému, CSc., doc. RNDr. Vlastislavu Dohnalovi, Ph.D. a RNDr. Martinu Svobodovi, Ph.D. za přečtení rukopisu textu a řadu cenných připomínek, které přispěly ke zkvalitnění výsledku. Dále děkujeme RNDr. Davidu Hokszozi, Ph.D., RNDr. Filipu Zavoralovi, Ph.D., RNDr. Jakubu Klímkovi, Ph.D., RNDr. Leu Galambošovi, Ph.D., Ing. Vladimíru Kyjonkovi, RNDr. Jakubu Lokočovi, Ph.D., Mgr. Jindřichu Mynarzovi a Lence Koskové Trískové za kontrolu vybraných kapitol textu a odborné konzultace k nim. V neposlední řadě pak patří velký dík recenzentům, doc. Ing. Michalu Krátkému, Ph.D. z Vysoké školy báňské – Technické univerzity Ostrava a RNDr. Jiřímu Maternovi, Ph.D. ze společnosti Seznam.cz. Děkujeme také sponzorům za významnou finanční pomoc. Autoři byli při přípravě knihy částečně financováni z Programu rozvoje vědních oblastí na Univerzitě Karlově (PRVOUK) č. 204-04/1204 (Irena Holubová).

Do knihy se promítají dlouholeté zkušenosti autorů z výuky kurzů zaměřených na zpracování dat, Big Data a NoSQL databáze. Látka byla Irenou Holubovou zpracována pro jednosemestrální kurz v informatické sekci MFF UK, který vznikl v roce 2012 a dnes je povinně volitelnou součástí magisterské výuky. Obdobný kurz vytvořil David Novák v roce 2014 na FI MU v Brně. V roce 2016 bude kurz v upravené podobě Irenou Holubovou vyučován i na FEL ČVUT. Studentům uvedených i obdobných kurzů bude kniha sloužit jako studijní opora.

Knihy obsahuje velké množství příkladů, které je možné nalézt i na webové stránce <http://www.ksi.mff.cuni.cz/bigdata>. Na stránkách naleznete i opravy případných chyb a další informace. Připomínky a dotazy ke knize můžete zasílat na adresu bigdata@ksi.mff.cuni.cz.

Přejeme vám příjemné čtení.

Autoři

Praha, Brno, Fryšava, Oldřichov v Hájích a Amherst, MA, USA, 25. září 2015

Část I.

**Pojem Big Data a principy
distribuovaného
zpracování dat**

1.

Úvod

*Big Data is like teenage sex:
everyone talks about it,
nobody really knows how to do it,
everyone thinks everyone else is doing it,
so everyone claims they are doing it...*

—Dan Ariely

Slovní spojení Big Data naznačuje, že budeme mluvit o datech, jež jsou velká. Tiše předpokládáme, že data jsou digitální data – žijeme přece ve 21. století. Otázkou ale zůstává, jak velká musejí být data, aby se z nich stala Big Data.¹ Formální, přesnou a všemi přijímanou definici nenajdeme. Společnost Gartner,² v oblasti informačních technologií uznávaná výzkumná a poradenská společnost se sídlem v USA, definuje Big Data jako „data, jejichž velikost (**v**olume), rychlost nárůstu (**v**elocity) a různorodost (**v**ariety) neumožňují zpracování pomocí doposud známých a ověřených technologií v rozumném čase“ [67]. Tyto tři základní vlastnosti bývají označovány jako „3 V“. Postupně k nim přibývají i další „V“, jako např. nejistá věrohodnost (**v**eracity) a vysoká hodnota (**v**alue) pro firmu, která je vlastní

¹ Stejně jako mnoho jiných pojmů z oblasti informačních technologií se ani pojem Big Data nepřekládá. Nebudeme ho tedy překládat ani my.

² <http://www.gartner.com>

[145], nebo limitovaná doba platnosti (**validity**) pro jejich využití a s tím související přechodná doba jejich nutného ukládání (**volatility**) [132].

Co tedy jsou Big Data? Jedná se o revoluci v IT, která znamená konec všech dosud užívaných nástrojů, nebo je to jen bublina, jež brzy splaskne? Přinesou nové aplikace a přístupy očekávané obrovské zisky a převratné vědecké výsledky, nebo sledujeme novodobou zlatou horečku? A odkud se berou Big Data?

Big Data se objevila především s příchodem nových technologií, služeb a jejich kombinací. Příkladem mohou být senzorové sítě nebo vědecké přístroje zkoumající přírodní jevy, sociální sítě nebo mobilní technologie a související aplikace. Tyto typy technologií a aplikací, s nemalou pomocí svých uživatelů, generují každou vteřinu obrovská množství dat, která potřebujeme efektivně uložit a účelně zpracovat. Myšlenku dobře vystihuje např. společnost IBM:³ „V závislosti na odvětví a organizaci zahrnují Big Data informace z interních a externích zdrojů, jako jsou transakce, sociální média, podniková data, senzory a mobilní zařízení. Firmy mohou tato data využívat, aby lépe přizpůsobily své výrobky a služby potřebám zákazníka, dále optimalizovaly provoz a infrastrukturu a/nebo našly zcela nové zdroje příjmů“ [45].

Podívejme se blíže na jednotlivá „V“ z uvedené charakteristiky společnosti Gartner. Hovoříme-li o *velikosti* dat, máme na mysli datové kolekce takových objemů, které nedokážeme uložit na jeden databázový server, ale potřebujeme jich několik desítek či stovek. Množství dat v Big Data kolekci typicky *narůstá* v čase velmi rychle, často exponenciálně. Navíc velké objemy přibývajících a obměňujících se dat potřebujeme zpracovávat velmi *rychle*. Dalším rozměrem problému je *různorodost* dat. Na rozdíl od klasických strukturovaných dat např. v relačních databázích (*relational database management systems*, RDBMS), v oblasti Big Data hovoříme o datech semi-strukturovaných (např. obecně textové dokumenty nebo data ve formátech XML nebo JSON) nebo zcela nestrukturovaných (např. multimediální data). Protože kolekce Big Data často vznikají z různých veřejně dostupných zdrojů, mohou se potýkat s nižší *věrohodností*. Například data získaná vytěžováním textů ze sociálních sítí nemůžeme považovat za tak konzistentní, úplná a přesná jako data z databázových záznamů v uzavřených firemních systémech.

1.1 Jak velká jsou Big Data?

Abychom mohli objektivně měřit v přesných číslech, je dobré získat měřítko pro velikost dat. Uvědomme si nejprve, že v současné době má pevný disk v běžném osobním počítači kapacitu řádově až několik terabajtů (TB), tj. 10^{12} bajtů (B).⁴ Zajímavé odhady týkající se Big Data pak uvádí společnost IBM [44]. Např. odhaduje, že v roce 2016 bude existovat 18,9 miliard internetových připojení, tj. 2,5 připojení na každou osobu na Zemi. V roce 2020 pak podle jejich odhadů bude 6 miliard lidí na zemi vlastnit mobilní telefon, denně vznikne 2,8 kvintilionů

³ <http://www.ibm.com>

⁴ Dříve předpony kilo-, mega-, giga- atd. odpovídaly násobkům 1 024, tedy 2^{10} . Dnes pro ně ale používáme speciální zkratky KiB, MiB, GiB atd.

$(2,8 \times 10^{18})$ bajtů dat⁵ a celkově bude na discích uloženo 40 zettabajtů dat (kde 1 zettabajt je 10^{21} bajtů, tedy ekvivalent miliardy pevných disků o velikosti 1 TB).

Častým zdrojem Big Data jsou sociální sítě. Dle zjištění společnosti IBM [44] např. uživatelé serveru YouTube⁶ každý měsíc shlédnou více než 4 miliardy hodin videa. Sít Twitter⁷ má v současné době 200 milionů uživatelů, kteří každý měsíc vyprodukují 400 milionů krátkých komentářů (*tweets*). A podle společnosti Zephoria [46] má Facebook⁸ (v době psaní této knihy, tedy jaro 2015) 900 milionů uživatelů aktivních každý den, kteří sdílí nějaký obsah více než 4,75 miliardy krát denně a na Facebook nahrají přes 300 milionů fotografií denně. V roce 2008 pro tyto účely využíval Facebook síť 10 000 serverů, v roce 2009 už to bylo 30 000 a odhady z roku 2012 mluví o více než 180 000 serverech.

Dalším zajímavým příkladem je obchodování na burze. Zde dnes namísto lidí obchodují převážně počítače využívající komplexní matematické algoritmy. Rychlost takových obchodů omezuje jen výkonnost hardwaru, tedy fyzikální vlastnosti použitých materiálů. Dle již uvedeného zdroje od IBM burza v New Yorku v současné době zachytí přibližně 1 TB dat během jednoho dne obchodování.

Obchodování na burze

Obrovský objem strojově generovaných dat, jako např. dat o obchodování na burze, přitom představuje problém nejen z prostého pohledu jejich ukládání a zpracování, ale především z pohledu jejich interpretace a analýzy. Dobrým příkladem může být projekt studia Stamen⁹ [136] [142], který vizualizuje data z jediného dne obchodování na burze NASDAQ¹⁰ (National Association of Securities Dealers Automated Quotations), po burze v New Yorku druhé největší svého druhu na světě.

Posledním pojmem, který zmíníme v souvislosti s příklady velkých dat, jsou *proudy dat* (*streams*). Jedná se o specifický typ zdrojů, z nichž data do aplikace přicházejí ve formě nikdy nekončícího rychlého toku dat, který typicky není možné opakovaně procházet. Data můžeme dočasně uložit, nicméně vzhledem k jejich neustálému nárůstu není možné uložit a zpracovávat celou množinu, ale pouze vybrané časové okno. S proudy dat se setkáváme např. v prostředí sítí – ať už senzorových, počítačových nebo mobilních.

⁵ Při překladu názvů velkých čísel si musíme dát pozor na to, kterou škálu měl autor na mysli. Krátká škála, používaná ve většině anglicky mluvících zemích, mění předpony názvů čísel s násobkem tisíc (tedy např. bilion je tisícinásobek milionu). Dlouhá škála, používaná v kontinentální Evropě, používá pro každou předponu vždy dvě přípony (milion, miliarda) a předpony mění s násobkem milion. Tedy např. kvintilion v krátké škále odpovídá 10^{18} , zatímco v dlouhé škále 10^{30} .

⁶ <https://www.youtube.com>

⁷ <https://www.twitter.com>

⁸ <https://www.facebook.com>

⁹ <http://stamen.com>

¹⁰ <http://www.nasdaq.com>

Zdroje Big Data

Z uvedeného je zřejmé, že klíčovou charakteristikou Big Data není oproti běžné představě ani tak celkový objem dat, ale právě trvalost a rychlost jejich nárůstu (velocity). Lidmi vytvářená data, jako jsou uváděné příklady sociálních sítí, přitom mají nějakou představitelnou, byť velmi vzdálenou mez; na celém světě je zkrátka pouze určitý počet lidí vybavených mobilním telefonem s fotoaparátem, připojením k internetu a účtem na Facebooku. Ale např. již datům získávaným z obchodování na burze tento praktický limit do značné míry chybí. Dalším z rapidně rostoucích zdrojů dat jsou dnes informace o provozu strojů samotných, tedy např. *logy* webových a e-mailových serverů, síťových směrovačů, webových API (Application Programming Interface) a aplikací, atd.

V současnosti a blízké budoucnosti budou přitom strojově generovaná data představovat mnohem větší objem celkově ukládaných a zpracovávaných dat. „Každých 30 minut provozu vygeneruje jediný motor letounu Boeing 10 TB dat. Jeden zaoceánský let čtyřmotorového letounu tak vygeneruje asi 640 TB dat. Vynásobíme-li to asi 25 tisíci lety, které se uskuteční každý den, uděláme si představu o obrovském množství dat, které existuje“. [143]

Dat získávaných ze senzorů přitom bude jen přibývat, ať už se bude jednat o měřiče zatížení vozovky zabudované v moderních dálnicích, nositelná zařízení v oblasti medicíny, RFID (Radio Frequency Identification) čipy v poštovních zásilkách nebo domácí spotřebiče propojené v tzv. *Internet of Things* (IoT). Ty všechny odesílají data prakticky neustále a často bez vědomí svého nositele či majitele.

S problémem rychlosti nárůstu je těsně spjat též problém *granularity*. Většina systémů dnes požaduje ukládat „surová data“ tak, jak přicházejí, protože nejsou dopředu známy všechny požadavky na jejich využití a stále častěji tak není možné „šetřit“ pomocí ukládání agregovaných dat.

1.2 Historie a vznik NoSQL databází

Donedávna nejpoužívanějším způsobem pro efektivní ukládání a správu dat byly tradiční databázové systémy, ať už relační, objektové, objektově relační, XML nebo další. Nejpoužívanějšími z nich jsou pak bezpochyby databáze relační, které vycházejí z jednoduchého, zato robustního matematického pojmu relace, který si většinou představujeme jako tabulku s řádky a sloupci. Všechny systémy tohoto typu mají mnoho společných rysů, jako je persistentní ukládání dat, řízení souběžného přístupu více uživatelů, dotazovací jazyky pro přístup k datům apod. Typicky jsou založeny na architektuře klient/server, která předpokládá, že máme veškerá data uložena na jednom výkonném uzlu (serveru), k němuž přistupují klientské programy. Pokud tedy potřebujeme uložit větší množství dat, zvýšíme diskovou kapacitu serveru. Pro zálohování nám obvykle stačí 1–2 další záložní servery (tzv. *zrcadla*) serveru primárního, na nichž jsou veškerá data duplikována. Jak je ale z vlastností Big Data zřejmé, pro jejich zpracování standardní databázové sys-

témy nestačí a potřebujeme přístupy nové. Ty jsou dnes souhrnně označovány jako *NoSQL databáze*, i když, jak si ukážeme později, nejedná se o tak jednoznačný termín jako v případě relačních nebo objektových databází.

Pojem „NoSQL“ pravděpodobně jako první použil Carlo Strozzi v článku [152] z roku 1998. Označil tak open source relační databázi, která nepodporovala klasický přístup k datům přes SQL [21]. V roce 2009 stejný pojem použil Eric Evans pro označení konference pro zastánce různých ne-relačních databází [82]. V současné době se jako NoSQL označují nové databázové systémy, které začaly vznikat začátkem nového tisíciletí pro podporu efektivního zpracování Big Data. Podle [81] se jedná o „novou generaci databázových systémů pro správu velkého množství dat, které jsou převážně ne-relační, distribuované, škálovatelné a podporují replikaci. Často jsou to databáze bez datového schématu, s jednoduchým rozhraním pro práci s daty a open source přístupem“.

V minulých desetiletích jsme mohli pozorovat vznik několika nových databázových konceptů, které s odstupem času můžeme vnímat spíše jako módní trendy, jež se příliš neprosadily. Např. objektové databáze začaly vznikat jako přirozenější úložiště pro struktury objektového programování. V dnešní době sice objektové programovací paradigma stále dominuje, ale používanost objektových databází je oproti relačním podstatně nižší.

Naopak postupný vznik NoSQL databází přichází jako odpověď na praktické potřeby společností internetové éry. NoSQL technologie jsou nyní buď vyvíjeny interně přímo pro potřeby společností jako je Google,¹¹ Amazon¹² či Facebook, nebo jsou rozvíjeny v režii open source komunit a relativně malých nově vzniklých firem. Vývoj NoSQL technologií probíhá většinou bez zásadnější účasti akademických komunit, ale tyto databáze často staví na solidních akademických výsledcích (např. v oblasti udržování konzistence dat v distribuovaných systémech). Často se jedná o teoretické výsledky publikované před desítkami let, pro jejichž širší využití dozrála doba až nyní, a to právě díky technologiím popisovaným v této knize.

Mezi hlavní uváděné výhody NoSQL databází patří především:

1. Flexibilní škálovatelnost: Zatímco klasické databázové systémy využívají *vertikální škálovatelnost*, tedy využívají výkonnější hardware, NoSQL databáze *škálují horizontálně*. Zpracování každé úlohy tyto systémy mohou distribuovat v rámci množiny uzlů (zvané *cluster*¹³), kterou je možné dle potřeby rozšiřovat přidáváním dalších uzlů.
2. Flexibilní datový model: NoSQL databáze nevyžadují určení žádného databázového schématu anebo je schéma dat volné. Dodržování tohoto schématu je často ponecháno na aplikaci a jeho změna neznamena pro databázi významnou zátěž.

¹¹ <http://www.google.com/about/company/>

¹² <http://www.amazon.com>

¹³ Pojem cluster může mít několik významů. V oblasti distribuovaného zpracování dat rozlišujeme především dva – konkrétní hardwarové řešení, kdy jsou jednotlivé výpočetní jednotky fyzicky umístěny v jednom nebo více stojanech a propojeny pomocí vysokorychlostních kabelů, popř. sdílející např. diskové pole, nebo obecně jakoukoli množinu síťově propojených počítačů. Pokud nebude řečeno jinak, budeme v textu nadále používat slovo cluster v tom druhém, obecnějším významu.

3. Orientace na efektivní čtení: NoSQL databáze vybízejí své uživatele k vytváření takových datových struktur, které budou nejlépe odpovídat často pokládaným dotazům. Každý takový dotaz je pak zpracován velmi rychle a systém může zvládat obrovské množství dotazů za jednotku času.
4. Ekonomický aspekt: Vzhledem k horizontální škálovatelnosti NoSQL databází neklademe na uzly, na něž ukládáme data, žádné speciální požadavky. Může se jednat o běžně dostupné, relativně levné počítače (*commodity hardware*).

Dynamické datové schéma

Přestože NoSQL databáze jsou v běžném pojetí téměř synonymní s „bezschémátovými“ (*schema-less*) databázemi, v praxi samozřejmě data vždy nějaké schéma mají, jen nemusí být přesně definováno. Spíše se jedná o meta informace k datům, ne o schéma ve svém klasickém pojetí. Např. článek má autora a datum publikace. Záznam v logu webového serveru obsahuje URL a návratový status. Na rozdíl od relačních databází NoSQL databáze často přenášejí zodpovědnost za správu schématu na aplikační logiku: přidáme-li do datového modelu článku nový atribut, musíme jej buď doplnit i k existujícím záznamům, což může být od určitého objemu dat ekonomicky nebo technicky vyloučené. Druhou možností je uzpůsobit aplikaci tak, aby se vyrovnala s *heterogenním* datovým modelem.

V praxi existuje ještě další řešení, které zcela opouští tradiční model „data mají jeden autoritativní datový model a jedno jediné úložiště“. Zdrojová data ukládáme v surovém formátu a dávkovým zpracováním je pravidelně obohacujeme, transformujeme nebo agregujeme do podoby vyžadované aplikací či aplikacemi. Zdrojová data jsou neměnná (*write-once* či *immutable*) a aplikace pracuje s odvozenými datovými modely, kterých může být několik. (Konkrétní příklad s počítáním přístupů na webové stránky viz první a druhá kapitola knihy od Nathana Marze [126].)

Hlavními výzvami pro NoSQL databáze jsou pak:

1. Zralost: I když jsou existující NoSQL databáze velmi efektivní, ještě nedosahují léty prověřené robustnosti, na kterou jsme zvyklí u relačních databází.
2. Uživatelská podpora: Výhodou, ale současně i nevýhodou NoSQL databází je jejich vznik v rámci tzv. *start-up* projektů a fakt, že jsou typicky open source. Díky tomu je sice jejich vývoj překotný, ale na druhou stranu nemáme vždy k dispozici poskytovatele s dobrým jménem a silným zázemím (opět srovnáme-li situaci se světem relačních databází).
3. Administrace: NoSQL databáze jsou často masivně distribuovanými systémy, což a priori znamená složitější instalaci a také následnou údržbu.
4. Standardizace přístupu k datům: Zatímco relační databáze nabízejí standardizovaný přístup pomocí jazyka SQL, v NoSQL světě si prakticky každá databáze vyvinula vlastní dotazovací jazyk a programátorské rozhraní (API). Zadávání složitějších dotazů často vyžaduje netriviální programátorské znalosti.

5. Experti: Na trhu je stále poměrně velký nedostatek odborníků v oblasti NoSQL databází.

Zdá se, že pro čím dál větší počet aplikací a IT společností vychází bilance uvedených výhod a nevýhod pozitivně, protože technologie popisované v této knize jsou v dnešní době již široce používány. Tím nemáme na mysli pouze internetové giganty, kteří stáli u zrodu těchto databází, ale také stovky společností uváděných v sekcích „Our Customers“ webových prezentací různých NoSQL databází.¹⁴ Záplava dat nutí čím dál více běžných firem zamýšlet se nad možnostmi jejich databázové vrstvy, často se pak poohlíží po technologiích popisovaných v této knize.

1.2.1 Konec relačních databází?

Jak je někdy chybně uváděno, pojem NoSQL neznamená „No to SQL“, neboli „Ne jazyku SQL“. Tyto databáze nemají nahradit relační (SQL) ani jiný typ databází. Jejich cílem je nabídnout řešení pro nové typy aplikací, pro něž stávající databázové systémy nebyly navrženy, a tudíž jim nevyhovují. NoSQL ale neznamená ani „Not only SQL“, čili „Nejen jazyk SQL“ nebo „Něco víc než SQL“. Například systémy Oracle DB¹⁵ nebo PostgreSQL¹⁶ podporují širokou škálu jiných principů než jen jazyk SQL, ale do skupiny NoSQL databází je neřadíme. Na druhou stranu existuje klasifikace NoSQL databází [81] na *pravé* (core), kterým bude primárně věnována tato kniha, a *nepravé* (non-core), kam můžeme zahrnout i všechny ostatní systémy, které nejsou založeny na relačním modelu, tedy např. objektové nebo XML databáze.

Příchod NoSQL databází neznamená konec tradičních relačních databází. Tyto systémy mají velké množství cílových aplikací, pro které je standardizovaný relační model vhodný. Desítky let vývoje přinesly velká množství efektivních technologií využívaných v relačních databázích – například techniky fyzické organizace dat na disku, vyhledávací indexy, optimalizace zpracování dotazů, implementace vyhledávacích operací a další. Nezanedbatelná hodnota existujících relačních databází je i v jejich prověřenosti časem, stabilitě, spolehlivosti a robustnosti, existenci standardizovaných rozhraní a dotazovacích jazyků pro komunikaci s nimi a v neposlední řadě také v záruce zachování silné konzistence dat. Těmito vlastnostmi se většina současných NoSQL úložišť pochlubit nemůže. Je ovšem pravděpodobné, že postupem času budou jednotlivé NoSQL databáze „dozrávat“ a tato situace se bude měnit.

V tabulce 1.1 na následující straně vidíme srovnání požadavků na zpracovávaná data, resp. předpokladů o těchto datech v prostředí databází relačních a NoSQL. Pochopitelně se v obou případech najdou výjimky a speciální případy, tabulku je tedy třeba chápat jako krátký shrnující přehled obecných tezí.

¹⁴ Viz například <http://basbo.com/about/customers/>, <http://planetcassandra.org/companies/>, <https://www.mongodb.com/who-uses-mongodb> nebo <http://neo4j.com/customers/>.

¹⁵ <https://www.oracle.com/database/>

¹⁶ <http://www.postgresql.org>

Tabulka 1.1: Relační vs. NoSQL databáze – předpoklady o datech a aplikaci

Relační databáze	NoSQL databáze
Integrita dat je zásadní.	Stačí, pokud je většina dat většinu času v pořádku.
Datový formát je konzistentní a dobře definovaný.	Datový formát nemusí být známý nebo konzistentní.
Předpokládáme dlouhodobé uložení dat.	Vzhledem k velkému množství dat často ukládáme pouze určité „časové okno“ (např. poslední měsíc, poslední rok).
Aktualizace dat jsou časté.	„Write-once/read-many“, tedy vložená data už typicky nejsou dále modifikována (nebo alespoň ne příliš často). Obvykle data neustále přibývají, aniž by byla modifikována. Již nepotřebné záznamy jsou pak smazány.
Předvídatelný (lineární) nárůst velikosti dat.	Nepředvídatelný (exponenciální) nárůst velikosti dat.
Nástroje pro dotazování dat umožňují přístup i ne-programátorům.	Typicky pouze programátoři píšící (implementují) zpracování dat.
Probíhají pravidelné zálohy dat.	Pro řešení výpadků je využívána replikace dat.
Přístup k datům zajišťuje jediný server.	Data jsou umístěna na více serverech, přistupujeme tedy ke clusteru uzlů.

Na relační, resp. NoSQL databáze bychom tedy měli nahlížet jako na dvě různé možnosti pro ukládání dat. Každá má své výhody a nevýhody a je určena pro jiný typ aplikací. Z obecnějšího hlediska používáme pojem *polyglotní persistence*, který znamená, že nastala doba, kdy máme mnohem více možností využívat různé typy úložišť pro různé typy úloh a dat – a to například i pro různé části jednoho systému [144].

Svět před relačními databázemi

Přestože se může zdát, že relační databáze tu byly „odjakživa“, a NoSQL systémy je přicházejí „nahradit“, relační databáze byly ve své době revoluční změnou v pohledu na data a vyzývatelem *tehdy* tradičních databází, jako byl např. IBM IMS, který využíval pro datový model hierarchickou strukturu, nebo síťové databázové systémy, jako byl např. systém IDMS.

1.3 O čem bude kniha

Naším cílem je představit čtenáři problematiku zpracování Big Data zejména pomocí NoSQL databází. Podíváme se na obecné základní principy, na různé druhy NoSQL databází včetně bližšího pohledu na konkrétní populární systémy i na teoretičtější pozadí využívaných metod. Kapitoly knihy jsou rozděleny do tří částí.

První část knihy je věnována oblasti Big Data a distribuovaného zpracování dat obecně. Konkrétně v kapitole 2 nejprve ukážeme, jaké datové formáty typicky využíváme v oblasti Big Data a jaké jsou jejich výhody a nevýhody. V kapitole 3 si vysvětlíme základní pojmy a techniky využívané při práci s Big Data, jako je distribuce, replikace, škálování, konzistence apod. A v kapitole 4 si blíže představíme také programový model MapReduce, který je často používaným principem při distribuovaném zpracování dat.

Druhá část knihy je věnována vlastním NoSQL databázím. Nejprve v kapitole 5 podrobněji rozebereme společné prvky databázových systémů tohoto typu, principy datového modelování v NoSQL světě a uvedeme nyní již poměrně standardizovanou klasifikaci NoSQL databází. V následujících čtyřech kapitolách se podrobněji seznámíme s hlavními třídami NoSQL databází, tedy systémy typu klíč-hodnota (kapitola 6), dokumentovými databázemi (kapitola 7), sloupcovými databázemi (kapitola 8) a grafovými databázemi (kapitola 9). Budeme se zabývat jejich hlavním zaměřením, vlastnostmi a technikami, pomocí nichž je zajištěna funkcionality těchto systémů. V každé z těchto kapitol typicky zvolíme jeden konkrétní systém využívaný v příkladech, kterými je text bohatě prokládán. Samozřejmě se budeme zamýšlet také nad tím, pro které typy úloh je daný typ databází vhodný a pro které méně.

V poslední části knihy se budeme věnovat pokročilejším problémům a přístupům v oblasti zpracování Big Data pomocí NoSQL databází. Abychom čtenáři netvrdili, že Big Data znamenají pouze databáze, úvodem si v kapitole 10 krátce vysvětlíme i další související pojmy, jako je analytické zpracování Big Data, jejich vizualizace, cloud computing nebo fulltextové vyhledávání. Dále se podíváme na oblasti jako je např. efektivní dotazování (kapitola 11) nebo transakce v distribuovaném prostředí (kapitola 12). Zaměříme se také podrobněji na oblast grafových NoSQL databází (kapitola 13), které se od ostatních tří typů výrazně liší, a představíme si také hybridní databáze, NewSQL databáze a další speciální typy databází pro Big Data (kapitola 14). V závěru se pokusíme shrnout aktuální stav oblasti a diskutovat předpokládaný další vývoj.

2.

Datové formáty

Možná přemýšlíte, proč kniha o databázích začíná kapitolou o datových formátech. Je to proto, že práce s moderními databázemi NoSQL znalost datových formátů vyžaduje. Na rozdíl od relačních databází, které staví na abstraktním relačním modelu dat, u NoSQL databází často formální datový model neexistuje nebo je tak jednoduchý, že vnitřní strukturu ukládaných dat řídí sama aplikace. Typicky k tomu používá právě některý existující široce rozšířený formát. Například databáze typu klíč-hodnota pracují s daty jako s nestrukturovaným binárním objektem a pro ukládání strukturovaných hodnot používají nejčastěji formát JSON. Ve světě dokumentově orientovaných databází jsou typickými kandidáty formáty JSON a XML.

Znalost datových formátů je důležitá i pro celkový návrh aplikací používajících NoSQL databáze. Většina aplikací dnes pracuje ve webovém prostředí. Klíčové části aplikace jsou implementovány jako služby, které jsou využívány dalšími částmi aplikace, například klientskou aplikací napsanou v Javascriptu a běžící přímo ve webovém prohlížeči. Jednotlivé části aplikace si data předávají opět v přesně definovaném formátu. V ideálním případě je datový formát pro komunikaci shodný s formátem pro ukládání, aby se omezila režie potřebná pro konverzi dat.

Nyní se stručně seznámíme s jednotlivými běžně používanými datovými formáty.