

Big Data i 5V

Nowe wyzwania w świecie danych

Krzysztof Goczyła



Wydział Elektroniki, Telekomunikacji i Informatyki
Politechnika Gdańska
[*kris@eti.pg.gda.pl*](mailto:kris@eti.pg.gda.pl)



Sopot, 10.09.2014

O czym będzie?

- ❑ Co to jest „Big Data”?
- ❑ 5V
- ❑ Problemy i metody przetwarzania Big Data
- ❑ Istniejące rozwiązania
- ❑ Podsumowanie



Co to jest „Big Data”?

Big Data (wielkie dane)

Zbiory informacji o bardzo dużej **objętości**, dużej **szybkości narastania** i dużej **różnorodności**, które wymagają nowych, nieklasycznych metod i narzędzi do ich **przechowywania**, **przetwarzania** i **analizowania** dla potrzeb podejmowania decyzji, odkrywania nowych zjawisk i optymalizacji procesów.

Szacowania:

- światowe zbiory cyfrowe osiągnęły 5 ZB ($= 5 \cdot 10^{21}$ B) (*International Data Corporation*)
- światowy przepływ danych w 2014 r osiągnie 75 EB ($= 75 \cdot 10^{18}$ B) (*IBM*)

Źródła Big Data

- ☐ Dane generowane przez użytkowników portali internetowych, w tym sieci społecznościowych
- ☐ Dane opisujące transakcje dokonywane poprzez Internet
- ☐ Dane naukowe (biologiczne, astronomiczne, pomiary fizyczne itp.)
- ☐ Dane generowane przez roboty przeszukujące Internet (*Web mining, Web crawling*)
- ☐ Dane grafowe obrazujące powiązania pomiędzy stronami WWW
- ☐ Dane pochodzące czujników i urządzeń podłączonych do Internetu (Internet of Things)
- ☐ ...

5 (=3+2) V

Volume – ogromna objętość danych, liczona w tera- i petabajtach, wykraczająca poza normy spotykane w bazach danych

Velocity – duża szybkość narastania danych, większa niż w tradycyjnych systemach

Variety – duża różnorodność struktury danych i ich treści

Veracity – wiarygodność i dokładność (prawdziwość) danych

Value – rzeczywista wartość dla analiz i procesów podejmowania decyzji

Rozproszenie danych

Rozmieszczenie danych na wielu komunikujących się ze sobą serwerach danych (komputerach fizycznych lub wirtualnych).

Zasady:

- serwery połączone są siecią
- serwery mogą funkcjonować autonomicznie (dopuszczamy możliwość izolowania fragmentów sieci)
- rozproszenie powinno być przezroczyste dla użytkownika

Replikacja danych

Sytuacja, w której w systemie rozproszonym te same dane znajdują się na wielu serwerach danych.

Problemy:

- spójność (*consistency*)
- spójność a dostępność (*consistency vs availability*)
- ostateczna spójność (*eventual consistency*)

Rozproszenie i replikacja (R+R)

Rozproszenie danych zwiększa dostępność (*availability*), efektywność przetwarzania (*efficiency*) i bezpieczeństwo dostępu do danych (*security*), a także zapewnia *poziomą skalowalność* (*horizontal scalability*), konieczną do obsługi Big Data.

Replikacja danych zwiększa dostępność (*availability*) i efektywność dostępu do danych (*efficiency*), a także polepsza ogólną niezawodność (*reliability*) systemu rozproszonego.

Problem: (twierdzenie CAP)

W systemie rozproszonym nie można jednocześnie zapewnić spójności i dostępności w każdym momencie funkcjonowania systemu.

Cechy charakterystyczne:

- ☐ brak schematu danych
(zasada *schema on read* zamiast *schema on write*)
- ☐ złożone, zagnieżdżone struktury danych
- ☐ brak normalizacji danych
- ☐ efektywność przetwarzania dużych danych drogą rozproszenia i zrównoleglenia operacji
- ☐ łatwość w dokonywaniu replikacji danych
- ☐ duża dostępność (*availability*)
- ☐ wysoka skalowalność horyzontalna
- ☐ brak deklaratywnego języka dostępu (brak możliwości złączeń, ...),
zamiast tego interfejsy programistyczne (Java, JavaScript, Ruby, Gremlin, ...)
- ☐ ograniczone zachowanie zasad ACID
- ☐ ograniczone możliwości indeksowania

Systemy typu *klucz/wartość*

Dane przechowywane są w postaci par *klucz/wartość*.
Wartość może być dowolnego typu, np. ciąg bajtów interpretowanych programowo. Klucz może być hierarchiczny.

Klucze	Wartości
...	...
pomorskie	<i>jakieś informacje o woj. pomorskim</i>
pomorskie/Gdańsk	<i>jakieś informacje o Gdańsku</i>
pomorskie/Gdynia	<i>jakieś informacje o Gdyni</i>
...	...
mazowieckie	<i>jakieś informacje o woj. mazowieckim</i>
mazowieckie/Warszawa	<i>jakieś informacje o Warszawie</i>
mazowieckie/Warszawa/Mokotów	<i>jakieś informacje o Mokotowie - dzielnicy Warszawy</i>
...	...

Przykłady baz *open source* tego typu: Riak, Redis, Voldemort

Systemy kolumnowe

Dane przechowywane są kolumnami a nie wierszami, tj. dane z danej kolumny są przechowywane razem. Dodawanie kolumn jest bardzo łatwe. Każdy wiersz może mieć inny zestaw kolumn, bez konieczności przechowywania wartości pustych (NULL) lub nieznanych (UNKNOWN).

	<i>klucze wierszy</i>	<i>rodzina kolumn "kolor"</i>	<i>rodzina kolumn "kształt"</i>	<i>rodzina kolumn "rozmiar"</i>
<i>wiersz</i>	"1"	"biały": "#F00" "zielony": "#F0F" "żółty": "#F00"	"kwadrat": "4" "trójkąt": "3"	
<i>wiersz</i>	"2"		"prostokąt": "4"	"a": "12" "b": "8"
<i>wiersz</i>	"3"	"biały": "#F00" "czarny": "#FFF"		

wiersze, klucze, rodziny kolumn, kolumny, wartości

Przykłady baz *open source* tego typu: HBase, Cassandra, Hypertable

Systemy dokumentowe

Służą do przechowywania dokumentów. Dokument ma identyfikator i treść o dowolnie złożonej, elastycznej strukturze dopuszczającej zagnieżdżanie innych dokumentów.

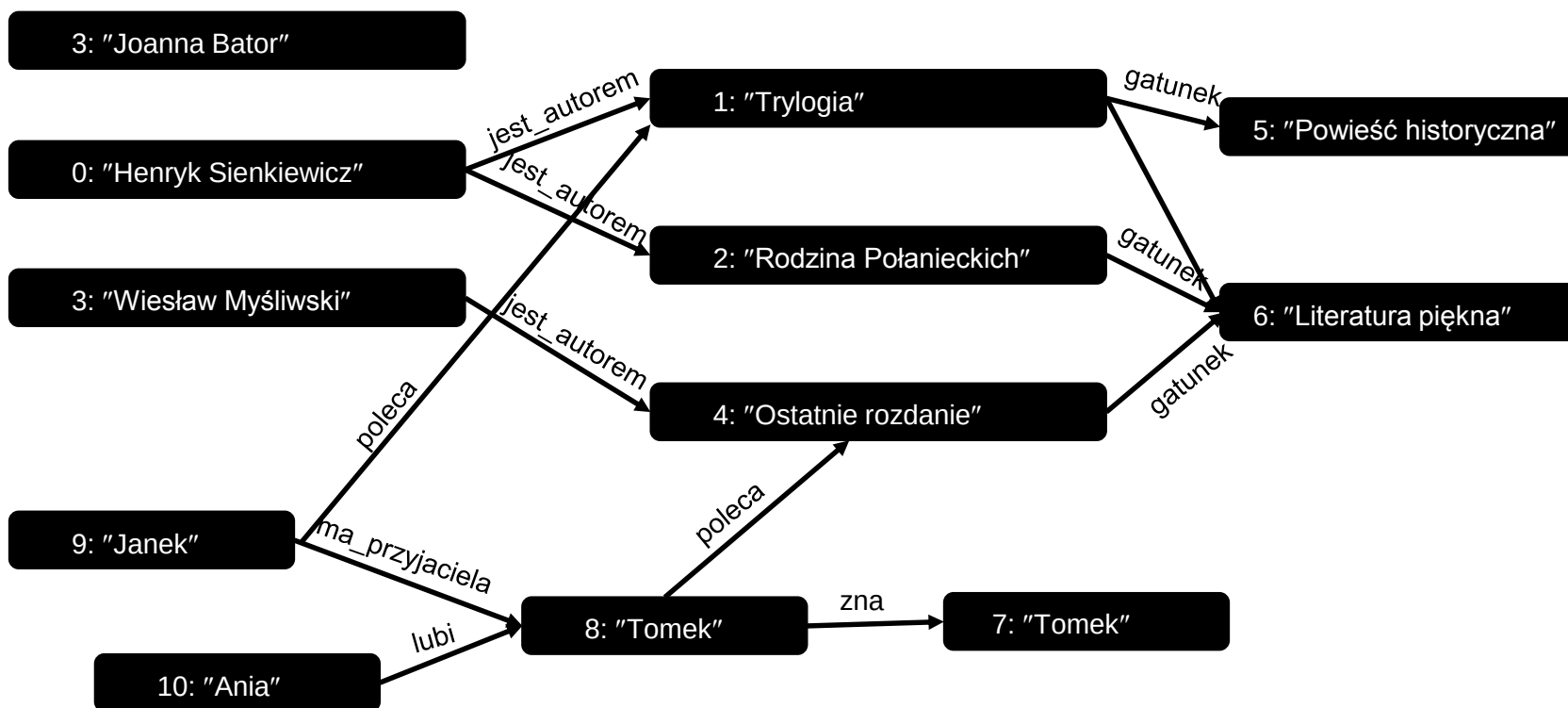


Wydruk dokumentu w formacie JSON

Przykłady baz *open source* tego typu: MongoDB, CouchDB

Systemy grafowe

Grafowa baza danych to zbiór węzłów i relacji pomiędzy węzłami (łuków). Węzły i łuki mogą mieć własności (atrybuty) postaci *klucz/wartość*.



Sieć społecznościowa w postaci grafowej bazy danych

Przykłady baz open source tego typu: Neo4J

Ogólna zasada:

Obliczenia (wykonywane równolegle na wielu maszynach)
pobierają wejściowe pary
klucz/wartość
i produkują inne pary wyjściowe
klucz/wartość.

Faza 1

Map: $(k1, v1) \rightarrow k2, v2$

Faza 2

Reduce: $(k2, \text{List}(v2)) \rightarrow (k2, \text{List}(v3))$

Potrzebna jest **warstwa pośrednicząca**, która zgrupuje wyniki fazy Map według tej samej wartości $k2$ i przekaże je do fazy Reduce.

MapReduce – Przykład

Dana jest duża kolekcja dokumentów. Oblicz liczbę wystąpień każdego słowa w tej kolekcji.

Map (String key, String value):

// key: nazwa dokumentu

// value: zawartość dokumentu

for each word w in value:

Emit (w, "1");

← Wynik: List(słowo, "1")

Reduce (String key, List values):

// key: słowo

// value: lista liczników

int result = 0;

for each v in values:

result += ParseInt(v);

Emit (key, AsString(result));

← Wynik: List(słowo, n)

MapReduce – Typowe zastosowania (1)

- **Poszukiwanie tekstu wg zadanego wzorca:**

W kolekcji dokumentów znajdź wiersze tekstu, które odpowiadają zadanemu wzorcowi.

- **Mierzenie popularności stron WWW:**

Dla każdej strony WWW z kolekcji adresów URL podaj liczbę dostępow do tej strony na podstawie logów.

- **Budowa grafu powiązań stron WWW:**

Dla każdej strony WWW z kolekcji adresów URL podaj listę stron, do których dana strona się odwołuje.

MapReduce – Typowe zastosowania (2)

➤ Budowa indeksu odwróconego:

Dla każdego słowa występującego w kolekcji dokumentów podaj listę dokumentów, w których dana słowo występuje.

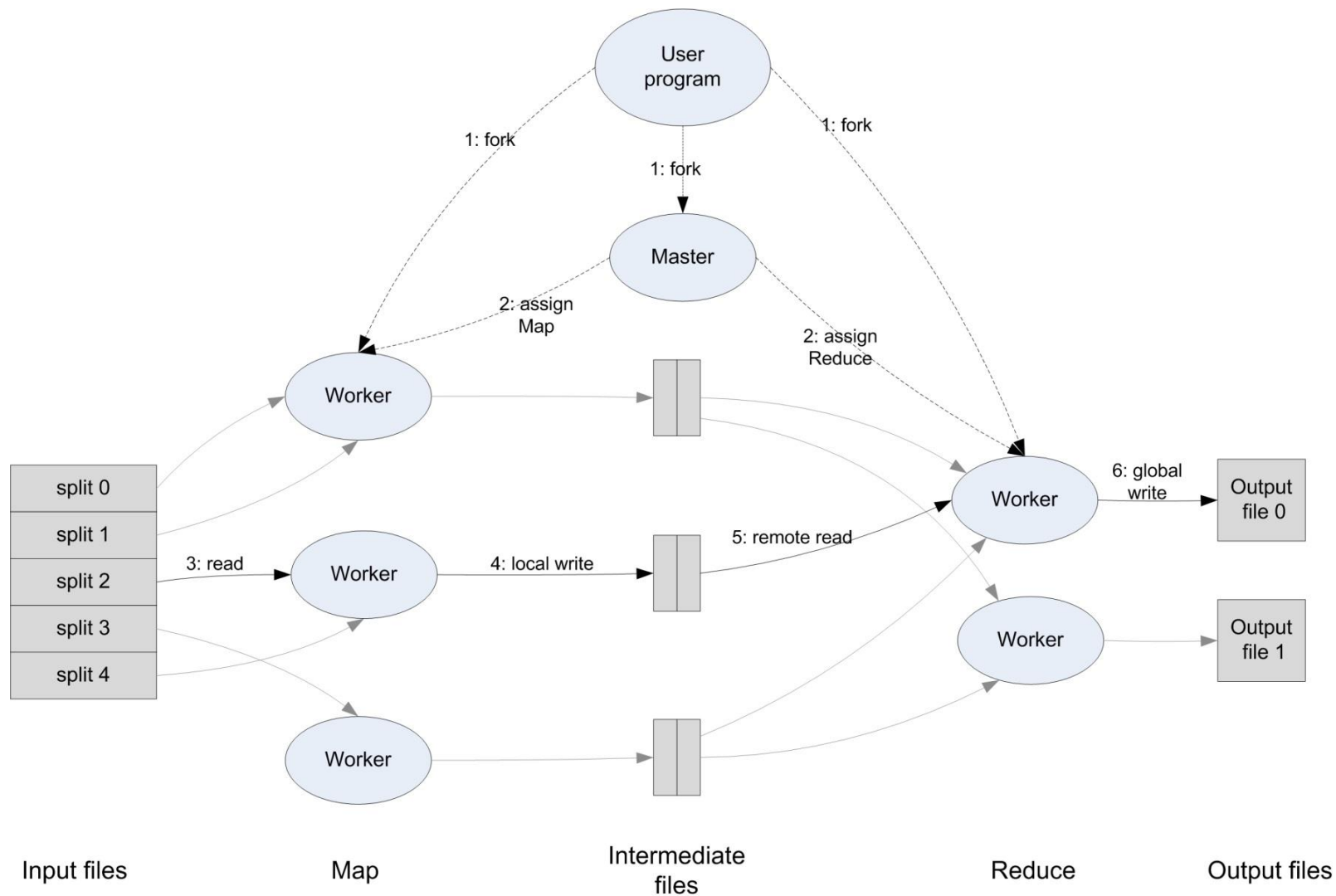
➤ Sortowanie:

Posortuj listę rekordów według wartości podanego klucza.

➤ Analiza częstotliwości słów:

Dla danego zbioru stron WWW danych adresami URL podaj częstotliwości słów występujących na danej stronie.

MapReduce - Realizacja



Cloud Computing (przetwarzanie w chmurze)

Paradygmat przetwarzania bazujący na wirtualizacji zasobów.

Wirtualizacja

Mechanizm dostępu do zasobów centrum przetwarzania (Data Center) poprzez oprogramowanie ukrywające fizyczne aspekty tych zasobów. Tworzy to wirtualną platformę przetwarzania, do której fizyczne zasoby mogą być przydzielane dynamicznie, zależnie od aktualnych potrzeb, bez wiedzy i ingerencji użytkownika.

Wirtualizowane zasoby:

- procesory
- pamięć dyskowa
- pamięć operacyjna
- sieć komunikacyjna

Modele dostarczania usług przetwarzania w chmurze:

➤ Infrastructure as a Service (IaaS)

Użytkownik otrzymuje sprzęt (procesory, pamięć operacyjną i dyskową, sieć komunikacyjną, tak jakby miał to u siebie.

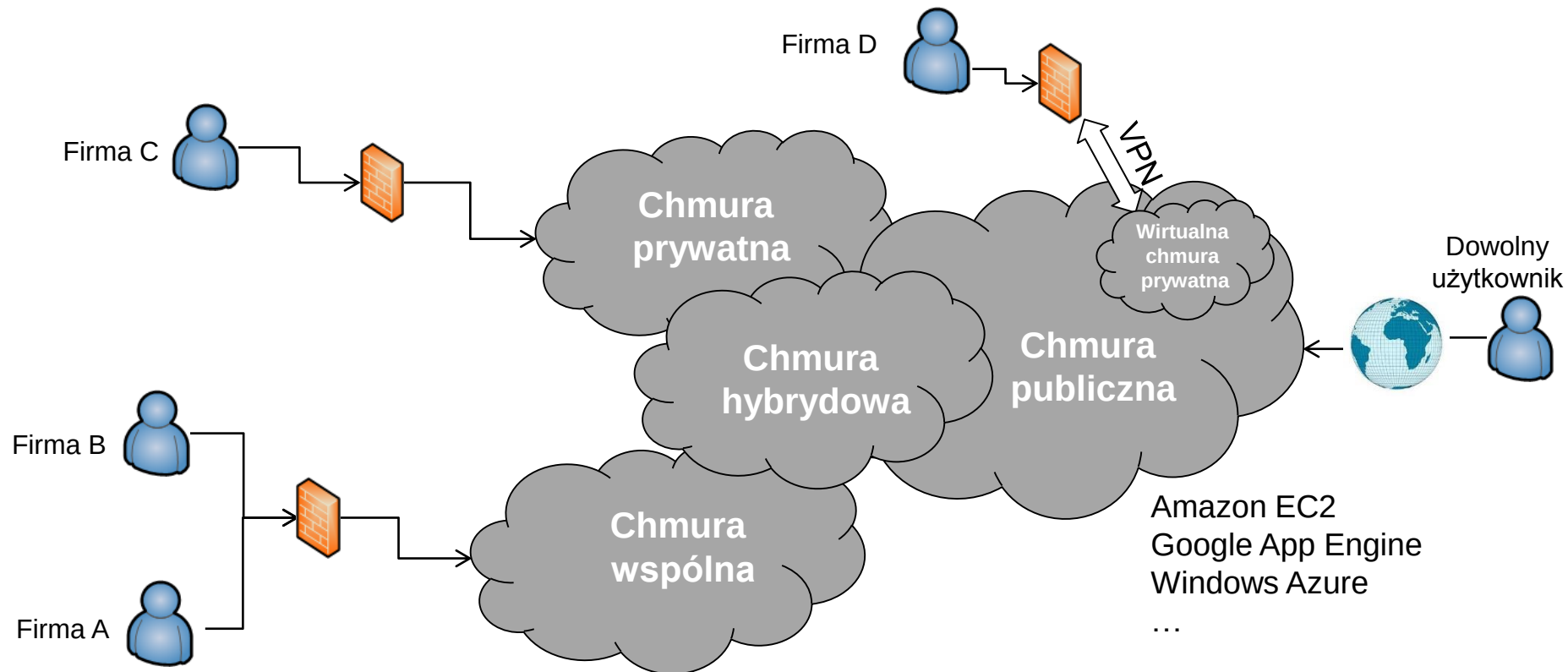
➤ Platform as a Service (PaaS)

Użytkownik otrzymuje system operacyjny i narzędzia programistyczne, wszystko na skalowalnej platformie.

➤ Software as a Service (SaaS)

Użytkownik ma dostęp do konkretnych aplikacji w ich najaktualniejszej wersji.

Cloud Computing



Modele biznesowe korzystania z usług w chmurze

Podsumowanie

- ❑ Big Data – obszerne dane o zmiennej strukturze, napływające szybko i nieustannie z różnych źródeł
- ❑ Wymagają nowych narzędzi do przechowywania, zarządzania i przetwarzania
- ❑ Odpowiedzią są wysoce skalowalne bazy NoSQL i algorytmy przetwarzania o wysokim stopniu rozproszenia i zrównoleglenia
- ❑ Problem najpoważniejszy – integracja, redukcja i interpretacja Big Data oraz wydobywanie z nich informacji i wiedzy

Dziękuję za uwagę!

