



Das SCAPE Projekt: Langzeitarchivierung und Skalierbarkeit

Teil 2: Anwendungsfälle an der Nationalbibliothek

Dr. Sven Schlarb

Österreichische Nationalbibliothek

SCAPE ½ Informationstag

05. Mai 2014, Österreichische Nationalbibliothek.

- Vorteil: Gute Lese-Performanz
Nachteil: Schlechte Performanz im Schreiben zufälliger Werte
- Schlechte Performanz bei große Tabellen mit vielen Spalten
- Effiziente Lastverteilung bei voller Verfügbarkeit bedeutet hohen Verwaltungsaufwand
→ Sharding: Aufwändig und wartungsintensiv
- Vertikale Skalierung → Teuere Hardware
- ÖNB: 2011 Versuch der Web-Archiv-Indexerstellung mit MySQL an unzureichender Performanz gescheitert!

- Elastizität → Einfache Erweiterbarkeit
- Hoher Schreibdurchsatz
- Konsistenz → Prüfsummen in HBase
- Effiziente zufällige Lesevorgänge
- Hohe Verfügbarkeit und Fehlertoleranz
- Atomare Lese- und Schreibmethoden → Parallele Datenmodifikation ohne Sperrungen
- Effiziente Bereichsanfragen → Gutes Leseverhalten bei sequentiellen Daten

- Kostenvorteil von Hadoop, da
 - meist auf preiswerter Hardware
 - Ohne Bindung an einen bestimmten Hersteller
 - Basierend auf Open-Source-Software

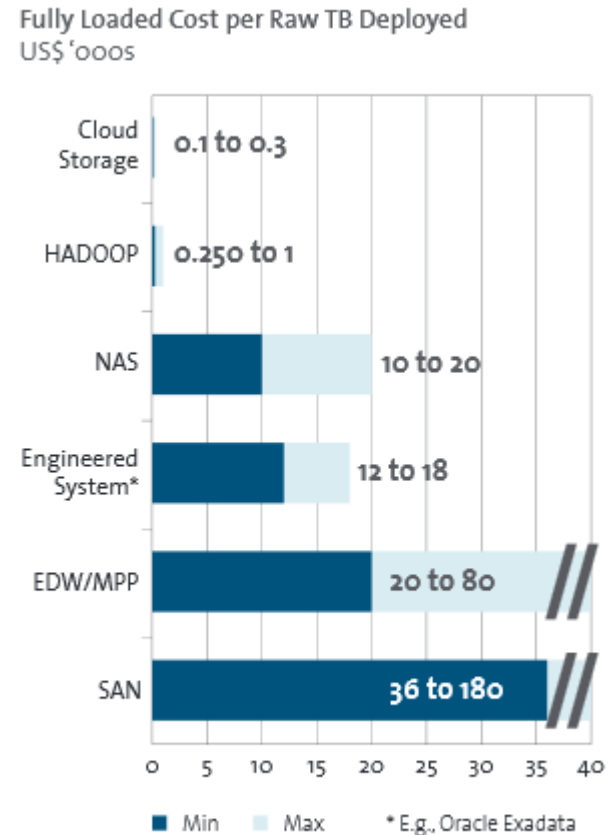


Abbildung 5: Kostenvergleich Hadoop versus Alternativen

Quelle: BITKOM Leitfaden Big-Data-Technologien-Wissen für Entscheider 2014, S. 39

Speichern und Verarbeiten getrennt



- Daten auf NAS-Speicher
- Müssen zur Verarbeitung in andere Server-Umgebung geladen werden
- Multi-Terabyte-Szenarien?



Speichern und Verarbeiten zugleich

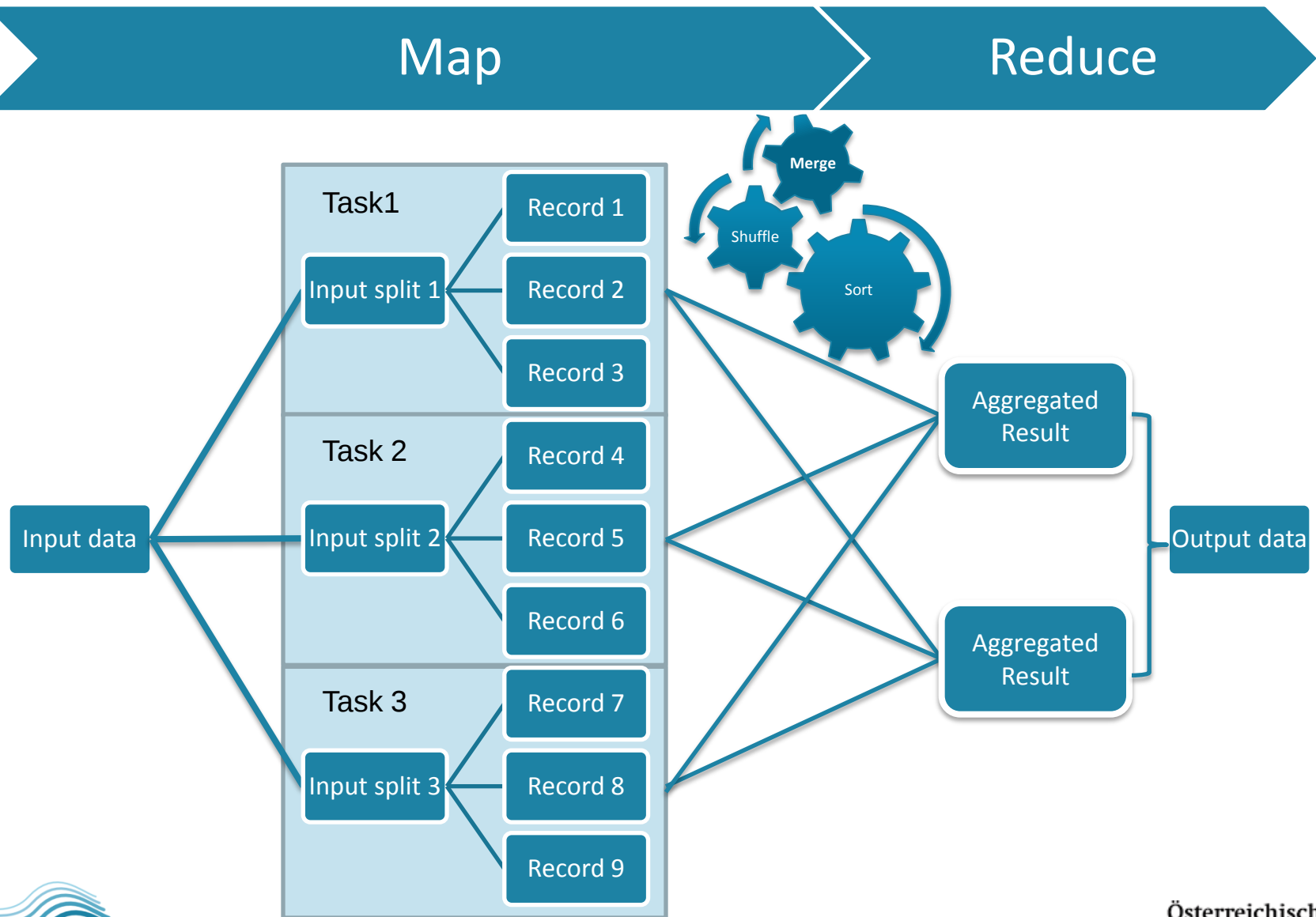


- Prozessoren bei den Daten
- Datenverarbeitung kann unmittelbar erfolgen

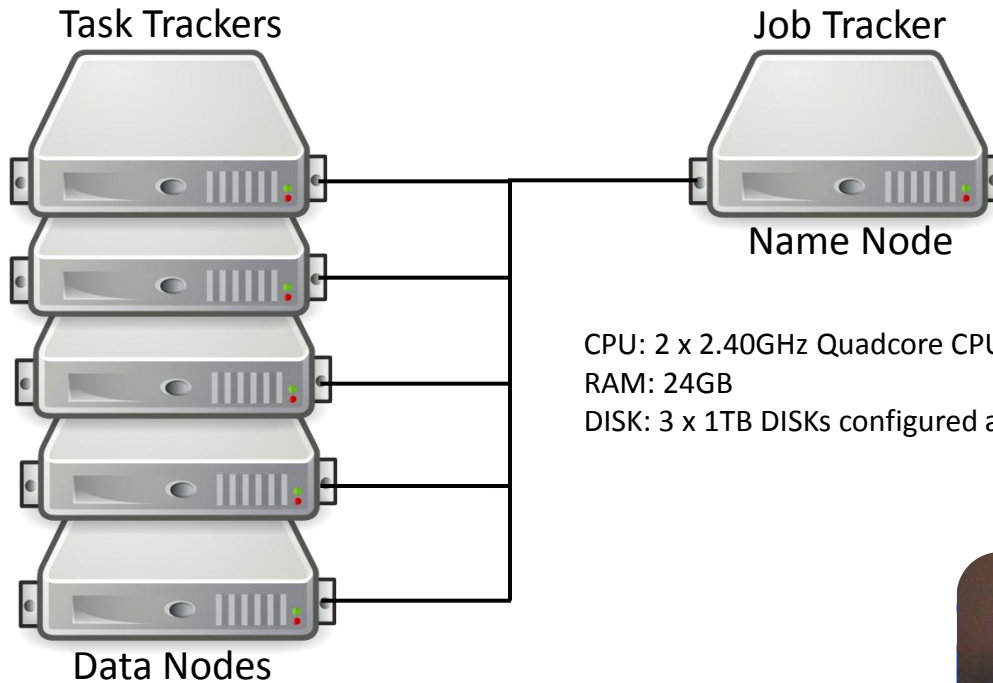


- Bei großen Datenmengen ist es meist einfacher die **verarbeitende Komponente zu den Daten** zu bringen als umgekehrt, die Daten zur verarbeitenden Komponente
- **Feingranulare Parallelisierung:** Die Ausführung der Datenverarbeitung findet auf den zur Verfügung stehenden Prozessorkernen statt
- **Hardware-bedingte Ausfälle** sind keine Besonderheit, sondern es gibt spezielle Vorkehrungen dafür
- **Redundanz:** Datenblöcke werden redundant gespeichert (Default: 3x) → Ausfallsicherheit, Flexibler Zugriff auf Daten
- **Daten-Lokalität:** Freier Knoten mit möglichst direktem Zugang zu Datenblock übernimmt die Verarbeitung

MapReduce/Hadoop auf einen Blick



Experimenteller Computer-Cluster



CPU: 2 x 2.40GHz Quadcore CPU (16 HyperThreading cores)
RAM: 24GB
DISK: 3 x 1TB DISKs configured as RAID5 (Redundanz) – 2 TB effektiv

CPU: 1 x 2.53GHz Quadcore CPU (8 HyperThreading)
RAM: 16GB
DISK: 2 x 1TB DISKs konfiguriert als RAID0 (Performance) – 2 TB effektiv

- Of 16 HT cores: 5 for Map; 2 for Reduce; 1 für Betriebssystem.

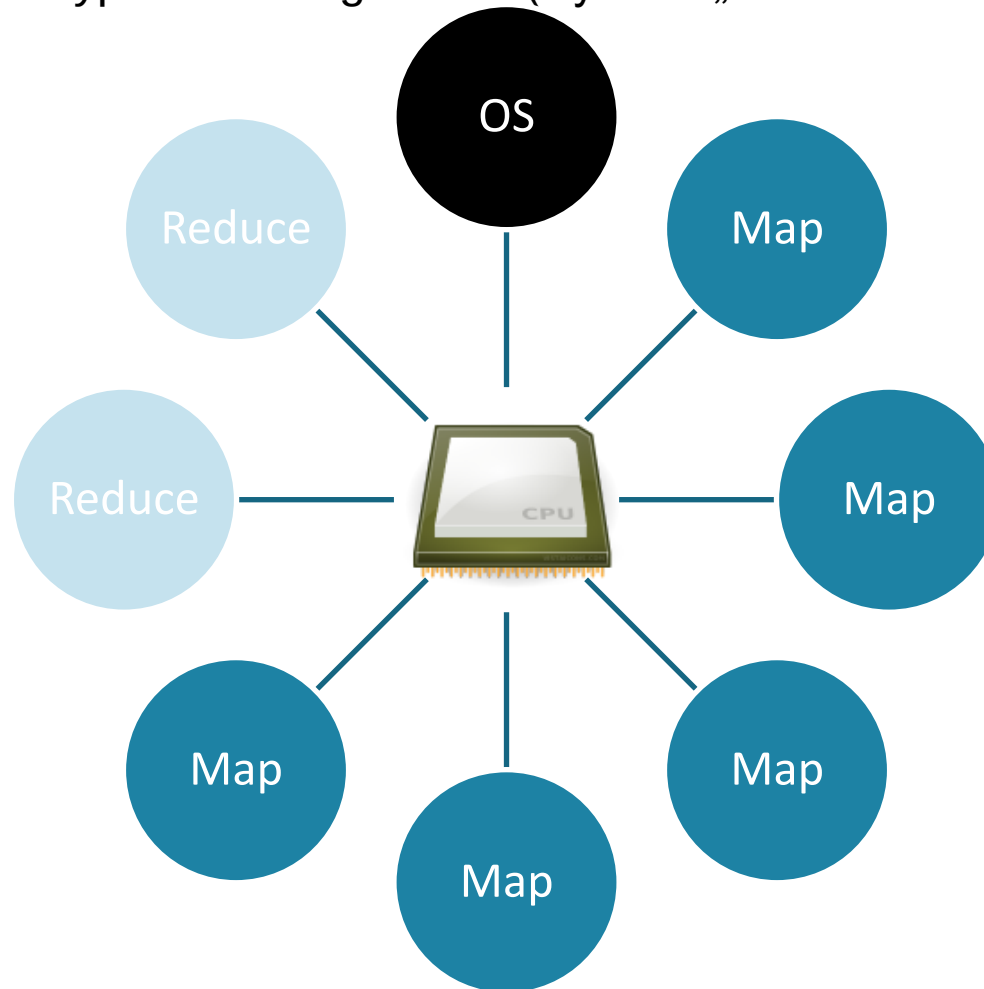
→ 25 Prozessorkerne für **Map** tasks

→ 10 Prozessorkerne für **Reduce** tasks



Die physische Sicht: Prozessor (CPU)

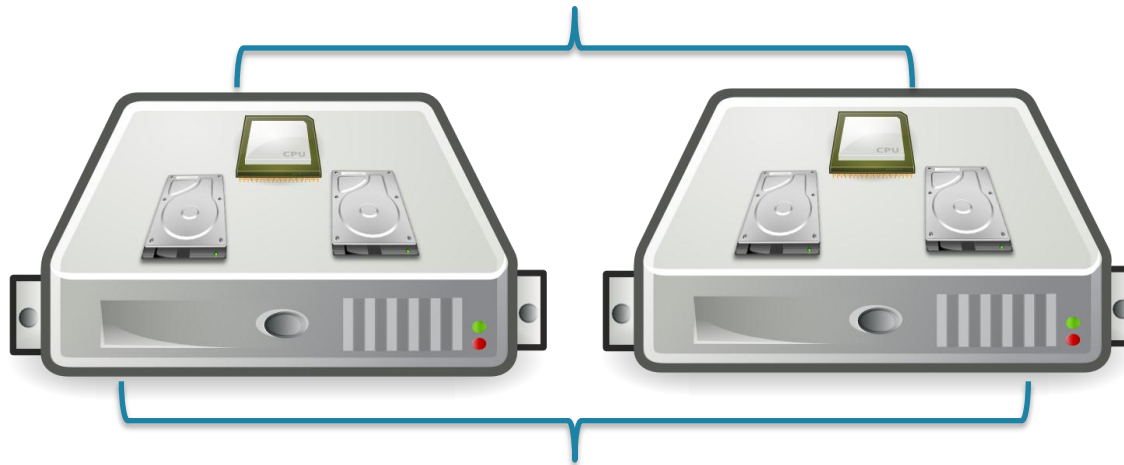
Beispielkonfiguration einer Quad-Core-CPU Pro Cluster-Knoten
4 physische Kerne
8 Hyperthreading-Kerne (System „sieht“ 8 Kerne)



Die physische Sicht: Cluster-Knoten

Verteilte Datenverarbeitung (MapReduce)

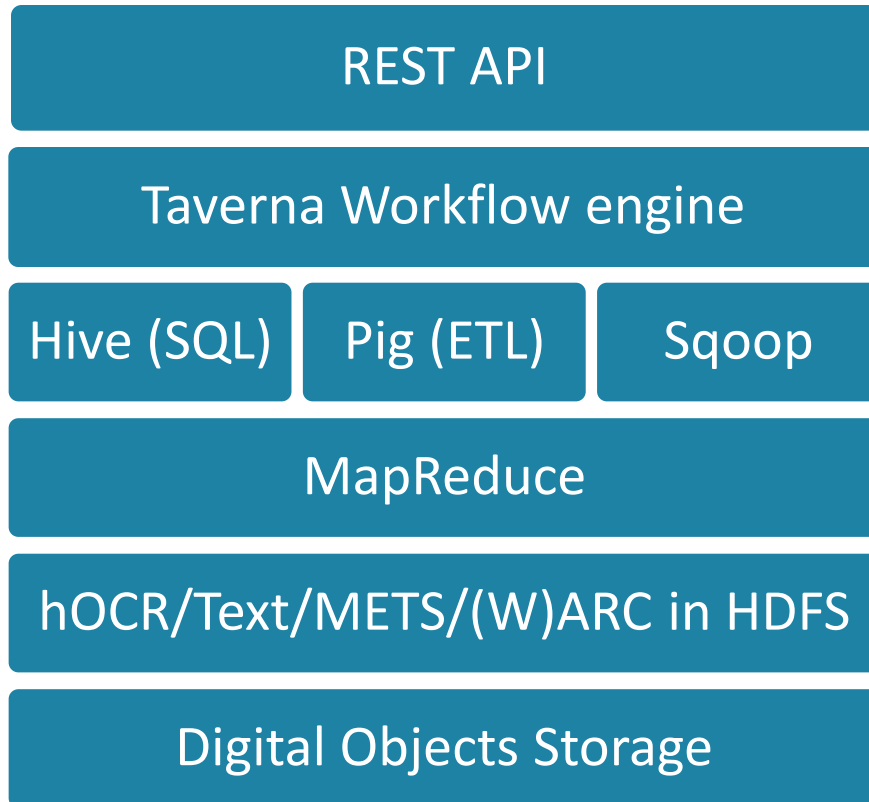
2 x Quad-Core-CPU:
10 Map(Parallelisierung)
4 Reduce (Aggregation)



4 x 1 TB Festplatten bei Redundanz 3:
1,33 TB effektiv (rein rechnerisch)

Verteilter Datenspeicher (HDFS)

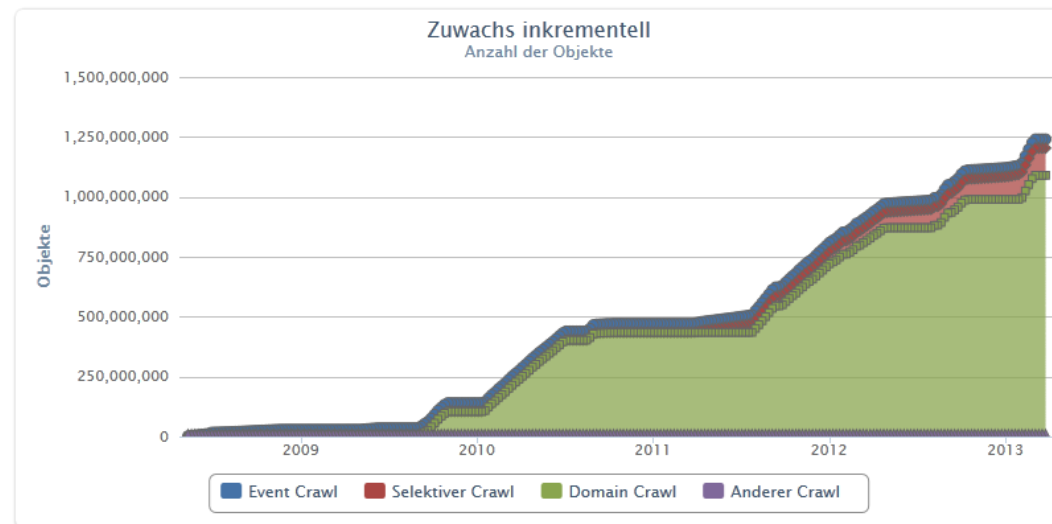
Hadoop = MapReduce + HDFS



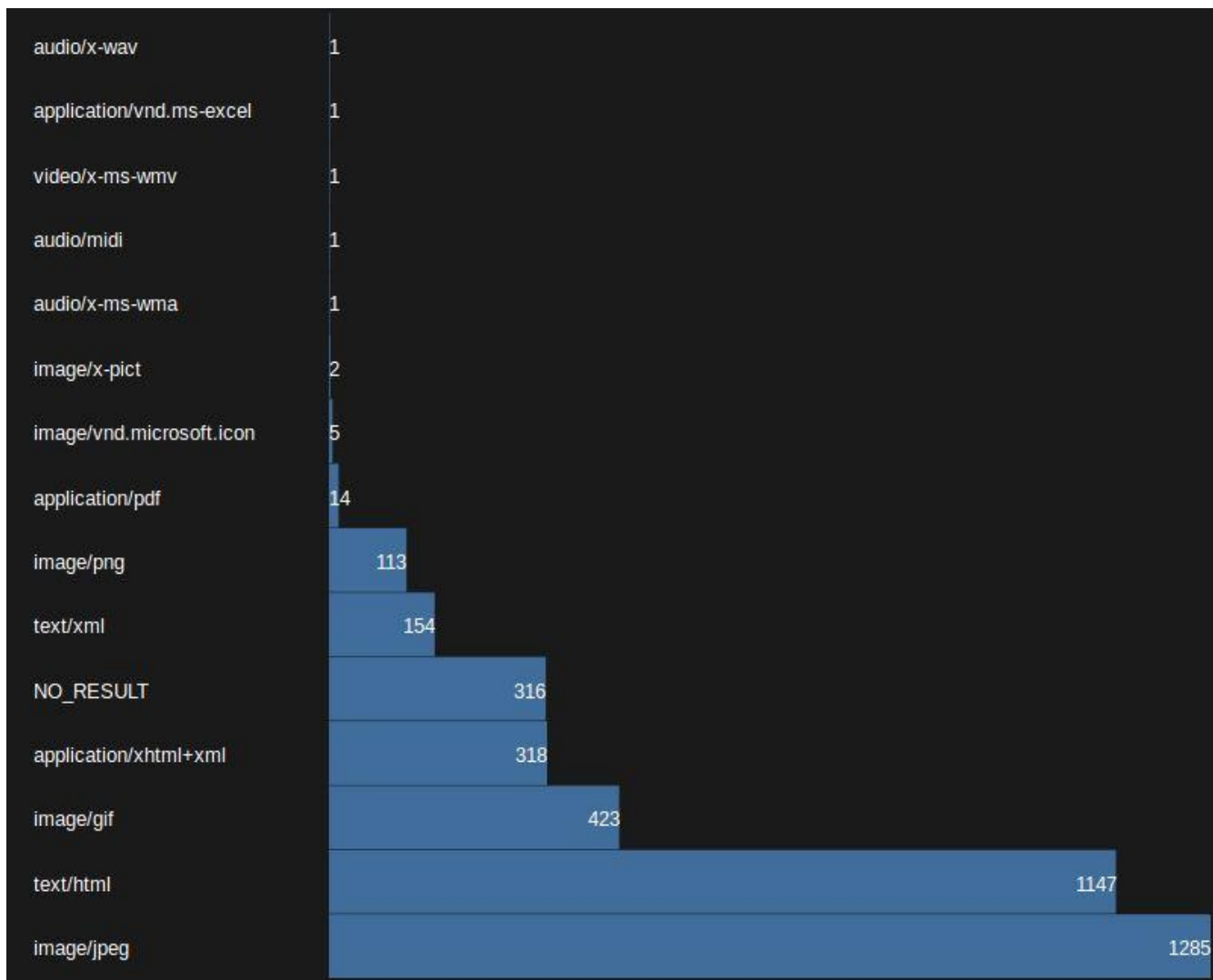
- Steuerung via REST API
- Workflow Engine für komplexe Jobausführung
- Hive als Abfrageschnittstelle – an (My)SQL-Syntax angelehnt
- MapReduce/Pig zum Laden und Transformieren v. Daten
- Sqoop für die DBMS-Integration
- „Kleine“ Objekte in HDFS
- „Große“ Objekte auf NAS

- Web-Archivierung
 - Dateiformat-Identifikation im Web-Archiv
- Austrian Books Online
 - Bild-Datei-Format-Migration
 - Vergleich verschiedener Buch-Derivate
 - MapReduce in der Qualitätssicherung digitaler Bücher

- Domain Harvesting
 - Gesamte **top-level-domain** **.at** alle 2 Jahre
- Selektives Harvesting
 - „Wichtige“ Websites die sich häufig ändern
- Event-basiertes Harvesting
 - Besondere Veranstaltungen und Events (z.B. Wahlen oder Sport-Events)
- Speicher: ca. 45TB
- ca. 1,7 Milliarden Objekte

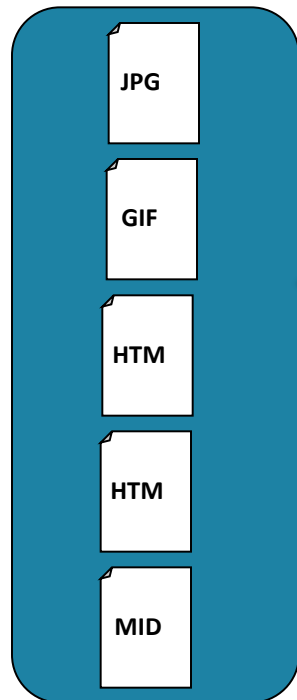


Dateiformat-Identifikation im Web-Archiv



Dateiformat-Identifikation im Web-Archiv

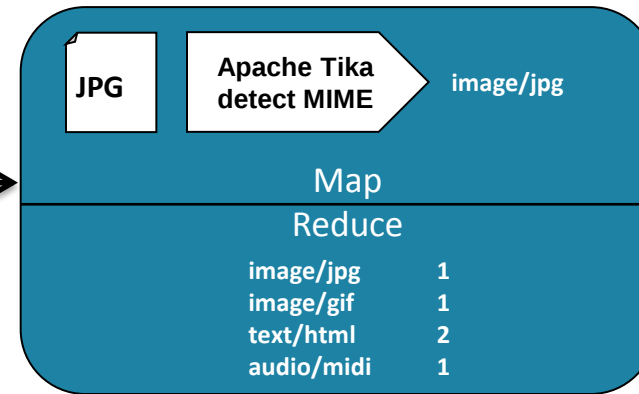
(W)ARC Container



(W)ARC InputFormat
(W)ARC RecordReader



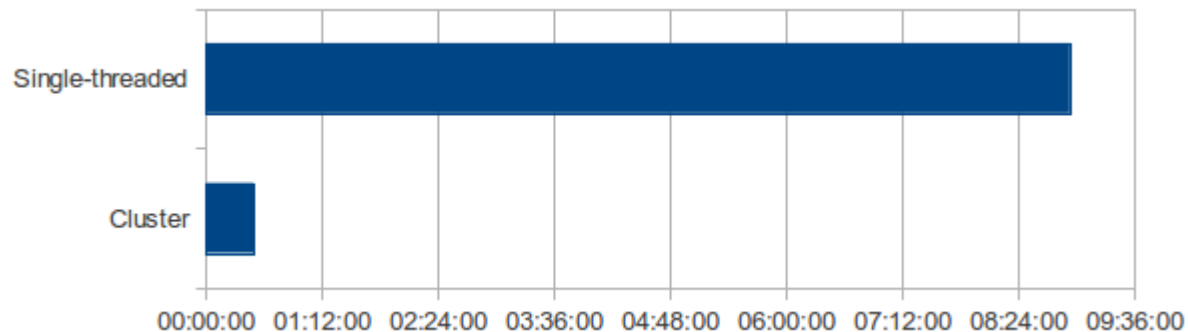
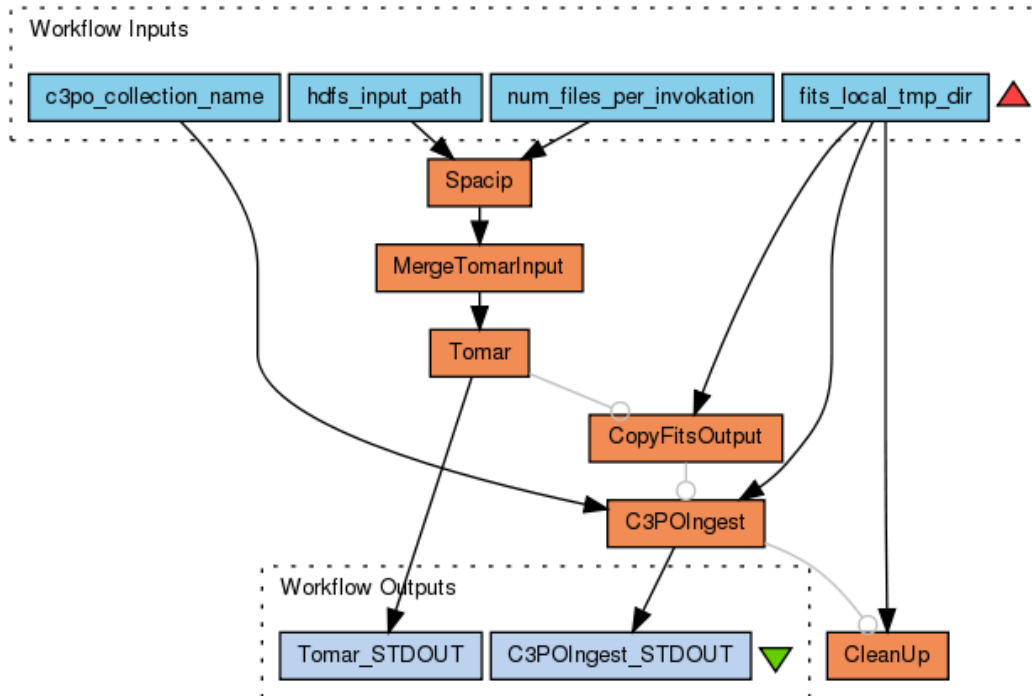
MapReduce



Software-Integration	Durchsatz(GB/min)
TIKA detector API in Map Phase	6,17 GB/min
FILE als Kommandozeilen-Applikation mit MapReduce	1,70 GB/min
TIKA JAR als Kommandozeilen-Applikation mit MapReduce	0,01 GB/min

Datenmenge	Anzahl der ARC-Dateien	Durchsatz(GB/min)
1 GB	10 x 100 MB	1,57 GB/min
2 GB	20 x 100 MB	2,5 GB/min
10 GB	100 x 100 MB	3,06 GB/min
20 GB	200 x 100 MB	3,40 GB/min
100 GB	1000 x 100 MB	3,71 GB/min

Extraktion von Datei-Eigenschaften im Web-Archiv



Screencast

- Partnerschaft mit Google
- Ausschließlich gemeinfreie Werke
- Zielsetzung ~ 600.000 Bände
 - ~ 200 Mio. Seiten
- ~ 70 Projekt-Teilnehmer
 - 20+ im Kernteam
- Aktuell ~ 170.000 Bände verfügbar
 - ~ 50 Mio pages
 - ~ 18 Milliarden Token im Volltext, der mit Hilfe automatisierter Texterkennung (OCR) generiert wurde



ADOCO (Austrian Books Online Download & Control)

Digitalisierung

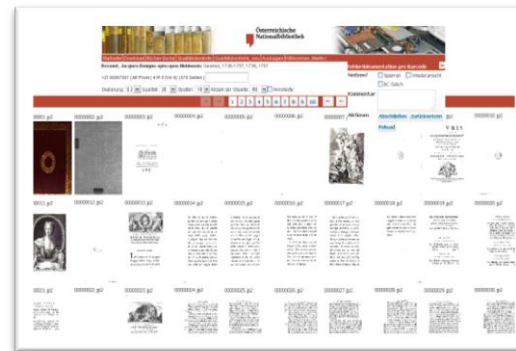
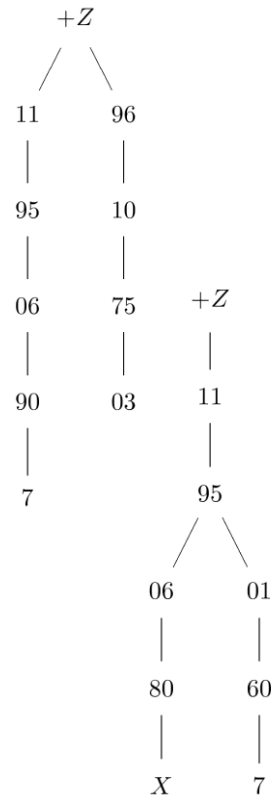
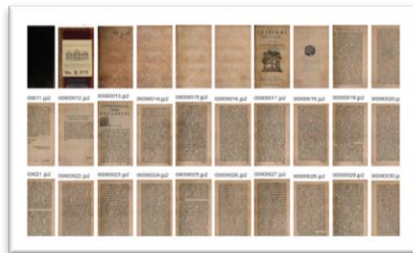
Download &
Speicherung

Qualitäts-
Kontrolle

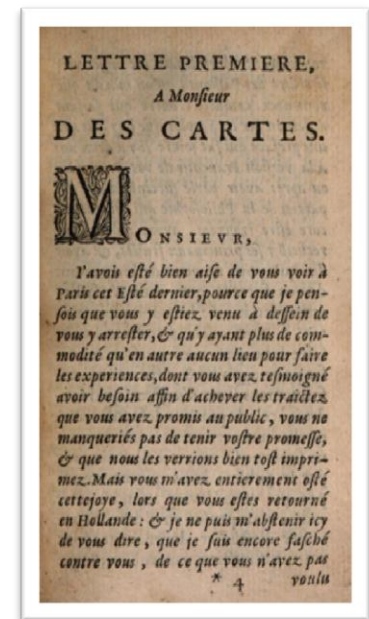
Zugang



Google
Public Private Partnership

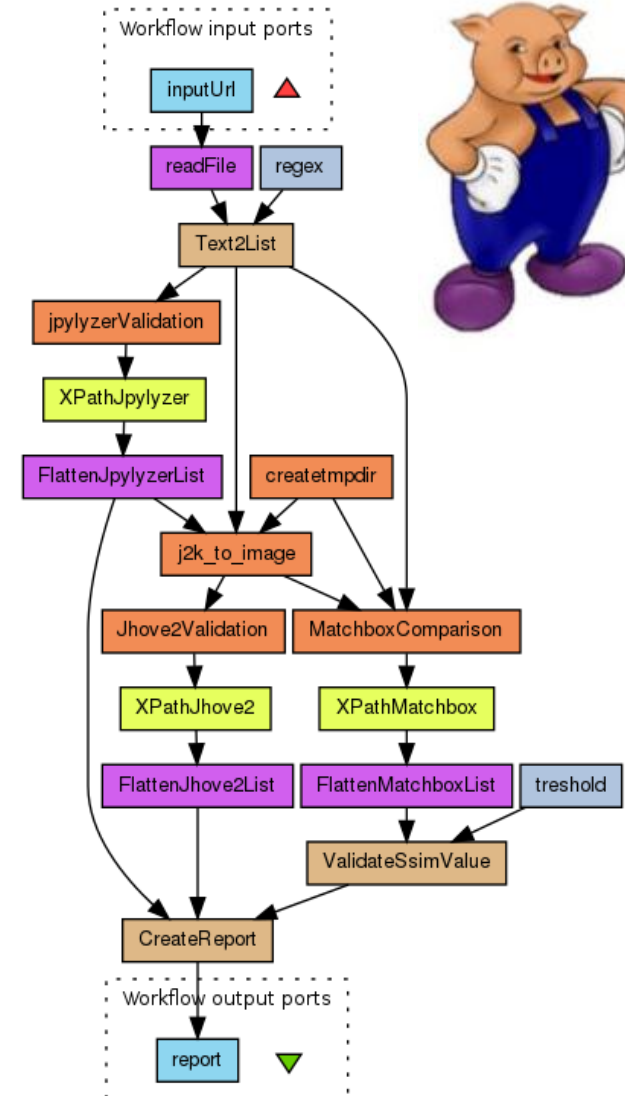


ADOCO



<https://confluence.ucop.edu/display/Curation/PairTree>

- Bildformat-Migration
 - TIFF nach JPEG2000
 - Speicherbedarf reduzieren
 - JPEG2000 nach TIFF
 - Risiko: Dateiformat nicht mehr ausreichend verbreitet, unterstützt und verwendet
- Herausforderungen:
 - Integration verschiedener Werkzeuge der Datei-Format-Validierung, Migration und Qualitätssicherung
 - Qualitätssicherung rechenintensiv (abhängig von Methode)



- Vergleich verschiedener Digitaler Buchversionen
 - Bilder wurden verändert
 - Bilder kommen von verschiedenen Scan-Quellen
- Herausforderungen:
 - Rechenintensiv
 - 170.000 Bände à ca. 400 Seiten
- SCAPE Software: Matchbox

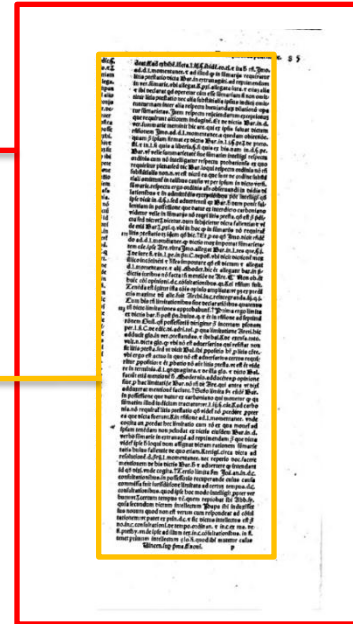
- Verarbeitung von 60.000 Bänden, ~ 24 Millionen Seiten
- Taverna Tool service (Job-Ausführung per SSH)
- Verschiedene Hadoop-Jobs
 - Hadoop-Streaming-API
 - Hadoop Map/Reduce
 - Hive
- Workflow auf myExperiment verfügbar:
<http://www.myexperiment.org/workflows/3105>
- Blogpost:
<http://www.openplanetsfoundation.org/blogs/2012-08-07-big-data-processing-chaining-hadoop-jobs-using-taverna>

Indikator für möglichen Textverlust

Korrekte Bildbeschneidung

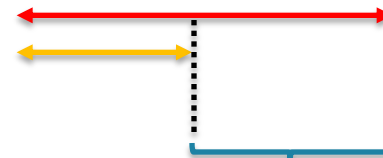
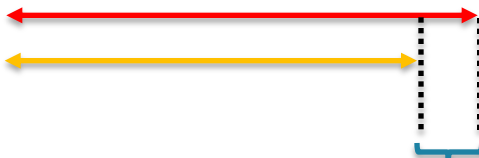


Fehlerhafte Bildbeschneidung



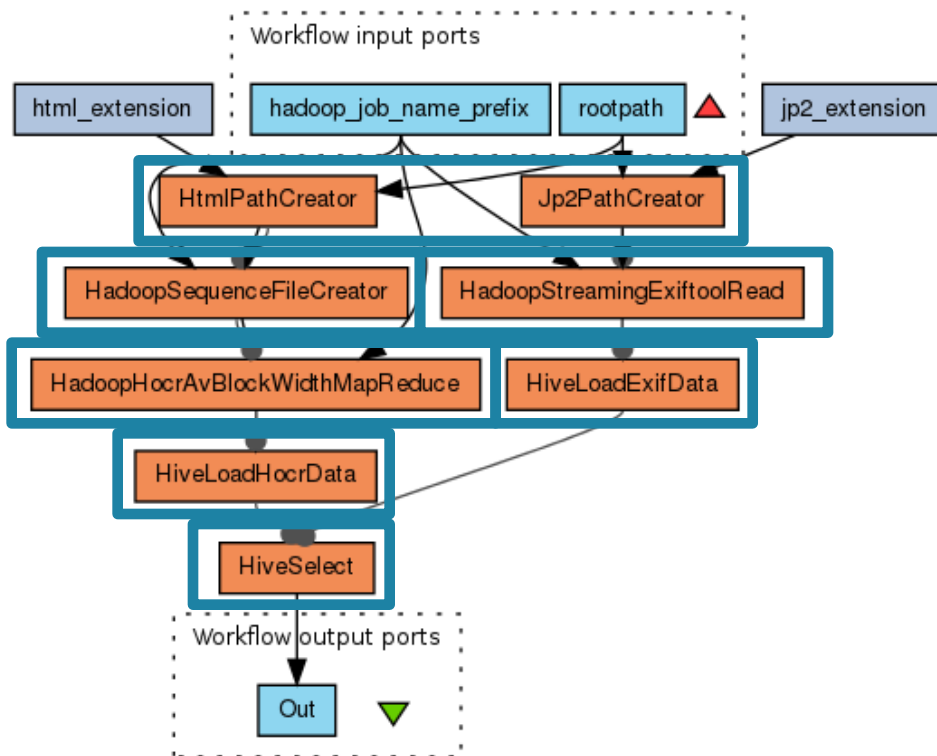
Bildbreite

Blockbreite



Annahme: „Signifikante“ Unterschiede zwischen durchschnittlicher Blockbreite und Bildbreite sind Hinweise auf möglichen Textverlust wegen Bild-Beschneidungsfehler.

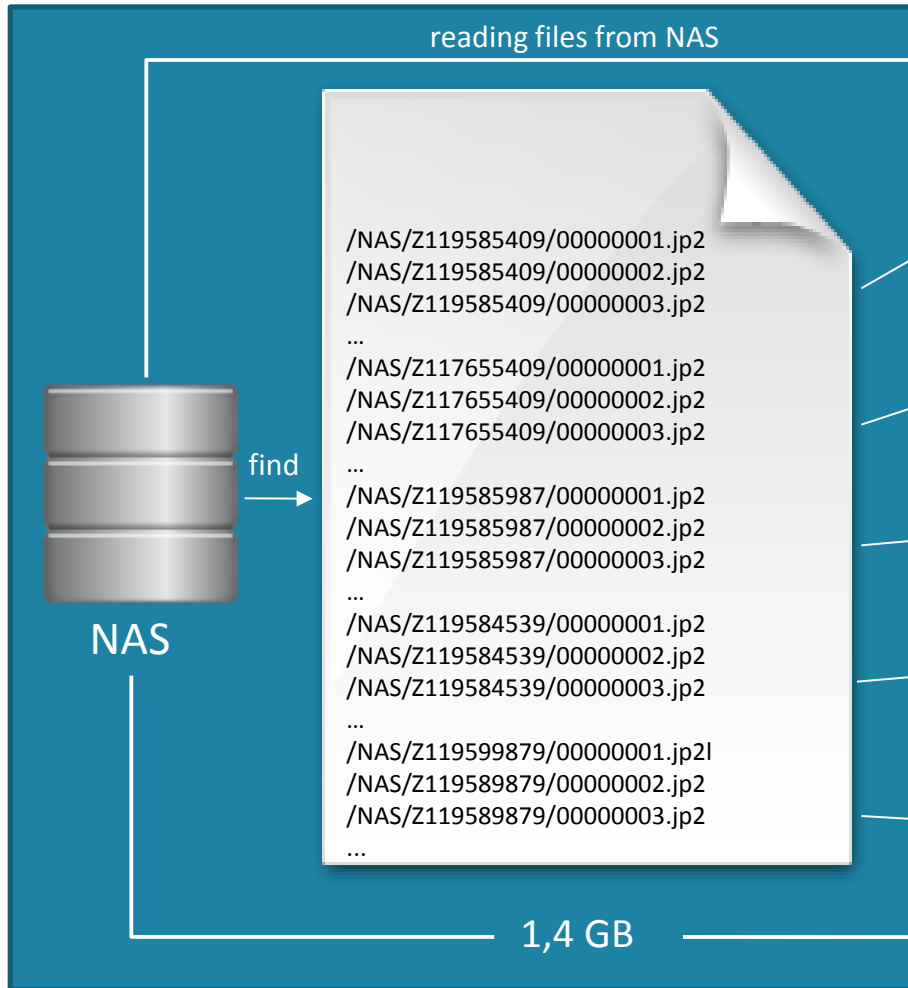
MapReduce in der Qualitätssicherung



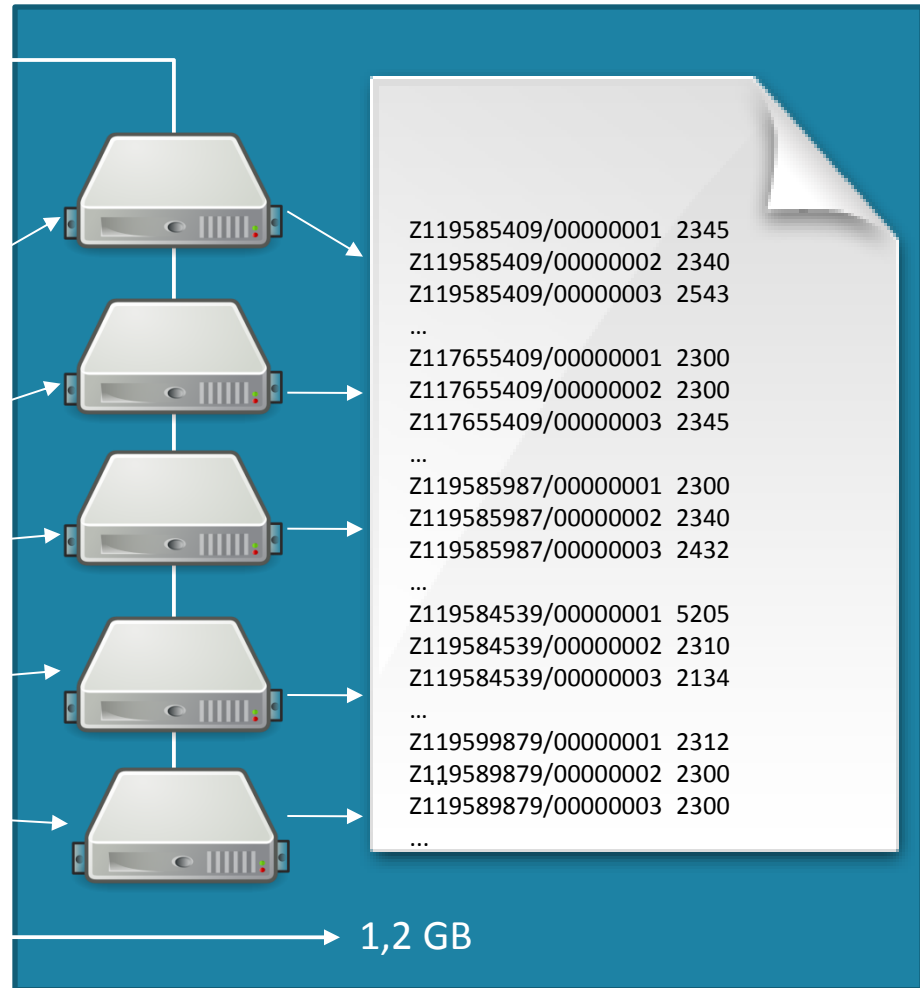
- Textdateien mit Pfaden zu Dateien erzeugen (JP2 & HTML)
- Bildmetadaten mit Exiftool lesen (Hadoop Streaming API)
- SequenceFile mit allen HTML Dateien erzeugen
- Durchschnittliche Blockbreite mit MapReduce berechnen
- Daten in Hive-Tabelle laden
- Test-Abfrage

Lesen der Bild-Metadaten

Jp2PathCreator



HadoopStreamingExiftoolRead



60.000 Bücher(24 Million Seiten): ~ 5 h

+

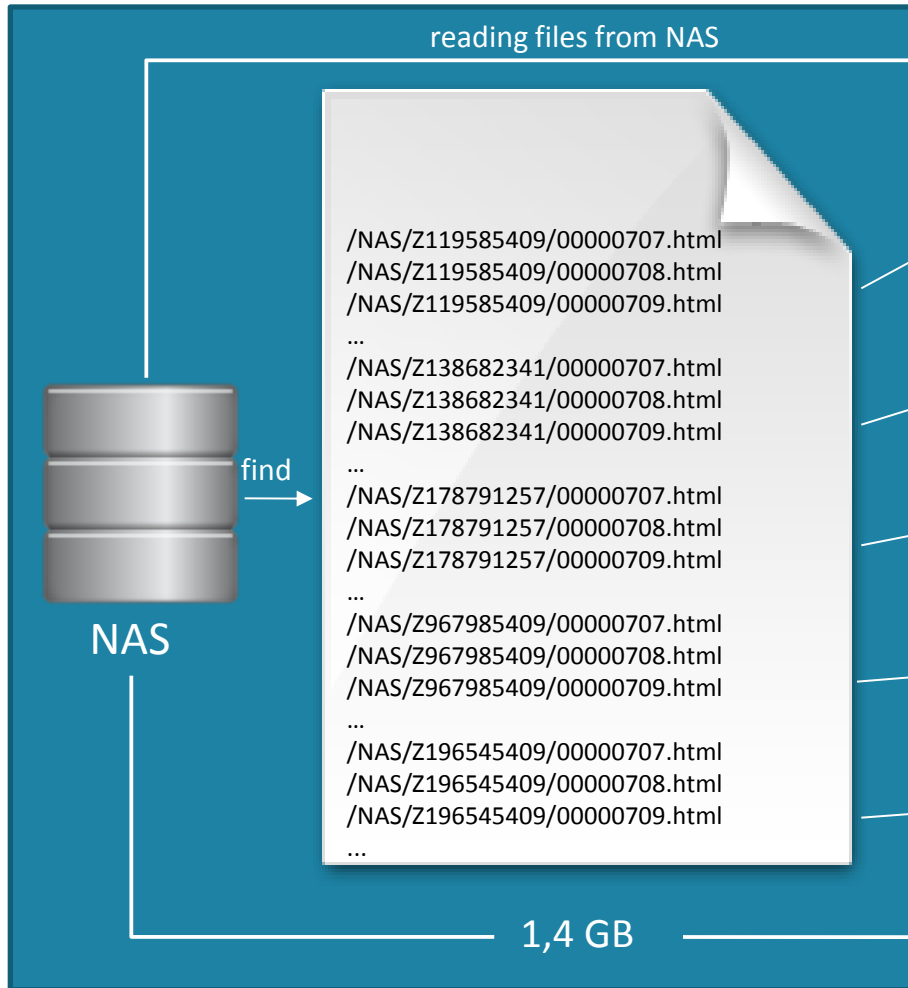
~ 38 h

=

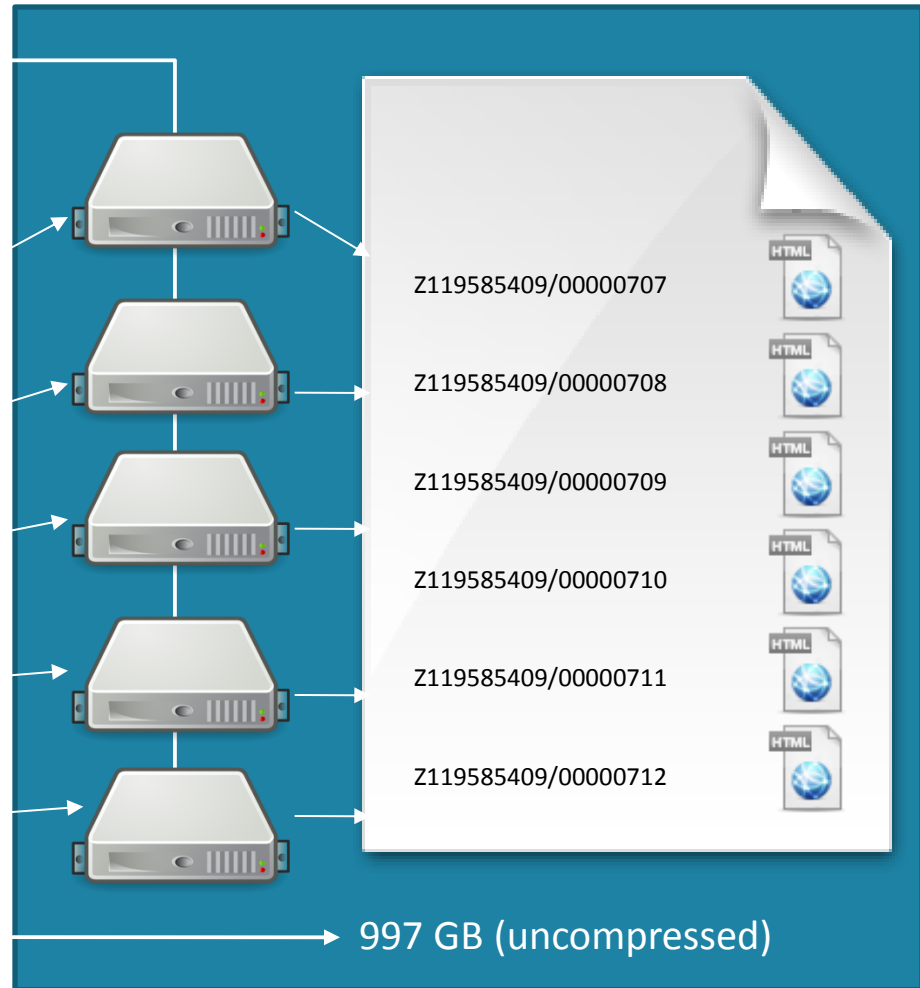
~ 43 h

Erzeugung eines SequenceFile

HtmlPathCreator



SequenceFileCreator



60.000 books (24 Million pages): ~ 5 h

+

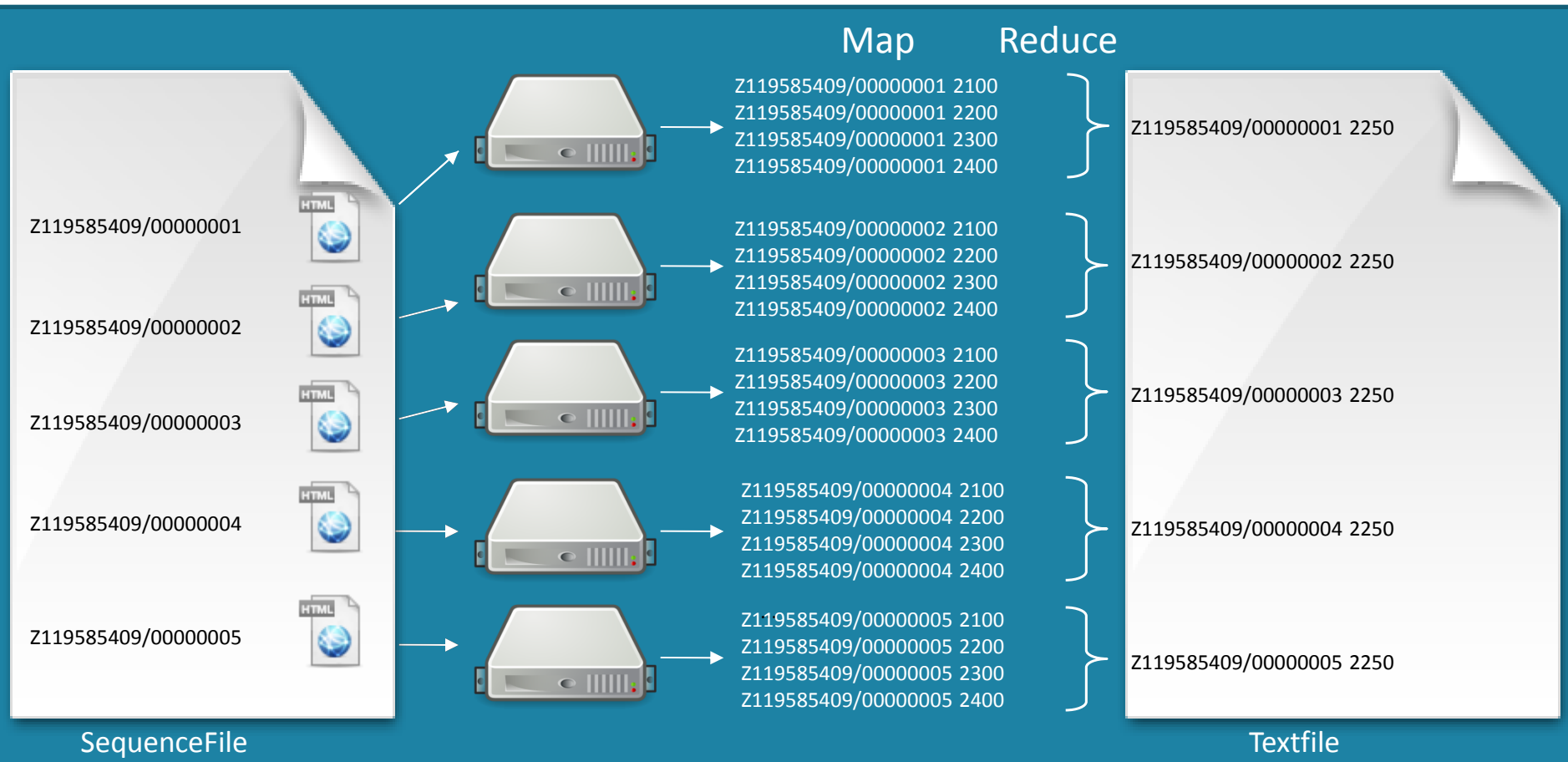
~ 24 h

=

~ 29 h

Berechnung der durchschnittlichen Blockbreite

HadoopAvBlockWidthMapReduce



60.000 books (24 Million pages): ~ 6 h

Laden der Daten (Hive)

HiveLoadExifData & HiveLoadHocrData

```
Z119585409/000000001 1870
Z119585409/000000002 2100
Z119585409/000000003 2015
Z119585409/000000004 1350
Z119585409/000000005 1700
```

```
CREATE TABLE htmlwidth
(hid STRING, hwidth INT)
```

htmlwidth

hid	hwidth
Z119585409/000000001	1870
Z119585409/000000002	2100
Z119585409/000000003	2015
Z119585409/000000004	1350
Z119585409/000000005	1700

```
Z119585409/000000001 2250
Z119585409/000000002 2150
Z119585409/000000003 2125
Z119585409/000000004 2125
Z119585409/000000005 2250
```

```
CREATE TABLE jp2width
(hid STRING, jwidth INT)
```

jp2width

jid	jwidth
Z119585409/000000001	2250
Z119585409/000000002	2150
Z119585409/000000003	2125
Z119585409/000000004	2125
Z119585409/000000005	2250

Abfragen (Hive)

HiveSelect

jp2width

jid	jwidth
Z119585409/000000001	2250
Z119585409/000000002	2150
Z119585409/000000003	2125
Z119585409/000000004	2125
Z119585409/000000005	2250

htmlwidth

hid	hwidth
Z119585409/000000001	1870
Z119585409/000000002	2100
Z119585409/000000003	2015
Z119585409/000000004	1350
Z119585409/000000005	1700

```
select jid, jwidth, hwidth from jp2width inner join htmlwidth on jid = hid
```

jid	jwidth	hwidth
Z119585409/000000001	2250	1870
Z119585409/000000002	2150	2100
Z119585409/000000003	2125	2015
Z119585409/000000004	2125	1350
Z119585409/000000005	2250	1700

- Für die digitalen Bibliothek bietet Apache Hadoop eine verlässliche Basis zum Aufbau einer kostengünstigen und skalierbaren IT-Infrastruktur
- Sorgfältige Auswahl der Werkzeuge aus dem Apache Ökosystem (weniger ist mehr!)
 - HBase, Hive, Pig, Oozie, etc.
- Mit neue Möglichkeiten kommen auch neue Fragestellungen
 - HDFS als Master- oder Bereitstellungsbereich?
 - Eigenen Cluster oder Infrastruktur mieten?
 - (Urheber-)Rechtliche Bedenken

Weiterführende Informationen

- Projekt-Website: www.scape-project.eu
- Github-Repository: www.github.com/openplanets
- Projekt-Wiki: www.wiki.opf-labs.org/display/SP/Home

Erwähnte SCAPE Werkzeuge

- SCAPE Plattform
 - <http://www.scape-project.eu/publication/an-architectural-overview-of-the-scape-preservation-platform>
- Jpylyzer – Jpeg2000 Validierung
 - <http://www.openplanetsfoundation.org/software/jpylyzer>
- Matchbox – Bildvergleich
 - <https://github.com/openplanets/scape/tree/master/pc-qa-matchbox>