



Integrating the Fedora based DOMS repository with Hadoop

Asger Askov Blekinge

State and University Library, Denmark

SCAPE Information Day

State and University Library, Denmark, June 25th 2014

Our Repositories

- Each **File** is stored in Bit Magasinet, our bit preservation storage system.
- Each **Record** is stored in DOMS and have have reference to the **File** in Bit Magasinet
- Can Hadoop be added to this setup?

Hadoop Data Locality

- **Rule 1:** The size of the Hadoop cluster should be independent of the size of the data storage
- The reading of data should happen from local disks. This prevents a central storage system from limiting the speed of the cluster
- With this restriction, the number of nodes in the cluster can keep growing
- Without, the cluster will reach a point where it will overload the central storage system.

Repositories (DOMS) and Hadoop

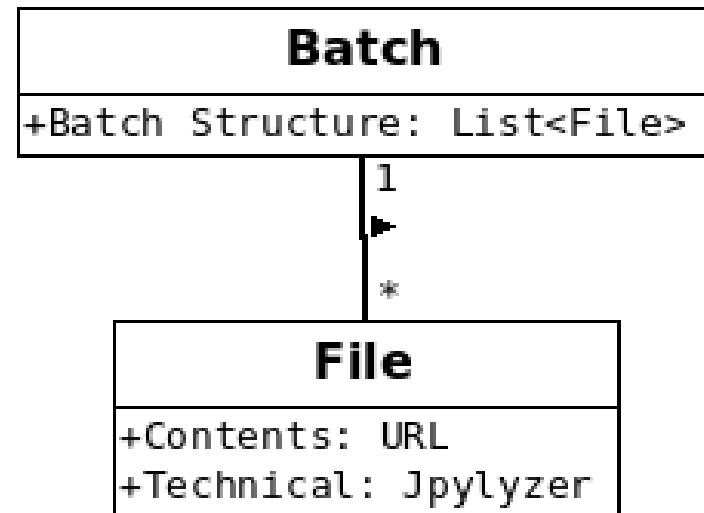
- Repositories, especially Fedora 3.x, are single headed. You cannot add more machines to the repository to increase the performance.
- If Hadoop accesses the repository directly, it will be limited to the speed of the repository.

- Hadoop provides its own bit archive system in the form of HDFS, which is integrated with the cluster
- We do not use this. We have built our own system instead, **Bit Magasinet**
- We can handle many more files because we use magnetic tapes, rather than disks.
- But: it requires us to request a number of files, which will then be made available for Hadoop.

- Hadoop does not play nice with DOMS or Bit magasin
- This state of affairs is not acceptable to us.
- Besides, it is a nice challenge ;)

How we do it in the Newspaper digitisation project

- Files are stored in Bit Magasinet
- One Batch Object
 - Batch object have list of files
- One Record per File



How we do it in the Newspaper digitisation project

- A Hadoop map/reduce job is split into two steps
 - Map, where the work on each “record” is performed.
 - Reduce, where the results are collated
- In the Map step, we run the tool on the file.
 - We have a lot of Map workers.
- In the Reduce step, we store the results in the repository.
 - We have only a few Reduce workers.

How we do it in the Newspaper digitisation project

- Retrieve the list of files from DOMS
- Request these files from Bit Magasinet
- Start Hadoop job on files
 - Map: Run Jpylyzer on each file (Many worker nodes)
 - Reduce: Store the results back in DOMS (Few worker nodes)
- This way, the actual work on the records is not connected to DOMS, and we can scale the cluster

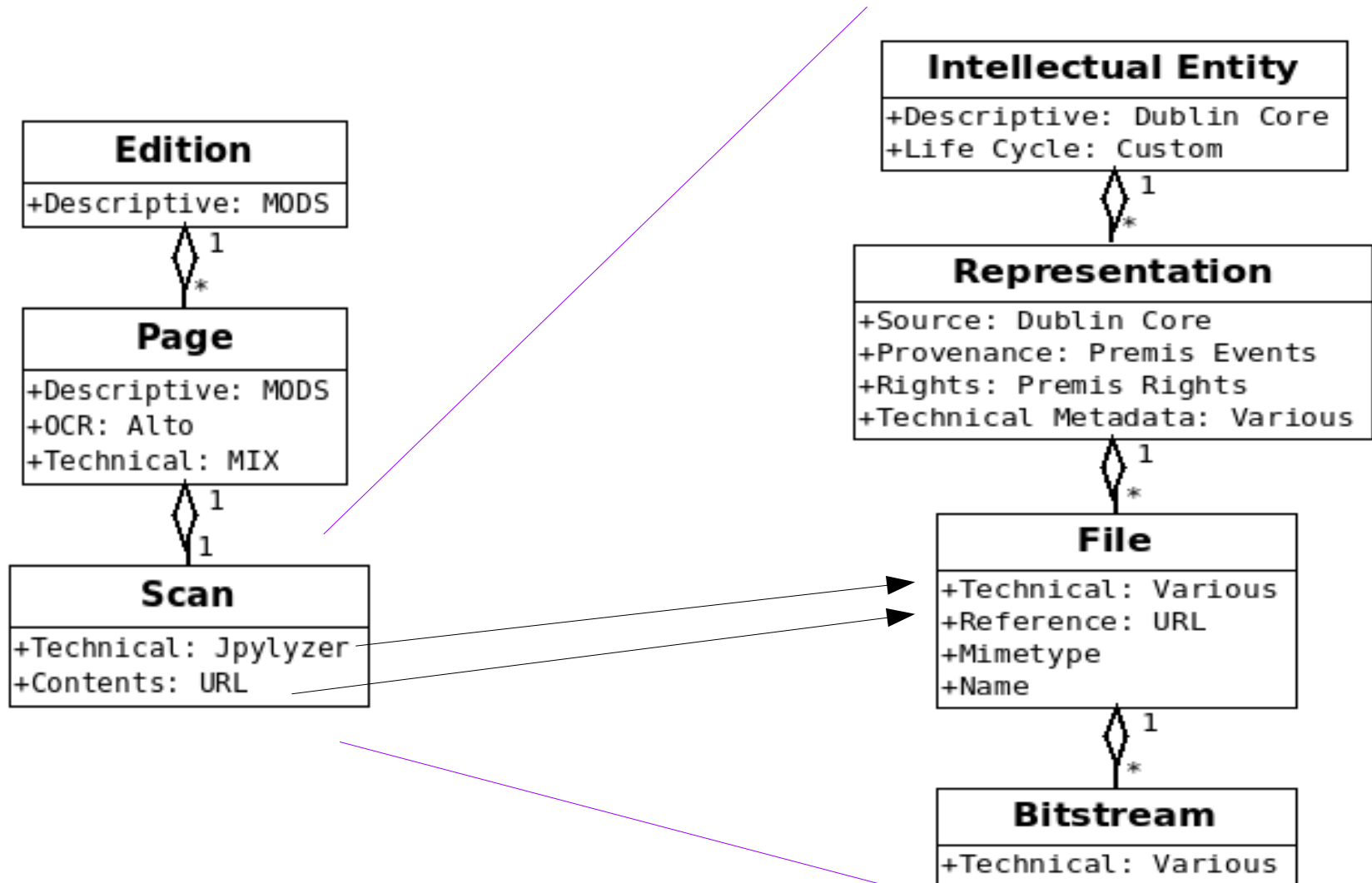
How we do it in SCAPE

- **Staging:** Retrieve the records from DOMS to an archive file
- **Hadooping:** Hadoop reads the records, work and writes new, updated records to the archive file
- **Loading:** Store the updated records in DOMS

Step 1 – Retrieve records

- SCAPE has devised a repository agnostic object format based on METS
 - github.com/openplanets/scape-platform-datamodel
- SCAPE has designed a generic repository REST interface
 - github.com/openplanets/scape-apis
- SB has implemented the SCAPE Repo API for DOMS
 - github.com/statsbiblioteket/scape-doms-data-connector
- We have implemented a client for the SCAPE Repo API
 - github.com/statsbiblioteket/scape-stager-loader

SCAPE Datamodel mapping



SCAPE Repository API

```
<mets:mets ID="scape-entity:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af" OBJID="scape-entity:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af" PROFILE="scape">
  <mets:metsHdr RECORDSTATUS="NEW"/>
  <mets:dmdSec ID="DMD-8c72c14d-475a-49a2-9f43-321732c4e7a2">
    <mets:mdWrap MDTYPE="OTHER">
      <mets:xmlData/>
    </mets:mdWrap>
  </mets:dmdSec>
  <mets:dmdSec ID="DMD-747421f1-fc0d-4c1d-896c-9087d43b5e10">
    <mets:mdWrap MDTYPE="OTHER">
      <mets:xmlData>
        <scape:versionMD version-number="1"/>
      </mets:xmlData>
    </mets:mdWrap>
  </mets:dmdSec>
  <mets:amdSec>
    <mets:techMD ID="TMD-scape-representation:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-SCAPE_REPRESENTATION_TECHNICAL">
      <mets:mdWrap MDTYPE="OTHER">
        <mets:xmlData/>
      </mets:mdWrap>
    </mets:techMD>
    <mets:techMD ID="TMD-scape-representation:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-scape-file:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-JPYLYZER">
      <mets:mdWrap MDTYPE="OTHER">
        <mets:xmlData>
          <!--
            +Technical: Various
            +Reference: URL
            +Mimetype
            +Name
          -->
        </mets:mdWrap>
      </mets:techMD>
    </mets:amdSec>
    <mets:fileSec>
      <mets:fileGrp>
        <mets:file ID="scape-file:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af" SEQ="0" ADMID="TMD-scape-representation:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-scape-file:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-JPYLYZER" MIMETYPE="image/jp2">
          <mets:Flocat xlink:href="http://bitfinder.statsbiblioteket.dk/newspapers/B400022028241-RT1_400022028241-1_1795-06-01_adresseavisen1759-1795-06-01-0006.jp2"
xlink:title="B400022028241-RT1_400022028241-1_1795-06-01_adresseavisen1759-1795-06-01-0006.jp2" LOCTYPE="URL"/>
        </mets:file>
      </mets:fileGrp>
    </mets:fileSec>
    <mets:structMap>
      <mets:div TYPE="Intellectual entity">
        <mets:div ID="scape-representation:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af" ADMID="TMD-scape-representation:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af-SCAPE_REPRESENTATION_TECHNICAL" TYPE="Representation" xlink:label="page-image-adresseavisen1759-1795-06-01-0007A">
          <mets:fptr FILEID="scape-file:uuid:1c0194a3-c5af-4b40-b140-5ac64cfa43af"/>
        </mets:div>
      </mets:div>
    </mets:structMap>
  </mets:mets>
```

File
+Technical: Various
+Reference: URL
+Mimetype
+Name

- Get Entity – GET /entity/<entityID>
- Update Entity – PUT /entity/<entityID>
- Create Entity – POST /entity/<entityID>
- And many more

- Checkout

```
java -jar scape-stager-loader.jar
```

```
--id_file=identifierFile.txt
```

```
--checkoutSequenceFile="test.seqfile"
```

```
checkout
```

- Commit

```
java -jar scape-stager-loader.jar
```

```
--commitSequenceFile="test.seqfile"
```

```
commit
```

Step 2: Hadoop reads and updates records

- The Hadoop job is started with the sequence file as input
- For each record in the sequence file
 - Read the record
 - Do work
 - Update the record in the sequence file with the result of the work

Step 3: Store the updated records in DOMS

- The hadoop job produces a sequence file
- For each record in the sequence file:
 - Read the record into memory
 - Any changed fields are updated in the corresponding DOMS objects

This way, the actual work on the records is not connected to DOMS, and we can scale the cluster independently from the repository