

# Bezpečnostní seminář

## **Big data**

TERADATA. ASTER

## Analytický Ecosystem

*RDBMS a MapReduce*

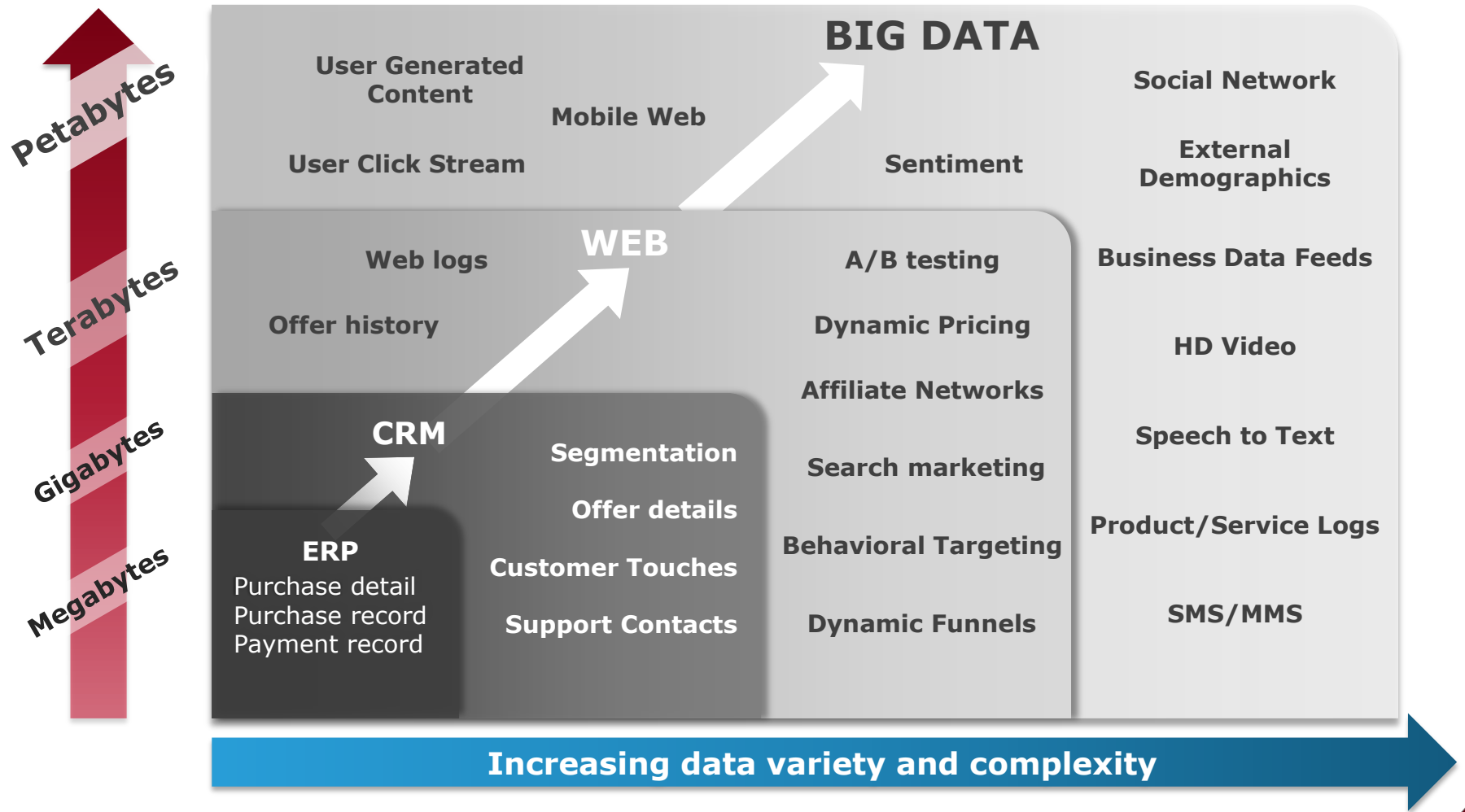
Luboš Musil  
Enterprise Architect

# Cíle prezentace

## **Představit koncepci Analytického Ecosystému**

- *Masivně paralelní databáze a Map reduce*
- *Discovery platforma*
- *Integrace systémů v Ecosystému*

# Big Data: Od transakcí k iteracím



# Klasické BI vs. Analytika v oblasti Big Data



**Business** determines what questions to ask

**Classic BI Method**  
Structured & Repeatable Analysis



**IT** structures the data to answer those questions

***"Capture only what's needed"***



**IT** delivers a platform for storing, refining, and analyzing **all data sources**

***"Capture in case it's needed"***

**Big Data Analytics**  
Multi-structured & Iterative Analysis



**Business** explores data for questions worth answering

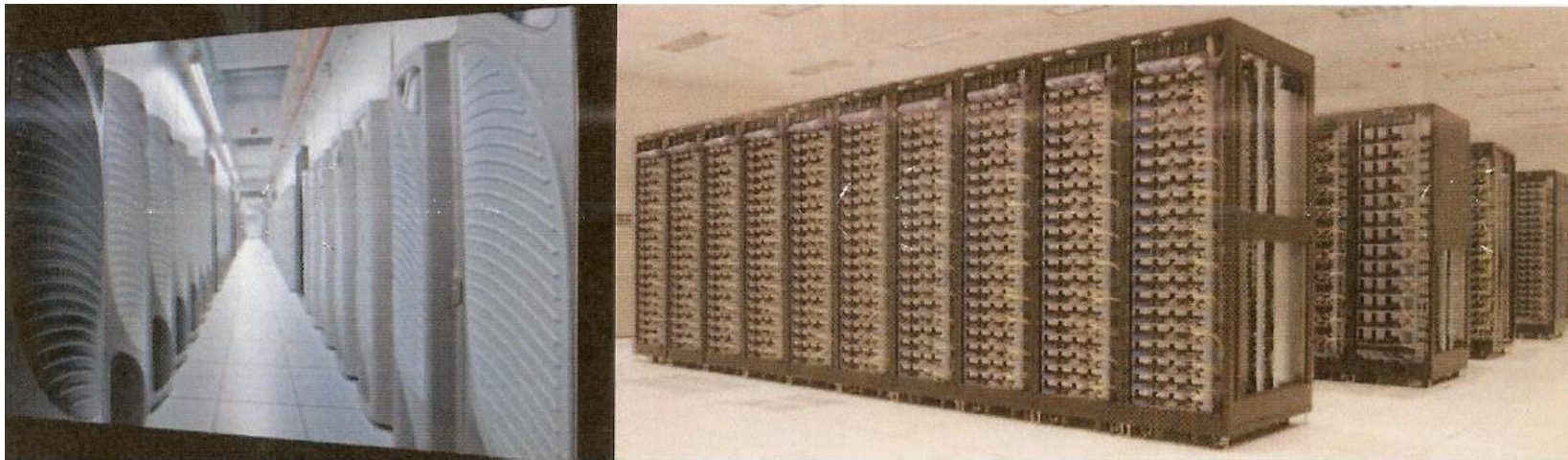
# Mapreduce reference: Yahoo Hadoop clusters

- Více jak 100,000 CPUs v 25,000+ počítačích s běžícím Hadoop
- Největší Hadoop cluster: 4000 nodů (boxů)
- 2\*4 jádrové CPU boxy s 4\*1TB disk a 16GB RAM
- Užito pro Ad Systems a Web Search
- Více jak 40% z Hadoop jobů Yahoo jsou Pig joby



# MPP RDBMS reference: eBay

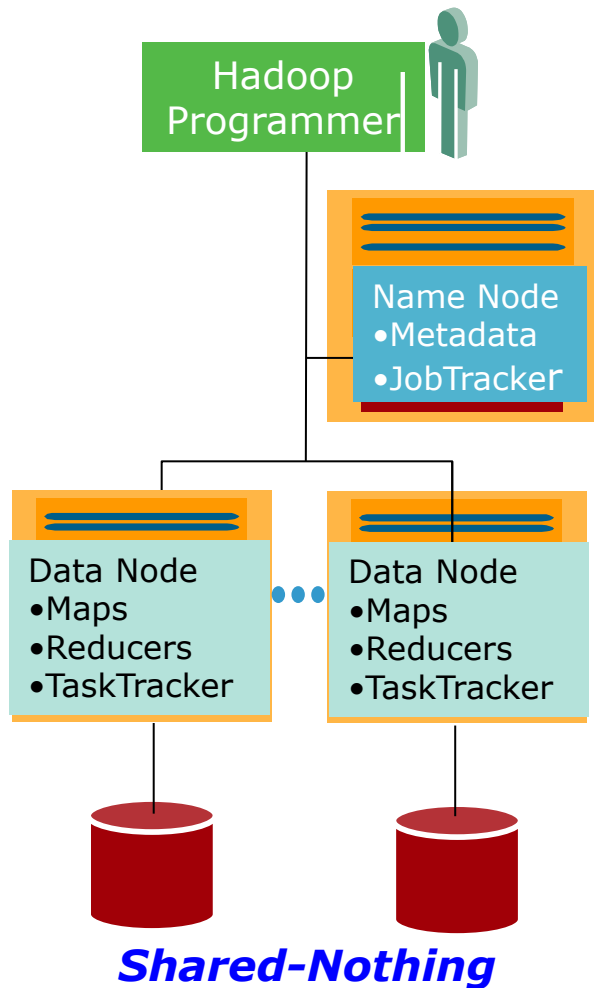
- Analytický systém pro analýzu Web log o velikosti 42 PB
- Využití MPP RDBMS Teradata



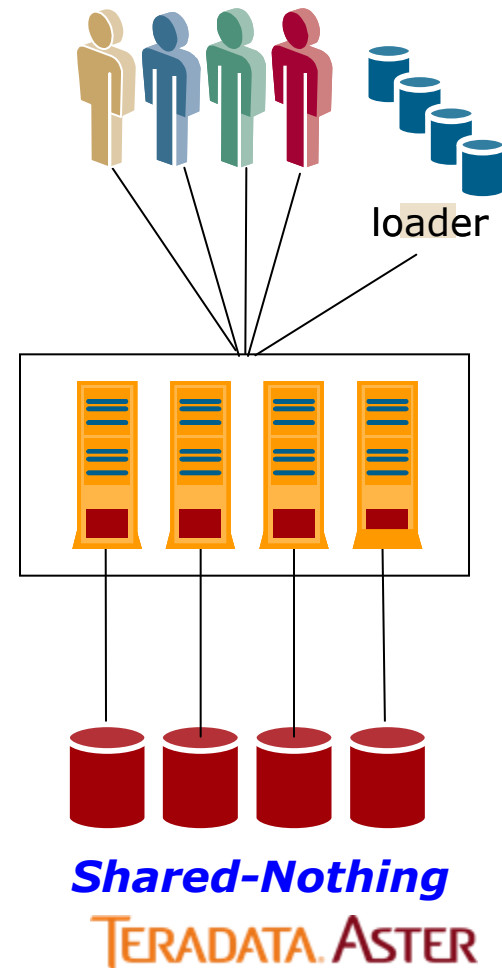
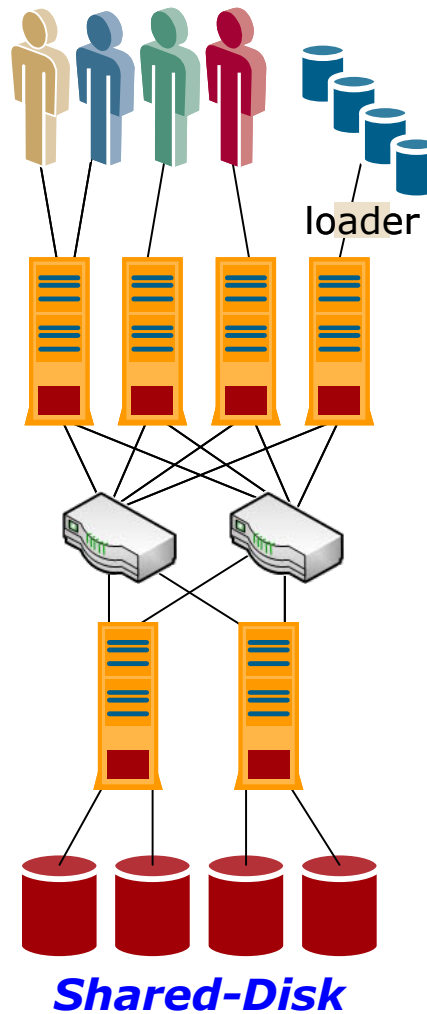
- **Similar structure of net-connected nodes**
- **Different approaches to processing and storage**
- **Complementary functionality enables a wider scope of analysis**

# Hadoop MapReduce vs. MPP RDBMS

## Hadoop MapReduce



## MPP RDBMS



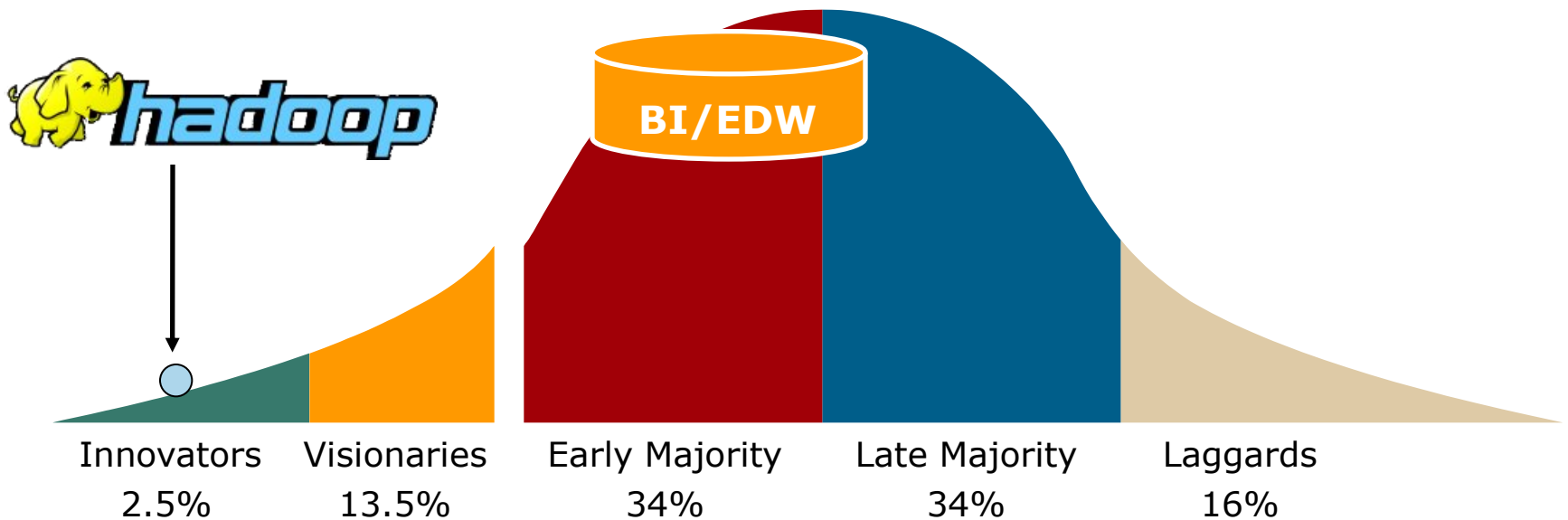
TERADATA.ASTER

# Srovnání základních vlastností

| Map Reduce / Hadoop                    | Data Warehouse                                                           |
|----------------------------------------|--------------------------------------------------------------------------|
| File system foundation                 | RDBMS foundation                                                         |
| Scale out to 1000s of nodes            | Scale out to 2048 nodes                                                  |
| Open source prices                     | Commercial vendor pricing                                                |
| Embryonic open source                  | Mature proprietary code base                                             |
| Numerous programming languages         | Java, C++, C, and a few others                                           |
| Batch processing                       | Batch, interactive, real time processing                                 |
| Programmer optimizes each job          | Cost based query optimization                                            |
| Single large fact table                | Integrated subject areas                                                 |
| Simple star schema like queries        | Unlimited query flexibility/composition                                  |
| First attempts at BI tools integration | Dozens of robust BI and ETL tools                                        |
| Complex multi-step processing          | Single step SQL process                                                  |
| Schema-less                            | External schema                                                          |
| Extensive text parsing functions       | Basic text parsing                                                       |
| Java programmer reporting              | End user interactive reporting                                           |
| 1-10 concurrent jobs per cluster       | 10s to 1000s of concurrent queries                                       |
| 100s of skilled programmers            | 100s of thousands of BI/EDW programmers                                  |
| Limited or no system management        | Extensive system management                                              |
| No security; LDAP only                 | Encryption, role based access control, privacy tools, single sign-on     |
| Fault tolerant data blocks             | Failover, fast recovery, checkpoints, redo logs, hot standby nodes, etc. |
| < 1% market adoption                   | 50-60% market adoption                                                   |

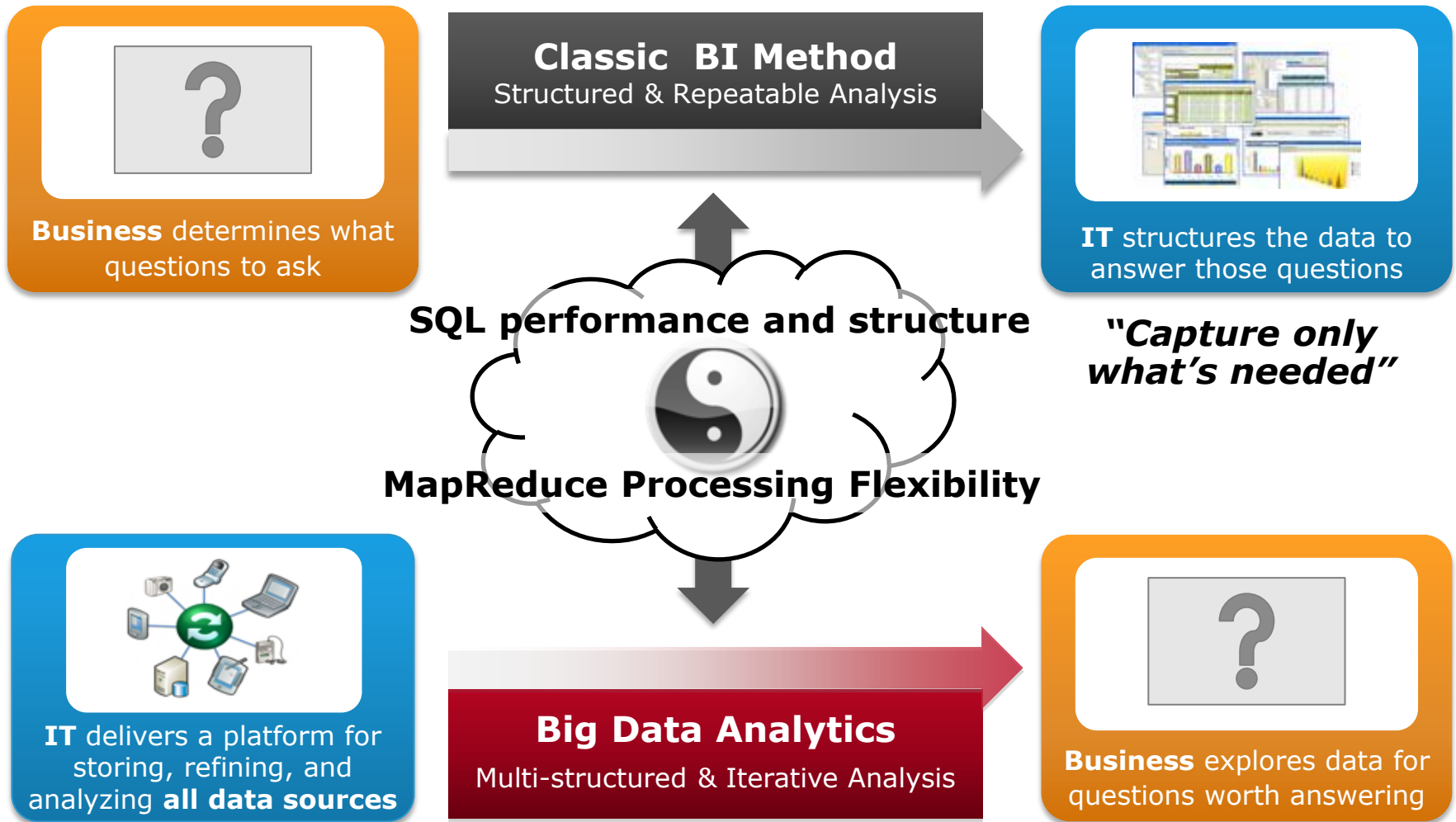
# Evoluce užití Map Reduce

- *Technologie využívající MapReduce framework se od roku 2011 dynamicky rozvíjí*
- *SMP RDBMS jsou nahrazovány MPP RDBMS*
- *Oba světy se začínají integrovat.....?*



# Unifikovaná Big Data architektura

Přemostění klasického BI a Big Data Analytického světa



*"Capture in case it's needed"*

# Proč Unifikovaná Big Data Architektura?

Zpřístupnit všem uživatelům analyzy všech typů dat společnosti



Java, C/C++, Pig, Python, R, SAS, SQL, Excel, BI, Visualization, etc.

**Discover and Explore**

**Reporting and Execution  
in the Enterprise**

**Capture, Store and Refine**

Audio/  
Video

Images

Docs

Text

Web &  
Social

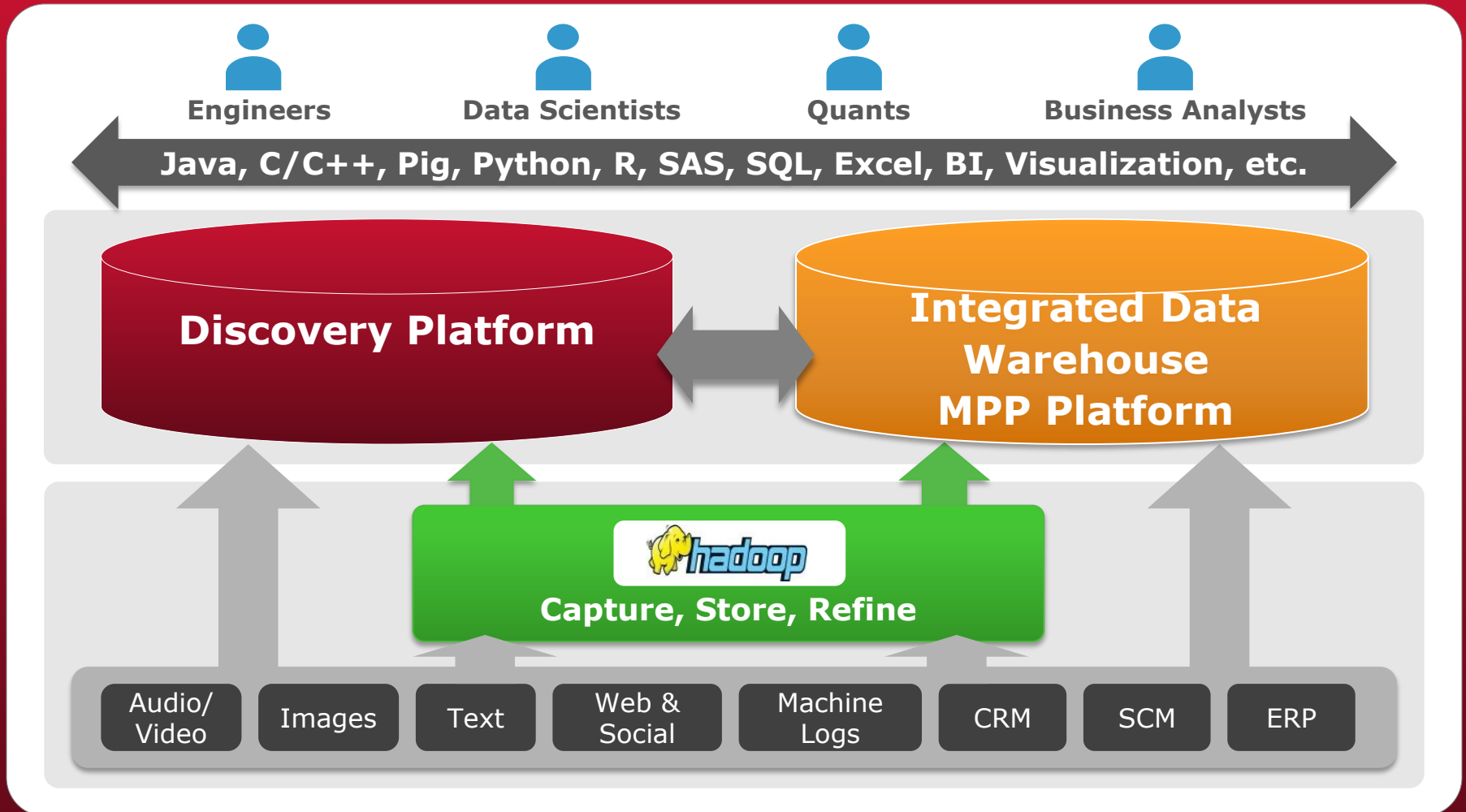
Machine  
Logs

CRM

SCM

ERP

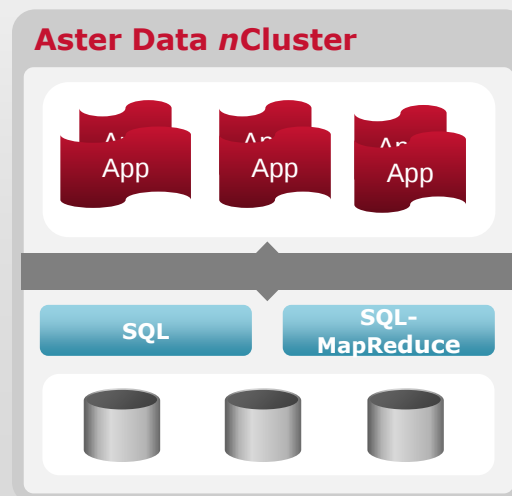
# Unifikovaná Big Data Architektura ve společnosti



# Discovery platform – SQL MapReduce

## Patentovaný Framework pro pokročilé analýzy, které je obtížné dělat v MPP RDBMS pomocí SQL

- Couples SQL (relational) with MapReduce (SQL-MapReduce) providing a new framework for rich analytics on diverse data (non-relational and relational).
- User code is installed in the cluster, and then it's invoked on database data from SQL. Execution is automatically parallelized across the cluster.
- Includes library of pre-packaged Analytic Modules (50+ currently) to speed analytics development (e.g. time-series, complex pattern/path, affinity, graph, data transformation, text, statistical...)
- Leverage existing investments in BI, ETL tools & resources
- Complete support for ANSI-standard SQL and MapReduce
- Supports custom analytics written in a variety of languages
- SQL-Hadoop integration via SQL-H™



# Discovery platform Aplikační portfolio

Aster data: Některé z 50+ analytických aplikací

## Path Analysis

Discover patterns in rows of sequential data

## Text Analysis

Derive patterns and extract features in textual data

## Statistical Analysis

High-performance processing of common statistical calculations

## Segmentation

Discover natural groupings of data points

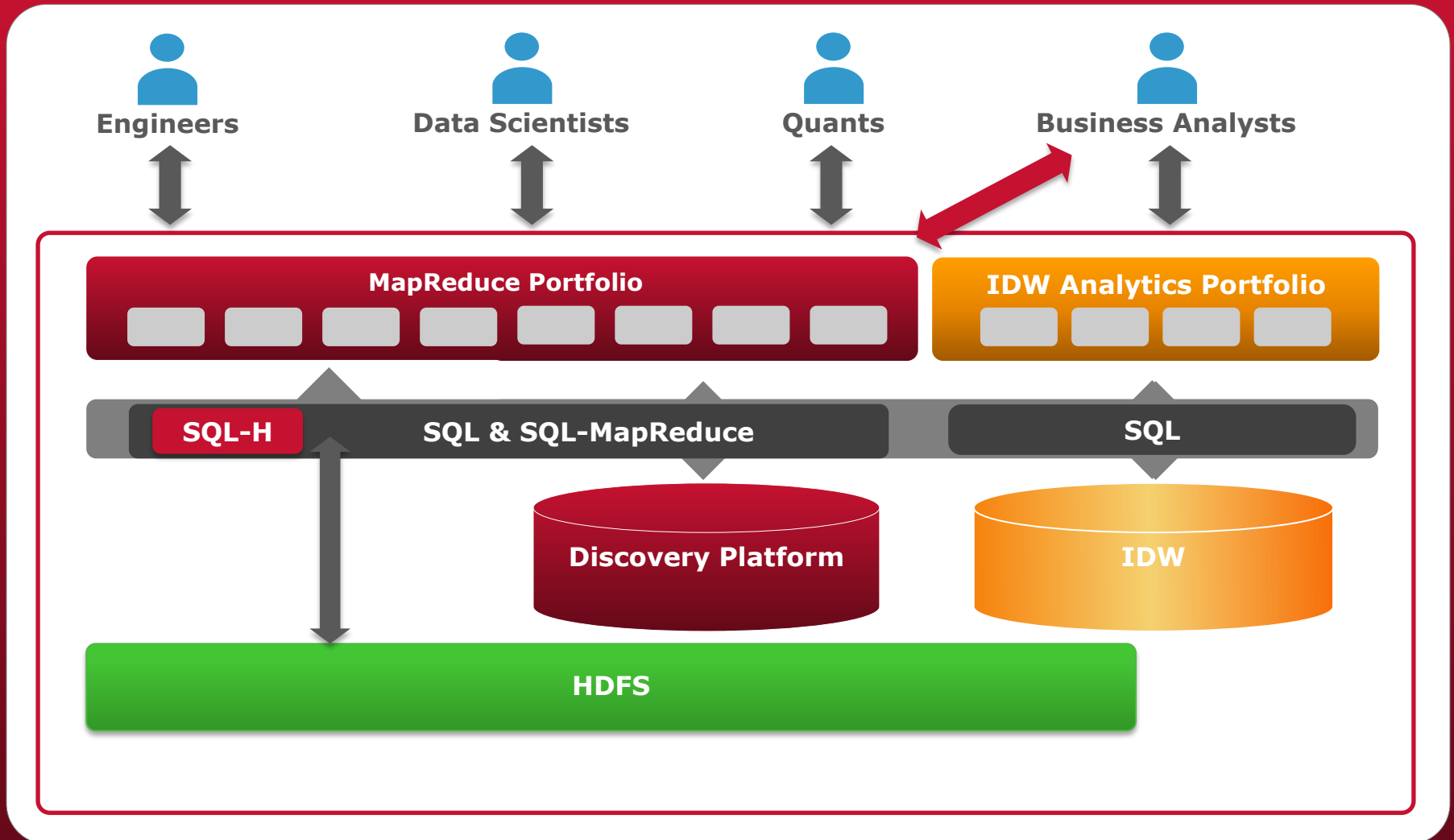
## Marketing Analytics

Analyze customer interactions to optimize marketing decisions

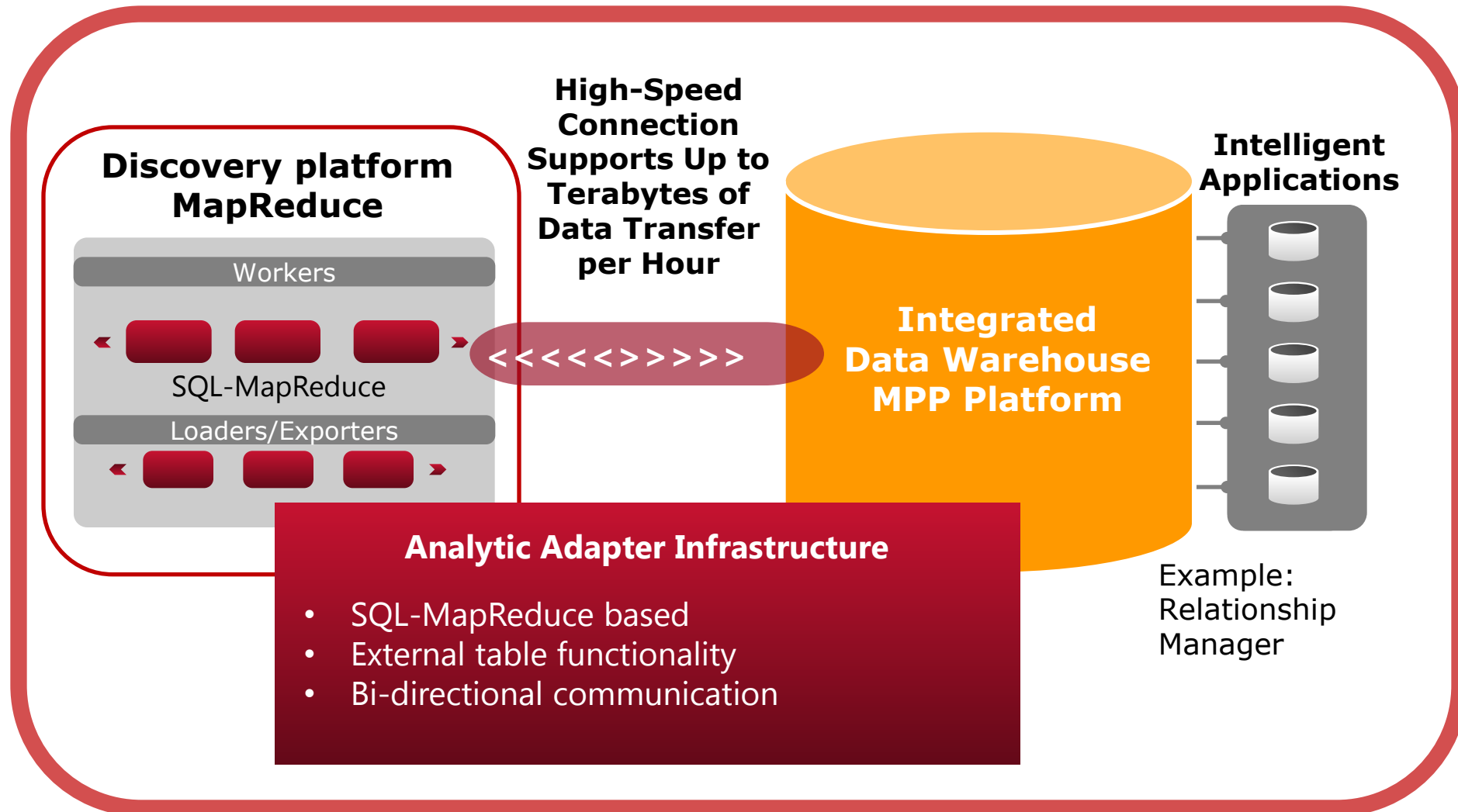
## Data Transformation

Transform data for more advanced analysis

# Discovery platform - konektor do hadoop

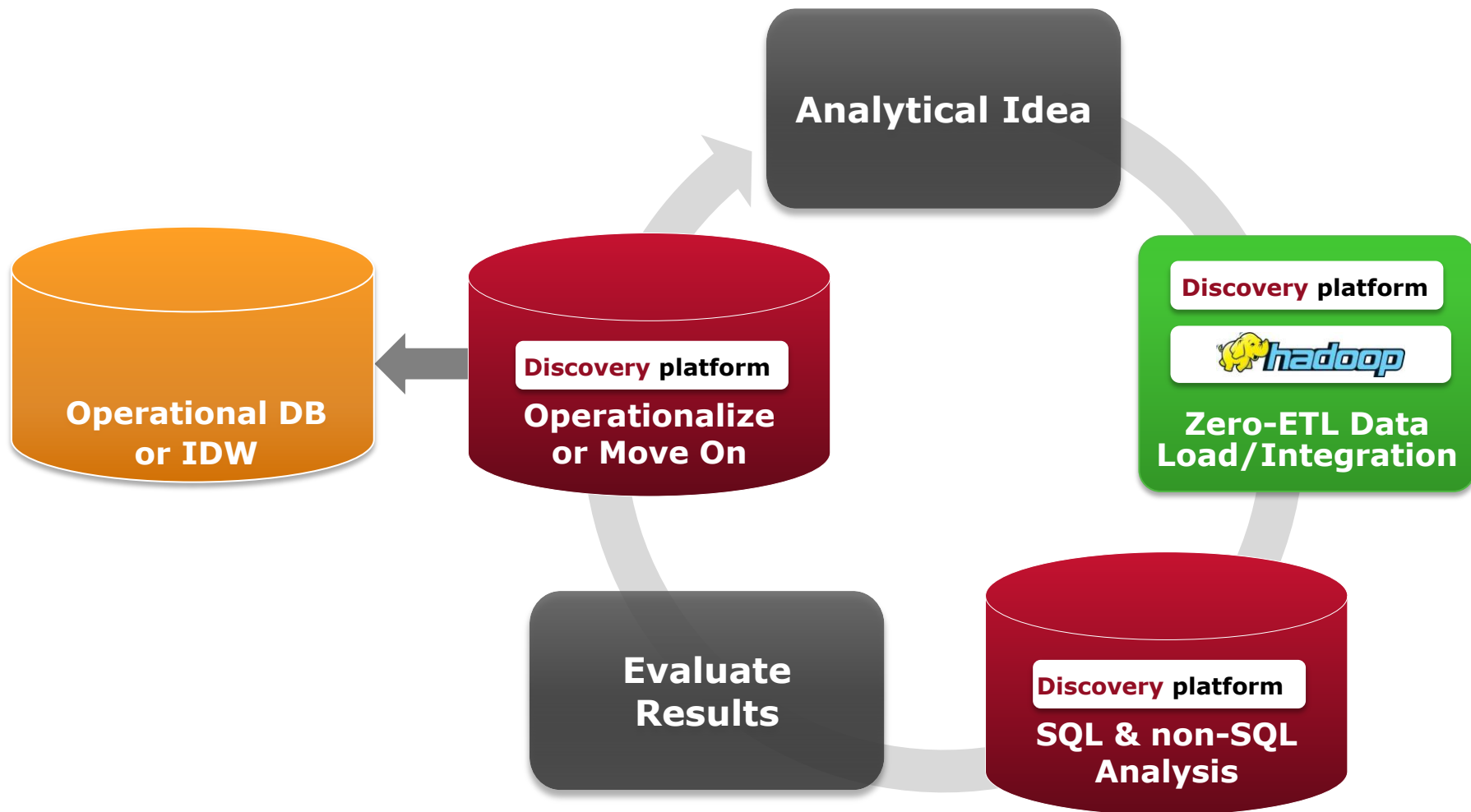


# Discovery platform - konektor do IDW



# Discovery cyklus

Nejefektivnější cesta jak získat business hodnotu z Big Dat



# Big Data analýzy – Technické rozdíly

*Různé datové typy potřebují různá schemata*

## Data that uses a **stable schema (structured)**

- Data from packaged business processes with well-defined & known attributes (e.g., ERP data, Inventory Records, Supply Chain records, ...)

## Data that has an **evolving schema (semi-structured)**

- Data generated by machine processes; known but changing set of attributes (e.g., Web logs, CDRs, Sensor logs, JSON, Social profiles, Twitter feeds, ...)

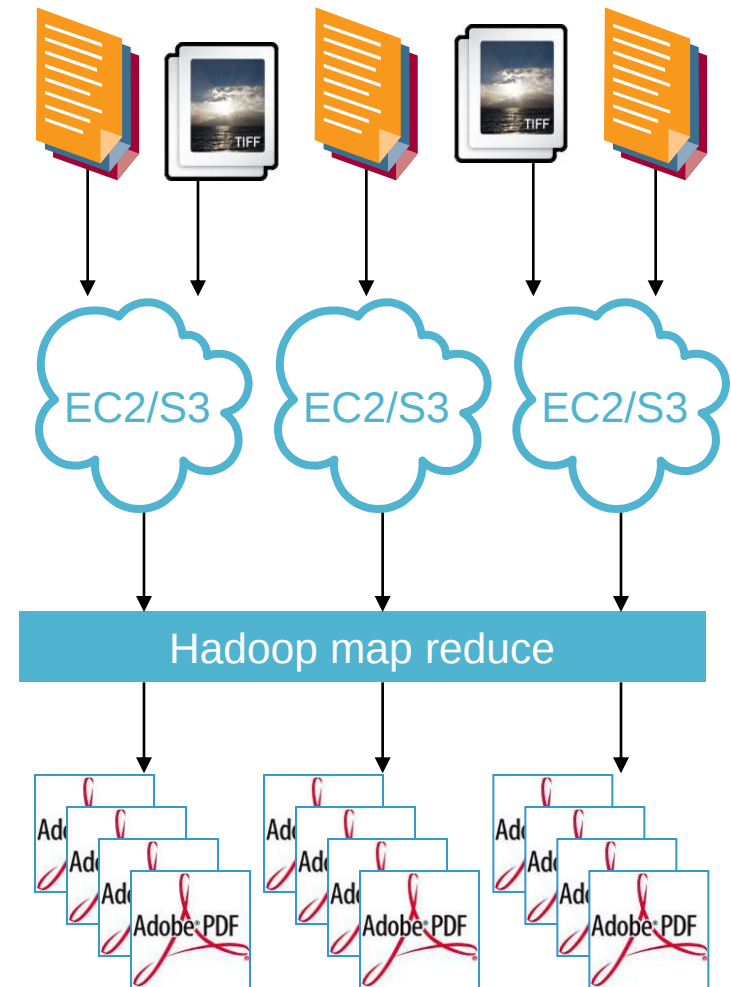
## Data that has a format, but **no schema (unstructured)**

- Data captured by machines with well-defined format, but no semantics (e.g., images, videos, web pages, PDF documents, ...)
- Semantics can be extracted from raw data by interpreting the format and pulling out required data (e.g., shapes from video, face recognition in images, logo detection, ...)
- Sometimes format data is accompanied by meta-data that can have (Stable Schema or Evolving Schema) – that needs to be classified and treated separately

# Příklad Formátu „No Schema“

## Image processing

- Millions of files or objects
- Find key object and transform it
  - Convert BMPs to JPGs
  - Convert DOCs to PDFs
  - Change NYC to New York City
  - Etc.
- ETL but data stays on node
- New York Times
  - Convert 11M articles to PDFs
  - Convert TIFFs in clouds
  - Use Amazon EC2 and S3 clouds
  - 4TB of articles → 1.5TB of PDFs



# Analytický workload

Unifikovaná Big Data Architektura musí podporovat daný workload optimálně

## Low cost storage and retention

- Retention of raw data in manner that can provide low TCO per terabyte storage costs
- Access in deep storage still required but not at same speeds as in a front line system

## Loading and refining

- **Load:** bring data into the system from the source system
- **Pre-processing / prep/ cleansing / constraint validation:** prepare data for downstream processing – e.g., fetch dimension data, record new incoming batch, archive old window batch, etc.
- **Transformations:** Convert one structure of data into another structure. This may require going from 3NF in relational to star/snowflake schema in relational, or going from text to relational, or going from relational to graph – I.e., structural transformations

## Reporting

- This is querying of what happened, where did it happen, how much happened, who did it

## Analytics (user-driven, interactive, ad-hoc)

- Relationship modeling that can be done via declarative SQL (e.g., scoring, basic stats)
- Relationship modeling done via procedural MR (E.g., model building, time series)

# Kdy použít jakou platformu?

Nejvodnější přístup je dle charakteru požadavku a dle datových typů

|                      | Low Cost<br>Storage &<br>Retention                                                                                                         | Loading and Refining                    |                 | Reporting | Analytics<br>(User-driven,<br>interactive) |
|----------------------|--------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------|-----------------|-----------|--------------------------------------------|
|                      |                                                                                                                                            | Data Pre-Processing,<br>Prep, Cleansing | Transformations |           |                                            |
| Stable<br>Schema     | <b>Financial analysis, ad-Hoc/OLAP</b><br><b>Enterprise-wide BI and Reporting</b><br><b>Spatial/Temporal</b><br><b>Active Execution</b>    |                                         |                 |           |                                            |
| Evolving<br>Schema   | <b>Interactive data discovery</b><br><b>Web clickstream, social feeds</b><br><b>Set-top box analysis</b><br><b>CDRs, Sensor logs, JSON</b> |                                         |                 |           |                                            |
| Format,<br>No Schema | <b>Image processing</b><br><b>Audio/video storage and refining</b><br><b>Storage and batch transformations</b>                             |                                         |                 |           |                                            |

\* Aster – example of Discovery platform

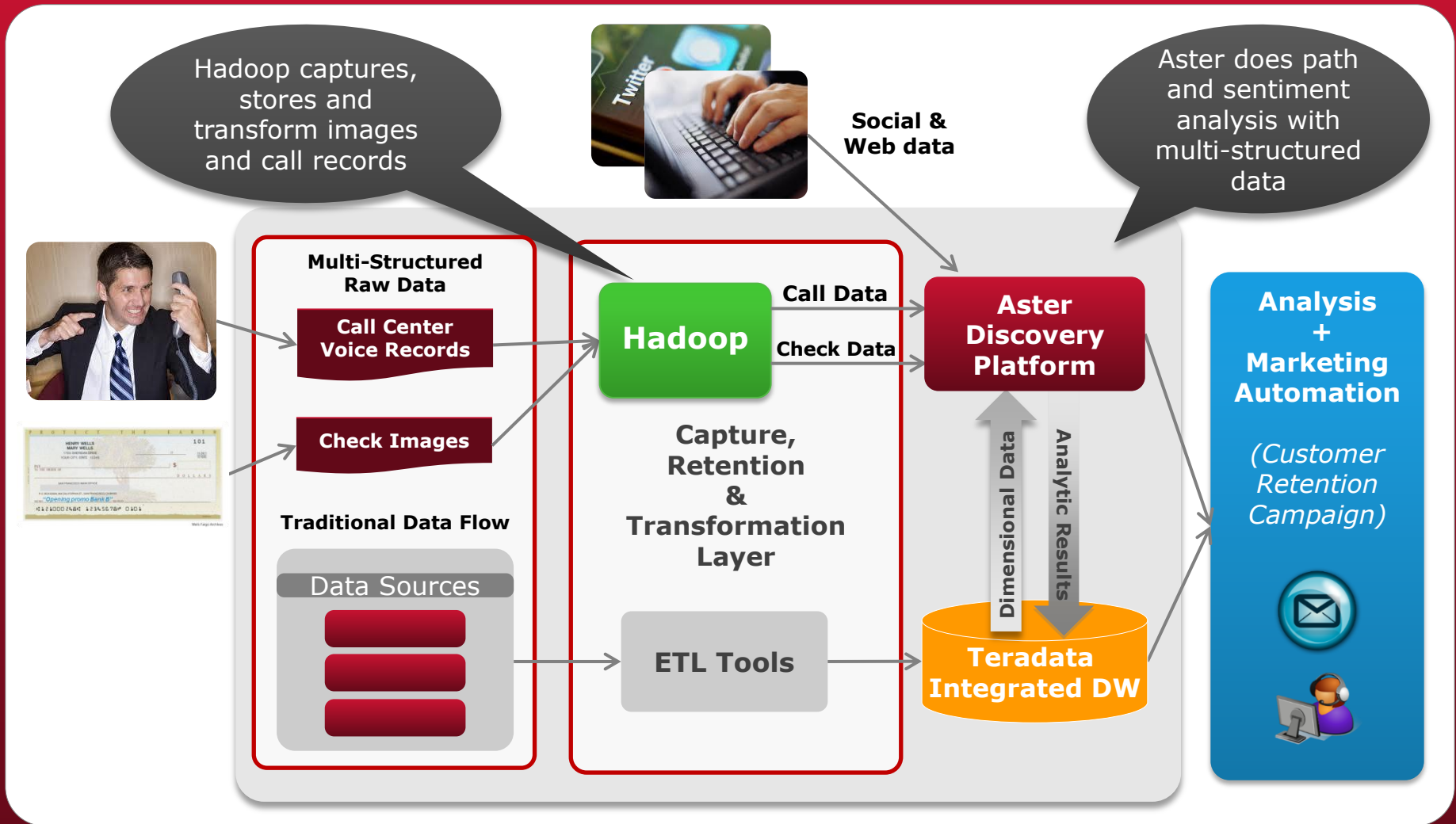
# Kdy použít jakou platformu?

Nejvodnější přístup je dle charakteru požadavku a dle datových typů

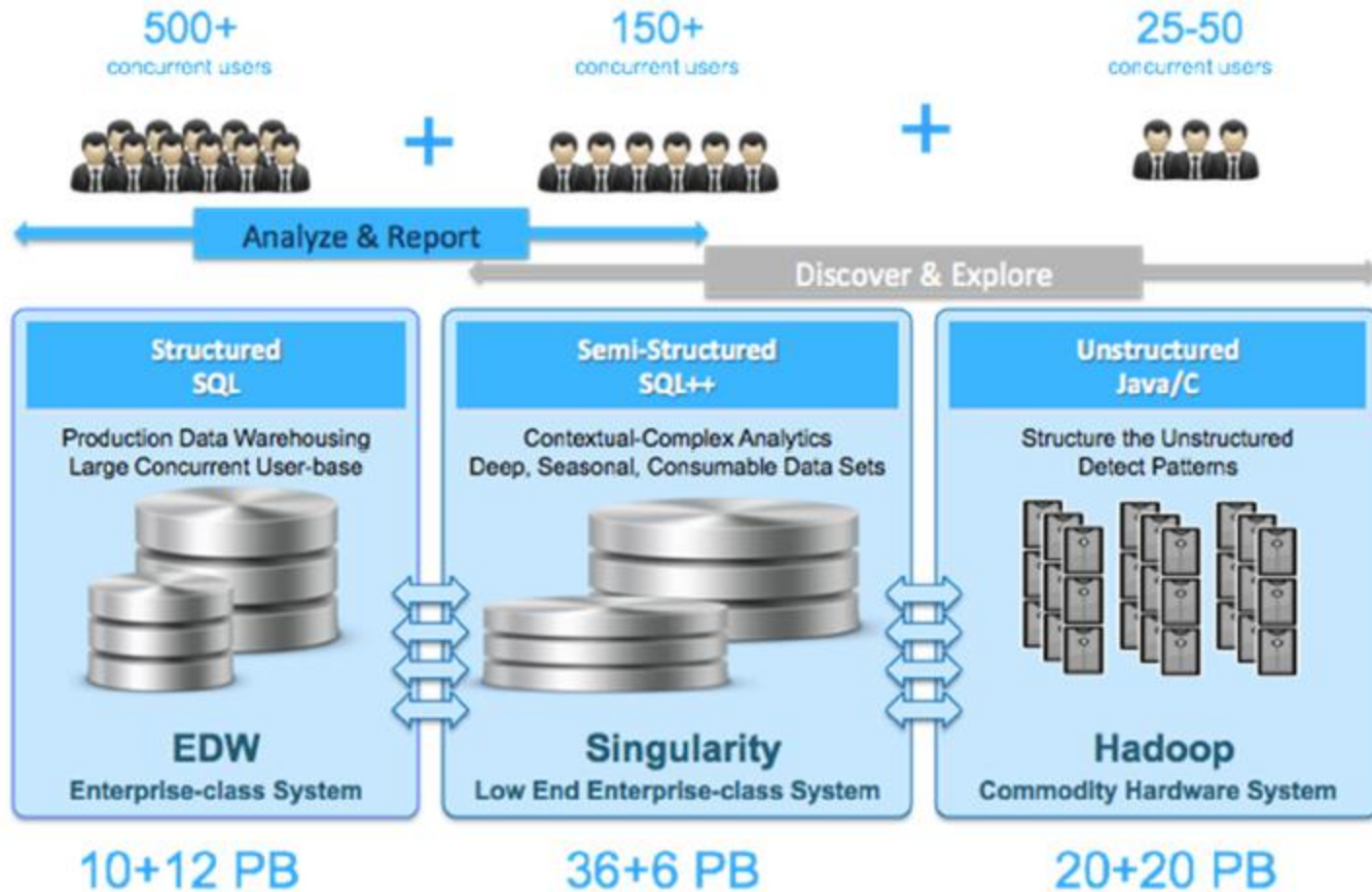
|                      | Low Cost<br>Storage &<br>Retention | Loading and Refining                    |                                             | Reporting | Analytics<br>(User-driven,<br>interactive) |
|----------------------|------------------------------------|-----------------------------------------|---------------------------------------------|-----------|--------------------------------------------|
|                      |                                    | Data Pre-Processing,<br>Prep, Cleansing | Transformations                             |           |                                            |
| Stable Schema        | MPP RDBMS /<br>Hadoop              | MPP RDBMS                               | MPP RDBMS                                   | MPP RDBMS | MPP RDBMS<br>(SQL analytics)               |
| Evolving Schema      | Hadoop                             | Aster* /<br>Hadoop                      | Aster*<br>(joining with<br>structured data) | Aster*    | Aster*<br>(SQL + MapReduce<br>Analytics)   |
| Format,<br>No Schema | Hadoop                             | Hadoop                                  | Hadoop                                      |           | Aster*<br>(MapReduce<br>Analytics)         |

\* Aster – example of Discovery platform

# Přesnější identifikace odchodu zákazníků (Churn)



# Analytický Ecosystém - eBay



# eBay si ověřila, že MPP RDBMS je vhodnější než MapReduce pro jejich Web analýzy

*"I talked with Oliver Ratzesberger and his team at eBay last week, who I already knew to be MapReduce non-fans. This time I added more detail.*

*Oliver believes that, on the whole, MapReduce is 6-8X slower than native functionality in an MPP DBMS, and hence should only be used sporadically. This view is based on part on simulations eBay ran of the Terasort benchmark. On 72 Teradata nodes or 96 lower-powered nodes running another (currently unnamed, as per yet another of my PR fire drills) MPP DBMS, a simulation of Terasort executed in 78 and 120 secs respectively, which is very comparable to the times Google and Yahoo got on 1000 nodes or more.*

*And by the way, if you use many fewer nodes, you also consume much less floor space or electric power."*

<http://www.dbms2.com/2009/04/14/eBay-thinks-mpp-dbms-clobber-mapreduce/>

TERADATA<sup>®</sup> ASTER