



DAU - Certified Data & Analytics Tester (CDAT) Syllabus

Version 1.1

Copyright Notice

This document may be copied in its entirety, or extracts made, if the source is acknowledged.

All CDAT syllabi and linked documents (including this document) are copyright of Data & Analytics United (hereafter referred to as DAU).

The material authors and international contributing experts involved in the creation of the CDAT resources hereby transfer the copyright to DAU. The material authors, international contributing experts and DAU have agreed to the following conditions of use:

- Any individual or training company may use this syllabus as the basis for a training course if DAU and the authors are acknowledged as the copyright owner and the source respectively of the syllabus, and they have been officially recognized by DAU. More regarding recognition is available via: <https://www.da-united.com/recognition>
- Any individual or group of individuals may use this syllabus as the basis for articles, books, or other derivative writings if DAU and the material authors are acknowledged as the copyright owner and the source respectively of the syllabus.

Thank you to the main authors

- Rogier Ammerlaan, Jaap de Roos and Armando Dörsek

Thank you to our Master Trainers who are supporting the Translation

- French: Jean-Luc Cossi
- Spanish: Juan Pablo Rios Alvarez

Thank you to the review committee

Alexis Jesús Herrera Colmenares, Ángel Rayo Acevedo, Arjan Brands, Aurelio Gandarillas Cordero, Beatriz López Botello, Benjamín Gillibrand H., Christine Green, Daniel Leo Lopez Romero, Egbert Bouman, Eibert Dijkgraaf, Emilie Potin-Suau, Fabiola Mero, Geoffrey Wemans, Héctor Ruvalcaba, Hermon Alfaro Fuentes, Huib Stoel, Jayapradeep Jiothis, Jean-Luc Cossi, Joan Gasull Jolis, José Díaz, Juan Pablo Rios Alvarez, Julie Gardiner, Kaan Sanli, Koos van Strien, Krishan Premcharan, Kyle Alexander Siemens, Márcia Araújo Coelho, Mario Alvarez Gómez, Mark Summers, Melissa Pontes, Miaomiao Tang, Miguel Angel de León Trejo, Nadia Soledad Cavalleri, Neriman Kocaman, Paula Castro, Petr Neugebauer, Dr. Ralph Elster, Ruth Margaret Florian Caipa, Santiago de Jesús Gonzalez Medellin, Samuel Ouko, Shreyansh Tewari, Silvie Schoenmaker, Dr. Srinivas Padmanabhuni, Thomas Cagley, Tim Koomen, Tobi Steenbakkers, Vanessa Islas Padilla, Vipul Kocher, Wim Decoutere

Revision History

Version	Date	Remarks
0.2	June 8 th , 2020	Initial BETA release
1.0	August 6 th , 2020	After review of review committee
1.02	May 14 th , 2021	Textual changes based on review comments
1.1	March 3 rd , 2025	Updated chapter on Data Quality, DTAP - minor changes

Table of contents

Table of contents	3
Business Outcomes	5
Learning Objectives/Cognitive Levels of Knowledge	5
Target Audience	6
Prerequisites	6
1 Introduction to Business Intelligence (BI) and Data & Analytics (DA)	7
1.1 Business Intelligence (LO-1.1)	8
1.2 Challenges faced when turning corporate data into information (LO-1.2)	8
1.3 Data Warehouse (LO-1.3, LO-1.4)	9
1.4 Data Mart (LO-1.5)	10
1.5 Data Lake (LO-1.6)	11
1.6 Data Analytics (LO-1.7, LO-1.8, LO-1.10)	12
1.7 Big Data (LO-1.9)	13
1.8 Data Mining (LO-1.11)	14
1.9 OLTP vs OLAP (LO-1.12, LO-1.13)	15
1.10 Reporting and Data Visualization (LO-1.14)	17
1.11 Data Modeling (LO-1.15)	18
1.12 The ETL Process (LO-1.16)	20
1.13 Internationalization and Localization (LO-1.17)	21
2 Data & Analytics Testing Strategy	24
2.1 Testing and Test Methodologies (LO-2.1, LO-2.2, LO-2.3)	24
2.2 BI & DA Testing in a traditional approach (LO-2.4)	25
2.3 BI & DA Testing in an agile approach (LO-2.5, LO-2.6)	26
2.4 Risk-based Testing in a BI & DA environment (LO-2.7, LO-2.8, LO-2.9, LO-2.10)	27
2.5 Testing Roles and Skills (LO-2.11, LO-2.12, LO-2.13)	27
2.6 Setting up a test-strategy and approach in a BI environment (LO-2.14)	29
3 Test Techniques	31
3.1 What are test techniques?	31
3.2 Different techniques (LO-3.1 to LO-3.7)	31
3.3 Mapping techniques on a BI & DA environment (LO-3.1 to LO-3.7)	32
4 BI Testing	33
4.1 Reports and Dashboards Testing (LO-4.1)	33
4.2 Testing of OLAP Cubes (LO-4.2)	35
4.3 Testing the Data Model (LO-4.3)	36
4.4 Testing in Data Mining (LO-4.4, LO-4.5)	37
4.5 Completeness Testing (LO-4.6)	37
4.6 Transformation Testing (LO-4.9)	39
4.7 E2E testing	40
5 Data Quality	42
5.1 Quality (LO-5.1)	42
5.2 Quality Characteristics (ISO/IEC 25010) (LO-5.2)	42
5.3 Data Quality Characteristics (ISO/IEC 25012) (LO-5.3)	43

5.4	<i>Data Profiling (LO-5.4)</i>	44
6	Environmental Needs	47
6.1	<i>Test environments according to DTAP (LO-6.1)</i>	47
6.2	<i>Multidimensional DTAP in the Data Analytics landscape (LO-6.2)</i>	48
6.3	<i>Impact of Security Standards on Analytics Testing (LO-6.3/6.4)</i>	50
6.4	<i>Impact of Privacy Regulations on Testing in the Analytics Environment (LO-6.5)</i>	50
6.5	<i>Encryption methodologies for Anonymization and Pseudonymization (LO-6.6)</i>	51
6.6	<i>Common Pitfalls using Production Data (LO-6.7)</i>	52
7	References	54

Business Outcomes

Business Outcomes (BOs) are a brief statement of what you are expected to have learned after the training.

BO-1	Understand the complexity and elements in a Business Intelligence (BI) & Data & Analytics (DA) environment
BO-2	Understand the specific aspects of testing in a BI & DA environment
BO-3	Have insights into the different skillsets required and the diversity of test roles/skills in a BI & DA environment
BO-4	Have insight into and the ability to apply different test techniques and know how they can be applied to a BI & DA environment
BO-5	Have insight into the different test aspects of a BI & DA environment compared to traditional testing, e.g., OLAP testing, Transformation testing, Completeness testing
BO-6	Understand the difference between quality attributes for system requirements and the specific quality attributes applied to data environments
BO-7	Understand and assess risk analysis and how to apply this in a BI & DA environment given the specific aspects of the environment
BO-8	Understand and assess the complexity of data and the importance of the use in DA
BO-9	Have knowledge about the complexity of DTAP environments (Development, Test, Acceptance and Production) in a BI & DA environment
BO-10	Understand the effects of privacy and data protection regulations for testing in DA projects

Learning Objectives/Cognitive Levels of Knowledge

To cover all the Business Outcomes (BOs), the Learning Objectives (LOs) contained in each chapter are defined. LOs are brief statements that describe what you are expected to know after studying each chapter, what can be examined for certification and are together with the BOs the training goals of this training.

The LOs are defined based on Bloom's modified taxonomy [BT1, BT2] as follows:

Definitions	K1 Remembering	K2 Understanding	K3 Applying
Bloom's definition	Exhibit memory of previously learned material by recalling facts, terms, basic concepts, and answers.	Demonstrate understanding of facts, and ideas by organizing, comparing, translating, interpreting, giving descriptions and stating main ideas.	Solve problems to new situations by applying acquired knowledge, facts, techniques, and rules in a different way.
Verbs (examples)	Remember	Summarize	Implement
	Recall	Generalize	Execute
	Choose	Classify	Use
	Define	Compare	Apply
	Find	Contrast	Plan
	Match	Demonstrate	Select
	Relate	Interpret	
	Select	Rephrase	

Target Audience

This certification program is designed for test engineers, test coordinators, managers, business analysts, data leads, data warehouse developers and BI-consultants.

Prerequisites

Recommended

- Basic knowledge of testing in general, e.g., ISTQB or TMAP.
- Basic knowledge of Database and Data Modeling principles, for instance BPMN or UML, ERD.

1 Introduction to Business Intelligence (BI) and Data & Analytics (DA)

In this introduction, terminology is defined about Business Intelligence, Data Warehouses, Data Analytics, Big Data, OLTP, OLAP, Data Marts, Reports and Dashboards.

Keywords

Business Intelligence, Data Warehouse, ETL, Data Marts, Reports, Dashboards, Analytics, Big Data, OLTP, OLAP, Three Vs, Modelling, Audit Trail, Star Schema, Snowflake Schema, Slicing and Dicing, Facts, Dimensions.

LO-1.1	K1	Recall the definition of business intelligence
LO-1.2	K2	Understand challenges turning data into information
LO-1.3	K2	Explain the concept of a data warehouse, including types of data
LO-1.4	K1	Recall the functional view and the 'how view' of a data warehouse
LO-1.5	K1	Recall the aspects of the data mart
LO-1.6	K1	Recall the description of a data lake
LO-1.7	K2	Explain the description of data analytics
LO-1.8	K1	Recall the three main types of analytics
LO-1.9	K2	Explain the big data concept by 3 Vs
LO-1.10	K2	Recall additional Vs associated with big data
LO-1.11	K1	Recall the definition of data mining
LO-1.12	K2	Explain and understand the difference between OLTP and OLAP
LO-1.13	K2	Explain cubes, facts, slicing and dicing
LO-1.14	K2	Understand what reports and dashboards are meant for
LO-1.15	K2	Recall the different modelling techniques associated with data warehouse design
LO-1.16	K2	Explain the ETL process and Source-to-Target mappings
LO-1.17	K2	Understand localization and internationalization
LO-1.18	K2	Understand the impact of (different) character sets on the BI & DA environment

In the Business Intelligence & Data Analytics (BI & DA) environment, different terms are being used for different parts of the model shown in Illustration 1. In this chapter, a common understanding of what we see as the world of BI & DA is described. This does not mean that other definitions do not exist, but, for the content of this syllabus, it is relevant that everybody gets the same understanding of the definitions used.

This chapter describes different aspects of a BI & DA environment. Illustration 1 will be used to explain the different parts of the model. It is important for a tester to understand the terminology and the environment.

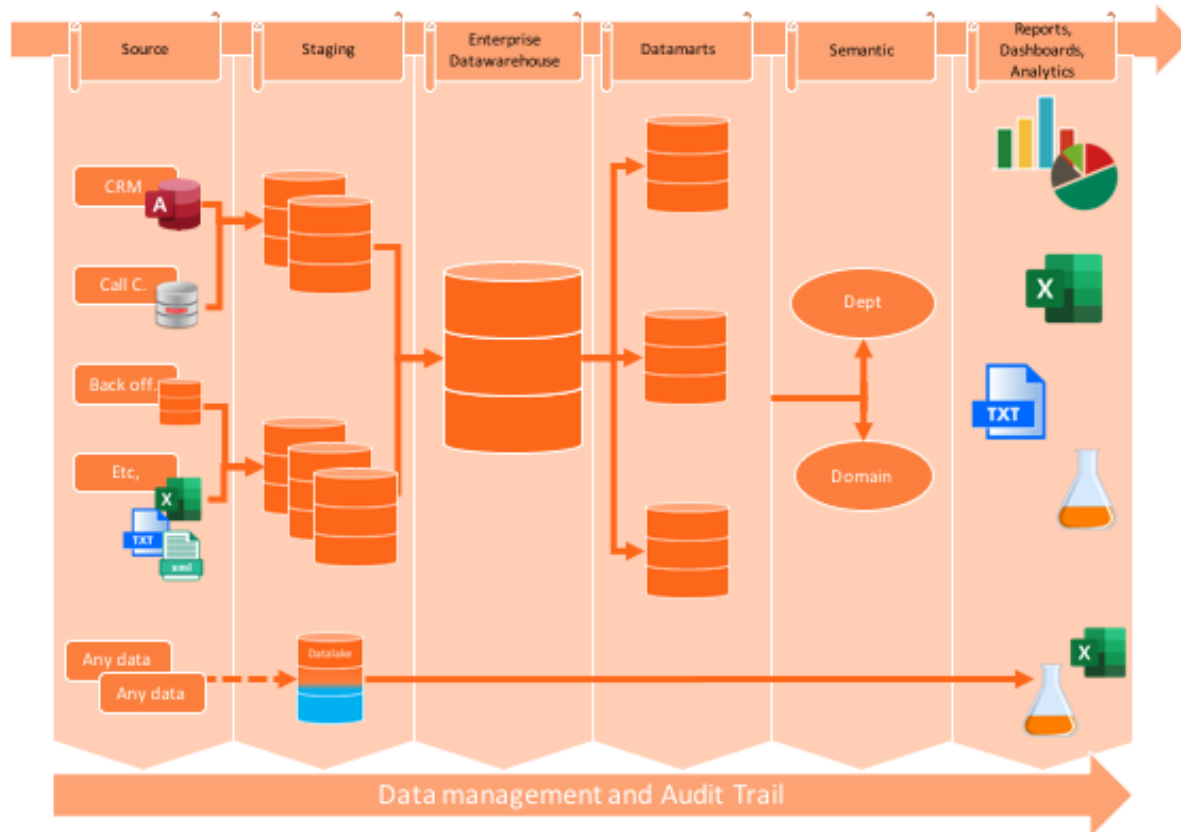


Illustration 1 - A typical BI environment

1.1 Business Intelligence (LO-1.1)

Business Intelligence (BI) comprises the strategy, methods and technologies used to gather, store, report and analyze business data to help people analyze data to make business decisions.

BI is further referred to as the way to visualize this analysis in a meaningful and clear representation. BI therefore refers to the “front end” technology in the DA environment. In the wider context, the concept of BI relates to the process of generating “knowledge” from data, including concepts like data warehousing. BI technologies like Tableau, Power BI, SAS, can handle large amounts of structured and sometimes unstructured data to help report the current state, but also to identify, develop, and otherwise create new strategic business opportunities.

1.2 Challenges faced when turning corporate data into information (LO-1.2)

The data which is needed to turn into meaningful information may come from several sources. Integrating these data to start analyzing may lead to challenges. The source systems are often built for different purposes than informing CEOs or feeding data scientists with useful data – in the right format and at the right moment.

For instance, the source systems may be running “on premise” (i.e. within the walls of our company) or somewhere “in the cloud” (at an external location, with a third party). Data may be delivered in batches, through pull systems like database queries or through push systems like messaging systems. Data may be received periodically (e.g., per batch run, daily or weekly) or continuously and in real time, like with feeds from energy meters. Some may store historical data; others may store only the current values. Companies may already be in possession of the data themselves, or they may need to acquire it from external parties (e.g. social media posts, weather reports, traffic information). The format may be proprietary or open, data quality may be taken care of by its owner or seen as a problem of data analysts.

To counter those differences and to make sure that there is no overload of the source systems with defined queries, a collection of data is stored in a different place so it can be used at the moment that is convenient. By adding timestamps to the data, historic information is introduced, which enables the analysis of the past periods and the spotting of trends.

1.3 Data Warehouse (LO-1.3, LO-1.4)

A data warehouse enables better BI & DA. Data warehouses are designed to overcome most of the challenges mentioned in the previous section.

The data warehouse can be found at the center of illustration 1. A data warehouse, also known as Enterprise Data Warehouse (EDW), stores data from several sources, either direct or via a staging area (an intermediate storage area used for data processing during the extract, transform and load process). It is a consistent way of storing (company) data, both current and historical, mainly associated with structured data. It consists of a data model regardless of the original source(s). The performance is optimized for processes storing new data. The data model of the data warehouse is created via different data modelling techniques, like: re-structuring data into Facts and Dimensions and de-normalization.

Definitions of a Data Warehouse

To get insight on a data warehouse, we present the (different) definitions of two important authors on the subject of business intelligence and data warehousing, namely W.H. (Bill) Inmon and R. (Ralph) Kimball.

W.H. (Bill) Inmon: “A data warehouse refers to a subject oriented, integrated, time-variant, non-volatile collection of data in support of management’s decision-making process”.

Subject Oriented

Creating a “subject oriented” collection of data results in the situation where all data regarding a subject, for instance “customer” or “product” is brought together, independent of its source.

Integrated

In an integrated collection of data, the same name is used for the same information (e.g., “customer-number” vs “client-id”), coding is harmonized (e.g., “male/female” vs “man/woman”), the same measurements are used (e.g., USD vs EUR and inches vs centimeters).

Time Variant and Non-Volatile

In operational systems, data is continuously added (inserted), edited (updated) and deleted, providing an “actual” or “current” view of the situation. A data warehouse is used to provide snapshots of the situation at any certain moment in time and *over* time.

This definition seems to stress the *writing* part of the data warehouse environment, as it describes data modeling choices that enable the fast and easy writing (inserts, updates, deletes) of data into the data warehouse.

Ralph Kimball refers to a data warehouse as “a copy of transaction data specifically structured for query and analysis”. This definition fits the “Business Intelligence” side very well, as it stresses the need of easy and fast querying (writing and reading) of the data that is made available in data marts (see section 1.4).

In section 1.11 more information on data modelling aspects can be found.

1.4 Data Mart (LO-1.5)

Data Mart (DM), mentioned in Illustration 1, is a subset of the Enterprise Data Warehouse (EDW). The data mart will hold a subset of the data in the EDW, focusing on a single subject area, and it is oriented towards a specific task or department/user group, enabling them to perform queries fast and in an easy manner. It presents data in a way that the end users will understand. It also holds summarized data, whereas the data warehouse will hold all details.

It is designed as a dimensional model (e.g. “star schema”, “snowflake”) and focuses on integration of information regarding a certain (single) subject area. Data marts can also tackle issues regarding privacy and security, for instance by showing less detailed information about individuals – or only to authorized colleagues.

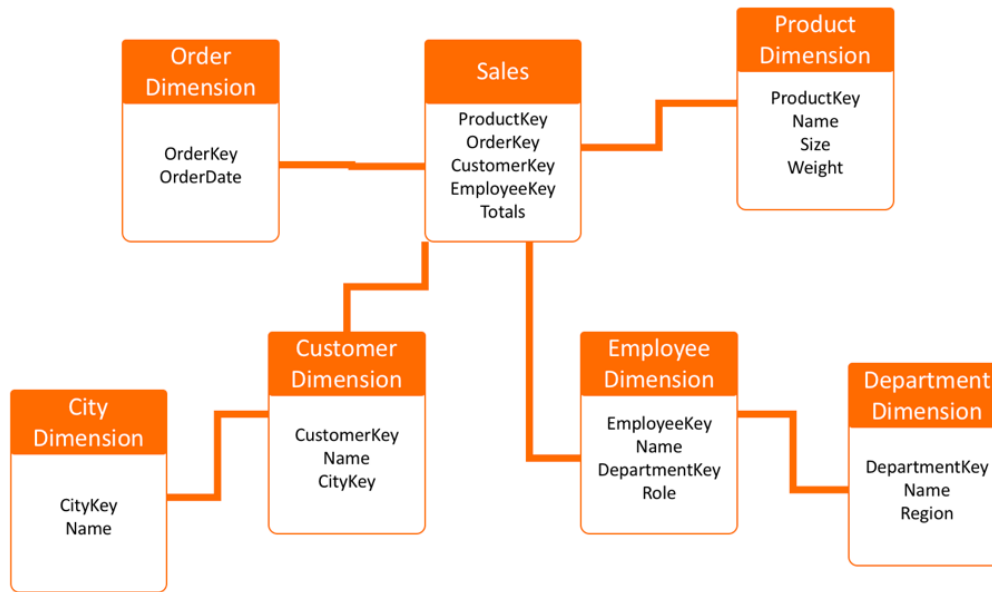


Illustration 2 - Example Snowflake Model

1.5 Data Lake (LO-1.6)

A data lake is a large repository (big data) of internal and external data in its natural, unstructured form as a raw data reservoir. In Illustration 1, it is illustrated on the left-hand bottom end. Raw data flows into a data lake in real time, and users can segregate, correlate, and analyze different parts of the data, based on their needs, or as a source for the ETL (Extract, Transfer, Load) process of a data warehouse.

In a data warehouse, data is often stored in tables (rows and columns). This requires that attributes of the data (e.g., data types, length) must be known in advance. Data lakes store (semi-structured) data types in their native format, without requiring the tables to be defined upfront. A data lake can include structured data from relational databases (rows and columns), semi-structured data (CSV, logs, XML, JSON), unstructured data (emails, documents, PDFs) and binary data (images, audio, video). Information should be tagged with a set of extended meta data tags (Campbell[CC1]).

Metadata is a description of a digital asset, “data about data”. These are terms associated with an object to describe it, for instance the date a record was created or added, the source where the information comes from, the type or category of the data. Metadata helps to organize digital assets and makes it therefore easier to find those specific assets within a group of objects.

Data lakes rely on low-cost storage options because of the volume of its contents. The infrastructure will need to fit the type of data stored in the data lake: large volumes of unstructured data. The core technology was based on the Apache Hadoop Ecosystem, an open-source software framework that distributes data storage and processing among commodity hardware located in on-premises

data centers. Hadoop includes a file system called HDFS that enables customers to store data in its native form.

As the cloud computing industry matured, object stores from large vendors introduced interim data lake solutions, such as Amazon Simple Storage Service (S3), Microsoft Azure Blob, and Google Cloud Storage.

Relational databases are often queried by use of SQL, Standard Query Language. As the data lakes are often stored in non-relational database systems or so-called NoSQL (“Not only SQL”) databases, testers are well advised to master skills associated with querying data in these specific systems too.

1.6 Data Analytics (LO-1.7, LO-1.8, LO-1.10)

Data analytics, which mainly takes place in the right-hand side of Illustration 1, i.e., “Reports, Dashboards and Analytics”, comprises the combined corporate and external data sources for the creation of analytical models. Based on these models, decisions can be made regarding groups and individual subjects. Data Analytics is strongly associated with “Big Data”.

A popular definition of analytics is “the extensive use of data, statistical and quantitative analysis, explanatory and predictive models, and fact-based management to drive decisions and actions” **[DA1]**.

There are three types of analytics: Descriptive Analytics, Predictive Analytics and Prescriptive Analytics (Davenport**[TD1]**).

1.6.1 Descriptive Analytics

Reports, dashboards, alerts, and scorecards are used to describe what has happened in the past. Descriptive Analytics may also be used to classify customers or other entities into groups that are similar on certain dimensions.

1.6.2 Predictive Analytics

Predictive Analytics is a type of analytics that uses data from the past to create models that help to make predictions. Today’s most used technique for creating these predictive models is machine learning. Typically, multiple variables are used to predict a particular *dependent* variable. For instance, using clients’ credit history to predict their ability to repay their loans in the future.

Predictive analytics uses a variety of quantitative techniques (such as optimization) and technologies (e.g., models, machine learning and recommendation engines) to specify the optimal behaviors and actions (Davenport**[DA1]**). It will tell you what to do, based on measuring cause and effect in test groups and control groups. E.g., a drug company tests its product on a group of test

subjects and gives a control group a placebo to verify the effect. If a statistically significant difference is discovered between the two, it will tell you which approach to take.

1.6.3 Prescriptive Analytics

In prescriptive analytics, data can be correlated, for instance where two or more variables (or data elements) have a mutual relationship or connection, or causation, that is the influence of causes on effects in a later state of the data.

Correlation vs Causation

Note that while causation and correlation can exist at the same time, correlation does not imply causation. Causation explicitly applies to cases where action A causes outcome B. On the other hand, correlation is simply a relationship between A and B. Action A relates to Action B – but one event does not necessarily cause the other event to happen.

1.7 Big Data (LO-1.9)

Big Data is a combination of a large amount of structured, semi-structured and unstructured data collected by organizations, that can be analyzed for predictive modeling and other advanced analytics applications. In a 2001 research report, META Group (now: Gartner) analyst Doug Laney characterized Big Data by using 3 Vs: Volume, Velocity and Variety.

1.7.1 Volume, Velocity, Variety (LO-1.9)

Volume

Volume refers to the amount of data that is created, stored, analyzed, and visualized. The sheer volume of the data is enormous, and a very large contributor to the ever-expanding digital universe is the Internet of Things (IoT) with sensors all over the world, in all sorts of devices creating data every second. And, of course, let's not forget all the social media including audio and video streams. Although big data does not equate to any specific volume of data, big data deployments often involve terabytes (TB), petabytes (PB) and even exabytes (EB) of data captured over time.

Velocity

Velocity refers to the speed at which data is created, stored, analyzed, and visualized. In the past, when batch processing was common practice, it was normal to receive an update from the database every night or even every week. Nowadays, real time data, streaming data, sensing data are used more and more. Because of IoT, velocity has become fundamental.

Variety

Variety refers to the many different formats of data. In the past, most data that was collected was structured data: it fitted neatly in columns and rows and in text files. Nowadays, data comes in many different formats: structured data, semi-structured data, unstructured data and even complex structured

data. The wide variety of data requires a different approach as well as different techniques to store all raw data (see: section 1.5, Data Lake) – also, to analyze and to effectively use the data.

1.7.2 Additional Vs (LO-1.10)

In the course of time, additional Vs were added to the big data definition, explaining important aspects of big data and a big data strategy that organizations cannot ignore. Four of these are: Veracity, Variability, Visualization and Value (Rijmenam[DF1]).

Veracity

Organizations need to make sure that both the data and the analysis are correct, so the information that is presented, is actually true.

Variability

Big data is extremely variable, i.e., the meaning of the data may differ from time to time, just like two words may have a different meaning. This will depend on the context and will be important in processes like sentiment analysis. Variability is often confused with variety, but while variety is concerned with the way data is stored and presented, variability is concerned with the actual *meaning* of the data.

Visualization

With the right analyses and visualizations, raw data can be put to use, otherwise it will remain essentially useless. Though perhaps not the most challenging part from a technical perspective, telling a complex story in a graph is not only very difficult, but also absolutely crucial, to present it in such a way that the recipients understand the message contained in the story.

Value

More and more organizations understand that data must be considered as an important asset. Of course, data itself is not valuable at all. Its value is in the analyses done on that data and how the data is turned into information and eventually turned into knowledge, which can be used for learning objectives or just to create more knowledge of a certain subject. Its value lies in how organizations will use that data and turn their organization into an information-centric company that relies on insights derived from data analyses for their decision-making.

1.8 Data Mining (LO-1.11)

Data mining consists in the process of discovering new patterns from large data sets involving methods at the intersection of artificial intelligence, machine learning, statistics, and database systems. This implies a discovery of trends and patterns in data – not by humans, but by the computer itself. (Davenport[TD1]). Data mining is most commonly associated with statistical and predictive models (see: Predictive Analytics). Data mining techniques use algorithms. Data mining mainly takes place in the phase Reports, Dashboards, and Analysis, as stated in the Illustration 1.

Some examples of data mining applications are:

- detection of fraud,
- predicting stock market price,
- analysis of customer behavior (e.g., in “Market Basket Analysis”).

1.9 OLTP vs OLAP (LO-1.12, LO-1.13)

The generic operational systems and the environments used in data analytics are different in several ways.

Operational systems, often also called Online Transaction Processing (OLTP), are used to support primary business activities like processing customer requests, invoice handling, processing lab results, planning courses and filling or filing tax forms (mainly ‘writing’ and ‘updating’ data). Whereas Online Analytical Processing (OLAP) is used to create value from the existing business data, by judging measures such as company results over several dimensions like Time, Pricing Categories, Regions and Customer Groups (mainly ‘reading’ data).

In the table below, you will find some of the differences between Online Transaction Processing (OLTP) and Online Analytical Processing (OLAP).

	OLTP System Online Transaction Processing (Operational System)	OLAP System Online Analytical Processing (Data Warehouse)
Source of data	Operational data; OLTPs are the original source of the data	Consolidation data; OLAP data comes from the various OLTP sources
Purpose of data	To control and run fundamental business tasks	To help with planning, problem solving, and decision support
Data insights	Reveals a snapshot of ongoing business processes, actual status	Multi-dimensional views of various kinds of business activities, time series
Inserts, Updates and Deletes	Swift inserts, updates and deletes initiated by end users	Periodic long-running batch jobs refresh the data, no physical deletion of data (but flagging “deleted” records)
Queries	Relatively standardized and simple queries returning relatively few records	Often complex queries involving joins, transformations, and aggregations
Processing Speed	Typically, very fast processing, focused on real time informing of end users and typically short batch windows	Depending on the amount of data involved and the complexity of the (batch) processing, loading data and querying for results may take many hours
Space Requirements	Can be relatively small if historical data is absent or archived	Larger due to the existence of redundant storage, aggregation structures and history data
Database Design	Highly normalized, with many tables	Typically, de-normalized with fewer tables, e.g., use of star and/or snowflake schemas for easy reporting
Backup and Recovery	As operational data is mostly critical to run the daily business, backups and restores are of the utmost importance	Instead of creating backups (and restores) re-loading the OLTP data may be sufficient as a recovery method, depending on the processing times and importance of certain analytics products

Table 1 - Differences OLTP vs OLAP

1.9.1 OLAP Cubes and Dimensional Reporting

Reports and dashboards are great for visualizations that present data in a uniform fashion. But when analytical users need to be able to explore data in an interactive way, clicking through regions, departments, turnover by month etc., the OLAP cubes come into play.

An OLAP cube, also known as multidimensional cube or hypercube, is a data structure which offers ways to display and sum large amounts of data and provide searchable access to any data points. The data can be rolled up, sliced, and diced as needed to handle all kinds of questions that arise when confronted with the data set. In contrary to most OLTP databases, which are not designed for analysis, OLAP databases are specialized databases that are designed to help extract this business intelligence information from the data.

In order to enable easy and fast querying, the underlying data model is often based on Kimball data modeling techniques. The star schemas and snowflake schemas (as opposed to Relational Modeling techniques) will hold Fact and Dimension tables which are used in the cube (see section 1.12: Data Modeling).

Due to its structure, the OLAP cube is often accessed by other means than SQL by multidimensional expressions (MDX¹). This is due to the fact that SQL is not sufficient enough with its technique, where records are inserted one by one. It is extremely slow, considering the vast volume of data to be loaded in the DWH. Several analytical tool vendors, like Tableau, Cognos, Microsoft Power BI and SAS, provide software to perform OLAP.

1.9.2 Key aspects in OLAP systems

Measures

Measures are the numeric values that users want to slice, dice, aggregate, and analyze; they are one of the fundamental reasons why you would want to build OLAP cubes using the data warehousing infrastructure. For instance, gross and net earnings, direct and indirect costs, the number of products sold.

Dimensions

Dimensions allow for the filtering, grouping, and labeling of data. For instance, Regions or Locations, Dates and Product Dimensions. The data can be presented in a format where the data is categorized naturally into these hierarchies and categories to allow a more in-depth analysis. Dimensions may also have natural hierarchies to allow users to “drill down” to further levels of detail. For instance, the date dimension that can be drilled down from Year to Quarter, Month, Week and Day.

¹ Multidimensional Expressions (MDX) is a query language for online analytical processing (OLAP) using a database management system. Much like SQL, it is a query language for OLAP cubes. It is also a calculation language, with syntax similar to spreadsheet formulas

Drill Down and Drill Through

“Drilling down” into the data in an OLAP cube, means that the user is analyzing the data at a different level of summarization. The level of detail of the data changes as the user drills down. Here terms like granularity (level of detail) are important, as the user is moving from summarized data to data with a narrower focus.

When users “drill through” data, they are requesting the individual transactions that contributed to the OLAP cube’s aggregated data at the lowest level of detail for a given measure value.

The difference between them is that a drill-down operates on a *predefined hierarchy* of data, for instance from Europe, to The Netherlands, to Amsterdam, within the OLAP cube. A drill-through goes directly to the lowest level of detail of data and will retrieve a set of rows from the data source that has been aggregated into a single cell.

Granularity

The most detailed unit of the data is a fact, a contract, invoice, spending, task, etc. Each fact might have a measure – an attribute that can be measured, such as: price, amount, revenue, duration, tax, discount, etc. The “grain” (or granularity) of a fact table states the level of detail in the fact table. A low level of granularity means that the table holds very detailed information – which can later be “rolled up”.

1.10 Reporting and Data Visualization (LO-1.14)

Visualizing data using charts, graphs, and maps is one of the most impactful ways to communicate complex data. Visualizing data for end users, is a very important aspect of the data analytics activities and require strong analytical and communication skills next to knowledge on the business and (financial) reporting.

Standards regarding reporting can be applied, for instance, those of the International Business Communication Standards (ICBS). They feature important aspects like “Conceptual Rules” ([Minto, **BM1**]). “Perceptual Rules” (Few[**SF**]), and “Semantic Rules” (Hichert[**RH**]).

Report

Report, according to Poonam Madan, is defined as “a document that presents information in an organized format for a specific audience and purpose” (Madan[**PM1**]). There are two types of reports: static and dynamic. Static reports are run immediately upon request, and then stored with their data and will not change anymore. Dynamic reports are created at runtime. Each time a dynamic report is run, it gathers the most recent data in the Data Warehouse. Only the report definition, which remains the same over time, is stored.

Dashboard

A dashboard is a visual display of the most important information needed to achieve one or more objectives; consolidated and arranged on a single screen so the information can be monitored at a glance.

Dashboards vs Reports

A report is a more detailed collection of tables, charts, graphs and text and it is used for a much more detailed, full analysis, while a dashboard is used for monitoring what is going on. The behavior of the pieces that make up dashboards and reports are similar, but their makeup itself is different. A dashboard answers a question in a single view, while a report provides information. To put it another way, the report can provide a more detailed view of the information that is presented on a dashboard.

1.11 Data Modeling (LO-1.15)

In order to store data in the environment in an orderly fashion, data modeling techniques will be used. Engineers and testers need to understand the reasons for applying certain techniques – and their effects.

In “regular” computer systems (see: OLTP), data is usually stored in a highly normalized fashion. Edgar Codd introduced rules for data modeling: Normal Forms. The purpose of normalization is to eliminate redundant data, therefore saving space and avoiding mistakes when updating and deleting data.

In the process of normalization, tables are divided into smaller tables, which are linked through relationships, based on the concept of primary keys (PK) and foreign keys (FK). This will usually lead to a great number of tables presented in an entity relationship diagram (ERD).

A drawback of highly normalized data tables is that *querying* the data for analysis purposes can be hard, time consuming and error prone. Therefore, data marts are usually designed as in a de-normalized² fashion, leading to logical data models (LDM) that are designed according to other design rules.

Well known designs for analysis purposes (OLAP, data marts) are star schemas or snowflake diagrams **[GUR]** which are multi-dimensional schemas, holding Fact and Dimension tables. Inside a data warehouse, design techniques like Data Vault modeling (Linstedt**[DL1]**) and Anchor Based modeling (Rönnback**[LR1]**) can be found.

1.11.1 Maintaining History: Slowly Changing Dimensions

A data warehouse is used to provide snapshots of the situation at any certain moment in time and *over* time. Therefore, data is almost never physically deleted from a data warehouse, but only flagged as such. Changes in data will be stored by adding new rows, reflecting the new situation. This requires design decisions, stating where and how to store the concept of time together *with* the data.

² Denormalization is a strategy used on a previously normalized database to increase performance. In computing, denormalization is the process of trying to improve the read performance of a database, at the expense of losing some write performance, by adding redundant copies of data or by grouping data.

When data is being copied from source systems to (the staging area of) the data warehouse, it is common to add time stamps to the data. By doing so, we will know at which moment the changes were known.

Not all changes will have the same effect in the data warehouse environment. Changes in data may reflect a change in the real world, like when a shop physically moves to another address. New data can also complete or correct situations, like adding actual values where NULLs used to be or correcting a typo in the Product Description.

In order to cope with changes in data over time, Kimball introduced the Slowly Changing Dimensions (SCD). We cover SCD Types 0 through 4 below.

Type 0 – No Changes

If SCD Type 0 is applied, dimensions never change. Changes are not allowed at all.

Type 1 – No History

With SCD Type 1, existing values are simply overwritten with the new data. The effect is that tables do not grow much, because no rows are added to handle the time aspects. The drawback of this is that the table will always only hold the current value for each attribute and any historical value of the data is lost.

Type 2 – Row Versioning

When using SCD Type 2, new records are added, encompassing the change, and flagging the old record as inactive (by using an “inactive” column and/or a start date and end date). This allows the fact table to continue to use the old version of the data for historical reporting purposes and to use the changed data in the new record to only impact the fact data from that point forward.

Type 3 – Previous Value column

With SCD Type 3, an extra column is added to store both the actual and the previous value of the column(s) you wish to be able to report on. When the data is updated, the existing value is “moved” to the column “previous value” and the new value is placed into the reportable column. This enables you to look back at what the value of the data was previously. Note that loading/updating data may become very difficult and error prone.

Type 4 – History Table

With SCD Type 4, current values are stored in the dimension table but all changes are tracked in a separate table. This provides an easy way to show the current data; selecting historic data may become more challenging. One of the reasons to introduce this SCD type can be database performance: this way the database table with “current” data can be kept small (less records).

1.12 The ETL Process (LO-1.16)

ETL stands for: Extract, Transform and Load (data). Data is *extracted* from the source systems, *transformed* in a usable format, and *loaded* into the (enterprise) data warehouse.

ETL is commonly associated with data warehousing projects but in reality, any form of bulk data movement from a source to a target can be considered ETL. ETL processes are seen, when the need arises to move bulk application data from one source to another for data integration or data migration purposes. ETL testing is a data-centric testing process used to validate that the data has been transformed and loaded into the target as expected [DG1].

Testing the ETL process may cover aspects like completeness, auditability, logging processing results, exit codes, processing times, re-starting of ETL processes, syntactical and semantical validations (are the right source fields moved to the right target fields and do they indeed mean the same, were the right tables joined) and are the conversions of dates, numbers and strings performed as expected.

ETL is documented by means of Source-to-Target Mappings (STM) (also known as: ETL mapping documents), which will often describe:

- Names of source and target tables,
- Names (and data types) of columns in the source and target tables,
- How to detect Inserts, Updates and Deletes,
- The way SCD are applied (adding History),
- Transformations of data.

Mapping Change Date	Source table	Source column	Data type, Length	Transformation	Defaults	Error Handling	Target table	Target column	Nullable	Primary Key	Data type, Length
MM/DD/YYYY	Name of source table	Column in source table from which data is extracted	Data type and length for the source system	Transformation information, lookups, functions, etc.	Value to use when source field is null	Used to document value, if, keys, comments, etc.	Name of target table	Target table column name	Whether a field can be null	Primary key field for target	Data type & Length for target column
SOURCE							TARGET				

Illustration 3 - Source to Target Mapping Header Example

The ETL process can be a very compute-intensive activity. Due to a shift to data lakes and cloud-based data warehousing solutions, in which the complete set of data is first moved to the cloud (storage) and then transformed, the abbreviation ELT is being introduced: Extract, Load and Transform. This stresses that the activity of Transformation is being performed where it makes most sense, i.e., not on premise but “in the cloud”.

1.13 Internationalization and Localization (LO-1.17)

There is a difference in adaptation of, for instance, a product, application, or document when it will be created for a local or international market. In the next sections, the differences will be explained.

1.13.1 Localization

According to the W3C internationalization site, localization refers to the adaptation of a product, application, or document content to meet the language, culture and other requirements of a specific target market (a local).

Amongst others, it can entail customization related to:

- Numeric, date and time formats,
- Use of currency,
- Keyboard usage,
- Collation³ and sorting,
- Symbols, icons and colors.

1.13.2 Internationalization

Internationalization is the design and development of a product, application or document content that enables easy (future) localization for target audiences that vary in culture, region, or language. Internationalization typically entails:

- Designing and developing in a way that removes barriers to localization or international deployment. This includes such things as enabling the use of proper encoding of characters.
- Providing support for features that may not be used until localization occurs. For example, using markup in your DTD⁴ (Document Type Definition) to support bidirectional text⁵, or for identifying language.
- Enabling code to support local, regional, language, or cultural-related preferences. For instance, date and time formats, local calendars, number formats and numeral systems, sorting and presentation of lists, handling of personal names and forms of address, etc.
- Separating localizable elements from source code or content, so that localized alternatives can be loaded or selected, based on the user's international preferences as needed.

1.13.3 Importance of Localization and Internationalization

³ Collation specifies how data is sorted and compared in a database.

⁴ A DTD defines the structure and the legal elements and attributes of an XML document.

⁵ A bidirectional text contains two text directionalities, right-to-left (RTL or dextrosinistral) and left-to-right (LTR or sinistrodextral). It generally involves text containing different types of alphabets, but may also refer to boustrophedon, which is changing text direction in each row.

Internationalization and Localization are important to Data Analytics engineers and testers. When collecting data from source systems, the way data is stored needs to be clear. Each interface to the data (both upon input as well as when retrieving the data) can influence the way the data is stored and/or presented. At processing, the ingested data⁶, handling internationalization and localization concepts is important, not only from a pure technical perspective, but also when looking at business rules: should customer “AEL Dörsek” be treated as a copy of “AEL Doersek” and “AEL D̄rsek”? Or are these to be treated as individual customers? This touches the subject of data quality (see chapter 5).

When presenting data, the correct implementation of localization is important: tools should be able to handle differences in currencies, thousands-separators, decimal separators, etc. Are characters being displayed correctly, and the non-Latin characters like Chinese, Japanese, Arabic, Greek, too?

1.13.4 Character Sets and Internationalization (LO-1.18)

In the course of history, different character sets were designed in order to store data. A few of these character sets are described below, including their main features.

ASCII

ASCII is a 7-bit encoding technique which assigns a number to each of the 128 characters that are most frequently used in American English. This allows most computers to record and display basic text. It does not include symbols frequently used in other countries, such as the British Pound sign or characters with the German Umlaut. ASCII is understood by almost all email and communications software.

ISO/IEC 8859 (Extended ASCII)

ISO/IEC 8859 is an 8-bit extension to ASCII developed by ISO (the international organization for standardization). Along with the 128 ASCII characters, it covers characters like the Euro sign (#128), British Pound (#163) and the American Dollar Cent (#162) symbols. Several variations of the ISO/IEC 8859/Latin standard exist, to cover different language families. For instance, Latin-1 for Western European languages, Latin-5 for Turkish, ISO/IEC 8859-5 for Cyrillic and ISO/IEC 8859-8 for Hebrew. Depending on the chosen standard, the stored bytes will present (partially) different characters.

UNICODE

Unicode is an attempt by ISO and the Unicode Consortium to develop a coding system for electronic text that includes every written alphabet in existence. It uses 8-, 16- and 32-bit characters (= multi byte characters) depending on the specific representation. This results in larger documents, because storing these characters takes up more space than ASCII or Latin-1 based texts.

Note: The first 256 characters of UNICODE are identical to Latin-1. Unicode will even store emojis (e.g., smileys) and most mathematical symbols.

⁶ Data ingestion is a process by which data is moved from one or more sources to a destination where it can be stored and further analyzed. The data might be in different formats and come from various sources, including RDBMS, other types of databases, S3 buckets, CSVs, or from streams.

1.13.5 Possible effects of mixing character sets

Amongst others, the following issues may occur when mixing character sets:

- Possible ambiguity in presentation and conversions as bytes may be presented in the wrong character set,
- Storage may not be sufficient when storing multi byte characters⁷; this issue is sometimes not foreseen at the design stage,
- When poorly implemented, the sort order of the bytes can be ambiguous: which byte of multi byte character should be read first and which one last (big endian, little endian)⁸,
- When using `STRING` functions, programming languages and database systems may have difficulty in selecting the correct character (selecting the second byte is not the same as selecting the second character),
- Storing special characters, like emoticons, may be supported by specific applications but not by the tools in which reports and dashboards are developed,
- Characters that appear the same to humans might not be recognized as such by computer systems (e $\neq$ e in different character sets for example).

⁷ A multi byte character is a character composed of sequences of one or more bytes. Each byte sequence represents a single character in the extended character set.

⁸ A big endian system stores the most significant byte of a word at the smallest memory address and the least significant byte at the largest. A little-endian system, in contrast, stores the least-significant byte at the smallest address. Computers store information in various sized groups of binary bits.

2 Data & Analytics Testing Strategy

This chapter will explain, the characteristics of testing, different test methods, risk analysis and discuss the differences and similarities of testing in a data & analytics environment and the skills of a tester.

Keywords

Testing, BI Testing, Test process, Waterfall, V-model, Scrum, Agile Testing, Risk Analysis, Test leader, (Reporting, ETL, Data Migration or Data Quality) Tester

LO-2.1	K1	Recall the definition of testing
LO-2.2	K1	Recall the test methodologies
LO-2.3	K2	Understand the different test processes
LO-2.4	K2	Explain BI testing
LO-2.5	K2	Explain Scrum and Agile testing
LO-2.6	K1	Recall the Agile BI checklist and BI best practices
LO-2.7	K2	Explain risk-based testing
LO-2.8	K3	Recall and apply a risk-based testing strategy
LO-2.9	K2	Explain the definition of risk identification
LO-2.10	K2	Explain how likelihood and impact Influence risk
LO-2.11	K1	Recall the definition of test leader and tester
LO-2.12	K1	Recall the tester's skills and additional skills for D&A testers
LO-2.13	K2	Explain the differences of skills for the reporting tester, ETL tester, data migration tester and data quality tester
LO-2.14	K3	Recall and set-up a test strategy and approach in a BI environment

2.1 Testing and Test Methodologies (LO-2.1, LO-2.2, LO-2.3)

There are different definitions of testing, one from the ISTQB syllabus:

"The process consisting of all lifecycle activities, both static and dynamic, concerned with planning, preparation and evaluation of software products and related work products to determine that they satisfy specified requirements, to demonstrate that they are fit for purpose and to detect defects." (ISTQB[IS1])

A common perception of testing is that it consists only of running tests, i.e., executing the software (mainly called 'dynamic testing'). Whereas this is *part* of testing, it is certainly not all of the testing activities. There are also test activities before and after test execution. There are several test methodologies that define a structured test process. These activities include: planning, monitoring and control, choosing test conditions, designing and executing test cases, checking results, evaluating entry and exit criteria, reporting on the testing process and system under test, and finalizing or completing closure activities after a test phase has been completed. Testing also includes reviewing documents (and/or source code) and conducting static analysis, so called as it forms part of static testing.

Both dynamic testing and static testing can be used as a means for achieving similar objectives and will provide information that can be used to improve both the system being tested and the development and testing processes.

Testing, if it is, for instance, in an app environment or a BI & DA environment, can have one or even more of the following objectives:

- Finding defects
- Gaining confidence about the level of quality
- Providing information for decision-making
- Preventing defects

The thought process and activities involved in designing tests early in the life cycle (verifying the test basis via test design) can help prevent defects from being introduced into code. Reviews of documents (e.g., requirements, use cases) and the identification and resolution of issues also help to prevent defects appearing in the code.

Different viewpoints in testing respond to different objectives. For example, in development testing (e.g., component, integration and system testing), the main objective may be “to detect as many failures as possible”, so that defects in the software are identified and can be fixed. In acceptance testing, the main objective may be “to confirm that the system works as expected”, to gain confidence that it has met the requirements. In other cases, the main objective of testing may be “to assess the quality of the software” (with no intention of fixing defects), to give information to stakeholders about the risk of releasing the system at a given time. Maintenance or regression testing often includes testing that no new defects have been introduced during the development of the changes. During operational testing, the main objective may be to assess system characteristics, such as reliability or availability.

Different approaches have been developed over the decades, but the most often used testing methodologies are:

- ISTQB®: Internationally most common and known. A test approach on how to setup a structured test process in different software development life cycles (agile, scrum, waterfall, or V-model approaches)
- TMap®, a test management approach on how to set up a structured test process in different software development life cycles (agile, scrum, waterfall, or V model approaches)
- Quality for DevOps Teams, a test approach based on the Knowledge of Body TMap® about testing in a DevOps environment

Different countries and organizations could also develop their own approaches, based on combinations of the above-mentioned methods, and the best practices of their environment.

2.2 BI & DA Testing in a traditional approach (LO-2.4)

BI & DA testing differs from a regular software testing project in various aspects. But even so, knowledge of testing methodologies is important to understand, plan, and organize a BI & DA testing project.

To define “testing”, we refer to the ISTQB®, (ISTQB[IS2]); however, specifically in the context of a Data, Analytics or Business Intelligence environment, an addition to the traditional test definitions is required.

“Business Intelligence & Data Analytics Testing is the process of validating the data, format and performance of the ETL, analytics, reports, subject areas and security aspects of the BI & DA projects.”

Main focus areas:

- Data at the source

- Data transformation
- Data loading
- Reporting

Because of these specific aspects of BI testing, it is important for testers to integrate them in the software development lifecycle model on the different test levels, in addition to the applicative changes on a product. For instance, data is an important part of the 'system' under test and therefore, testing of data (and possibly data quality) is part of functional testing. Integration testing can apply to the flow of data from source to target system. All elements of the BI & DA environment that are in scope of the project should be mapped on the several phases of the development life cycle.

For instance, data transformation could be part of the component test, where the development team validates the individual components developed for the data transformation, whereas a possible separate test team will validate the functional business rules of data transformation on system test level.

For each aspect of the BI test, a test approach needs to be chosen on the several test levels. In practice, there could be more, fewer, or different levels of development and testing, depending on the BI environment.

Software work products (such as data transformation documents, business rules but also business scenarios or user stories, use cases, requirements specifications, design documents and code) produced during development are often the basis of testing. Therefore, implementing reviews of these documents can help improve their quality. So, do not focus only on the execution of testing with, for instance, running transformations, but also on improving the quality of the description of the rules, to prevent wrong interpretations.

2.3 BI & DA Testing in an agile approach (LO-2.5, LO-2.6)

Within Agile, many agile frameworks are developed. Scrum is one of the agile frameworks that is most commonly used. Scrum is founded on empirical process control theory, or empiricism. Empiricism asserts that knowledge comes from experience and making decisions based on what is known. Scrum employs an iterative, incremental approach to optimize predictability and control risk. Three pillars uphold every implementation of empirical process control: transparency, inspection, and adaptation (Schwaber[SG1]).

Agile testing is an iterative test approach based within the context of the agile way of working. There are many definitions of agile testing (C. Kaner, M. Bolton, J. Bach and many more), but a good explanation is mentioned in section 1.4 of the Agile United-Certified Practitioner in Agile Testing (AU-CPAT) Syllabus. Some other interesting reads are Lisa Crispin's two books, *Agile Testing* (Crispin[JL1]) and *More Agile Testing* (Crispin[JL2]).

Before implementing agile testing and using the BI development and testing process in an agile environment, the following checklist could help with the thought sustaining the implementation process:

- Start with the business information needs to provide context for scope,
- Iterations should be time-boxed,
- Stress on data discovery through the requirements and design phase,
- Validate the BI Architecture and get approval on the proof of concept,

- Data validation and verification⁹ should be completed for each development iteration,
- Use flow charts or diagrams to explain the BI process along with some documentation,
- Any change that will be deployed to production should be thoroughly tested in regression environment,
- Environment needs for the different DevOps teams/ CI/CD pipelines,
- Lean solutions to add like Monitoring, tests in production, Build-in Quality (especially in a D&A environment where Dataflows can be unpredictable),
- Have a formal change control; this will minimize the risk as all changes have to be approved before it goes into production.

2.4 Risk-based Testing in a BI & DA environment (LO-2.7, LO-2.8, LO-2.9, LO2.10)

A risk is something that might lead to a future negative consequence. It can be defined, based on the probability that something occurs and the impact when that events happen. Risk-based testing is an approach for testing that is trying to reduce the level of risks on the product by informing stakeholders of the status of these risks. For a description of risk-based testing, we will refer to (ISTQB[IS1]).

A tester is able to create a risk analysis using a risk-based approach. During risk analysis, quality characteristics models like ISO/IEC 25010 are used. Most of these models are focusing on the quality of the product itself, which is fine, but for BI & DA environments, additional quality characteristics should be applied. There are specific information and data quality characteristics that are very valuable to be considered. In chapter 5, a more detailed description of these quality characteristics is given.

A tester who is involved in setting up a test strategy and risk-based testing approach should understand these characteristics and specific risks in a BI & DA environment, and how risk-based testing supports BI testing. For instance, localization issues, or format issues (e.g., amounts) between different data sources or countries.

When a tester works in an agile environment, it is expected that, during the sprint planning meeting, the tester thinks about risks and risk-based testing. Therefore, the tester should be capable to apply and understand risk-based testing, and, in this case, specifically in a BI & DA environment.

2.5 Testing Roles and Skills (LO-2.11, LO-2.12, LO-2.13)

The ISTQB covers the roles of tester and test leader (ISTQB[IS2 par 5.1.2]) and describes their tasks in a generic manner. While the organization of testing in a BI & DA environment is quite similar, its landscape is complex and testing this requires a great variety of skills.

Think of skills in the following areas:

- Querying (e.g., SQL, NoSQL)

⁹ Data verification is a way of ensuring the user types in what he or she intends, in other words, to make sure the user does not make a mistake when inputting data. Validation is about checking the input data to ensure it conforms with the data requirements of the system to avoid data errors.

- Infrastructure (e.g., system engineering)
- Data Modelling
- Data Visualization
- Data Integration and Reporting Tools
- Test Data generation
- Data Quality
- (Inter)nationally acknowledged legal frameworks regarding Privacy, e.g., General Data Protection Regulation (EU) 2016/679 (GDPR).

When addressing roles and responsibilities, using generic terms such as “tester” or even “data warehouse tester” can be dangerous, as its meaning may differ in various organizations and projects. The following roles are examples of how they can be defined more precisely:

1. Reporting Tester
2. ETL Tester
3. Data Migration Tester
4. Data Quality Tester

Depending on the organization and type of project, different knowledge and skills will be required from the tester. The spider graph in Illustration 4 provides a strong *suggestion* on the knowledge levels needed for each role and (project) environment. The descriptions listed below will highlight some of the key skills which characterize each role.

Reporting Tester

The Reporting Tester (or: “Reports Tester”, “BI Tester”) will focus mainly on the end user product, interacting with the developer of reports and dashboards and with the business. This tester will excel in knowledge of business standards, subject matter expertise and industry knowledge (for instance, health care, finance or marketing). Most reports testers can be found close to the business.

ETL Tester

In general, ETL Testers will have knowledge of logical database design, ETL-processing tools and their querying skills will enable them to create comparisons between source and target tables, incorporating aggregations and transformations according to (general) design rules and technical designs like Source-to-Target Mappings. ETL testers will often be knowledgeable in the underlying infrastructure (databases) so they can design test cases accordingly. ETL testers will also focus on robust processes and on the availability of the system/data. Most ETL testers can be found close to the developers.

Data Migration Tester

Like the ETL tester, Data Migration Testers will need to show that data is moved to a different system in a correct and complete fashion. Their target may be a different technology and/or from on premise to a cloud-based solution. This requires additional knowledge about storage of data, data types, performance

characteristics, etc. As the migration is often planned as a one-off activity (one migration after a few trials), robustness is often of lesser interest. Data migration testers are likely to be found in task forces that perform similar data migrations for different customers, using elaborated re-usable scenarios.

Data Quality Tester

Should data quality be explicitly mentioned as an important aspect, then Data Quality (DQ) testers can be introduced. DQ testers will excel in measuring the expected and perceived data quality in the system, at source systems (OLTP) or the Data & Analytics landscape. In order to analyze results, such roles require some knowledge of the industry, of the behavior of the underlying infrastructure and of the tools used to perform the analysis. In the field, Data Stewards may be asked to perform this testing role, along with their other activities.

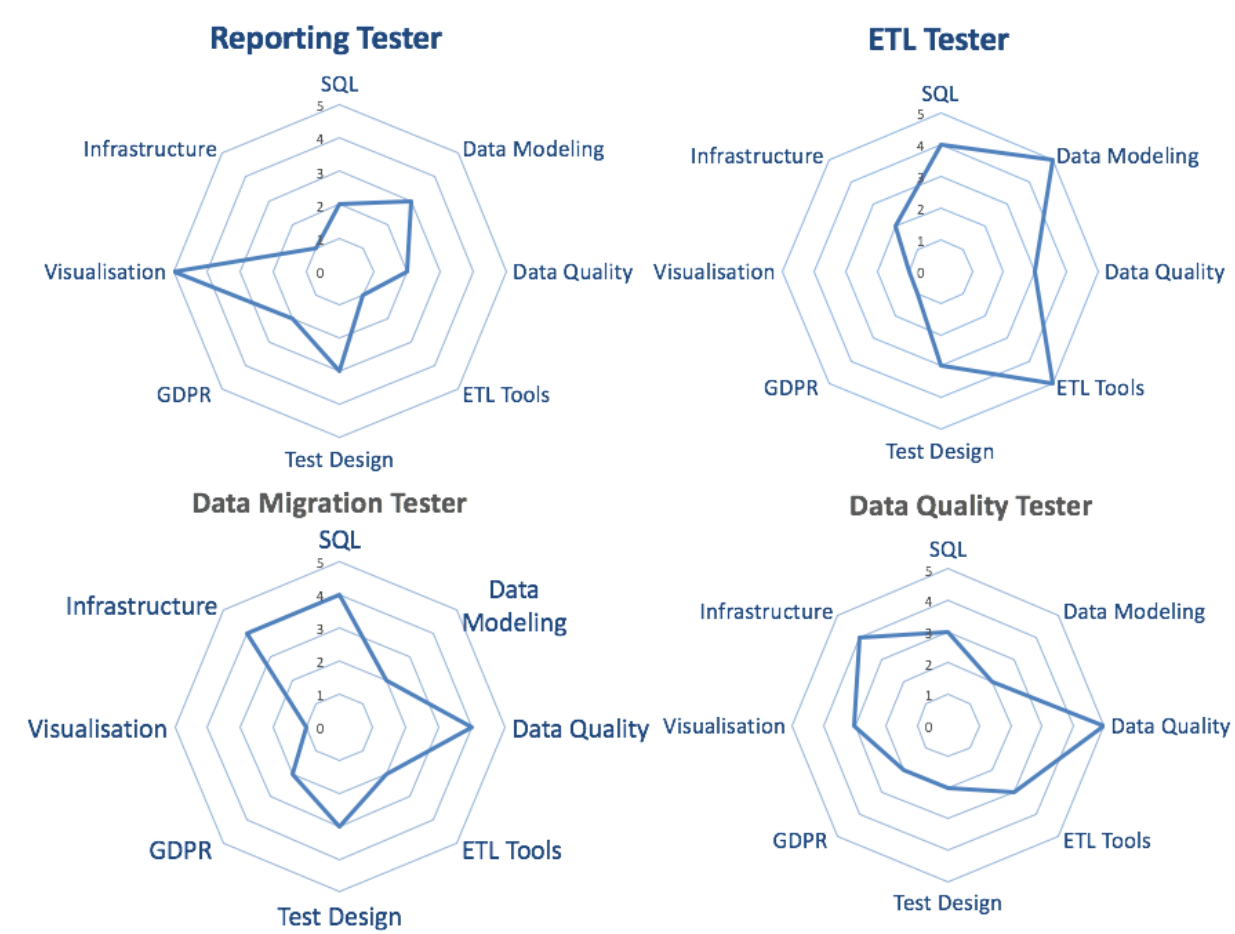


Illustration 4 - Testing roles and their specific skills (Source: DAU, 2020)

2.6 Setting up a test-strategy and approach in a BI environment (LO-2.14)

Because a BI & DA environment is composed of many different aspects, like transformation, complex infrastructure, reports and a diversity of database solutions, it is important for a test manager (or tester,

when there is no test manager) to think about setting up a test strategy and approach to have a good understanding of the different test aspects of testing in a BI & DA environment. You should set up a test strategy and approach based on the BI & DA environment. This should include an explanation of aspects like functional and nonfunctional testing, and other BI or DA aspects based on the specific areas in the Illustration 1. In this approach, elements of test techniques, quality attributes (non-functionals and data), what needs to be tested and to which extent (risk-based) testing will be executed.

3 Test Techniques

In this chapter, we will summarize the test techniques that can be useful for a tester to apply.

Keywords

Syntactic test, Semantic test, Equivalence partitioning, Boundary testing, Decision table test, Data combination test, Coverages

LO-3.1	K3	Explain the Syntactic test technique and be able to apply it
LO-3.2	K3	Explain the Semantic test technique and be able to apply it
LO-3.3	K3	Explain the Equivalence Class Partitioning test technique and be able to apply it
LO-3.4	K3	Explain Boundary Testing and be able to apply it
LO-3.5	K3	Explain the Decision Table test technique and be able to apply it
LO-3.6	K2	Explain the Data Combination test technique and be able to apply it
LO-3.7	K2	Explain why different coverages are required within test techniques

Several books explain test techniques in a very detailed way. Especially for experienced testers, this information can be common knowledge. From the perspective of a BI analyst or –consultant, this information could be rather new. For this last group, we refer to (Koomen[**TM1**]) and (ISTQB[**IS2**]) to have a broad understanding of how these techniques work and can help you in a BI & DA environment.

3.1 What are test techniques?

Test techniques are methods used to derive test cases from source documentation in a structured way. Test techniques are useful to help testers have a good understanding of the source documentation, find defects in the source documentation, and to set up test cases in a structured and most efficient way. Test techniques and combining techniques help testers setup and define test cases, so testing can be executed in the most sufficient and efficient way.

3.2 Different techniques (LO-3.1 to LO-3.7)

A lot of different techniques are available worldwide. There are several test methodologies, such as (Koomen[**TM1**]) and (ISTQB[**IS2**]), that describe some of those techniques, like:

- Syntactic Testing (SYN)
- Semantic Testing (SEM)
- Equivalence Partitioning (EP)
- Boundary Value Analysis (BVA)
- Decision Table Testing (DTT)
- Elementary Comparison Testing (ECT)
- Data Combination Test (DCT)
- Checklists-Based Testing (CHK)
- Pairwise Testing (PT)
- Experience-Based Techniques, like Error Guessing (EG) and Exploratory Testing (ET)
- Coverages ([**ST1**]), like:
 - Statement Coverage (SC),

- Decision Coverage (DC),
- Condition Decision Coverage (CDC),
- Multiple Condition Coverage (MCC).

For a detailed description and the usage of these techniques, we refer to (Koomen[**TM1**]), (ISTQB[**IS2**]), (Veenendaal[**TT1**]) and (Black[**ST1**]). It does not mean that other available techniques are not useful. The most important for a tester is to be knowledgeable of all kinds of techniques that are useful in a BI & DA environment and able to practice and apply these techniques. The techniques mentioned above are useful in all kinds of areas of software development. In chapter 4, the test techniques described could be useful in specific areas of a BI & DA environment.

3.3 Mapping techniques on a BI & DA environment (LO-3.1 to LO-3.7)

As Illustration 1 - A typical BI environment shows, there are different tests techniques defined on the areas of the model. The table below describes a common use of techniques applied to those different areas. This does not mean that, depending on the situation, techniques could be used on all other areas; this table rather focuses on the typical techniques used for a certain area:

Area	Typical test techniques used to test the area	Typical documentation used as source for testing
Databases (DWH, Staging, data marts, data lake)	Syntax, Semantic	functional specifications (of source systems), ERD, users' subject matter expertise (SME)
OLAP	Syntax, Semantic, Decision Tables, Coverages, Pairwise, Checklists, Business Process Test	Business processes, functional and technical specifications
Reports/Dashboards	Syntax, Semantic, Checklists, BVA, EP	Functional- and technical specification document, mockups, company standards
Data Models	Experience Based, Checklist	Logical Data Models (in modelling tools)
ETL	Syntactical, Semantic, Equivalence Partitioning, Boundary Value Analysis, Coverages, Pairwise, Decision Table	Business rules, Source Target Mapping
Internationalization/Localization	Syntax, Semantic, Checklist	Source Target Mapping, Functional Specification Document, Standards (e.g., ISO), SME
Data Quality	Data Profiling, Experience Based	(Production) Data, Data Dictionary, (company and industry) Standards, SME

Table 2- Mapping techniques on a BI & DA Environment

This table shows the typical techniques that are used and that can be used in different areas of the BI & DA environment. Coverages, for example, consists of a variety of different approaches that could result in less or more test cases. Depending on the risk of the subject area or the information that the data is used for, the tester can decide to use more extensive coverage. This table is therefore helpful in the process of setting up a risk-based test strategy and to define the testing approach.

4 BI Testing

This chapter defines the terminology used for testing in a BI & DA environment like reports, data mining, and testing transformations and how to apply test techniques.

Keywords

BI Reports, Queries, Dashboards, Self Service BI, OLAP, Data Modeling, Data Mining, Data environment, ETL, Transformations, Aggregations, Conversions, STM, Business rules, Completeness testing, Minus/Except, Intersect

LO-4.1	K2	Explain the testing of BI Reports and Dashboards, including its place in the data environment
LO-4.2	K1	Recall OLAP, Slicing and Dicing
LO-4.3	K1	Recall the description of Data Modeling and normalization
LO-4.4	K1	Recall the description of Data Mining, including Testers' Tasks
LO-4.5	K1	Recall common linguistic confusions among testers and data scientists
LO-4.6	K2	Recall the aspects of completeness testing
LO-4.7	K2	Recall common ways to compare datasets, by using visual aids, file comparison tools, database comparison tools and SQL statements
LO-4.8	K3	Explain Comparing Datasets using database SET operators MINUS/EXCEPT and INTERSECT and be able to apply them
LO-4.9	K2	Recall the aspects of transformation testing

In a BI & DA environment, test types **[IS1]** can be executed on several test levels (ISTQB**[IS1]**). Alongside the more traditional test types like performance etc., there are some specific ones in a BI & DA environment:

- Reports and Dashboard Testing
- Testing of OLAP cubes
- Testing of Models
- Testing in Data Mining
- Completeness Testing
- Transformation Testing
- E2E Testing

4.1 Reports and Dashboards Testing (LO-4.1)

In chapter 1, reports and dashboards are introduced as a way to visualize data, to provide insights to our business colleagues. When discussing testing of reports, think about queries (ad-hoc), reports, dashboards, self-service BI, OLAP (see section 4.2) and data mining (see section 4.4).

In general, testing reports and dashboards leans heavily on the use of checklists **[ST2]**. These checklists will focus, amongst others, on the following aspects:

- Access
 - Who is authorized to see the report and who is not? Is any Role-Based Access (RBA) applied?

- Can certain fields be viewed by some users only, for instance, based on his/her role? Should these values be (partially) scrambled or hidden?
 - Can the report be saved and shared with other users? With authorized users only or with anyone?
 - Is the report accessible (and still readable) on mobile devices?
 - Can the report be printed and is its lay-out as expected?
- Field mappings
 - Are all fields mapped to the right fields in the source (database)?
 - Are the naming conventions as expected?
 - Are the transformations and aggregations as expected?
 - Are the values logically sorted and/or according to the design?
 - How should records with NULL values be treated?
- Visuals
 - Is the lay-out (incl. font, color, etc.) as expected, when looking at company standards?
 - Are the widths of columns sufficient to store the longest names and greatest numbers in the source – and what if the numbers are totaled?
 - Are graphs implemented as expected, including the right colors, captions, and legends?
 - Can visuals be understood by the color blind?

Many of the aspects mentioned above can be considered as syntactic and semantic testing. Benefits of using checklists are, amongst others, its adaptability to the customer's situation and learnability: in general, they can be easily taught to and applied by new colleagues.

4.1.1 Data for testing Reports

Most reports will eventually get used by business users, so, for acceptance testing, it is important to provide test data that can be recognized by users and therefore is similar to production data. When the report is being developed on an existing Data Mart and personal data is involved, it may be best to have (a sample of) the Data Mart pseudonymized (see section 6.4 regarding privacy regulations).

If a data mart does not yet exist, it is recommended to have the report built upon a *virtual* data mart, created from (database views on) data in the source systems and/or the data warehouse. Because of the work required during this phase, performance may be poor, but it will allow for a great deal of the necessary functional testing. Performance testing will be performed in a later stage, on the real data mart. The development process of the virtual data mart itself will probably lead to finding design flaws – which can then be tackled early.

At the stage of systems testing, when testing technical aspects such as mappings and formatting, it may be best to define a small synthetic test set.

4.1.2 Query Skills for Testing Reports and Dashboards

Testing reports and dashboards requires skills relating to the presentation tool itself, and skills for querying the data in the underlying systems (e.g., the data mart). In general, as data marts are designed for “easy” querying, this requires some basic query skills like SQL. The tester has to be able to compare the presented values to the database entries. Usually, no complex JOIN operations need to be performed in this part of the DA environment.

4.1.3 Testing Ad Hoc Queries

The business will sometimes require an *immediate* answer based on company data that was not (yet) disclosed in an orderly fashion (e.g., through reports or dashboards). This is when Ad Hoc Queries may be used. Ad Hoc Queries are often (SQL) database queries written by developers, creating a one off product that provides the necessary data to the end user.

As the source systems may not be designed for reporting purposes, the result will depend greatly on the knowledge of the underlying data and the way source systems and business processes work. Also, as the query will often not be required to be future-proof, the coding will likely to be more complex (and harder to read). Extensive (SQL) coding skills from the developer and the tester(s) are thus needed.

Ad hoc queries can be tested on at least two levels. The technical level: is the query selecting the right data (using the correct joins, tables etc.). The functional level: is the query giving the right data that the user asks for (is it the right information delivered). During design, development and testing, frequent contact between the subject matter expert, developer and tester is heavily recommended.

4.2 Testing of OLAP Cubes (LO-4.2)

4.2.1 Drill Down, Drill Through, Slicing and Dicing

OLAP Software offers some common functionalities, that serve to drill down, drill through and to slice and dice the data. By comparing results of testing queries (by ways of SQL, MDX) with the presentation in the OLAP Cube, testers can validate that correct values are presented. Aggregations are particularly important validations: will regions, product groups, timelines be shown in a correct manner? What is the lowest grain of the values, and should this be shown? What is the result of renaming country names, what will happen when very large and very small sets exist? Are dates and timelines (SCD) and codes/decodes implemented in a correct manner? What are default values for NULL?

More information about this subject is available, for example, Kimball group provides a series of flaws that emerge in data modeling, amongst others: grain, codes, hierarchy handling, date dimension issues, surrogate key issues, SCD strategies **[KIM1]**. SQL skills needed to test OLAP Cubes will include basic aggregation skills such as SUM and GROUP BY functions.

4.2.2 Self Service BI

Tools that are used for OLAP and reporting belong to the group of the “Self Service BI (SSBI)” tools, i.e. tools that provide querying features for (business) end users. As these tools usually come as “off the shelf” products, testing skills for such software are required as well. Labels, meanings, and references in data sets for “SSBI” must be as transparent as possible for end users’ purposes: where developers and testers tend to read data definitions, many end users will not invest time in learning the exact details and differences between the presented data objects.

4.2.3 Performance Testing of OLAP Cubes

Copying data from relational databases to in-memory OLAP cubes, or loading the data, will take up time. The time needed may increase due to the number of users busy on the system, the amount of data being loaded and the time of day (e.g. when other batch processes are running). Therefore, performance tests

may be introduced to validate that the cubes can be loaded within the available time window. Questions may arise on *who* can load cubes (administrators, end users?) and *when* can cubes be loaded (batch window, ad hoc?). Also, requesting data from the cube will take up time. Testing response times of OLAP cubes resembles performance testing of (generic) front-end tools. For this purpose, often common test execution tools for front ends can be applied. Tools like Tableau even have built in performance measurement.

It is a good practice to have database administrators and performance engineers involved early, who assure the efficiency of the queries and database schemes. A lot of performance issues can thus be prevented.

4.3 Testing the Data Model (LO-4.3)

4.3.1 Testing the Data Model

Static testing of the data model may consist of reading design decisions and reviewing logical data models. It is important to recognize “implied” design rules when modeling rules are applicable, such as dimensional modeling¹⁰ and/or data vault modeling¹¹.

When Logical Data Models (LDM) are implemented on the actual infrastructure, designers and developers may choose to apply Physical Data Models that are different from the LDM. Effects of changes should be considered when writing test plans. Especially adding Indexes¹² and Partitions¹³ may influence performance of reading and writing the data.

Testers and designers may create test scenarios focused on the growth of the data sets and their performance in the near future (see section 4.3.2).

- When starting with an empty data warehouse, are empty tables allowed?
- Are scripts needed to create a valid starting point for loading the initial load and the following increments?
- How much time does the initial load take?
- Will the system provide information (warnings, error logs) when unexpected loading patterns and occur, for instance a “wrong loading order” of data sets?
- Who has access to this audit trail or logs?
- Does it hold any personal identifiable data (see section 6.4)?

4.3.2 Testing Implementation of History (Data Lineage)

Data lineage states where data comes from, where it is going, and which transformations are applied to it as it flows through multiple processes. It helps understand the data life cycle and its time aspects. It is one of the most critical pieces of information from a metadata management point of view (Allen[MA1]).

¹⁰ Dimensional Modeling is a data structure technique optimized for data storage in a Data warehouse

¹¹ Data vault modeling is a database modeling method that is designed to provide long-term historical storage of data coming in from multiple operational systems

¹² An index is a method to track the performance of some group of assets in a standardized way

¹³ A partition is a section of a storage device, such as a hard disk drive or solid-state drive

In both static and dynamic testing, testers should stress testing time aspects:

- When loading data into the data model, how will the order of loading data sets influence the results? Should a certain source be loaded before another? Can we already start with loading a certain table or should another loading job be completed, first?
- Are Dimension Tables loaded before the Fact Table?
- Are the correct SCD types (as documented in the Source-to-Target mappings) applied?
- Will starting and ending timelines be correct? Will data not “disappear” or be stored twice?
- Can some tables be changed periodically, for example, only once a month or year?
- When loading multiple sets at the same time, do timing issues occur?
- Can time zones negatively impact the loading strategy?
- Are time windows for loading data tailored to the customers’ situation? i.e. when the system starts loading at 12 PM, loading will take up 8 hours’ time, but first reports need to be run at 7 AM, what will happen?

4.4 Testing in Data Mining (LO-4.4, LO-4.5)

As stated in chapter 1, data mining is leaning heavily on the automated discovery of patterns in data and will require in-depth knowledge of statistics and data science.

The data science process partially resembles software development; in these regards, data and analytics testers can successfully contribute to the development of data mining deliverable, for example, by reviewing designs, challenging assumptions, testing the underlying infrastructure.

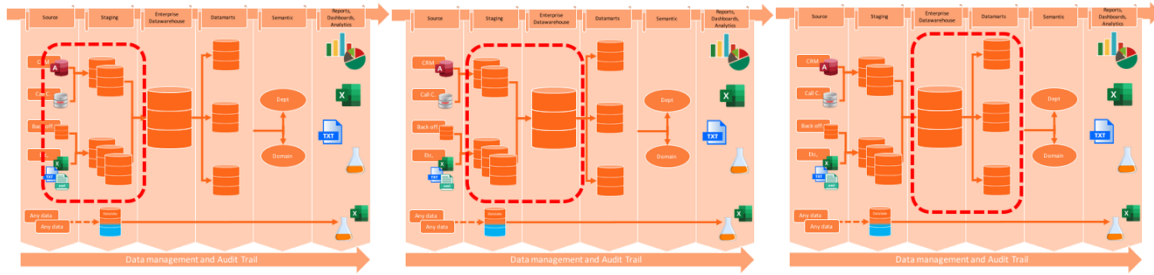
4.4.1 Beware of sector specific terminology

Due to different naming conventions in the world of data science and software development, linguistic confusion among testers, developers and data scientists may arise. In data science terminology, “performance of an algorithm” refers to the quality of its predictions, not to its speed. Also, a “test set” has a different meaning to a tester than to a data scientist. The latter will speak of training sets, validation sets and a test set or holdout data set, in the process of creating and testing the (machine learning) algorithms.

For a tester, regression will mean that errors are being introduced when new versions of the software are deployed. For a data scientist, regression is a statistical term, used to describe the strength of the relationship between two or more variables.

4.5 Completeness Testing (LO-4.6)

Completeness testing is required to ensure that all relevant records are transferred from the source to the target and that the contents (format, precision) of each record (row) and value (column) is still correct. Other ways to verify completeness may include checking the audit trail or executing other reconciliation procedures. This is relevant to both the ETL-process and the Reporting activities.



One of the methods for completeness testing is to count records, for instance, by the SQL `SELECT COUNT` statement. In most cases, this is definitely not sufficient: counting will only prove that the number of records is as expected, not that the values in the cells (records) are correct – or that data inadvertently shifted over records.

4.5.1 Common ways to compare data sets (LO-4.7)

To test the ETL process, data sets need to be compared frequently. Depending on the size (volume) and type of data, different approaches and tools may be applied. In general, testing needs to validate that all relevant rows and columns are being copied and that no data is either lost or wrongly duplicated.

4.5.1.1 File Comparison Tools

File comparison tools will support testers when loading two files and reporting in which rows or columns differences occur. For example, by visually showing this (on-screen) to the tester or by generating a log with differences. Using file comparison tools may become difficult in a big data environment, due to, for example, long loading times, insufficient memory space or lengthy (visual) analysis by human testers. For smaller data sets, like reference tables and sets up to a few hundred rows, using file comparison tools is recommended, as the tools are easily available.

4.5.1.2 Data Quality Tools featuring Comparison Functionalities

Data quality tools (see chapter 5) often provide features which enable file/table comparisons. Of course, this is not the only function of tools, they also inform their users about many aspects of the data quality, like its frequencies and “odd” situations in the data sets.

4.5.1.3 Database-oriented comparison

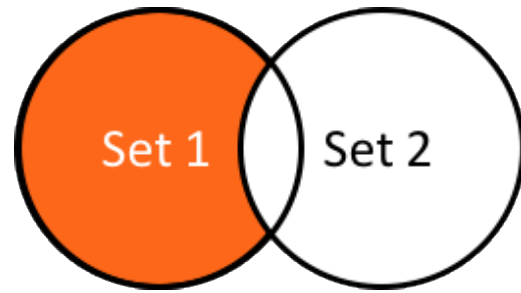
Data sets that are stored in (relational) databases can be compared using SQL statements, in order to verify that all the rows were copied, and all the content is in the expected place. The most basic SQL statements, unfortunately, often require lengthy and complex JOIN constructs. The time and the needed skills may not be sufficiently available. Some database systems will offer specific SET operators that render comparison easier (see section 4.5.2).

4.5.2 Compare datasets with database SET operators (LO-4.8)

Some database environments support the MINUS (or EXCEPT) and INTERSECT queries. Those queries (officially called SET operators) can facilitate a much easier comparison of source vs target tables. The process is used to execute a source MINUS target query and a target MINUS source query. Check whether there are any duplicates – and missing rows.

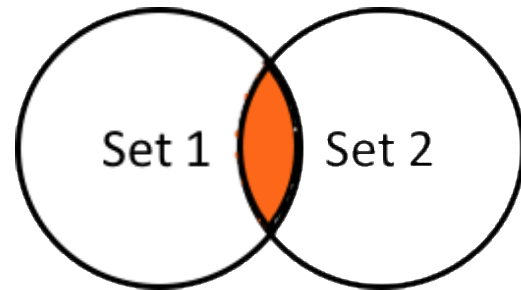
The MINUS operation combines the result of two SELECT statements and returns distinct rows which belong to the first set of results only. A MINUS or EXCEPT query looks as follows:

```
SELECT column1 [, column2 ]  
FROM table1  
[WHERE condition]  
EXCEPT  
SELECT column1 [, column2 ]  
FROM table2  
[WHERE condition]
```



The SQL INTERSECT clause/operator is used to combine two SELECT statements but returns rows only from the first SELECT statement that are identical to a row in the second SELECT statement. This means INTERSECT returns only common rows returned by the two SELECT statements. The basic syntax of INTERSECT is as follows:

```
SELECT column1 [, column2 ]  
FROM table1  
[WHERE condition]  
INTERSECT  
SELECT column1 [, column2 ]  
FROM table2  
[WHERE condition]
```



MINUS queries must conform to the following rules: first, the number of columns in the selected lists of both queries must be the same. Second, the order of the columns and their types must be comparable. With very large (broad) tables, it can be efficient to compare the most relevant columns only, for example, by creating VIEWS on the tables first, and then performing the MINUS queries on these VIEWS.

Note that in the standard MINUS query implementation, MINUS queries need to be executed twice (Source-to-Target and Target-to-Source). This doubles execution time and resource utilization. MINUS queries are processed on either the source or the target database, which can draw significantly on database resources (CPU, memory, and hard drive read/write). In most cases, as the environments have to be set up to provide for these kinds of queries, it should not be an issue. To make sure no issue arises, always contact your DBA in order to avoid performance issues.

4.6 Transformation Testing (LO-4.9)

In chapter 1, the Source-to-Target Mapping document was discussed, which contains a description of the required transformations. For transformation testing, it is important to understand the different forms of transformations that can be described in the STM.

4.6.1 Examples of Transformations

An example of a simple transformation in the ETL processes is the concatenation of values. Columns in the data mapping document are sometimes combined in a single target column. In this case, testers will need

to validate at least that this simple concatenation will function properly, even when NULLs or empty values exist in the source table and upper and lower case characters are used.

Aggregations can be filed under Transformations as well. Some common aggregations are:

- periodic values, for instance Hour > Day > Week > Month > Quarter > Year.
- geographical values, for instance City > Province > Country > World region.

Some other forms of transformation are decoding conversions for the integration of data, or easier readability for end users:

- M/F and 0,1,2 into: Male, Female, Unknown
- EUR, €, Euro into: EUR
- 1.000.000,= / 1,000,000.- into: 1000000,00
- 01-10-2012, 1-OCT-2012 into: 20121001000000.

A tester is required to have in-depth knowledge about the datatypes available in the database system and the data integration tools used. Look for tricky transformations and point out risky differences in character code settings. When designing test cases, data profiling can be performed in order to discover which actual values in the source tables exist, that need to be transformed in the design and the implementation of the ETL.

4.6.2 Applicable Test Techniques

Depending on the type of transformation performed, different techniques could be used. When data is transferred from column A in source to column B in target without any transformation, semantic testing will be a helpful technique. When there is a business rule with a transformation (for example source value "SALES" to target value "S"), then several techniques can be used, for instance Equivalence Partitioning, Boundary Value Analysis, Data Combination or Decision Table Test are all good techniques that can be applied. Which one you choose will depend on the level of risk identified during the risk analyses. You can expect that a decision table can result in more test cases than an equivalence partitioning. Also, the coverage techniques from (Koomen[**TM1**]) can be applied here.

4.7 E2E testing

E2E testing is an abbreviation of end to end testing. By this, we mean a test that starts at the absolute beginning of a process, the trigger or cause, until the processes get to the finished (or: end) state. This can be applied in a BI & DA environment as well. The process starts with extracting data from the source, flowing through the staging, data warehouse and data marts to the end result: reports, dashboards, files or external systems. E2E testing in a BI & DA environment can therefore be compared with a full integration test from start (extracting the data) up and until the end results. Therefore, this test is sometimes called an Integration Test as well. Please note that integration tests *between two components* also exist, and that this ambiguity can lead to issues when discussing what is being tested.

Typical aspects that are tested during an E2E test are infrastructure, interfaces, and consistency, traceability, and completeness of the data flows.

When using multiple data sources, setting up an E2E test may become laborious and error prone, as all parts of the tested landscape must be in sync, including the (test) data that is being used. Timing can

become an issue, e.g., when “customers” do not yet exist in one test system while they are already being referenced in another. This will require significant attention from business subject matter experts, test data specialists and engineers.

5 Data Quality

In this chapter, the terminology concerning Data Quality is defined and explained, as well as how to measure the level of Data Quality.

Keywords

Quality, ISO/IEC 25010, ISO/IEC 25012, data profiling, SQL, outlier detection, datatypes, big data

LO-5.1	K1	Understand that there are different definitions of and perspectives on Quality
LO-5.2	K2	Recall the main quality characteristics according to ISO/IEC 25010 and explain the shortcomings when it comes to Data Quality
LO-5.3	K2	Recall the definition of Data Quality according to DAMA and apply each of data quality characteristics to a given data set
LO-5.4	K1	Recall the definition and reasons of Data Profiling
LO-5.5	K1	Recall the most common (single column) data profiling tasks
LO-5.6	K1	Identify basic SQL statements for Data Profiling
LO-5.7	K3	Apply data profiling tasks to an example data set incl. using the correct terminology
LO-5.8	K1	Recall the benefits of using Data Profiling Software
LO-5.9	K1	Explain the impact of Big Data on the Data Profiling process

5.1 Quality (LO-5.1)

There are many definitions of quality; the following bulleted list contains four of them. To a certain extent, all of them explain what quality is, in their context. It is important to understand the different viewpoints of the different definitions, as this is also the authors' view on quality in the context of data.

- William Deming, a US-American engineer, statistician, professor, author, lecturer, and management consultant has defined quality as follows: *"Good quality means a predictable degree of uniformity and dependability with a quality standard suited to the customer"*.
- The American Society for Quality Control Standards Committee defined quality as: *"The totality of features and characteristics of a product or service that bear on its ability to satisfy stated or implied needs"*.
- The European International Standards Organization defined quality in the ISO/IEC 9000 as: *"The degree to which a set of inherent characteristics, of a product or service, fulfill requirements"*.
- Joseph Juran built theories and methods of Total-Quality Management and defined quality as: *"Fitness for use. Those product features which meet the needs of customers and thereby provide product satisfaction. Freedom from deficiencies"*.

5.2 Quality Characteristics (ISO/IEC 25010) (LO-5.2)

The international standard ISO/IEC 25010 [ISO] provides us with a model focused on the quality of computer systems and software. The standard describes product quality using the following eight main characteristics:

1. Functional suitability
2. Performance efficiency
3. Compatibility
4. Usability

5. Reliability
6. Security
7. Maintainability
8. Portability

Data and Analytics systems can be regarded as a form of computer systems and software too, so the characteristics will help us to analyze the (desired) quality level of the system. Depending on the focus, for instance extracting data from sources vs. reporting, other quality attributes (see section 5.3) can be more useful, see also section 2.4 on Risk Based Testing.

However, establishing the level of data quality may prove difficult, using the ISO/IEC 25010 model. Data quality may be seen as part of “Functional Suitability”, when considering its sub-characteristics: Functional Completeness, Correctness and Appropriateness. However, it does not (yet) provide us with sufficient means to measure data quality hands on, which forces Data Analytics professionals to explore other options for measuring data quality.

5.3 Data Quality Characteristics (ISO/IEC 25012) (LO-5.3)

In order to enable clear communication about data quality, we need to distinguish between various data quality characteristics. ISO/IEC provides this by means of the ISO/IEC 25012 in which the data quality model outlines fifteen DQ characteristics [I12] for data quality:

1. Accuracy
2. Completeness
3. Consistency
4. Credibility
5. Currentness
6. Accessibility
7. Compliance
8. Confidentiality
9. Efficiency
10. Precision
11. Traceability
12. Understandability
13. Availability
14. Portability
15. Recoverability

5.3.1 Testing of Data Quality

In order to provide a clear report on the data quality (and its influence on the system, End-to-end) it is important to select the most relevant characteristics and to define a level of required quality. This is typically done in the product requirement analysis (PRA) or shortly after the first analysis of source systems data.

Setting up measures to test the quality and executing these tests repetitively will help inform both business and development teams about potential problems that may be difficult to find afterwards, downstream, in the environment. Typically, after this phase, important choices will be made regarding the robustness of the system, such as to what extent the code will need to be able to process unexpected data and when

it is confronted with erroneous data, will it log such situations and continue processing or abort processing altogether? Depending on the choices made, this may lead to *adding* test cases with synthetic test data testing the robustness with “illegal” situations (as these faulty situations may not yet be available in the production systems).

With the ISO/IEC 25024 standard for measurement of data quality, the International Organization for Standardization provides measures including associated measurement methods and quality measure elements for the quality characteristics in the data quality model [I24].

5.4 Data Profiling (LO-5.4)

Data Profiling is the process of examining the data available in an existing data source [...] and collecting statistics and information about that data (Johnson[TJ1]).

Data profiling will provide important insights to designers, developers, and testers in the Data Analytics domain. Checking whether values really “never occur”, what are the actual domains in which values fall, what seem to be the character settings in source systems, etc. These will lead to better specifications and/or more and improved test cases.

Testers will apply data profiling in order to:

- “Shift Left”
- Facilitate better modeling
- Facilitate better data integration (ETL)
- Improve and extend test cases
- Improve coverage levels of tests
- Enable testing of error handling and audit trails

Other phases will benefit from data profiling, such as:

- Efficient database design
- Data clean(s)ing process
- Reverse engineering

5.4.1 Common Data Profiling Tasks (LO-5.5)

A summary of common data profiling tasks is presented below:

Category	Task	Description
Cardinality	Num-rows	Number of rows
	Lengths	Measurements of value length (min, max, median, average)
	Nulls	Percentage of NULL values
	Distinct	Number of distinct values (also called: Cardinality)
	Uniqueness	Distinct Values/Number of Rows
Distribution	Histogram	Frequency histograms
	Constancy	Most frequent value / number of rows
	n-tiles	Grouping numeric values in equal groups, e.g. Quartiles
	1st Digit	Check for Benford’s Law
Patterns, Data Types and Domains	Basic Type	Generic data type, e.g., numeric, alphabetic, alphanumeric, date/time
	Data Type	Database specific datatypes, e.g., VARCHAR, DATE, CHAR(10)
	Size	Maximum of digits in a numeric value (precision)
	Decimals	Maximum of decimals in decimal types (precision)
	Patterns	Value Patterns, e.g., “Aa9...”

	Class	Semantic, generic data type, e.g., Code, Indicator, Text, Date/Time
	Domain	Classification of a semantic domain, e.g., credit card, first name, city

Table 3 - Common Data profiling tasks

Data profiling will look into the way that data is distributed in the system. Data profiling in order to discover errors and questionable situations can be performed either by a designer or by a tester. The effect of a data profiling assessment is that design decisions can be adjusted, additional test cases can be created, and small but effective test sets can be derived from the complete (production) data set.

5.4.2 Basic SQL statements for Data Profiling (LO-5.6/5.7)

When performing data profiling on a relational database, basic and manual SQL statements can be used.

The following two examples show common SQL queries¹⁴:

```

/* INDIVIDUAL UNIQUE VALUES */
SELECT COUNT (Products) FROM TABLE T_Inventory
GROUP BY X;

```

```

/* PROPORTION OF UNIQUE VARIABLES */
SELECT CAST(COUNT(DISTINCT Products) AS DECIMAL) /COUNT(X)
FROM TABLE T_Inventory;

```

In these examples, only one variable/column at a time are being profiled – and only for its frequency. In the BI & DA environment, many data stores exist, each with dozens to hundreds of tables, and each table can contain many columns. Data profiling may become a very time-consuming task. BI tools may be used (as they often are available in the BI & DA environment) to help out, or dedicated data profiling tools may be introduced (see section 5.4.4).

As a tester, practical experience with some data profiling tasks or tools is useful to have a good understanding of how a tester can benefit from data profiling. Of course, manual querying on a database can help, some tools can speed up this process and give quick insights and information about databases. The information gathered with the tool can help the tester understand the data in the database, and support with developing test cases. Even with some basic data profiling knowledge, it is possible to identify potential problem areas with, for example, transformation rules.

5.4.3 Benefits of using Data Profiling Software (LO-5.8)

Using specific tools to perform data profiling activities has some advantages:

Less manual work required

Data profiling software can help render the data profiling process more efficient. Data profiling software can be used to automate certain profiling tasks, enabling the team to run data profiling processes whenever they are needed without becoming a burden to the team.

¹⁴ In these examples we use Microsoft SQL (T(ransact)SQL) and not Oracle SQL.

In-depth profiling

Whereas humans might be inclined to save time by skipping (time consuming) profiling actions, tools will run their program and provide all required data profiling information.

Repository of re-usable checks

Using Data Profiling suites will likely lead to a central repository of checks. These can be re-used in tests and even in code, not only in profiling activities. Having all the checks in one place is also beneficial for the quality of the checks themselves as these can be centrally reviewed whenever business rules are added or changed.

5.4.4 Impact of Big Data on Data Profiling (LO-5.9)

The rise of Big Data influences the process of Data profiling. Looking at the three Vs of Big Data (see: *Big Data – LO1.9*) the following impact is seen:

Volume

The sheer volume of data can become a challenge as systems may not be able to profile data that is arriving in real time and in large volumes. Profiling data that is stored in extremely large data stores may become ineffective and inefficient (the cost is too high, the results come in too late).

Velocity

If data is entering the system at high velocity, the overall profile of data may change faster than can be tracked by standard data profiling tools and principles.

Variety

Being used to profiling texts stored in columns and rows (tables) in relational databases, new ways of storing data leave us with challenges: the data itself is different (e.g. audio, video, x-rays) or formatted in a different way (e.g., XML, JSON). This forces us to be creative and find new ways to profile this data.

6 Environmental Needs

In this chapter, the challenges and complexity of using test data and multiple test environments for data & analytics testing is discussed, including privacy and standard compliance rules in the modern world.

Keywords

DTAP, test environment, test tools, compliance, data privacy, GDPR, ISO/IEC 27001, anonymization, pseudonymization, de-personification

LO-6.1	K1	Recall the function of a DTAP structure (in the SDLC)
LO-6.2	K2	Explain the necessity to align different DTAP environments (e.g., for visualization (BI), model development (analytics), data processing (data warehouse), infrastructure and source systems)
LO-6.3	K1	Recall the compliancy standard for Information Security ISO/IEC 27001
LO-6.4	K2	Recall and explain the necessity for separation of operational and development and test systems according to information security standards (ISO/IEC 27001)
LO-6.5	K2	Recall and explain the necessity to separate the operational and development and test systems according to data privacy law and regulations (GDPR)
LO-6.6	K1	Recall the different encryption methodologies for anonymization, de personification and pseudonymization
LO-6.7	K2	Explain the benefits and pitfalls of using production data vs test data

6.1 Test environments according to DTAP (LO-6.1)

In a phased approach for the software development lifecycle (SDLC), it is common to have separate environments that support the various goals of developing (coding), systems testing, acceptance testing and running production. To enable this, a multi-tiered environment needs to be set up and isolated.

DTAP is an approach used in software development and -testing that describes a 4-tier environment, where each letter of the acronym D-T-A-P stands for a specific environment:

Environment	Target Audience
Development	Developers
Test	Test Engineers
Acceptance	(Testing) End Users
Production	End Users

Table 4 - DTAP Environment and Target Audience

The set of environments used for a DTAP cycle is often called a DTAP street.

The Development (D) environment is mainly used by developers to create and (unit) test their own programs or components. Developers often use their own environment to test their own work (unit test).

When the developer reaches the exit criteria or the agreed minimum quality level, the software is transferred to the Test (T) environment, where independent testers often run their tests. This is mostly a standardized environment, that resembles more the actual production environment. In this environment, technical features may have been added to enable testing activities, such as stubs and drivers. Testers who work in a T-environment are mostly responsible for creating their own test data.

Once testers have reached the exit criteria or the agreed minimum quality level, the software will be transferred to the Acceptance (A) environment, in which acceptance testers such as business users (and sometimes the system manager/administrator) will test the product to validate that it is fit for purpose and get confidence that it meets their expectations. Some may refer to this environment as the QA (Quality Assurance) environment.

After successful acceptance, testing the software will be transferred to the Production (P) environment. P it is mainly used by the business and functional maintenance and is hardly ever used for testing activities.

The actual number of environments can be greater, depending on the need for developers and testers. For instance, a development team may use multiple development environments in parallel, all deploying to one “main” testing environment.

Variations

Some customers will decide to use a 3-tier environment, e.d. Development, Quality Assurance (QA) and Production instead of a 4-tier environment. Others may want to extend the DTAP street with dedicated environments for experimental (Lab) or education (Training) purposes.

6.2 Multidimensional DTAP in the Data Analytics landscape (LO-6.2)

A Data Analytics landscape is not only about coding of ETL processes and providing information products such as dashboards. A Data Analytics environment also depends heavily on the data flowing through the data analytics landscape and on the technical infrastructure it is running on. Taking these facts into account results in a so-called multi-dimensional DTAP. It is important to structure and manage this properly and not underestimate this task.

Data sources

Data in the Data Analytics landscape is flowing in from various source environments (or: “primary systems”). Those sources will probably have their own DTAP environments. Linking the DTAP of a data analytics environment to the DTAP’s of these primary systems will require making arrangements between parties and configuring interfaces.

Infrastructure

We need to build, configure and test the hard- and software on which the Data Analytics landscape is deployed. This infrastructure will probably have its own DTAP environment. You may think of physical assets such as servers and networking technology. Database Management Systems (DBMSs) and software for data scientists and business analysts will also be upgraded or replaced on a regular basis.

Data Models

When data is stored in tables, the specifications of these tables will change over time. This may require its own DTAP-street. Teams that are developing end user products (e.g. reports) based on the structure and content of these tables will need to “plug in” on the correct version of the tables.

Information Products

Data scientists and business analysts will be developing solutions for end users in their own version of the DTAP-street.

Combining the various DTAP-environments leads to a complex matrix or multi-dimensional DTAP:

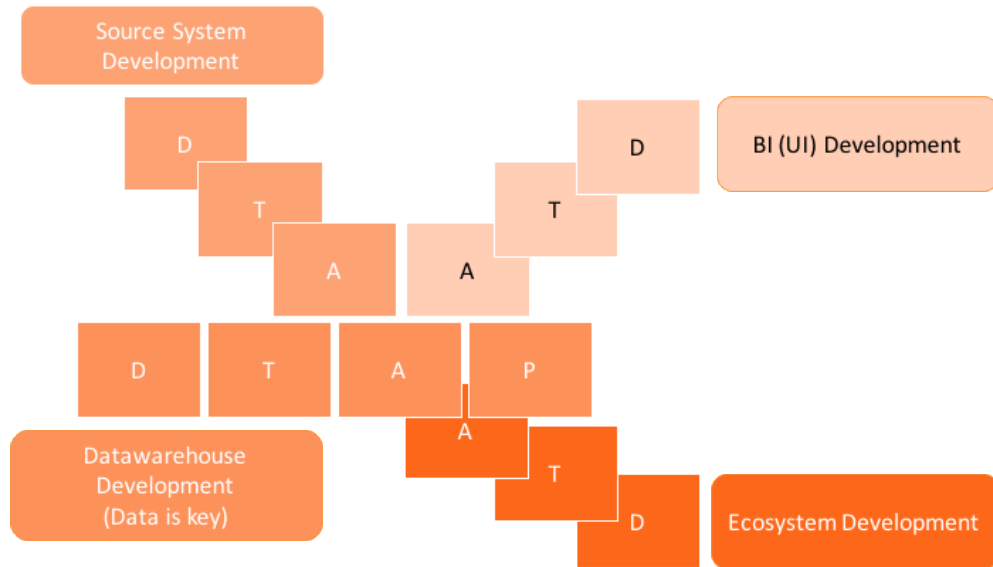


Illustration 5- A typical BI multi dimensional test environment

The complexity depends on the scope, organization, project, etc. There are a lot of factors that need to be considered and, when the complexity is such that all the test environments need to be integrated to each other, a central coordination activity would be a good suggestion.

Some aspects that are very important to consider and are arguments to have a structured approach and central coordination:

- Identical date/time on all environments,
- Data synchronization,
- Time travelling,
- Availability,
- Privacy,
- Identity & Access,
- Synchronization or anonymization.

Besides, depending on the needs for tests, different requirements for the test environment are applicable. The following are some examples of the specific important areas which can result in specific requirements:

- Data Processing Test
 - Full or partial copies of source data (profiling),
 - Need for synthetic data,
 - Access to or copies of source systems/target systems,
 - Need for the right identity & access.
- Model development and testing
 - Statistics and predictive modelling: need for (full) production data sets.
- Reports testing
 - Need for “real-life” data sets,

- Development of reports separate from data mart development,
 - Availability of various devices, in use by end users,
 - Need for the right identity & access.
- Infrastructure testing
 - Hardware vs. software,
 - Upgrading operating systems and analytical tools,
 - Storage, processing, networking bandwidth,
 - Need for bare metal and/or virtual systems for multiple teams,
 - Local (on premises) vs. cloud-based infrastructure
 - Availability of various end user devices.

A data and analytics infrastructure often requires substantial investments and, for instance, data warehouse appliances are not bought in numbers, like generic hardware (e.g. blades), virtual systems or cost-effective and scalable cloud-based solutions. Therefore, the DTAP strategy needs to be addressed as soon as the first investment plans are made for the BI & DA environment and the test (infrastructure) management needs to be consulted.

6.3 Impact of Security Standards on Analytics Testing (LO-6.3/6.4)

Testing (as well as development and change management) require clearly written information security policies. The security standard ISO/IEC 27001 emphasizes clear rules and policies for the handling of information assets and the engineering process, without stating how this is to be achieved by the organization. ISO/IEC 27002 provides ways to comply with these requirements.

First, organizations must clarify which data is sensitive and how to handle it. Second, they must explain how the organization buys or creates *secure* software in a *secure* development area. Third, they must state how they deploy software into production without any risk for production stability.

Under ISO/IEC 27001 the organizations need to enforce the policies and create evidence for auditing purposes. In other words, a supporting organization needs to be established and employees need to be continuously educated and motivated to act accordingly (Haller[KH1]).

6.4 Impact of Privacy Regulations on Testing in the Analytics Environment (LO-6.5)

Privacy is stated as one of the fundamental rights in the Universal Declaration of Human Rights (United Nations). Privacy regulations, like the General Data Protection Regulation (EU) 2016/679 (GDPR) in Europe or the US Privacy Act in the USA, lead to new requirements to the (use of) the (test) environment. It is very likely that your analytics environment will store personal identifiable information (PII) which is subject to privacy laws and regulations.

These regulations also lead to additional tests that need to be performed:

- The “right to be forgotten” gives individuals the right to ask organizations to delete their personal data throughout the whole the BI & DA environment, including data warehouses,
- User logging and monitoring: who is using the (personal identifiable) data?
- Encryption techniques.

6.4.1 The Right to be Forgotten – Deleting Data

According to European privacy regulations, data processing parties need to purge (delete) any data that is no longer needed –when a customer puts in an explicit request to purge their data, it is referred to as “the right to be forgotten”. This applies to both the development process and the production stage.

Make sure that all the (personal identifiable) information that was used during the development of analytical products is successfully deleted. If development is outsourced, it is recommended to incorporate this obligation into any contracts between the parties involved.

During the production stage, data needs to be deleted when it is no longer needed by the organization or when a customer requires its deletion. Correct and *complete* deletion of data from production systems, including test systems and records, backups, and data warehouse tables, needs to be tested.

6.4.2 User Logging and Monitoring

In OLTP systems, it is common (and often required) to have systems in place that record and monitor the activities of its users, in order to have a proper audit trail. Because of the nature of the data in the BI & DA environment (often corporate-wide and sensitive data), this requirement is applicable here as well. The activities of (end) users of the analytics products need to be recorded (logged) in order to enable their monitoring and analysis by privacy and/or data protection officers. This will require additional tests to be designed and executed, focused on logging and monitoring.

6.5 Encryption methodologies for Anonymization and Pseudonymization (LO-6.6)

For testers as well as for most data scientists, there is often no need to identify individual subjects in the data by using personal identifiable information, such as email addresses, social security numbers, credit card details and full face photographs. According to privacy regulations, this data should not be made available, and the data should be de-identified. By encrypting or removing personally identifiable information from data sets, the people whom the data describe remain anonymous.

6.5.1 Anonymization

Data anonymization has been defined as “a process by which personal data is irreversibly altered in such a way that a data subject can no longer be identified directly or indirectly (for instance by using additional data sources)”. In environments such as data warehouses, where a lot of data is available, there is a risk that subjects can be *indirectly* identified– so extra caution is required.

6.5.2 Pseudonymization

The process of obscuring data with the ability to re-identify it later is also called pseudonymization. By contrast to anonymization, pseudonymization is defined as “the processing of personal data in such a way that the data can no longer be attributed to a specific data subject without the use of additional information.” (Article 4(5) of the GDPR).

Pseudonymization is a data management and de-identification procedure by which personally identifiable information fields within a data record are replaced by one or more artificial identifiers, or pseudonyms.

A single pseudonym for each replaced field or collection of replaced fields makes the data record less identifiable while remaining suitable for data analysis and data processing.

Testers will need to inquire if measures need to be taken for anonymization/pseudonymization in the BI & DA environment (for themselves and/or for their co-developers and analysts) and will need to establish the correct operation of the encryption techniques. In case a tester has doubts or is not sure about privacy regulations or anonymization and pseudonymization with regards to testing, they should contact the respective data protection officer (DPO) as they will provide the necessary information.

6.5.3 Complexity of anonymization and pseudonymization in a BI & DA environment

Be aware that, in general, anonymization and pseudonymization can already be very hard to achieve in the right way in a traditional development project. In a BI & DA environment though, having to deal with many sources, both internal and external, requiring consistency between the sources, anonymization and pseudonymization can become very hard. Data sets are related to one another, and keys need to match. Testers in a BI & DA environment will require, depending on objectives, consistent data sets with data that are consistent to a certain degree. Creating such data sets is a highly knowledgeable and time-consuming task.

6.5.4 Privacy-preserving Data Mining

In generic data-oriented environments, rules and regulations stress that the data holding PII should be processed in such a way that it cannot be attributed to a specific data subject without using additional information. In data mining, technology is actively linking data (from several sources) and discovering links between them, creating the risk of re-identifying subjects. In order to counter this, algorithms and techniques were established under the name of Privacy-Preserving Data Mining (PPDM) that enable data mining without breaching data privacy. The moment of actively applying these techniques can vary between the time of collection (e.g. by adding noise) and the moment of presentation (for instance, reducing accuracy, suppression, and generalization) (Menders[MEN]).

6.6 Common Pitfalls using Production Data (LO-6.7)

Testing Analytics systems using production data sets have their pros and cons. Obvious advantages are the availability of data and its “recognizability” to the business user. Production data may show the existence of situations that “should not occur according to the design”. It does contain, however, the following disadvantages, among others:

Processing Time

At the stage of development and testing functionality, it is recommended to use small data sets that provide all designed test cases, so that it can be run fast and frequently, and because analyzing errors is easier when done on a small sample. When using (full) production sets, running time may cost many hours or even days. Therefore, this approach may be harmful to a fast development and testing process.

Little Variety

The production set may be large (volume), but its variety may be very limited. It may not at all be useful to test the processing of certain characters, error situations. A large part of the business rules may remain untouched.

Privacy Issues and/or Fines

Production data sets with personal identifiable information need to be prepared (by ways of anonymization/pseudonymization, see section 5.5.) for testing purposes. Analyzing data sets for personally identifiable information (PII) may take up precious time. On the other hand, if anonymization/pseudonymization is poorly executed, companies run the risk of non-compliance with rules and regulations, privacy violations (for instance GDPR or local equivalent regulations) and data breaches are severely fined.

7 References

Reference	Source
[BM1]	"The Pyramid Principle: Logic in Writing and Thinking (Financial Times Series)", Prentice Hall, ISBN 9780273710516, Barbara Minto (2008)
[BR1]	"Testing Embedded Software", Addison-Wesley, ISBN: 9780321159861, Broekman and Notenboom (2003).
[BT1]	"A Taxonomy for Learning, Teaching, and Assessing: A Revision of Bloom's Taxonomy of Educational Objectives", L. Anderson, P. W. Airasian, and D. R. Krathwohl (Allyn & Bacon 2001).
[CC1]	"Top Five Differences between DataWarehouses and Data Lakes", C. Campbell on Blue-Granite.com (November 2015).
[DA1]	"Competing on analytics", Harvard Business School Press, EAN: 9781633693722, T.H. Davenport and J.G. Harris (2017).
[DF1]	"Why The 3V's Are Not Sufficient To Describe Big Data", M. van Rijmenam on Datafloq.
[DG1]	"ETL Testing", Datagaps website, https://www.datagaps.com/data-testing-concepts/etl-testing/ .
[DL1]	"Data Vault Series 1 – Data Vault Overview", D. Linstedt (2002), https://tdan.com/data-vault-series-1-data-vault-overview/5054 .
[EB1]	"SmarTEST", Egbert Bouman (2008).
[GUR]	"Star Vs Snowflake Schema: Key Differences", on Guru99.com, https://www.guru99.com/star-snowflake-data-warehousing.html#4 .
[ISO]	ISO/IEC 25010:2011, Systems and software engineering — Systems and software Quality Requirements and Evaluation (SQuaRE) — System and software quality models
[I12]	ISO/IEC 25012:2008, Software engineering — Software product Quality Requirements and Evaluation (SQuaRE) — Data quality model
[I24]	ISO/IEC Standard 25024:2015 - Systems and software engineering — Systems and software Quality Requirements and Evaluation (SQuaRE) — Measurement of data quality
[IS1]	ISTQB Glossary Testing (v3.1, 2018).
[IS2]	ISTQB CTFL Syllabus 2018 v3.1.
[JL1]	"Agile Testing, A practical guide for testers and agile teams", L. Crispin and J. Gregory (2009)
[JL2]	"More Agile Testing, Learning Journeys for the Whole Team", L. Crispin and J. Gregory (2014).
[KIM1]	"Fistful of flaws", M. Ross (2003), https://www.kimballgroup.com/2003/10/fistful-of-flaws/ . Link accessed 7-JUN-2020.
[KH1]	"Testing experience", K. Haller (2014).
[LR1]	"Anchor Modeling in the data warehouse", L. Rönneck (2007) (www.anchor modeling.com).
[MA1]	"Multi-Domain Master Data Management", DOI, ISBN 978-0-12-800835-5, M. Allen and D. Cervo (2015)
[MEN]	"Privacy-Preserving Data Mining: Methods, Metrics and Applications", Mendes, Ricardo & Vilela, Joao. (2017). IEEE Access. PP. 1-1. 10.1109/ACCESS.2017.2706947.
[PM1]	"Language proficiency in English", Agarwal publication, ISBN 9789385872280, Madan Poonam (2016–2017), p138.
[RH]	"International Business Communication Standards", Rolf Hichert, Jürgen Faisst, IBCS Version 1.1, (2017)
[SF1]	"Show me the numbers", Stephen Few, 2nd edition, (2012).
[SG1]	"The Scrum guide", K. Schwaber and J. Sutherland (2017).
[ST1]	"Foundation of Software Testing-ISTQB Certification", Cengage Learning, ISBN-10: 1473764793, D. Graham, R. Black and E. van Veenendaal (2019).
[ST2]	"Testen von Data-Warehouse- und Business-Intelligence-Systemen", Stauffer et al (2013), pages 82 and 192.
[TD1]	"Enterprise Analytics", Thomas H. Davenport (2013), pages 12-14.
[TJ1]	"Encyclopedia of Database Systems, chapter Data Profiling", Springer, ISBN:978-0-387-49616-0, T. Johnson (2009).
[TM1]	"TMAP Next for result driven testing", Sogeti, ISBN: 9789072194961, Koomen, Aalst, Broekman and Vroon (2014).
[TT1]	"Test Techniques for the Test Analyst", E. van Veenendaal (2018).