

The Linked Data Visualization Model

Josep Maria Brunetti
GRIHO, Universitat de Lleida
josepmbrunetti@diei.udl.cat
<http://griho.udl.cat/>

Sören Auer
AKSW, Universität Leipzig
auer@uni-leipzig.de
<http://aksw.org/>

Roberto García
GRIHO, Universitat de Lleida
rgarcia@diei.udl.cat
<http://griho.udl.cat/>

Jakub Klímek
Charles University, Prague
klimek@ksi.mff.cuni.cz
<http://xrg.cz/>

Martin Nečaský
Charles University, Prague
necasky@ksi.ms.mff.cuni.cz
<http://xrg.cz/>

ABSTRACT

In the last years, the amount of semantic data available in the Web has increased dramatically. The potential of this vast amount of data is enormous but in most cases it is very difficult for users to explore and use this data, especially for those without experience with Semantic Web technologies. Applying information visualization techniques to the Semantic Web helps users to easily explore large amounts of data and interact with them. In this article we devise a formal Linked Data Visualization Model (LDVM), which allows to *dynamically* connect data with visualizations. In order to achieve such flexibility and a high degree of automation the LDVM is based on a visualization workflow incorporating analytical extraction and visual abstraction steps. We report about our comprehensive implementation of the LDVM comprising a library of generic visualizations that enable both users and data analysts to get an overview on, visualize and explore the Data Web and perform detailed analyses on Linked Data.

Categories and Subject Descriptors

H.5.2 [User interfaces]: GUIs, Interaction styles

Keywords

Semantic Web, Linked Data, Visualization, Interaction

1. INTRODUCTION

In the last years, the amount of semantic data available on the Web has increased dramatically, especially thanks to initiatives like Linked Open Data (LOD). The potential of this vast amount of data is enormous but in most cases it is very difficult *and* cumbersome for users to visualize, explore and use this data, especially for lay-users [9] without experience with Semantic Web technologies. Visualizing and interacting with Linked Data is an issue that has been recognized from the beginning of the Semantic Web (cf. [12]). Ap-

plying information visualization techniques to the Semantic Web helps users to explore large amounts of data and interact with them. The main objectives of information visualization are to transform and present data into a visual representation, in such a way that users can obtain a better understanding of the data [6]. Visualizations are useful for obtaining an overview of the datasets, their main types, properties and the relationships between them.

Compared to prior information visualization strategies, we have a unique opportunity on the Data Web. The unified RDF data model being prevalent on the Data Web enables us to bind data to visualizations in an *unforeseen* and *dynamic* way. An information visualization technique requires certain data structures to be present. When we can derive and generate these data structures automatically from reused vocabularies or semantic representations, we are able to realize a largely automatic visualization workflow. Ultimately, we aim to realize an ecosystem of data extractions and visualizations, which can be bound together in a dynamic and unforeseen way. This will enable users to explore datasets even if the publisher of the data does not provide any exploration or visualization means. Most existing work related to visualizing RDF is focused on concrete domains and concrete datatypes.

The *Linked Data Visualization Model* (LDVM) we propose in this paper allows to connect different datasets with different visualizations in a dynamic way. In order to achieve such flexibility and a high degree of automation the LDVM is based on a visualization workflow incorporating analytical extraction and visual abstraction steps. Each of the visualization workflow steps comprises a number of transformation operators, which can be defined in a declarative way. As a result, the LDVM balances between flexibility of visualization options and efficiency of implementation or configuration. Ultimately we aim to contribute with the LDVM to the creation of an ecosystem of data publication and data visualization approaches, which can co-exist and evolve independently. Our main contributions are in particular:

1. The adoption of the Data State Reference Model [7] for the RDF data model through the creation of a formal Linked Data Visualization Model, which allows to *dynamically* connect data with visualizations.
2. A comprehensive, scalable implementation of the Linked

Data Visualization Model comprising a library of generic data extractions and visualizations that allow to obtain an overview on, visualize and explore the Data Web in real-time.

3. An evaluation of our LDVM implementation using a benchmark consisting of 16 visualization tasks, which demonstrates that the breath of visualization possibilities can be substantially increased without users being required to write source code or have extensive domain knowledge.

The remainder of this article is organized as follows: Section 2 discusses related work. Section 3 introduces the Linked Data Visualization Model. Section 4 describes an implementation of the model and Section 5 presents its evaluation with different datasets and visualizations. Finally, Section 6 contains conclusions and future work.

2. RELATED WORK

Related work can be roughly classified into tools supporting Linked Data visualization and exploration as well as data visualization models in general.

Linked Data Visualization and Exploration. Exploring and visualizing Linked Data is a problem that has been addressed by several projects. Dadzie and Rowe [9] present the most exhaustive and comprehensive survey to date of existing approaches to visualising and exploring Linked Data. They conclude that most of the tools are designed only for tech-users and do not provide overviews on the data.

Linked Data browsers such as *Tabulator* [2] or *Explorator* [1] allow users to navigate the graph structures and usually display property-value pairs in tables. They provide a view of a subject, or a set of subjects and their properties, but not any additional support getting a broader view of the dataset being explored. *Rhizomer* [4] provides an overview of the datasets and allows to interact with data through Information Architecture components such as navigation menus, breadcrumbs and facets. It also provides visualizations such as maps and timelines. However, it needs to precompute some aggregated values. *DERI Pipes*¹ is an engine and graphical environment for general Web Data transformations and Mashup. Although it provides an end user GUI for the user to enter parameter values and browse the results, it is not intended for lay-users and requires SW expertise.

Graph-based tools such as *Fenfire* [13], *RDF-Gravity*², *IsaViz*³ provide node-link visualizations of the datasets and the relationships between them. Although this approach can help obtaining a better understanding of the data structure, in some cases graph visualization does not scale well to large datasets [11]. Sometimes the result is a complex graph difficult to manage and understand [14]. There are also *JavaScript libraries* for visualizing RDF. *Sgvizler*⁴ renders the results of SPARQL queries into HTML visualizations such as charts, maps, treemaps, etc. However, it requires SPARQL knowledge in order to create RDF visualizations. *Exhibit*⁵ helps users creating interactive sites with advanced text searching and filtering functionality. It also provides

Tool	Overview	Detail View	Automation	Lay users	Specific visualizations	Collaboration
Tabulator	-	✓	-	✓	M, T	-
Explorator	-	✓	✓	-	-	-
DERI Pipes	-	✓	-	-	-	✓
Fresnel	-	✓	-	-	-	-
Exhibit	-	✓	-	✓	M, T, C	-
Fenfire	-	✓	-	-	G	-
Sgvizler	✓	✓	-	-	C, G, T, M, R	-
Rhizomer	✓	✓	✓	✓	M, T, C	-
LODVis	✓	✓	✓	✓	R, M, E, B	-
Payola	✓	✓	-	-	C, E, O	✓

Table 1: Comparison of generic Linked Data visualization approaches. (Specific visualizations: M - Map, T - Timeline, C - Chart, B - Bubble Chart, O - Circles, R - Treemap, G - Graph, E - Tree.)

visualizations such as maps, timelines or charts. However, it requires a domain-expert to configure the different facets. Moreover, the data must to be stored in a concrete JSON structure and can not work directly on SPARQL endpoints. Other tools are restricted to visualizing and browsing concrete domains, e.g. *LinkedGeoData browser* [16] or *map4rdf*⁶ for spatial data or *FoaF Explorer* for FOAF profiles. Table 1 shows a summary of generic Linked Data visualization approaches. Most existing tools make it difficult for non-technical users to explore linked data or are restricted to concrete domains. Very few of them provide generic visualizations for RDF data combined with high-automation and overview visualizations. Consequently, it is still difficult for end users to obtain an overview of datasets, comprehend what kind of structures and resources are available and what properties resources typically have and how they are mostly related with each other.

Data Visualization. Regarding data visualization, most of the existing work is focused on categorizing visualization techniques and systems [10]. For example, [5] and [8] propose different taxonomies of visualization techniques. However, most visualization techniques are only compatible with concrete domains or data types [6]. Chi’s *Data State Reference Model* [7] defines the visualization process in a generic way. It describes a process for transforming raw data into a concrete visualization by defining four data stages as well as a set of data transformations and operators. This framework provides a conceptual model that allows to identify all the components in the visualization process. In Section 3 we describe how this model serves as a starting point for our Linked Data Visualization Model.

3. LINKED DATA VISUALIZATION MODEL

In this section we present the Linked Data Visualization Model (LDVM) by first giving an overview and then formalizing the key elements of the model.

3.1 Overview of LDVM

We use the *Data State Reference Model* (DSRM) proposed by Chi [7] as a conceptual framework for our *Linked Data Visualization Model* (LDVM). DSRM describes the visualization process in a generic way. Our LDVM is an adaptation of this generic model for the specifics of the visualization of RDF and Linked Data. The main difference is that in certain parts, LDVM works solely with RDF data model for increased automation while DSRM is generic in each of its

¹<http://pipes.deri.org/>

²<http://semweb.salzburgresearch.at/apps/rdf-gravity>

³<http://www.w3.org/2001/11/IsaViz>

⁴<http://code.google.com/p/sgvizler/>

⁵<http://simile-widgets.org/exhibit3/>

⁶<http://oegdev.dia.fi.upm.es/map4rdf/>

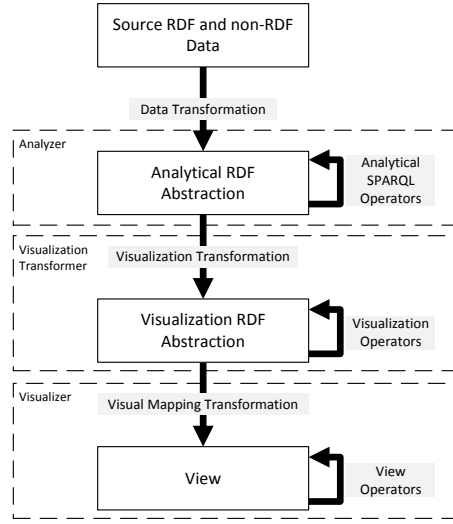


Figure 1: High level LDVM overview.

parts and does not constraint the applied data models. We also extend DSRM with three additional concepts – *analyzers*, *transformers* and *visualizers*. They denote reusable software components which can be chained to form an LDVM instance. Figure 1 shows an overview of the LDVM.

LDVM resembles a pipeline starting with raw source data (not necessarily RDF) and results with a visualization of the source data. It is organized into four stages that source data needs to pass through:

1. *Source RDF and non-RDF data*: raw data which can be RDF or adhering to other data models and formats (e.g. XML, CSV) as well as semi-structured or even non-structured data (e.g. HTML pages or raw text).
2. *Analytical abstraction*: extraction and representation of relevant data in RDF obtained from source data.
3. *Visualization abstraction*: preparation of an RDF data structure required by a particular visualization technique; the data structure is based on generic data types for visual analysis proposed by Shneiderman [15] (i.e., 1D, 2D, 3D or multi-dimensional data, temporal data (as a special case of 1D), tree data, or network data)
4. *View*: creation of a visualization for the end user.

Data is propagated through the LDVM pipeline from one stage to another by applying three types of transformation operators:

1. *Data transformation*: transforms the raw data represented in a source data model or format into a representation in the RDF data model; the result forms the base for creating the analytical RDF abstraction.
2. *Visualization transformation*: transforms the obtained analytical abstraction into a visualization abstraction.
3. *Visual mapping transformation*: maps the visualization abstraction data structure to a concrete visual structure on the screen using a particular visualization technique specified using a set of parameters.

There are also operators within the stages that allow for in-stage data transformations:

1. *Analytical SPARQL operators*: transform the output of the data transformation to the final analytical ab-

- straction (e.g. aggregations, enrichment from LOD).
2. *Visualization operators*: further refine the visualization abstraction data structure (e.g., its condensation if it is too large for transparent visualization).
3. *View operators*: allow a user to interact with the view (e.g., rotate, scale, zoom, etc.).

3.2 LDVM Stages

Let us now look at each stage of LDVM in more detail.

Source RDF and non-RDF Data Stage. The first stage considers RDF as well as non-RDF data sources. This is important because many data sources on the current Web do not (yet) provide data represented in the RDF data model. The data transformation transforms the source data to an RDF representation which forms a base for creating an analytical abstraction. In case of RDF data sources, the transformation can be a simple identity mapping. If the source RDF data does not have a suitable structure for the following analysis, the transformation can be a sequence of one or more SPARQL queries which map the source data to the required structure. In case of non-RDF data sources, it is a transformation that converts the source data model to the RDF data model. For example, it can transform a given CSV file to an RDF representation. Since data extraction is a vast research field on its own, we will not consider non-RDF data sources in this paper and refer the reader to [17] for a survey on knowledge extraction approaches.

Analytical RDF Abstraction Stage. The output of the second stage (*analytical RDF abstraction*) is produced by applying a sequence of various in-stage analytical SPARQL operators on the RDF output produced by the data transformation. We call the sequence *analyzer* (see Figure 1). Our goal is to enable users to reuse existing analyzers (possibly created by other users) for analyzing different datasets. We want to enable a user to find analyzers which can be applied for analyzing a given data set and, vice versa, to find datasets which may be analyzed by a given analyzer. Therefore, it is necessary to be able to decide whether an analyzer can be applied on a given dataset, i.e. whether the analyzer is *compatible with the dataset*. We formalize the notion of compatibility later in Section 3.4.

Visualization Abstraction Stage. An analytical abstraction is not a suitable data structure for visualization. We want to ensure that a visualization tool is reusable for different analytical abstractions. Building specific visualizers for particular analytical abstractions would not enable such reuse. This is because each visualization tool visualizes particular generic characteristics captured by the analytical abstraction. For example, there can be a tool which visualizes tree structures using the TreeMap technique or another tool which visualizes the same structures using the SunBurst technique. And, another tool may visualize 2-dimensional structures on Google Maps. An analytical abstraction may contain encoded both tree structure as well as 2-dimensional structure. All three mentioned tools can be applied to the analytical abstraction as well as on any other abstraction which contains the same structures encoded. Therefore, we need to transform the analytical abstraction into a more generic data structure based on Generic Visualization Data Types (GVDTs), listed in Table 2. The GVDTs are inspired

RDF Vocabulary	Data Type	Visualization Tool
xsd:int, dc:subject,... (count)	1D	Histogram
wgs84:lat, geo:point,...	2D	Map
visko:3DPointPlot,...	3D	3D Rendering
qb:Observation, scovo:Item,...	Multidimensional	Chart
xsd:date, ical:dtstart,...	Temporal	Timeline, Calendar,...
rdfs:subClassOf, skos:narrower,...	Tree	Treemap, SunBurst,...
foaf:knows,...	Network	Graph,...

Table 2: Generic visualization data types.

by Shneiderman [15] data types. This structure is then what is visualized by visualization tools. This transformation is ensured by the visualization abstraction stage of LDVM. Its output (*visualization abstraction*) is produced by a component which performs visualization transformation followed by a set of in-stage visualization operators. We call it *visualization transformer*. To facilitate associating visualization tools to analytical abstractions, we provide mappings from GVDTs to visualization tools but also to RDF vocabularies. This way, it is possible to associate input analytical RDF abstractions to GVDTs and then provide or recommend the more suitable visualization tools to deal with them. The mappings are summarized in Table 2. For example, an observation modelled using the the Data Cube Vocabulary (QB) for dimensional data and part of an analytical abstraction can be modelled using the generic Multidimensional datatype and then visualised by any charting tool, which will be based on this visualization abstraction. A visualization transformer can be reused for extracting visualization abstractions from various analytical abstractions. The only requirement is that *the transformer must be compatible with the analytical abstraction*. This is similar to the compatibility of analyzers. It is formalized later in Section 3.4.

View Stage. The output of the last stage (*view*) is produced by a component called *visualizer*. A view is a visual representation of a visualization abstraction on the screen. A visualizer performs *visual mapping transformation* which may be configured by a user using various parameters (e.g., selection of visualization technique supported by a visualizer, colors, shapes, proportions, etc.). The user can also manipulate the final view using the view in-stage operators (e.g. zooming, moving, etc.). A visualizer can be reused for visualizing various visualization abstractions which use GVDTs supported by the visualizer. In other words, *the visualizer must be compatible with a visualization abstraction*. Again, it is formalized in the same way as for visualization transformers and analyzers in Section 3.4. Let us note that the introduced analyzer, visualization transformers and visualizers can all exploit the benefits of Linked Data. As we have already mentioned, it is necessary to ensure compatibility between analyzers and RDF datasets, between visualization transformers and analytical RDF abstractions, and between visualizers and visualization RDF abstractions. As we show later in Section 3.4, the notion of compatibility is based on ontologies. For example, an analytical RDF abstraction adhering to a given ontology may be processed only by a visualization transformer which supports this ontology. However, Linked Data enables to exploit, e.g., equivalence mappings between ontologies. Therefore, a visualization transformer is able to process all analytical abstractions adhering to the supported ontology or any other ontology mapped to the supported one. Linked Data also enables analyzers, visualization transformers and even visualizers to enrich the pro-

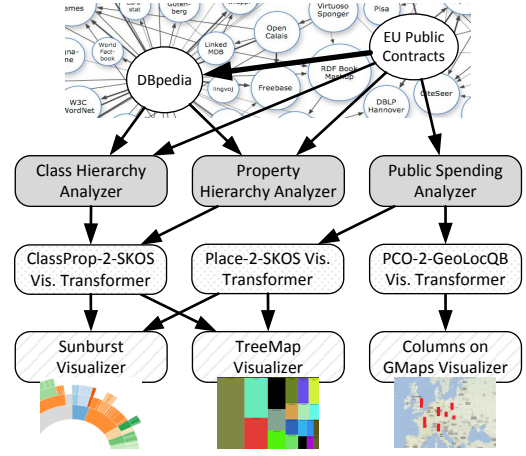


Figure 2: Sample analyzers and visualizers using LOD.

cessed data with other data from the LOD cloud. For example, a visualization transformer may exploit links to the LOD cloud to get GPS coordinates of the processed geographical objects.

3.3 Sample LDVM Pipeline Instance

Figure 2 shows several sample instances of LDVM. There are three analyzers available to users. Two of them can work with any data source and extract its class hierarchy or property hierarchy (expressed with `rdfs:subClassOf` or `rdfs:subPropertyOf` properties, respectively). The third one, *public spending analyzer*, analyzes the amount of money spent on public contracts of public authorities grouped by EU regions at various levels (countries, regions, municipalities etc.). It requires the data to be represented according to the *Public Contracts Ontology*⁷. The output contains EU regions represented as instances of class `s:AdministrativeArea` from `http://schema.org` name space. Regions are organized into a hierarchy using `s:containedIn`.

Further, there are three visualization transformers. *ClassProp-2-SKOS Visualization Transformer* transforms a class or property hierarchy to a visualization abstraction adhering to tree GVDT modeled with SKOS. Each node of the tree is associated with the name of the class or property, respectively, and the number of instances. *Place-2-SKOS Visualization Transformer* transforms a hierarchy of places expressed using `http://schema.org` constructs to tree GVDT modeled using SKOS as well. Each node of the tree is enriched with the name of the respective region and the money spent in the region. The last *PCO-2-GeoLocQB Visualization Transformer* transforms the analytical abstraction containing EU regions each complemented with the measured money spent to 2-dimensional GVDT modeled with QB. The dimensions are GPS longitude and latitude of the capital cities of each region. The items of the 2-dimensional structure are the capital cities with their names and the money spent. It is not necessary that the GPS coordinates are in the analytical abstraction. The transformer enriches the abstraction with GPS coordinates from the LOD cloud (DBpedia or other

⁷<http://purl.org/procurement/public-contracts#>

source) using links which must be present.

And, finally, there are three different visualizers. *Sunburst Visualizer* shows a tree GVDT modeled using SKOS as a sunburst view. Similarly, *TreeMap Visualizer* shows a tree GVDT modeled with SKOS as a TreeMap view. *Columns on GMaps Visualizer* shows a 2-dimensional GVDT modeled with QB with two dimensions representing GPS coordinates on a Google Map. Each 2-dimensional item with a label and number is represented as a column placed on the given coordinates. The height of the column depends on the number. The column is enriched with the label.

Various instances of LDVM pipelines can be created using these analyzers, visualization transformers, and visualizers. A sample instance may combine *Class Hierarchy Analyzer*, *ClassProp-2-SKOS Visualization Transformer*, and *TreeMap Visualizer* and can be applied on *DBpedia*. The data flows through the particular LDVM stages are:

- *Source Data*: raw RDF statements from *DBpedia* including the *DBpedia* ontology.
- *Analytical RDF Abstraction*: RDF statements stating a label, children and number of instances for each class in the *DBpedia* ontology.
- *Visual RDF Abstraction*: RDF statements adhering to SKOS representing a tree GVDT of *DBpedia* classes and their hierarchy.
- *View*: a TreeMap visualization of the *DBpedia* class hierarchy created by the visualizer.

The data is transformed between various stages as follows:

- *Data Transformation*: since we consider an RDF data source this transformation is an identity function.
- *Visualization Transformation*: transforms the class hierarchy to a SKOS representation.
- *Visual Mapping Transformation*: a function that maps the SKOS hierarchy to a TreeMap

Finally, the following in-stage operators are applied within each stage:

- *Analytical RDF Abstraction*: SPARQL queries which compute the numbers of instances of each class and represent the numbers as RDF statements with classes as subjects (the SPARQL queries form the analyzer).
- *Visual Abstraction Operators*: if there are more sibling nodes with 0 assigned, they are condensed to a single node.
- *View*: zoom in, zoom out.

3.4 Formalization

As we have discussed in the previous section, the core concept of LDVM are reusable components, i.e. analyzers, visualization transformers, and visualizers. They are software components which consume RDF data via their input interfaces. An *analyzer* consumes source RDF data. A *visualization transformer* consumes an analytical RDF abstraction. A *visualizer* consumes a visualization RDF abstraction. The goal of this section is to formally introduce the concept of compatibility of a component with its input RDF data. This formalization can then be easily applied on analyzers, visualization transformers, and visualizers. Given the formalization, we are then able to decide whether a given analyzer can be applied on a given RDF dataset. Similarly, we can decide whether a visualization transformer can be applied on a given analytical abstraction, etc.

Our approach is based on the idea to describe the expected input of a LDVM component with an *input signature*. The signature comprises an ontology, which describes the classes and properties expected in the processed RDF data, and a set of SPARQL queries which further restrict the data a component is able to process. The input signature is then compatible with the RDF data when the ontology of the signature is semantically compatible with the ontology of the RDF data and the SPARQL queries of the signature are evaluated on the RDF data to a non-empty result. Our approach is, intentionally, simple. As we show later, we define the semantic compatibility of two ontologies only on the base of simple equivalence and sub-type-of mappings between their classes and properties. Our rationale is to provide a simple and lightweight solution, which allows to resolve the compatibility without complex reasoning. In addition, we allow to specify more complex rules of compatibility using SPARQL queries defined by the input signature. While the ontologies can be used to resolve the static compatibility at design-time of a LDVM pipeline instance, SPARQL queries can be evaluated at run-time.

To define the concept of compatibility based on input signatures we first introduce a formal representation of an ontology. We do not require the definition to cover all aspects of ontologies formally. We only need to know the classes and properties proposed by the ontology and, optionally, RDF statements which define their equivalence and sub-type-of mappings to classes and properties of other ontologies.

DEFINITION 1 (ONTOLOGY). *An ontology $\mathcal{O}$ is a triple $(\mathcal{C}, \mathcal{P}, \mathcal{M})$ where $\mathcal{C}$ is a set of classes, $\mathcal{P}$ is a set of predicates, and $\mathcal{M}$ is a set of RDF statements (mappings of the classes and properties to other ontologies). Both classes and predicates are specified using their unique URIs.*

EXAMPLE 1 (SAMPLE ONTOLOGY). *Let us demonstrate the definition on two simple ontologies $\mathcal{O}_{pc}$ and $\mathcal{O}_{vcard}$. $\mathcal{O}_{pc} = (\{pc:Location\}, \{geo:long, geo:lat\}, \mathcal{M}_{pc})$. $\mathcal{M}_{pc}$ contains the following statement:*

- *$pc:Location$ $rdfs:subClassOf$ $geo:SpatialThing$*
As we can see, the ontology contains a class for a geographical location and two properties representing GPS coordinates. Moreover, the ontology contains an equivalence mapping of its class to a class in another ontology. $\mathcal{O}_{vcard} = (\{vcard:Location\}, \{vcard:longitude, vcard:latitude\}, \mathcal{M}_v)$ where $\mathcal{M}_v$ contains following statements:
- *$vcard:Location$ $rdfs:equivalentClass$ $geo:SpatialThing$*
- *$vcard:longitude$ $rdfs:equivalentProperty$ $geo:long$*
- *$vcard:latitude$ $rdfs:equivalentProperty$ $geo:lat$*

This ontology definition allows us to define the concept of compatibility of two given ontologies. We first need to define compatibility of two classes and compatibility of two properties.

DEFINITION 2 (CLASS COMPATIBILITY). *Let $\mathcal{O}_s = (\mathcal{C}_s, \mathcal{P}_s, \mathcal{M}_s)$ and $\mathcal{O}_t = (\mathcal{C}_t, \mathcal{P}_t, \mathcal{M}_t)$ be two ontologies. Let*

$\mathcal{M} = \mathcal{M}_s \cup \mathcal{M}_t$. Let $\mathcal{C} = \mathcal{C}_s \cup \mathcal{C}_t \cup \{\text{classes used as subjects or objects in statements from } \mathcal{M}\}$. A class $C_t \in \mathcal{C}_t$ is compatible with a class $C_s \in \mathcal{C}_s$ iff

- $C_t = C_s$, OR
- $\exists C \in \mathcal{C}$ s.t. C_t is compatible with C and
 - $(C_s, \text{rdfs:equivalentClass}, C) \in \mathcal{M}$,
 - $(C, \text{rdfs:equivalentClass}, C_s) \in \mathcal{M}$, or
 - $(C_s, \text{rdfs:subClassOf}, C) \in \mathcal{M}$.

A class C_s is compatible with class C_t when they are the same class or there exists a chain of equivalence or sub-class-of mappings between both classes. Therefore, the relation being compatible with between two classes is reflexive and transitive but not symmetrical.

DEFINITION 3 (PROPERTY COMPATIBILITY). Let $\mathcal{O}_s = (\mathcal{C}_s, \mathcal{P}_s, \mathcal{M}_s)$ and $\mathcal{O}_t = (\mathcal{C}_t, \mathcal{P}_t, \mathcal{M}_t)$ be two ontologies. Let $\mathcal{O}$ denote the union of both. A property $P_t \in \mathcal{P}_t$ is compatible with a property $P_s \in \mathcal{P}_s$ iff

- $P_t = P_s$, OR
- $\exists P \in \mathcal{P}$ s.t. P_t is compatible with P and
 - $(P_s, \text{rdfs:equivalentProperty}, P) \in \mathcal{M}$,
 - $(P, \text{rdfs:equivalentProperty}, P_s) \in \mathcal{M}$, or
 - $(P_s, \text{rdfs:subPropertyOf}, P) \in \mathcal{M}$

Compatibility of properties is defined analogously to compatibility of classes. Again, the relationship is a reflexive and transitive but not symmetrical. Based on the compatibility of classes and properties we can define the compatibility of two ontologies. The definition says that an ontology $\mathcal{O}_t$ is compatible with $\mathcal{O}_s$ iff for each component of $\mathcal{O}_t$ (i.e. each class or property) there is a compatible property in $\mathcal{O}_s$.

DEFINITION 4 (ONTOLOGY COMPATIBILITY). An ontology $\mathcal{O}_t = (\mathcal{C}_t, \mathcal{P}_t, \mathcal{M}_t)$ is compatible with an ontology $\mathcal{O}_s = (\mathcal{C}_s, \mathcal{P}_s, \mathcal{M}_s)$ iff

- $(\forall C_t \in \mathcal{C}_t)(\exists C_s \in \mathcal{C}_s)(C_t \text{ is compatible with } C_s)$
- $(\forall P_t \in \mathcal{P}_t)(\exists P_s \in \mathcal{P}_s)(P_t \text{ is compatible with } P_s)$

As we can easily prove, again ontology compatibility is reflexive and transitive but not symmetrical.

EXAMPLE 2 (COMPATIBLE ONTOLOGIES). In our previous example, class `vcard:Location` from the ontology $\mathcal{O}_{vcard}$ is compatible with the class `pc:Location` from the ontology $\mathcal{O}_{pc}$. Also, properties `vcard:longitude` and `vcard:latitude` from $\mathcal{O}_{vcard}$ are compatible with `geo:long` and `geo:lat` from $\mathcal{O}_{pc}$, respectively. Therefore, according to the previous definition, the ontology $\mathcal{O}_{vcard}$ is compatible with $\mathcal{O}_{pc}$.

We further need to formalize a dataset. Basically, a dataset comprises a set of RDF statements, an ontology which describes the structure of the RDF statements and, optionally, their semantics in a form of mappings to other ontologies.

DEFINITION 5 (DATASET). A dataset $\mathcal{DS}$ is a pair $(\mathcal{D}, \mathcal{O})$ where $\mathcal{D}$ is a set of RDF statements and $\mathcal{O}$ is an ontology with classes and properties used as types in $\mathcal{D}$.

EXAMPLE 3 (EU PUBLIC CONTRACTS DATASET). Note the datasets depicted in Figure 2. The dataset labeled EU

Public Contracts is formally expressed as a pair $\mathcal{DS}_{pc} = (\mathcal{D}_{pc}, \mathcal{O}_{pc})$. The ontology $\mathcal{O}_{pc}$ has been presented in a simplified way in previous examples. Note, that the original ontology is more comprehensive and complex. $\mathcal{D}_{pc}$ is a set of RDF statements (27×10^6) which represent the data about public contracts tendered in EU.

We now define *input signatures*. Each LDVM component has an input signature which comprises an ontology and a set of SPARQL queries. Both are optional. The ontology describes the basic required structure of the data which may be processed by the component. The SPARQL queries specify further restrictions on the data.

DEFINITION 6 (INPUT SIGNATURE AND COMPATIBILITY). An input signature $\mathcal{S}$ is a pair $(\mathcal{O}, \mathcal{Q})$ where $\mathcal{O}$ is an ontology and $\mathcal{Q}$ is a set of SPARQL queries. $\mathcal{O}$ specifies the basic required structure of the input RDF data. $\mathcal{Q}$ specify further more complex restrictions. We say that $\mathcal{S}$ is compatible with a dataset $\mathcal{DS} = (\mathcal{D}_{DS}, \mathcal{O}_{DS})$ iff

- $\mathcal{O}$ is compatible with $\mathcal{O}_{DS}$
- each Q in $\mathcal{Q}$ returns a non-empty result when executed on $\mathcal{D}_{DS}$

EXAMPLE 4 (ANALYZER COMPATIBILITY WITH A DATASET). Based on Figure 2 we demonstrate how compatibility of PCO-2-GeoLocQB Visualization Transformer with the analytical abstraction created by Public Spending Analyzer can be checked. The visualization transformer has an input signature $\mathcal{S}_{cgmap} = (\mathcal{O}_{vcard}, \mathcal{Q}_{cgmap})$. The analytical abstraction created by the analyzer on the base of the EU Public Contracts Dataset is a dataset $\mathcal{DS}_{spend} = (\mathcal{D}_{spend}, \mathcal{O}_{pc})$. To check the compatibility we first need to check the compatibility of $\mathcal{O}_{vcard}$ with $\mathcal{O}_{pc}$ (as performed in Example 2). Then we execute queries in $\mathcal{Q}_{cgmap}$ against $\mathcal{D}_{spend}$. The queries allow for checking a more detailed compatibility. Our sample query checks whether there are instances of `vcard:Location` in the dataset which have not only the values of `vcard:longitude` and `vcard:latitude` but also other properties which are the same for all the instances and have integer values, which can then be used for rendering by a visualizer.

```

1 SELECT ?property WHERE {
2   ?location a vcard:Location ;
3   vcard:longitude ?longitude ;
4   vcard:latitude ?latitude ;
5   ?property ?value .
6   FILTER (datatype(?value) = xsd:int)
7   { SELECT COUNT(?l) AS ?total WHERE
8     { ?l a vcard:Location } }
9 } GROUP BY ?property ?total
10 HAVING COUNT(?location)=MIN(?total)

```

If the query returns a non-empty result then the visualization transformer is compatible with the analytical abstraction created by the analyzer and the visualizer can be used to visualize the output of the analyzer.

Let us note that we could replace the ontology with corresponding SPARQL queries. However, ontologies allow us to resolve the basic compatibility statically by comparing both ontologies. SPARQL queries must be evaluated dynamically on the current data of the dataset. Considering only

SPARQL queries to specify the input signature would prevent us from resolving a part of the compatibility statically. This could be a discomfort for users who built a LDVM pipeline instance. For them, static compatibility evaluation is beneficial.

4. IMPLEMENTATION

Based on LDVM, we implemented a comprehensive two-level prototype. The two levels correspond to two levels of visualization detail from Shneiderman’s visual information seeking mantra: “*overview first, zoom and filter, then details on demand*” [15]. Figure 3 shows the general architecture of our LDVM implementation.

Step 1 corresponds to the LDVM *source RDF and non-RDF data* stage, in which a user can enter SPARQL endpoints and select graphs to visualize. The LDVM implementation may probe the selected SPARQL endpoints to determine which analyzers can be offered to the user in the next step. The probing is based on the formalism of compatibility introduced in Def. 6. In *step 2* the user selects an analyzer from the offered list. (S)he can also customize the analyzer or create a completely new one. *Step 3* performs the data analysis specified by the analyzer. The analysis is sent to the server, executed and the LDVM *analytical RDF abstraction* is returned. This involves querying of the specified SPARQL endpoints, which may involve SPARQL query caching mechanism for performance optimization, and synthesis of data gathered from the individual SPARQL endpoints involved. In *Step 4* the user chooses a visualizer to visualize the obtained analytical abstraction. First, the user chooses what structure encoded in the analytical abstraction should be visualized. This corresponds to choosing a suitable *visualization transformer*. Then, he chooses a suitable *visualizer* to visualize the *visualization abstraction* created by the transformer. Note that it is possible to incorporate the selection of the visualization transformer into the selection of the visualizer. In that case, the user chooses only the visualizer which already contains a suitable visualization transformer. This is not so flexible and reusable as if they were separated but it is more comfortable for the user in certain cases. The optional *step 5* involves optional customization of the specified visualization (e.g. shapes, colors, etc.). In other words, the user specifies the *visual mapping transformation* in this step. Our LDVM implementation allows for ontology-driven visual mapping transformation. For example, if we are visualizing data about public contracts, we may want to represent individual contracts (instances of class **Contract** from the Public Contracts Ontology) as file icons, etc. Finally, in *step 6* our LDVM implementation runs the chosen visualizer (i.e. creates a *view*) and presents it to the user. The user can then execute LDVM *view operators*, which correspond to zooming, scrolling etc. The arrow from step 6 to step 4 represents a possible change or customization of the used visualizer over the same data. The arrow from step 6 to step 2 represents a change of level of detail. This may be the zoom-in from Shneiderman’s mantra from the *overview* level to the *detail on demand* level. Also, the change of detail can be in the opposite direction, i.e. from a detailed view to its context in the overview. In our case, this means transition between our two parts of our LDVM implementation.

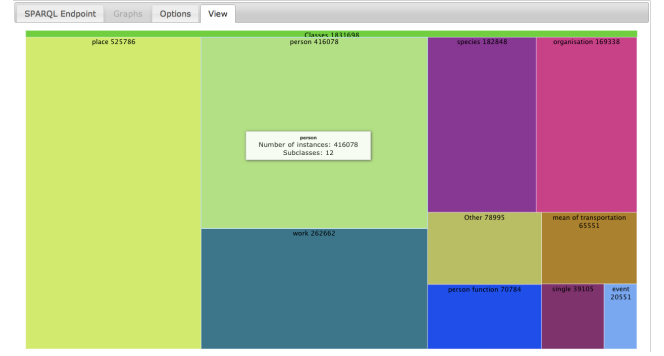


Figure 4: LDVM implementation LODVisualization showing the DBpedia class hierarchy as a treemap.

*LODVisualization*⁸ implements our architecture on the overview level. It allows its users to explore and view the Data Web through different visualizations on the *overview* level of detail. These visualizations allow users to obtain an overview of RDF datasets and realize what the data is about: their main types, properties, etc. *Payola*⁹ is a general framework for analysis and visualization of RDF data and in contrast to *LOD visualization* focuses on the *details on demand* part of Shneiderman’s mantra. It allows its users to specify detailed analyses of a selected dataset and inspect the results using various visualizations. We allow users to switch between the levels of detail as we pass the necessary information between the two parts. In the sequel, we describe the two parts of our prototype in greater detail.

Overview – LODVisualization. LODVisualization allows to connect different datasets, different data analysis and different visualizations according to the LDVM in a dynamic way. These visualizations allow users to obtain an overview of RDF datasets and realize what the data is about: their main types, properties, etc. In LODVisualization users can enter or select a SPARQL endpoint and select the graphs to visualize (first step of LDVM model). Then, the compatibility between analyzers and datasets is checked in order to determine which of them are available (second step). Once the analyzer has been executed, the results are stored into a visual abstraction, which corresponds to the third step of the model. Finally, users can visualize the results using different visualizers depending on their compatibility (fourth step). Our implementation includes analysis such as the class hierarchy, property hierarchy, SKOS concepts hierarchy, properties connecting two classes, etc. The results can be visualized using techniques such as tables, treemaps, charts, maps, etc. Since being based on the LDVM, LODVisualization is easy to extend with additional analysis and visualization techniques. Figure 4 shows a treemap with the DBpedia class hierarchy generated with LODVisualization. LODVisualization is developed using *Google App Engine* (GAE), a cloud computing platform for developing and hosting web applications on Google’s infrastructure. The frontend is developed using HTML, CSS and JavaScript, while the backend is implemented in Python. Most visual-

⁸<http://lodvisualization.appspot.com/>

⁹<http://live.payola.cz>

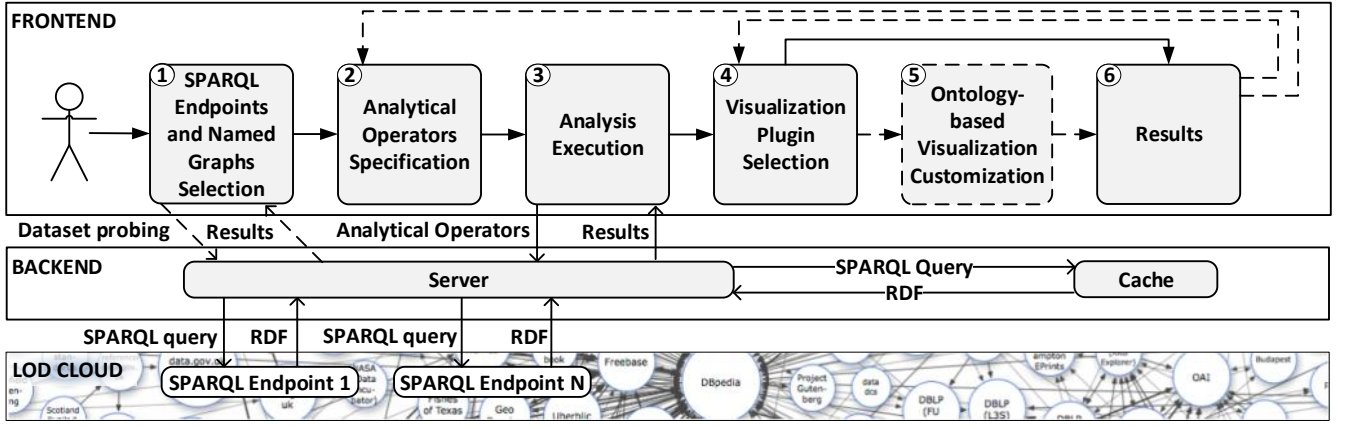


Figure 3: High-level Linked Data Visualization architecture.

izations are created using the *JavaScript Infovis Toolkit*¹⁰ and *D3.js* [3]. *Google Maps* and *OpenStreetMaps* are embedded for geospatial visualizations. Data is transferred between SPARQL endpoints, as well as server and client using JSON. LODVisualization is compatible with arbitrary SPARQL 1.1 compatible endpoints. The server provides a cache infrastructure implemented using GAE Datastore and Blobstore. The cache substantially increases performance, scalability and reliability of Linked Data visualizations by preventing the execution of the same queries when users aim to visualize the same data extraction with different visualization techniques. LODVisualization was integrated with Payola to provide users with easy transition between the two levels of visualization details. When exploring a dataset, there is a link provided that transfers all relevant information (SPARQL endpoint, graphs, classes, properties) to Payola and lets users continue analyzing data in Payola.

Details on demand – Payola. The workflow of Payola consists of two main parts. *First*, a user creates an analysis. An analysis starts with SPARQL endpoints selection and then consists of individual connected *analytical SPARQL operators*, which have an RDF graph on their input, do a transformation and send the transformed RDF graph to their output. An example can be seen on the left side of Figure 5. Payola is a framework, so the concrete operator functionality and usability is the responsibility of plugins creators. We created a library of default operators that represent typical parts of SPARQL queries. These include selection of resources of a certain type, selection of properties that are interesting to the user and filters for resources which satisfy SPARQL filter conditions. However, more complex operators are possible (e.g. an ontological filter which selects resources and properties being instances of a given ontology). These operators are integrated into SPARQL queries and executed on the selected SPARQL endpoints. Once an analysis is defined, it can be executed in a *second* part creating the analytical abstraction, which can be then visualized using different visualizers. The analytical abstraction in Payola is always an RDF graph. Concrete visualizations are again implemented via visualization plugins. We created a library of simple visualizer plugins that include various graph and chart visualizations. An example can be seen on

the right side of Figure 5. While the graph visualization is generically applicable, the chart visualization has certain requirements on the input RDF graph (i.e. LDVM visualization abstraction). It has to contain resources of a certain type, each one equipped with a label and property with a value. When the input graph meets this criteria, i.e. it is a visualization RDF abstraction compatible with the chart visualizer, it can be visualized as a chart. The vision for Payola is that for each problem domain, a domain expert creates custom LDVM analyzers for that domain and also custom LDVM visualizers appropriate for that domain. For the public procurement domain, for example, we created a visualization customization for the graph visualization for public contracts where individual contracts are visualized as green circles with a bobby pin inside. Payola allows users to download the resulting RDF data in their preferred RDF serialization. Payola enables users to share their visualizations and to embed them into external web pages. Payola is written in Scala and part of its build process is creation of the client code in HTML, CSS and JavaScript. It is compatible with arbitrary SPARQL endpoints. Payola was integrated with LODVisualization to allow an easy transition from the details on demand visualization level back to the overview level. Whenever a SPARQL endpoint and a graph is selected in Payola, a link to LODVisualization enables the user to explore the selected dataset on the overview level.

5. EVALUATION

Our evaluation consists of three parts. First, we evaluated our implementation with end users performing several tasks. Second, we have compared our prototype with the other approaches described in the Related Work section. Finally, we evaluated the performance of the critical overview tasks in LODVisualization.

User evaluation. We asked our test users to perform several tasks using our prototype LDVM implementation that consists of two levels of detail, *overview* implemented by LODVisualization and *details on demand* implemented by Payola. Expert users, who know the RDF and SPARQL, can create visualizations on both levels of detail and implement them in the tools directly. However, our goal was to show that Linked Data visualization does not have to require such expertise. This is consistent with our vision where expert

¹⁰<http://thejit.org/>

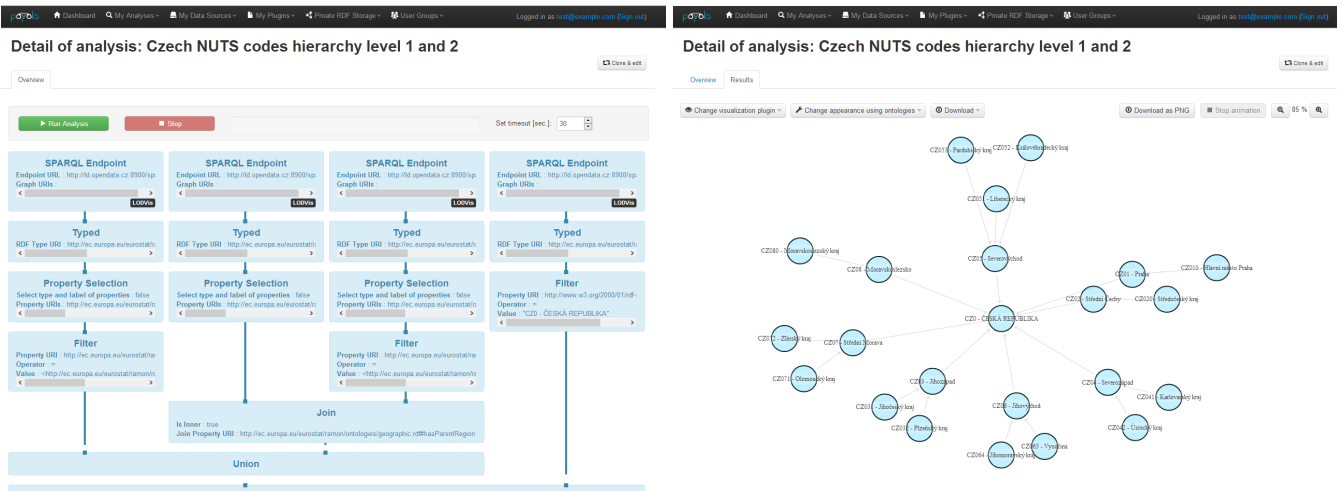


Figure 5: Payola screenshots: On the left analysis specification, on the right graph visualization of the analysis result.

users create complex configurable visualizations with user-friendly parameters, which are used as black boxes by lay users. A lay user should be able to use LDVM implementations without knowledge of programming and SPARQL. Therefore, we created some configurable sample analyzers and visualizers and we let lay users use and customize them to perform various tasks. We asked 16 users to assess the difficulty of the tasks on a scale from 0 – not able to solve, 1 – difficult to 5 – easy and to record the time required to complete each task. The vast majority of the users did not have any experience with the tools, was not instructed before the experiment and had no opportunity to interact with the LDVM team. The following scenarios and tasks were defined for the user evaluation.

Scenario 1: Generic tasks. First, we asked users to perform the following overview tasks with DBpedia using LOD-Visualization:

1. five most frequently used classes
2. five most frequently used properties
3. five least frequently used classes
4. five least frequently used properties
5. most generic classes in the class hierarchy
6. most generic concepts in the SKOS concept hierarchy
7. five instances of class `dbo:Village` with the highest indegree
8. five instances with the highest outdegree

Next, we prepared two more complex scenarios. A general one using DBpedia and a domain specific one regarding public procurement.

Scenario 2: DBpedia. The overview level tasks are:

1. Explore city and country classes in the class hierarchy.
2. Properties connecting cities and countries.
3. Person classes in the class hierarchy and all subclasses.
4. Properties connecting persons and cities.

The details on demand level tasks are:

5. Cities with a population of more than 2 million and their countries. Use the prepared analysis in Payola

and visualize using Gravity visualizer.

6. Cities in the Czech republic with a population of more than 50.000 and soccer players born there. A similar analysis is available in Payola showing cities in Germany with population more than one million and artists born there. Clone this analysis and change it accordingly. Use the circle visualizer.

Scenario 3: Public procurement. Here, only details on demand level tasks are performed:

1. Really expensive Czech public contracts (i.e. cost more than 20 billion CZK). Use the prepared analysis in Payola. Use the triple table visualizer.
2. Hierarchy of Germany NUTS regions of level 1 and 2. Clone and change a similar analysis in Payola for Czech NUTS regions. Use the Gravity visualizer.

The results are summarized in Figure 6. Users were not able to perform tasks 1.2, 1.4 and 1.6. They knew how to perform these tasks with LODVisualization but the SPARQL queries to obtain the property hierarchy and the SKOS hierarchy timed out. These analyses can be executed correctly over smaller datasets, but they do not scale yet to large datasets such as DBpedia. The average easiness rating for overview tasks is 3.67 and 2.82 for details on demand tasks. Complex tasks such as 2.6 and 3.2 were difficult for users and still required much time. However, tasks 2.5 and 3.1 were easier and quicker to solve. Overall, 12 out of the 16 partially quite complex tasks can be solved in reasonable time (only three task require more than 3 minutes).

Comparison with other tools. We compared our LDVM implementation with other tools described in the Related Work section. We determined whether it is possible to perform the visualization tasks using these tools. Table 3 shows the results. With Tabulator, Fenfire or Fresnel it is not possible to perform any of the tasks because they only allow to display a concrete resource or a set of resources and their properties. Explorator can perform details on demand tasks by combining subsets of resources and filtering them. However, it is not possible to perform any overview task. Some of the details on demand tasks can be performed with Exhibit

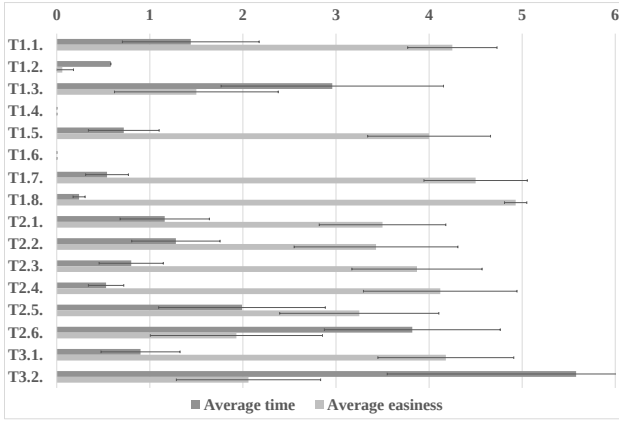


Figure 6: Average easiness (0-5, higher is better) and time in minutes (lower is better) including standard deviations for each task (N = 16)

and Rhizomer through their facet navigation. Rhizomer also provides class and SKOS concepts overviews that allow to perform some of the overview tasks. Finally, using Semantic Pipes or Sgvizler it is possible to perform all tasks, but SPARQL and programming skills as well as extensive domain knowledge are necessary. In summary, this evaluation shows, that our LDVM implementation allows users to solve the by far largest variety of visualization tasks without programming or domain knowledge.

Tool	1.1	1.2	1.3	1.4	1.5	1.6	1.7	1.8	2.1	2.2	2.3	2.4	2.5	2.6	3.1	3.2
Tabulator	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Explorer	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
D. Pipes	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Fresnel	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Exhibit	-	-	-	-	-	-	-	-	-	-	-	-	✓	-	✓	✓
Fenfire	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Sgvizler	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Rhizomer	✓	-	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	-	-	✓
LODVis	✓	-	✓	-	✓	-	✓	✓	✓	✓	✓	✓	-	-	-	-
Payola	-	-	-	-	-	-	-	-	-	-	-	-	✓	✓	✓	✓

Table 3: Visualization capability comparison with other tools.

The performance perspective. In order to demonstrate the feasibility to perform overview tasks in real-time, we evaluated the performance of our LODVisualization implementation with different datasets, data extractions and vi-

ID	Dataset	Data Extraction	Visualization Configuration	Execution time (s)	
				w/o cache	w/ cache
1	DBpedia	Class hierarchy	Treemap	3.58	0.87
2	DBpedia	Class hierarchy	Crop Circles	3.35	0.76
3	DBpedia	SKOS concept hierarchy	Treemap	48.79	16.96
4	DBpedia	Spatial data	Google Maps	2.05	0.37
5	Dbpedia	Spatial data	OpenStreetMaps	2.03	0.79
6	Wine Ontology	Property hierarchy	Indented Tree	1.95	0.57
7	Wine Ontology	Property hierarchy	Space Tree	2.45	0.70
8	LinkedMDB	Class hierarchy	Treemap	3.28	0.61
9	WWW 2011	Class hierarchy	Crop Circles	2.53	0.75
10	AGRIIS – FAO	SKOS concept hierarchy	Treemap	28.72	8.54

Table 4: Timings for 10 combinations of datasets, data extractions and visualization configurations.

sualizations. Table 4 shows a summary of the evaluation results. Timings for each concrete visualization were averaged from 10 execution cycles without cache and 10 execution cycles with cache. Despite using some of the largest datasets available on the Data Web, most of the visualizations can be generated in real-time (<5s rendering time) and the use of the cache further reduces the execution time substantially (<1s rendering time). Creating different visualizations for the same data extraction takes a similar time. This is due to the fact that most of the execution time can be attributed to the data transformation. The visual transformation timing depends on the size of the results to process. In the same way, the visual mapping transformation depends on the number of items to visualize. This number is particularly high in experiments #3 and #10, with a large hierarchy of SKOS concepts. However, it is important to note that the execution times depend mainly on the complexity of the SPARQL queries as well as on the availability of SPARQL endpoints or servers.

6. CONCLUSIONS AND FUTURE WORK

We presented the Linked Data Visualization Model (LDVM) that can be applied to rapidly create visualizations of RDF data. It allows to connect different datasets, different data extractions and different visualizations in a dynamic way. When applying this model, developers and designers can obtain a better understanding of the visualization process with data stages, transformations and operators. The LDVM offers user guidance on how to create visualizations for RDF data. We have implemented the model in two complementary applications which allow to create generic visualizations for RDF. In our implementations we provide visualizations that support the overview and details-on-demand tasks proposed by Shneiderman. These visualizations are useful for obtaining a broad view of the datasets, their main types, properties and the relationships between them.

The LDVM is the first step on a larger research agenda aiming at largely automatizing the visualization of semantic data. In future work we plan to focus on complementing the Linked Data Visualization Model with an ontology. A visualization ontology can help during the matching process between data and visualizations, capture the intermediate data structures that can be generated and choose the visualizations more suitable for each data structure. Regarding the implementations of the model, we plan to extend the library of data analyzers and visualizations. Having more data analyzers and visualizations will facilitate the creation of an ecosystem of data publication and data visualization approaches, which can co-exist and evolve independently. We aim to integrate these components into a general dashboard for visualizing and interacting with Linked Data through different visualizations as well as other exploration and authoring components.

Acknowledgments

We would like to thank our colleagues from the GRIHO, XRG and AKSW research groups for their helpful comments and inspiring discussions during the development of LDVM. This work was partially supported by a grant from the European Union’s 7th Framework Programme provided for the project LOD2 (GA no. 257943).

7. REFERENCES

- [1] Experimenting with explorer: a direct manipulation generic rdf browser and querying tool. In *WS on Visual Interfaces to the Social and the Semantic Web (VISSW2009)*.
- [2] T. Berners-Lee, Y. Chen, L. Chilton, D. Connolly, R. Dhanaraj, J. Hollenbach, A. Lerer, and D. Sheets. Tabulator: Exploring and analyzing linked data on the semantic web. In *3rd Int. Semantic Web User Interaction WS*, 2006.
- [3] M. Bostock, V. Ogievetsky, and J. Heer. D3: Data-driven documents. *IEEE Trans. Visualization & Comp. Graphics (Proc. InfoVis)*, 2011.
- [4] J. Brunetti, R. Gil, and R. Garcia. Facets and pivoting for flexible and usable linked data exploration. Crete, Greece, May 2012.
- [5] S. K. Card and J. Mackinlay. The structure of the information visualization design space. In *IEEE Symp. on Information Visualization*, INFOVIS '97.
- [6] S. K. Card, J. D. Mackinlay, and B. Shneiderman. *Readings in Information Visualization: Using Vision to Think*. Academic Press, London, 1999.
- [7] E. H. Chi. A taxonomy of visualization techniques using the data state reference model. In *IEEE Symposium on Information Visualization 2000*, INFOVIS '00, Washington, DC, USA, 2000. IEEE.
- [8] E. H. H. Chi, P. Barry, J. Riedl, and J. Konstan. A spreadsheet approach to information visualization. In *IEEE Symposium on Information Visualization '97*.
- [9] A.-S. Dadzie and M. Rowe. Approaches to visualising linked data. *Semantic Web*, 2(2):89–124, 2011.
- [10] M. C. F. de Oliveira and H. Levkowitz. From visual data exploration to visual data mining: A survey. *IEEE Transactions on Visualization and Computer Graphics*, 9:378–394, 2003.
- [11] F. Frasincar, R. Telea, and G.-J. Houben. Adapting graph visualization techniques for the visualization of rdf data. In *Visualizing the Semantic Web*, 2006.
- [12] V. Geroimenko and C. Chen, editors. *Visualizing the Semantic Web*. Springer, 2002.
- [13] T. Hastrup, R. Cyganiak, and U. Bojars. Browsing linked data with fenfire, 2008.
- [14] D. Karger and M. Schraefel. The pathetic fallacy of RDF. Position Paper for SWUI06, 2006.
- [15] B. Shneiderman. The eyes have it. In *IEEE Symposium on Visual Languages*, 1996.
- [16] C. Stadler, J. Lehmann, K. Höffner, and S. Auer. Linkedgeodata: A core for a web of spatial open data. *Semantic Web Journal*, 2011.
- [17] J. Unbehauen, S. Hellmann, S. Auer, and C. Stadler. Knowledge extraction from structured sources. In S. Ceri and M. Brambilla, editors, *SeCO Book*, volume 7538 of *Lecture Notes in Computer Science*, pages 34–52. Springer, 2012.