

Rialto: A Knowledge Discovery Suite for Data Analysis

Giuseppe Manco^{a,*}, Pasquale Rullo^b, Lorenzo Galucci^c, Mirco Paturzo^c

^a*ICAR-CNR, Via Bucci 41c, 87036 Rende (CS) - Italy*

^b*Dip. di Matematica e Informatica, Università della Calabria, Via Bucci 30b, 87036 Rende (CS) - Italy*

^c*Exeura S.r.l., Via P.A. Cabrai, 87036 Rende (CS) - Italy*

Abstract

A Knowledge Discovery (KD) process is a complex inter-disciplinary task, where different types of techniques coexist and cooperate for the purpose of extracting useful knowledge from large amounts of data. So, it is desirable having a unifying environment, built on a formal basis, where to design and perform the overall process. In this paper we propose a general framework which formalizes a KD process as an algebraic expression, that is, as a composition of operators representing elementary operations on two worlds: the *data* and the *model* worlds. Then, we describe a KD platform, named Rialto, based on such a framework. In particular, we provide the design principles of the underlying architecture, highlight the basic features, and provide a number of experimental results aimed at assessing the effectiveness of the design choices.

Keywords: Knowledge Discovery Process, Data Mining, Business Analytics Platforms

1. Introduction

Knowledge Discovery (KD) is the process of extracting novel, potentially useful, and ultimately understandable patterns in data (Han and Kamber, 2001). A KD process involves several steps, including data acquisition, data pre-processing,

*Corresponding author

Email addresses: `manco@icar.cnr.it` (Giuseppe Manco), `rullo@mat.unical.it` (Pasquale Rullo), `lorenzo.gallucci@exeura.eu` (Lorenzo Galucci), `mirko.paturzo@exeura.eu` (Mirco Paturzo)

5 data mining and data post-processing. In general, each step requires a variety of
6 algorithms and techniques, such as SQL queries to manipulate data in the pre-
7 processing phase, different learning methods in the data mining step, or model
8 visualization metaphors in the post-processing phase.

9 Though the main efforts in the KD community have been so far devoted to
10 the development of efficient mining algorithms and techniques, some standard-
11 ization initiatives, aimed at solving the KD problem as a whole, have been pro-
12 posed. As an example, the Cross Industry Standard Process (CRISP) method-
13 ology (Shearer, 2000), where data mining is set in the larger context of KD, is
14 considered a *de facto* standard methodology used by industry data miners. A
15 more formal attempt to define a general and unified framework is the 3-World
16 model proposed in (Johnson et al., 2000). Here, a KD process is seen as a
17 "multi-step process where the input of one mining operation can be the output
18 of another" and a data mining algebra is proposed.

19 Thus, in the perspective of a unified approach, a data mining system should
20 provide support to the entire KD lifecycle within a unique, integrated environ-
21 ment, providing suitable features to enable the construction of a KD analysis
22 as an interactive and an incremental process. To this end, there is a number of
23 requirements that an integrated KD environment should meet, including:

- 24 • capability to gather data from diverse data sources, such as databases,
25 Excel, CSV files, etc. (*data acquisition*)
- 26 • availability of a wide variety of techniques to reduce data preprocessing ef-
27 forts (nearly 80% of the time that an analyzer spends on analytics projects
28 is during preprocessing)
- 29 • richness of mining algorithms and techniques
- 30 • availability of effective model and data visualization tools, for decision
31 making and informed data manipulation, respectively
- 32 • *extensibility*, i.e., open architecture to be easily customizable with new
33 specific tasks.

34 Additional requirements that a KD system should fulfill in order to cope
35 with real-world applications include:

- 36 • *Usability*: a KD process can be extremely complex with several interacting
37 components. In order to efficiently design and build such a process, the
38 supporting KD environment should provide a GUI allowing for a simple
39 and interactive creation of the analysis process, obtained by assembling a
40 number of elementary tasks according to a given logical flow. Furthermore,
41 it should enable the user to visually explore data and models.
- 42 • *Multiple data types*: while classical KD approaches mainly deal with re-
43 lational data, the need to deal with unstructured data is lately growing
44 (Sebastiani, 2002). Indeed, textual data represent the large majority of
45 stored digital information and is growing at a much faster rate than busi-
46 ness quantitative data. As a consequence, text mining applications have
47 been steadily increasing in the last few years (see, e.g., (Vilares et al.,
48 2015), (Hristovski et al., 2008)). Besides texts, other data types like time
49 series (chung Fu, 2011), graphs (Borgelt, 2009), streams (Gaber et al.,
50 2005; S., 2005) have lately gained increasing popularity.
- 51 • *Big Data*: a KD tool should be able to efficiently deal with mass-memory
52 resident data and provide mechanisms to scale up to large real-world do-
53 mains.

54 Today, a large number of KD tools exists, with both commercial and open
55 source licenses, e.g., (Hall et al., 2009; Berthold et al., 2009). KD systems are
56 diverse in design and implementation. Many commercial KD tools available in
57 the market result from the integration of data mining methods into commercial
58 statistical tools, e.g., SPSS and SAS, while others, such as Oracle Data Mining,
59 result from the integration of data mining solutions into business intelligence
60 products. A few tools, such as KNIME (Berthold et al., 2009) and IBM SPSS
61 Modeler (formerly Clementine), were natively designed to provide full support
62 to the KD lifecycle, giving the possibility of executing the entire KD process

63 within a unique, integrated environment, thus eliminating the need for context
64 switching. Surveys on KD tools can be found in (Mikut and Reischl, 2011; Chen
65 et al., 2007).

66 In this paper we present Rialto, a KD platform assisting the designer dur-
67 ing the entire KD lifecycle. Rialto relies on a visual interaction which enables
68 the construction of a KD process as a workflow representing a composition
69 of elementary operators on two worlds: the *data* and the *model* worlds (the
70 2W Model). Thus, a workflow consists of nodes that are either tables (data), or
71 models or tasks (operators). Tables are unnormalized relations, as *Event Collec-*
72 *tions* for managing complex data types, such as texts, itemsets, etc., are allowed.
73 Operators are implemented as plug-ins of the platform. In order to facilitate
74 the work of the analyst, Rialto offers a broad variety of pre-packaged plug-ins
75 covering all phases of the KD process (from data acquisition to model deploy-
76 ment). In addition, new *ad hoc* Java-based plug-ins can be created whenever
77 needed (*extensibility*). To this end, Rialto provides an integrated development
78 environment which enables and simplifies code generation.

79 Currently, Rialto works on both relational data and texts. Data mining is
80 supported by plug-ins for classification (including decision trees, Naive Bayes,
81 rule induction), regression, clustering and association mining. As far as texts
82 is concerned, text-specific plug-ins are available for preprocessing, classification
83 (Complement Naive Bayes (Rennie et al., 2003), Olex (Rullo et al., 2009) ,
84 OlexGa (Pietramala et al., 2008) and Max Entropy (Nigam et al., 1999)), and
85 topic detection based on the Latent Dirichlet Allocation (LDA) model (Blei
86 et al., 2003). Ongoing work is devoted to extensions to deal with data streams
87 as well as spatio-temporal data.

88 Rialto is also equipped with plug-ins allowing the analyst to inspect both
89 data and models, thus taking advantage of descriptive statistics, explanatory
90 charts, and models visualizations metaphors.

91 Finally, Rialto exhibits features, such as the *bridge* mechanism, for the ef-
92 ficient and transparent manipulation of large amounts of data. Efficiency and
93 scalability are ensured by a workflow execution engine relying on parallel exe-

94 cution strategies along with effective buffering policies and management of data
95 in main memory.

96 In short, the contributions of the paper can be summarized as follows.

- 97 • We propose an algebraic model, namely the 2W model, for specifying data
98 mining queries and, more in general, a KD process. The 2W Model is a
99 variant of the 3W model which emphasizes the underlying philosophy of
100 combining several forms of knowledge and the cooperation among solvers
101 of different nature, in a more general setting.
- 102 • We introduce an effective architecture which implements the 2W model
103 and visually models data mining queries as interactions among the two
104 worlds. Rialto relies on two key aspects: closure and extensibility. These
105 aspects are directly derived from the specification of the 2W model, and
106 they are the essential tools to achieve the combination of deductive infer-
107 ences along with inductive mechanisms and statistical methods. However,
108 Rialto also exhibits some important aspects beyond those: some architec-
109 tural choices are specifically introduced to enable the optimized execution
110 of data mining queries. In addition, a strong emphasis is given to visual
111 inspection and support for the choice of the optimal data/model manipu-
112 lation operator.

113 The paper is organized as follows. In Section 2 we provide a description of
114 the 2W Model. In Section 3 we give an overview of Rialto. In Section 4 we
115 describe the architecture of Rialto. In Section 5 we show two case studies based
116 on Rialto. In Section 6 we evaluate Rialto against a number of both functional
117 and architectural requirements, and perform a number of experimental analyses
118 aimed at assessing the effectiveness of the proposed architecture. In Section 7
119 we describe a number of features we are currently working on and, finally, in
120 Section 8 we provide some concluding remarks.

121 2. The Foundation: Modeling KD Processes as Algebraic Processes

122 There is a sort of standard framework upon which complex data analyt-
123 ics applications can be approached (Manco et al., 2008). Notably, analytical
124 questions posed by the end user need to be translated into several tasks such
125 as choose analysis methods, prepare the data for application of these methods,
126 apply the methods to the data, and interpret and evaluate the results obtained.
127 The interaction of these tasks requires combining several forms of knowledge and
128 the cooperation among solvers of different nature: we need to support deduc-
129 tive inferences along with inductive mechanisms, in conjunction with statistical
130 methods. Historically, there has been an effort to express this combination
131 in a unified framework (Imielinski and Mannila, 1996): either by integrating
132 data mining algorithms with the underlying database systems (Chaudhuri et al.,
133 2002; Giannotti et al., 2004); or by providing ad-hoc extensions of SQL (Imielin-
134 ski and Virmani, 1999; Wang et al., 2003; Ortale et al., 2008). We can devise
135 however a unifying picture where we consider the Knowledge Discovery process
136 as an algebraic process, with some primitive data types and instances which can
137 be manipulated by means of a set of basic operators.

138 The impact of such an algebraic framework can be summarized in two main
139 aspects which we expect the framework will support. First, in an effort to
140 combine reasoning and mining, it is important to treat results of data mining
141 on a par with data objects which can be further manipulated. Second, suitable
142 combinations and compositions of known mining tasks can enable more powerful
143 mining computations and findings. In essence, the principle upon which these
144 aspects can be enabled is the *closure* capability.

145 The current literature has studied the foundational aspects of this algebraic
146 view of the KD process. In particular, the 3-Worlds Model (3W in the following)
147 (Johnson et al., 2000; Calders et al., 2006) introduces and develops this abstract
148 view as a mathematical structure where three entities, namely data, intensional
149 models and extensional models, can interact by means of algebraic operators.
150 In the following we will describe a variant of the 3W, which we refer to as the

151 2-World Model (2W). In this model, a KD process can be specified as the
 152 interaction between *two* apparently separate worlds, the *data world* (D-world)
 153 and the *model world* (M-world).

154 The D-world consists of a set of entities to be analyzed, with their properties
 155 and mutual relationships. Following (Johnson et al., 2000; Calders et al., 2006),
 156 we assume that the D-world D is made of all possible nested relational tables
 157 that can be defined over an infinite attribute set $\mathcal{A} = \{A_1, \dots, A_m, \dots\}$, where
 158 each A_i is associated with a domain $dom(A_i)$.

159 The M-world consists of the intensional representations of the patterns re-
 160 lating the elements of the D-World, expressed through specific languages (rules,
 161 mathematical properties, etc.). In (Johnson et al., 2000; Calders et al., 2006),
 162 M is populated by region tables, i.e., relational tables where attributes can
 163 only be associated with the special *region* domain. In particular, a region can
 164 be described in terms of a description of its members, i.e., region membership
 165 criteria. The simplest criteria can be devised as constraints over a subset of
 166 attributes in $\mathcal{A}$. For example, $A \leq 5 \wedge B > 6$ intentionally represents the set
 167 of all tuples t defined over a relation $r[A, B]$ such that $t[A] \leq 5$ and $t[B] > 6$.
 168 Thus, an example model can be defined over the scheme given by the attributes
 169 $\{\text{Condition}, \text{Classification}\}$, and a possible instance is represented by the
 170 table r defined as:

Condition	Classification
$A \leq 5 \wedge B > 6$	$C = 1$
$A < 5$	$C = 0$

172 It is easy to see how the above table encodes a decision tree, where each row
 173 represent a path and in particular the **Condition** attribute represents the path,
 174 whereas the **Classification** attribute represents the label associated with that
 175 path.

176 Within the 2W, we deliberately choose a more general setting, and we do
 177 not specify directly models as region tables. Rather, we assume that a model
 178 m is any mathematical structure characterized by

- 179 • two schemas $S, S' \subset \mathcal{A}$ and $S \subset S'$; intuitively, S is the set of attributes
180 upon which m is built, and $S' \setminus S$ is the set of attributes describing the
181 result of the application of m to a tuple t over S ;
- 182 • an *entailment* $\vdash_S$ operator such that, for each tuple t defined over S ,
183 $m \vdash_S t$ if t satisfies m ;
- 184 • an *extension* operator τ such that, for each t entailed by m , $\tau_m(t) = t'$,
185 where t' is a tuple defined over S' , such that $t'[S] = t$.

186 We notice that the above two operators subsume the basic properties that mod-
187 els should have: the capability of bounding tuples (entailment), and the capa-
188 bility to annotate tuples according to these bounds (extension).

189 The table r illustrated above is a special case of this formalization, where
190 $S = \{A, B\}$ and $S' = \{A, B, C\}$. Also, the entailment $\vdash$ holds if there is any
191 region in **Condition** which is satisfied by t , and $\tau_r(t)$ extends t with the value
192 of the region associated with **Classification**. For example, if $t = \langle 3, 6 \rangle$ then
193 $\tau_r(t) = \langle 3, 6, 0 \rangle$ because the entailment is satisfied by the second row of r , and
194 the corresponding **Classification** region is $C = 0$. In the following we refer
195 to $\langle S, S', \vdash_S, \tau \rangle$ as the schema of a model.

196 Thus, in our formalization a 2W database is a tuple $\langle \mathbf{D}, \mathbf{M} \rangle$, where $\mathbf{D}$ is a set
197 of relations in D , and $\mathbf{M}$ is a set of models in M . A knowledge discovery process
198 can hence be formalized as a transition from a 2W database $\langle \mathbf{D}, \mathbf{M} \rangle$ to another
199 one $\langle \mathbf{D}', \mathbf{M}' \rangle$, through a number of suitable operators. In particular, the closure
200 between the two worlds D and M is achieved by means of the following set of
201 classes of operators (see Figure 1):

- 202 • *Filtering* operators are self-injecting relationships. They take a set of enti-
203 ties as input and produce a new set of entities, generating data from data
204 (for pre-processing purposes) and models from models. Thus, a generic
205 data filter operator is of the form

$$\phi_D : D^n \rightarrow D^m$$

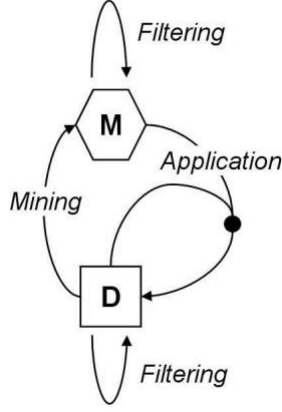


Figure 1: The 2W Model

206

while a generic model filter operator is of the form

$$\phi_M : M^n \rightarrow M^m.$$

- *Mining* operators relate data entities to model entities, thus enabling transitions from the D-world to the M-world. A generic mining operator (representing a mining algorithm) is of the form

$$\mu : D \rightarrow M.$$

207

The result of the application of μ to a data entity d is a model, which describes a pattern holding over d .

208

- *Application* operators can be seen as a sort of reverse mining operators, generating a data entity from the application of a model to another data entity, i.e.

$$\alpha : D \times M \rightarrow D$$

209

Applying a model m to a data entity d essentially means making the intensional properties of m explicit in extensional form: e.g., by associating each tuple in a table with the most likely target class according to the model, or by enumerating the frequent patterns appearing within the

210

211

212

213 tuple. The only requirement is that schemes are compatible, i.e., given
 214 $d \in D$ with scheme X and $m \in M$ with scheme $\langle S, S', \vdash, \tau \rangle$, then $S \subseteq X$.

215 Now, every data mining query can be represented by means of a composition
 216 of the above operators, and a KD process as a transition of 2W databases
 217 through data mining queries. For an instance, a typical learn-and-test process
 218 can be represented as the transition

$$\begin{aligned}
 &\langle \{d\}, \emptyset \rangle \\
 &\quad \rightarrow \langle \{d, d_{Train}, d_{Test}\}, \emptyset \rangle \\
 &\quad \rightarrow \langle \{d, d_{Train}, d_{Test}\}, \{m\} \rangle \\
 &\quad \rightarrow \langle \{d, d_{Train}, d_{Test}, d_c\}, \{m\} \rangle,
 \end{aligned} \tag{1}$$

where each transition is determined by the application of some of the above
 described operators, notably:

$$\begin{aligned}
 \langle d_{Train}, d_{Test} \rangle &:= \phi_D^{(split)}(d) \\
 m &:= \mu(d_{Train}) \\
 d_c &:= \alpha(m, d_{Test})
 \end{aligned}$$

219 Here, $\phi_D^{(split)}(d)$ represents a data filtering operator which splits d into two parti-
 220 tions. In turn, m is the model learned by the execution of a mining algorithm μ
 221 over d_{Train} (e.g., a decision tree), and d_c is the data entity obtained by applying
 222 the model m on the test set d_{Test} .

223 As another example, a process inducing a decision tree DT from a data set d ,
 224 by using the learning algorithm μ^{DT} , and then transforming DT into a set of
 225 rules, is synthetically represented by the expression $\phi_M^{DT2Rules}(\mu^{DT}(d))$, where
 226 $\phi_M^{DT2Rules}$ is a suitable model filtering operator.

227 Compared to the 3W model, the 2W model exhibits some substantial similar-
 228 ities and differences, as shown in table 1. First, the Model world does not entail
 229 an extensional representation in our simplified view. As a matter of fact, the
 230 extensional representation of a model corresponds to a set of tuples associated
 231 with the regions they belong to. Since we implicitly embed these associations

	$3W$	$2W$
<i>Data</i>	Nested relational tables	Nested relational tables
<i>Models</i>	Region tables	Any mathematical structure with entailment and extension
	Both intensional and extensional representation	No extensional representation
<i>Data operators</i>	Nested relational algebra	Generic (any function $D^n \mapsto D^m$)
<i>Model operators</i>	Dimension algebra	Generic (any function $M^n \mapsto M^m$)
<i>Mining operators</i>	Regionize, loop	Generic (any function $D \mapsto M$)
<i>Model to data</i>	Populate	Generic (any function $D \times M \mapsto D$ with compatible schemes)

Table 1: Comparison of 3W and 2W.

232 within the entailment and extension operators, we don't need to differentiate be-
 233 tween intensional and extensional representation. The differentiation, although
 234 useful from a theoretical point of view, is of little interest from a practical point
 235 of view.

236 Second, we regard the operators ϕ_D , ϕ_M , μ and α as black boxes, whereas
 237 the 3W world only allows a pre-specified set of basic operators, upon which
 238 to derive more complex operators by means of compositions. In particular,
 239 the operators ϕ_D in 3W represent all the nested relational algebra operators,
 240 and ϕ_M is any operator composing and decomposing regions. Also, the only
 241 mining operators in (Calders et al., 2006) are the *regionize* and *loop* operators.
 242 These operators are useful from a theoretical point of view, since they allow to
 243 investigate the expressive power of the overall algebra: that is whether any data
 244 mining algorithm can be expressed by means of a combination of this minimal
 245 set of operators. However, it is extremely unpractical to rely on such operators
 246 from an application viewpoint. This is the main reason why we decided to keep
 247 the operators at an abstract level, as functionals within the two worlds. The key
 248 point in the 2W is that objects in different worlds can interact, and that there

249 is a way of mapping an object within D to an object within M , and viceversa.

250 Finally, another substantial difference, is the choice of the mining models.
251 As we already mentioned, the 3W focuses on region tables, and regions are only
252 expressible as sets of linear inequalities. As Calders et al. (2006) declare, “*this*
253 *limitation means that the results of some data mining operations might not be*
254 *expressible, as they require more complex mathematical objects. Straightforward*
255 *examples are, e.g., non-linear regression methods, support vector machines, and*
256 *clustering methods resulting in non-linear regions, etc’ [...]. Therefore, the de-*
257 *scription of the output of some data mining tasks requires more sophisticated*
258 *constraints than the ones used in the paper”*. By contrast, in the 2W we choose
259 to describe models in a more general setting, i.e., as mathematical structures
260 equipped with the entailment and extension operators, the former providing
261 the capability of bounding tuples, and the latter to annotate tuples according
262 to these bounds. In this respect, the 2W model should be able to express data
263 mining queries involving different models.

264 Rialto was designed with the 2W philosophy in mind, and the architectural
265 choices are aimed at injecting the underlying algebra within a visual and inter-
266 active environment. In the following we explore the basic features of such an
267 environment.

268 3. An Overview of Rialto

269 The framework we propose is a graphical environment for designing and
270 managing KD processes as algebraic expressions based on the 2W Model. These
271 are defined by means of workflows, whose nodes represent either *tables*, *models*
272 or *tasks*:

- 273 • *Tables* represent entities of the D-world. A table is a set of tuples, char-
274 acterized by attributes. Attribute domains can be either primitive or
275 complex data types. In particular, an attribute can be an *Event Collec-*
276 *tion* (see Section 4.2.1), used to efficiently handle complex objects, such as
277 itemsets, sequences, time series, texts, etc. In general, Event Collections

278 can be used for the efficient storing of sparse data sets (i.e., with only a
279 few nonzero features).

- 280 • *Models* represent entities of the M-world.
- 281 • *Tasks* represent the algebraic operators of the 2W Model - thus, we have
282 tasks of type *filtering*, *mining* and *application*.

283 The screenshot of an example learn-and-test workflow is shown in Figure 2.
284 It essentially represents a KD process similar to the one shown in equation 1
285 above. This workflow starts with an *acquisition* task from a CSV file (a kind
286 of filtering task), whose data is imported within Rialto and stored in table *Ac-*
287 *quiredDataTable*. Each row in this table represents a pre-classified example.
288 Then, a number of *pre-processing* steps are performed. In particular, the above
289 table is first submitted to a *filter* node in charge of managing *null* values. The
290 resulting table *NoMissingValuesTable* is then processed by another filter node,
291 namely, *Discretize*, in charge of discretizing the values of attribute *Age*, so gen-
292 erating *DiscretizedAgeTable* (note that both the above filters transform tables
293 into tables). Now data is ready for *mining* and, then, the EC4.5 algorithm (an
294 extension of C4.5 (Quinlan, 1987)) is run by using the holdout method. To this
295 end, the input table *DiscretizedAgeTable* is split into a training set and a test
296 set (*HoldOut Testset*). The former is used to induce the model (a decision tree,
297 in this case) represented in the workflow by the hexagonal shaped node. The
298 latter (i.e., the test set) is used to test the accuracy of the induced model. This
299 is done by using the *application* task Prediction Join (visually, the model node
300 and the test table *HoldOut Testset* are fed to the prediction join node). The
301 resulting table *Predicted*, enabling comparison of the predicted labels with those
302 of the ideal classification, is finally *exported* as a CSV file.

303 As mentioned above, every task node in a workflow represents an operator
304 of the 2W Model (filtering, mining, application), and is implemented as a Java
305 *plug-in* (so any transformation on tables and models can be performed, as no
306 predefined operators are in the 2W Model). However, in order to make it easier
307 for the analyst the construction of workflows, Rialto offers a broad variety of

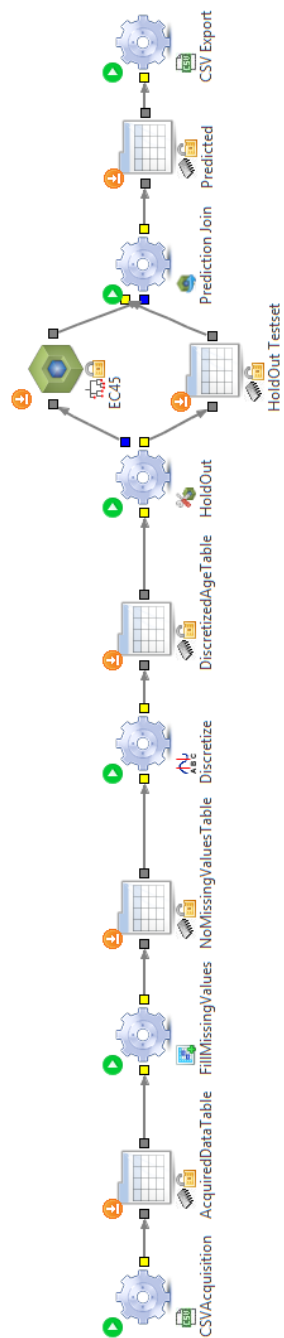


Figure 2: An example learning workflow

pre-built plug-ins covering all phases of the KD process. For an instance, for the purpose of data acquisition, Rialto provides plug-ins for different data types and format, namely, Arff, CSV, Xls, texts, and tables from relational databases through JDBC. Among the pre-processing plug-ins, SQL filters can be used to create tables from other tables. Data mining is supported by a number of plug-ins for classification, regression, clustering and association mining. Since Rialto works also on textual data, text-specific plug-ins are available for pre-processing (feature extraction, feature selection, concept recognition, etc.), text classification (Complement Naive Bayes (Rennie et al., 2003), Olex (Rullo et al., 2009) , OlexGa (Pietramala et al., 2008) and Max Entropy (Nigam et al., 1999)), and topic detection based on the Latent Dirichlet Allocation (LDA) model (Blei et al., 2003).

One peculiarity of Rialto is that, in the logic of an open and extensible architecture, not only tasks of the 2W Model are implemented as plug-ins, but *every* functionality in Rialto is implemented as a plug-in. For instance, Rialto offers several plug-ins (not appearing in workflows - for this reason called *behind-the-scene*, as opposed to tasks that are *on-the-scene* plug-ins) whereby the analyst can effectively and easily examine and inspect both data and models, taking advantage of descriptive statistics, explanatory charts, models visualization metaphors.

One of the fundamental behind-the-scene plug-in types in Rialto is that of *bridges*. A bridge is a component in charge of supporting the mapping of the table logical primitives into their physical level, thus providing a decoupling mechanism between logical data and its physical storage. This way, the bridge mechanism makes plug-ins independent from table implementations. A bridge is either a main memory bridge, used for "small" tables, or a mass memory bridge, the latter relying on a relational DBMS. The bridge mechanism ensures both transparency and efficiency of data retrieval and storing.

Table 2 provides a summary of the different classes of plug-ins. In its current release, Rialto features over 100 pre-packaged plug-ins.

This notwithstanding, new *ad hoc* plug-ins may be needed from time to time

Table 2: Summary of the different classes of plug-ins

Plug-in type	Note
Task	Implements 2W Model operators - Appears in workflows
Model visualizer	Applies to models - Does not appear in workflows
Data visualizer	Applies to tables - Does not appear in workflows
Statistical tool	Applies to tables - Does not appear in workflows
Bridge	Applies to tables - Does not appear in workflows

339 to fulfill the specific needs of the application at hand, e.g., an acquisition task
340 to get data from a legacy system, or a new model visualizer. At this aim, Rialto
341 provides an exhaustive Java-based application programming interface allowing
342 for quick access to the main methods. This way, a user can easily implement,
343 say, a new model induction algorithm, a new data visualization metaphor or a
344 new bridge, and store it in the plug-in archive of Rialto.

345 Finally, efficiency and scalability of the Rialto platform are ensured by a
346 workflow execution engine relying on parallel execution strategies along with
347 effective data handling policies. Rialto provides support for three kinds of paral-
348 lelism: parallel execution of independent workflow branches, parallel table scan
349 and pipelining. Data handling relies on techniques for the efficient management
350 of data in main memory and effective buffering strategies.

351 4. System Architecture

352 The general architecture of the system is shown in Figure 3. Here, we can
353 observe two basic modules, namely, GUI (Graphical User Interface) and CORE.
354 The former implements the required functionalities for visual workflow compo-
355 sition, as well as data/model explorations. The latter is responsible for the
356 efficient execution of workflows.

357 Besides these two modules, there are several other components that can be
358 *plugged* into the system. Data and model *visualizers* are bound to tables and
359 models, respectively, and are aimed at performing data and model explorations.

360 *Bridges* are bound to tables, and are in charge of implementing a decoupling
 361 mechanism between logical data and physical memory (see Section 4.2.4 for
 362 more details). *Tasks* are implementations of the 2W Model operators.

363 The *extensibility* is provided through Rialto SDK, a software development
 364 kit allowing users to implement, build and integrate their own plug-ins into
 365 their repositories. Plug-ins in Rialto are based on Java and can be developed
 366 through the Rialto integrated development environment (IDE) which enables
 367 and simplifies code generation. The current release of Rialto includes over 100
 plug-ins, covering all phases of the KD process.

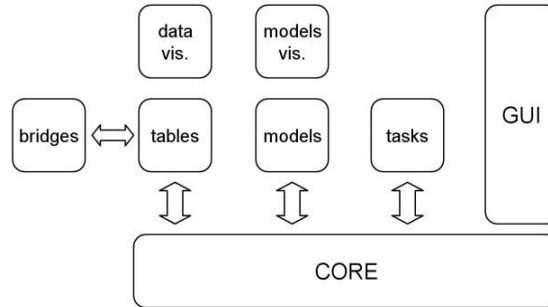


Figure 3: Rialto Architecture

368
 369 We next analyze each architectural component in turn, starting from the
 370 GUI.

371 4.1. Graphical User Interface

372 The basic functionalities of the user interface are aimed at modeling KD
 373 analyses by means of workflows. To this purpose, Rialto implements a visual
 374 interaction metaphor for representing data, models and tasks, along with their
 375 interactions. In addition, Rialto provides suitable features to enable the con-
 376 struction of a KD analysis as an interactive and an incremental process, e.g.,
 377 statistical visualization facilities to inspect, during pre-processing, data prop-
 378 erties, as well as tools for models inspection and performance measurement

(like classification accuracy, or statistical relevance) during post-processing. In

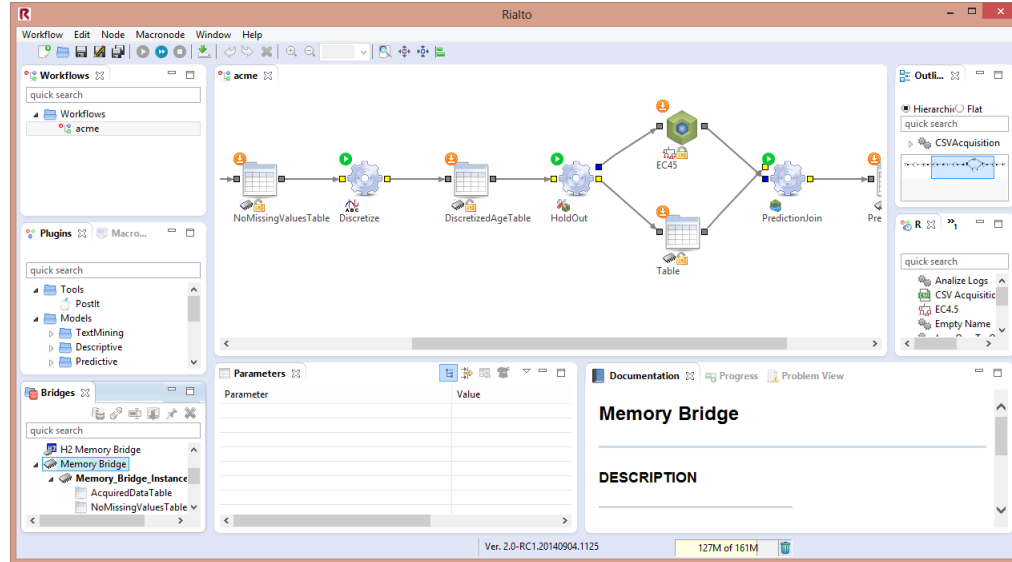


Figure 4: Workbench

379

380 particular, the user interface provides three different perspectives:

- 381 • *workflow perspective*, used to create, configure and connect nodes;
- 382 • *data perspective*, which enables exploratory analyses;
- 383 • *model perspective*, which allows inspecting and visualizing models.

384 4.1.1. Workflow Perspective

385 Figure 4 shows a screenshot of the Rialto workbench. The Workflow Edi-
 386 tor Area (WEA) is the heart of the workbench. Here, a workflow is designed
 387 by exploiting a graphical editor, where a user can create, position, move and
 388 connect nodes. Nodes can also be moved around in the WEA with a drag-and-
 389 drop operation. Additional functionalities available in the Workflow Perspective
 390 are: the *workflows navigator* allowing the user to navigate and display on the
 391 WEA the workflows which are currently open; the *plug-ins view* which lists all

the available nodes that can be instantiated in a workflow through drag-and-drop operations; the *macronodes view*, which lists all available macronodes (i.e., pieces of workflows); the *bridge view* allowing to select a bridge to be attached to a given table (recall that each table has a bridge associated - see Section 4.2.4 for more details).

The screenshot shows a software window titled 'acme' with a tab 'NoMissingValuesTable'. It displays a table with four columns: '0-AGE (integer)', '1-WORKCLASS (nominal)', '2-EDUCATION (nominal)', and '3-MARITAL_STATUS (nominal)'. A context menu is open over the '1-WORKCLASS' column, showing options: 'Show Nominal Values', 'Order A-Z', 'Order Z-A', 'Add Column Filter', 'Hide', and 'Remove Column Filters'. The table contains 22 rows of data.

	0-AGE (integer)	1-WORKCLASS (nominal)	2-EDUCATION (nominal)	3-MARITAL_STATUS (nominal)
0	37	Private	Masters	Married-civ-spouse
1	37	Private	Masters	Married-civ-spouse
2	49	Private	9th	Married-spouse-absent
3	52	Self-emp-not-inc	HS-grad	Married-civ-spouse
4	52	Self-emp-not-inc	HS-grad	Married-civ-spouse
5	30	State-gov	Bachelors	Married-civ-spouse
6	25	Self-emp-not-inc	HS-grad	Never-married
7	32	Private	HS-grad	Never-married
8	43	Self-emp-not-inc	HS-grad	Divorced
9	54	Private	HS-grad	Separated
10	54	Private	HS-grad	Married-civ-spouse
11	49	Private	HS-grad	Married-civ-spouse
12	30	Federal-gov	HS-grad	Married-civ-spouse
13	22	State-gov	Some-college	Married-civ-spouse
14	19	Private	HS-grad	Married-AF-spouse
15	49	Private	HS-grad	Separated
16	18	Private	HS-grad	Never-married
17	50	Federal-gov	Bachelors	Divorced
18	47	Self-emp-inc	HS-grad	Divorced
19	43	Private	Some-college	Married-civ-spouse
20	35	Private	Assoc-voc	Married-civ-spouse
21	35	Private	Assoc-voc	Married-civ-spouse

Figure 5: Table View

4.1.2. Data Perspective

In the Data Perspective, the analyst can get insights into the elements of the D-world, i.e., tables (see Figure 5). To this end, he/she can perform exploratory analyses by means of *ad hoc* visualization and statistical tools as well as simple data manipulation tasks (such as ordering or filtering values). The data perspective can be accessed by double-clicking on a table node.

4.1.3. Model Perspective

Similarly to data, models can also be explored by clicking on the respective nodes in a workflow. Model visualizers allow for a visual inspection of the

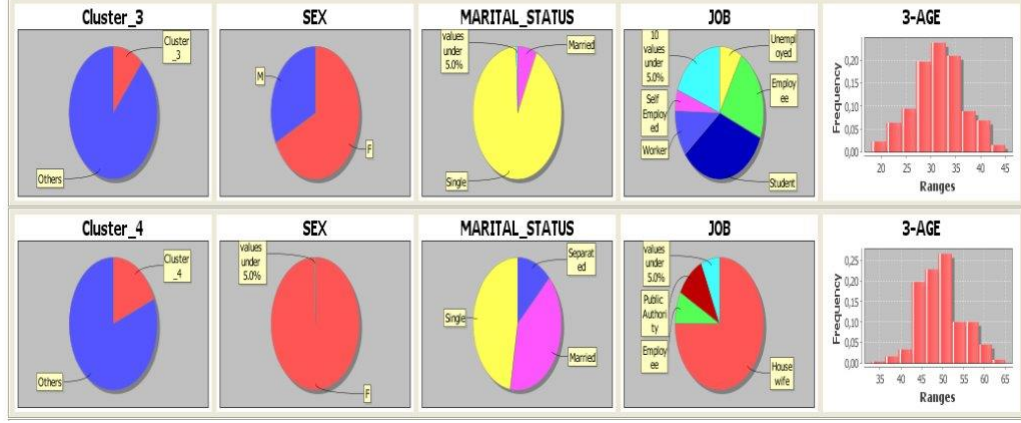


Figure 6: Clusters representation

underlying model in order to understand its features. An example visualization for a clustering model is shown in Figure 6, where the properties of Cluster_3 and Cluster_4 are represented in terms of population frequencies (leftmost pie charts), and attribute value distributions (sex, marital.status, job, 3-age).

4.2. Tables

Data in Rialto is stored within tables. From an abstract point of view, a table represents a set of entities characterized by properties. Properties are associated with a data type, which can be simple (nominal, ordinal, numeric) or complex (string, *event collection*). Nominal properties are characterized by an enumerable set of admissible values, which exhibit no order. By contrast, ordinal properties exhibit a standard order, and they are internally represented as integers. Finally, numeric properties can be exploited for representing results of measurements of real-life phenomena, and are implemented as double. Strings are exploited in Rialto to represent information that is not relevant to the analysis: in practice, a string is any sequence of bytes that are not manageable by the core system.

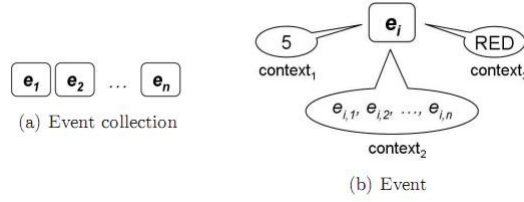


Figure 7: Event Collections

4.2.1. Dealing with complex data: Event Collections

Modeling of complex properties is usually mandatory in real-world applications, which typically require to deal with complex structures, e.g., itemsets, texts, sequences, time series. Within Rialto, we introduced a data type which can subsume several complex properties: the Event Collection (EC). An EC is essentially a set of events, where each event is a complex structure characterized by a set of predefined contexts, and a context is associated again to a data type. Thus, ECs in principle embody a recursive structure (a table within a table), where events correspond to sub-entities, and contexts represent the related sub-properties. Figure 7 illustrates the above ideas: each event has three contexts, namely, *context1* which is numeric, *context2* which is an EC, and *context3* which is nominal. It is easy to see how an EC can be used to model, e.g., itemsets (which are essentially event collections where events expose a single nominal context), time series (where events are characterized by a timestamp and a value) or even textual data. In the latter case, a document can be mapped to an EC of terms. Textual events can be characterized by contexts such as n-gram length, frequency within the document, and so on.

4.2.2. Manipulating Tables

Following the philosophy of an open architecture, tables can be accessed and managed by means of system APIs. In practice, a table can be created by specifying the metadata information and the underlying storage manager (or

443 bridge, explained later in this section). The following code fragment specifies a
444 table with 4 attributes, and then generates a row which is added to the table:

Code Fragment 1: Table specification

```
445  
446  
447 TableMetadataBuild metadataBuilder =  
448     new TableMetadata.Builder();  
449 metadataBuilder.add(new StringAttribute("Name"));  
450 metadataBuilder.add(new DoubleAttribute("Height"));  
451 metadataBuilder.add(new IntegerAttribute("Age"));  
452 metadataBuilder.add(new NominalAttribute("SEX", "F", "M", "F"));  
453 final TableMetadata tableMetadata = metadataBuilder.build();  
454 Table table = bridge.generateTable(tableMetadata);  
455 final Row newRow = table.getNewRow();  
456 newRow.setValue(0, "John");  
457 newRow.setValue(1, "1.74");  
458 newRow.setValue(2, 30);  
459 newRow.setValue(3, "M");  
460 table.appendRow(newRow);  
461
```

462 We can notice from the example that the library functions for creating and
463 manipulating tables are quite straightforward.

464 4.2.3. Table Visualization and Statistical Tools

465 Besides the pre-defined tools, new user-defined visualization and statistical
466 tools can be built over tables thank to the open architecture of Rialto. For an
467 instance, a statistic relative to a table is specified by instantiating its evaluation.
468 Specifically, a statistic implements the following Java interface:

```
469  
470 public interface Statistic extends RialtoPlugin {  
471     void check(TableMetadata)  
472         throws RialtoException;  
473     void evaluate(Table);  
474     Object getValue();  
475 }  
476
```

477 The first method is aimed at checking the compatibility of the statistic with the
478 metadata associated with the underlying table. The second method computes

479 the statistic, which can be fetched by using the third method. Notice how the
 480 interface exploits the TableMetadata, Table and Value components from the
 481 core, which are detailed later in this paper. As an example, Code Fragment 2
 482 is a simplified version of the Average statistic.

Code Fragment 2: the Average statistic

```

483
484 public class Average extends ParametersOwner implements Statistic {
485
486     @ParameterName("Attribute index")
487     @ParameterAttributeIndexDecorator(acceptableTypes = { AttributeType.DOUBLE,
488         ↪ AttributeType.INTEGER })
489     @ParameterMandatoryDecorator
490     public final PluginParameter<Integer> attributeIndex = createParameter();
491
492     private transient double value;
493
494     @Override
495     public void check(TableMetadata metadata) throws RialtoException {
496         for (final Attribute attribute : metadata) {
497             if (AttributeType.DOUBLE == attribute.getAttributeType()
498                 || AttributeType.INTEGER == attribute.getAttributeType())
499                 return;
500         }
501         throw new RialtoException("Table has no numeric attributes.");
502     }
503     @Override
504     public void evaluate(Table table) {
505         int count = 0;
506         value = 0;
507
508         for (final Row row : table) {
509             value += row.getValue(attributeIndex.getValue());
510             count++;
511         }
512         value = value / count;
513     }
514
515     @Override
516     public Object getValue() {
517         return value;
518     }
519 }
520

```

521 Within Code Fragment 2, we can see how parameters are specified within
 522 plug-ins by means of the Java Annotation technology. The same technology is
 523 used to map the specific user interface components to the parameters: in our
 524 case, the annotation concerning acceptable types forces the interface to allow
 525 the sole selection of either integer or double attributes.

526 Table charts can be embedded within Rialto in a similar fashion, by providing
 527 classes which implement the following interface:

```

528 public interface Chart extends Visualizer {
529     void check(Metadata) throws RialtoException;
530     JComponent getChart();
531     DataCharts getDataChart();
532     void setTable(Table);
533 }
534
535 
```

536 In practice, a chart in Rialto is an object capable of providing both a DataChart
 537 and a JComponent. The first element represents the relevant summary informa-
 538 tion which can be extracted from the associated table, whereas the JComponent
 539 is a Java Swing Object implementing the visualization metaphor for the infor-
 540 mation contained within DataChart. As an example, the DataChart associated
 541 with a Pie Chart subsumes a table with the frequency of each value relative
 542 to the attribute under consideration, and the JComponent represents the panel
 543 where the Pie is built starting from the DataChart. One key feature of the Data
 544 Perspective is the capability to ease the filtering process by materializing data
 545 manipulations, which are directly performed within the data perspective, in a
 546 visual fashion.

547 4.2.4. Table Bridges

548 Bridges are used to store tables and guarantee efficient and transparent
 549 access to their data. In practice, a bridge can be regarded as a decoupling
 550 mechanism between logical data and its storage, thus making tasks independent
 551 of the physical implementation of tables. Every table in a workflow has a bridge
 552 attached to it and, vice versa, each bridge may have one or more attached tables.
 553 The scenario is exemplified in Figure 8, where three different bridges are used

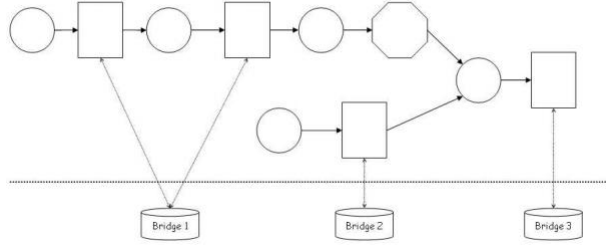


Figure 8: Bridges within a workflow

554 behind the scene to physically store the contents of the four tables available
 555 within the workflow. Each bridge is implemented as a plug-in of the platform.

556 Looking at Rialto from a data perspective, we can see that bridges create
 557 a two-level architecture: a first, abstract logical level where tables and models
 558 are semantically represented and manipulated. A second, physical level where
 559 tables are physically hosted in *ad hoc* structures. Rialto supports both main
 560 memory and mass memory bridges, thus making it possible to handle data sets
 561 of very large size.

562 A main memory bridge is essentially a component which provides a physical
 563 representation of a table as a matrix. Thus, a table bridged on main memory is a
 564 collection of rows represented as arrays of double values, where each value repre-
 565 sents either a numeric type or an index to a complex (nominal, string, EC) value,
 566 hosted in a specialized data structure. With reference to the Code Fragment
 567 1 shown in Section 4.2.2, the *bridge.generateTable* method creates one such a
 568 table in main memory. The operations on the row access the array representa-
 569 tion and store the values accordingly. By contrast, a mass memory bridge relies
 570 on a relational DBMS and exploits relational tables for storing and accessing
 571 data. Current mass memory bridges include both commercial DBMSs, namely,
 572 Oracle and MS-SQLServer, and Open DBMSs, namely, Postgres, MySQL and
 573 H2. Again, with reference to the data manipulation primitives of Section 4.2.2,
 574 operations on the table are translated into SQL statements, which are mediated

by an appropriate buffering strategy. Within a bridge, data structures and indexing methods are used to guarantee efficient implementation of the primitives for accessing and manipulating table components. For example, the access to an Event Collection (which can be seen as a non-normalized relation) can be optimized by relying on *ad hoc* structures based, for example, on inverted index structures.

The internal bridge behavior is exposed via easy APIs, enabling “fluent” declaration of complex filtering/sorting/randomization/etc. operations and, at the same time, optimal execution with respect to resources from the underlying storage medium. As an example, suppose that you want to manipulate table *customerTable*, by (a) selecting only rows having attribute *age* lower than 30, (b) excluding rows with one or more event in event collection *debts*, and (c) sorting the resulting rows by descending values of attribute *income*. In Rialto this is expressed as in the following code snippet:

```

589 Table customerTable = ...
590 TableMetadata metadata = customerTable.getMetadata();
591 IntegerAttribute age = metadata.getAttribute("age");
592 EventCollectionAttribute debts = metadata.getAttribute("debts");
593 DoubleAttribute income = metadata.getAttribute("income");
594 TableView myView = customerTable.having(age).lessThan(30)
595     .having(income).missing().sortBy(descending(income));
596

```

Here, we build a *TableView* object, which can be seen as a “normal” table in a workflow, except it is not materialized.

4.3. Models

Semantically, a model is an mathematical structure which subsumes intentional properties of a set of objects. Within the 2W, the objects are identified through entailment, whereas the properties are described through extension. Essentially, such properties can be induced by means of a specific mining algorithm, and can also be externalized by characterizing an entity according to

whether such properties hold for it. Models can be categorized as follows: Descriptive Models, which include Clustering and pattern models, and Predictive Models, which include classification, regression and outlier models. In Rialto, a model exhibits a specific signature as follows:

- A table schema, which characterizes the set of objects to which the model itself applies;
- A build operator, which implements a specific induction strategy;
- A PMML representation, which can be used to instantiate the model itself.

The hierarchy of all possible models which can be instantiated is shown in Figure 9. Here, we can see that each node (representing a type of model) provides a number of methods which actually represent the extension property, and that can be exploited by the application tasks in order to apply a model to a table (as detailed later in the section).

4.3.1. Model Visualizers

Besides the pre-defined visualizers, new additional visualization tools can easily be integrated by exploiting an abstract representation of model features based on an extended PMML representation. That is, models are described by using an appropriate XML representation, and visualizers take this representation as input to produce the desired visualization. This way, new visualizers can be created (as plug-ins of Rialto), independently of the algorithms producing a specific model. The general interface that a custom plugin should implement is similar to the Chart interface, with the only difference that the checking is relative to a model rather than a table:

```

public interface ModelVisualizer extends Visualizer {
    void check(Model) throws RialtoException;
    JComponent generateComponent();
}

```

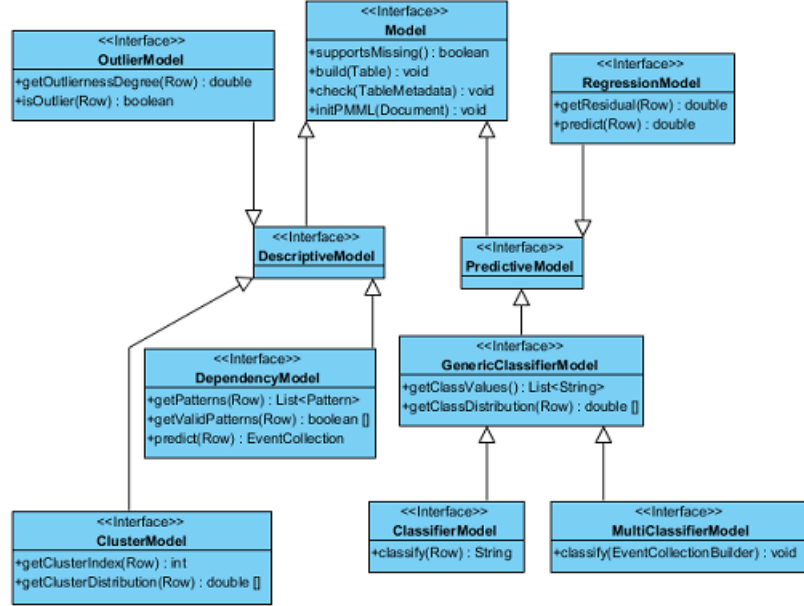


Figure 9: Model Hierarchy

4.4. Tasks

As we have seen in the previous sections, the algebraic operators of the 2W Model appear in a workflow as nodes of type *tasks*, that are implemented as Java plug-ins. Next we give some details about the different types of tasks.

Filtering tasks. They allow to acquire, export and modify tables and models. More specifically:

- *Acquisition* for importing tables or models from external sources. For example, we can exploit a CSV Acquisition task in order to acquire the content of a CSV file within a Rialto table, or a PMML representation of a decision tree into a model.
- *Export* for exporting tables and models to external destinations. For example, a CSV export task can pour the content of a table into a CSV

647 file, and a PMML model export can produce an XML file containing the
648 abstract representation of a given model.

649 • *Table filtering*, for modifying tables into new tables with derived infor-
650 mation. Example filtering tasks are selections of rows or columns, trans-
651 formation of values, joining of multiple tables according to some criteria.
652 SQL operations are supported natively within a SQL Filter, which per-
653 forms a specified SQL query on the input table to produce a table where
654 the result of the query is stored.

655 • *Model filtering* used to perform model post-processing. For example, a
656 tree-based model can be transformed into a rule-based model.

657 ***Mining tasks*** relate tables to models. Mining tasks are responsible for
658 invoking the *build()* method of the Model class (see top node in the hierarchy
659 of Figure 9) according to a specific building strategy, namely, by simple mining,
660 hold-out or cross-validation.

661 ***Application tasks*** enrich a table of entities with the properties which can
662 be derived from the given model. By default, we devise two main application
663 tasks, namely *Prediction Join* and *Description Join*. The basic idea is that a
664 table and a model can only join if the underlying schemas are compatible (that
665 is, the schema of the model is a subset of the schema of the table). The Pre-
666 diction/Description distinction depends on the operators which can be used to
667 enrich each entity with new properties. For example, a prediction join can be
668 applied to a table and a classifierModel produces a new table where the meta-
669 data associated with the table is enriched by a further nominal attribute, which
670 corresponds to the class which can be computed by the model on the tuples of
671 the table. By contrast, a dependency join applies to a table and a Dependen-
672 cyModel returns a table with a series of binary attributes, where each attribute
673 corresponds to the pattern associated with the model, and a true/false value
674 which denotes whether the pattern holds for the specific row under considera-
675 tion.

676

677 Again, the open architecture enables the definition of *ad hoc* tasks which can
678 implement user-specific operations. Structurally, the signature of a task requires
679 that a Task specifies some input and output ports, the metadata associated with
680 each port and the implementation of the *execute()* method. Code Fragment 3
681 describes the code snippet for an example task that takes two tables as input,
682 and returns as output a new table resulting as the merge of the input tables.

Code Fragment 3: An example Task

```

683
684 public class MergeTables implements Task {
685
686     @Input(name = "First Table")
687     @Order(0)
688     private Table firstTable ;
689
690     @Input(name = "Second Table", optional = true)
691     @Order(1)
692     private Table secondTable;
693
694     @Output(name = "Output Table")
695     @Order(0)
696     private WriteableTable outputTable;
697
698     @PreconditionCheck
699     public void check() {
700         if (secondTable != null
701             && !firstTable.getMetadata() != secondTable.getMetadata()) {
702             throw new RialtoException("Metadata differ on the input tables");
703         }
704     }
705
706     @Override
707     public TableMetadata[] getAllTargetMetadata() {
708         return new TableMetadata[] {
709             firstTable.getMetadata().createBuilder().build()
710         }
711     }
712
713     @Override
714     public void execute() {
715         for (Row row: firstTable) {
716             Row newRow = outputTable.getNewRow();
717             newRow.copyFrom(row);
718             outputTable.appendRow(newRow);

```

```

718     }
719     for (Row row: secondTable) {
720         Row newRow = outputTable.getNewRow();
721         newRow.copyFrom(row);
722         outputTable.appendRow(newRow);
723     }
724 }
725 }
726 }

```

727 4.5. Core

728 The Core component provides three basic functionalities, namely, *Workflow*
729 *Execution*, *Data Handling* and *Parallelization*.

730 **Workflow execution.** Rialto includes an engine for the efficient execution
731 of workflows. The workflow execution layer implements scheduling policies and
732 buffering strategies. The latter are aimed at “on the fly” processing of data
733 streams.

734 **Data Handling.** Rialto includes techniques for the efficient management
735 of data in main memory (under both main memory and mass memory bridges,
736 in the latter case relative to the portions of data which is buffered into main
737 memory). To this end, lossless data compression techniques, namely the LZ
738 (Lempel-Ziv) algorithms (Ziv and Lempel, 1977), are used.

739 **Parallelization.** Rialto provides support for three kinds of parallelism:
740 parallel execution of independent workflow branches, parallel table scan and
741 pipelining. Parallel execution of workflow branches is quite straightforward,
742 and it is enabled by static analysis of the workflow graph. Within parallel table
743 scan, large tables are automatically partitioned and each packet is assigned to
744 a different processor. The underlying paradigm here is that of processing single
745 packets and then merging the produced results. To this purpose, Rialto provides
746 a native object, namely *TaskExecutionSupport*, which can be exploited within
747 tasks or statistics. This object autonomously implements the most suitable
748 policy for processing a table, provided that a native method for processing single
749 rows is provided. Code Fragment 4 shows a naive task implementing the copy
750 of an input table to an output table. We can observe two components here: an

751 alternate version of the *execute* method, which uses a *TaskExecutionSupport*
 752 object, and a local object implementing a *RowProcessorBuilder* interface. The
 753 latter basically defines a method for processing an input row and producing an
 754 output row out from it. Within the *execute* method, a request is issued to the
 755 *taskExecutionSupport* object to scan the input table, process all rows using the
 756 *RowProcessor* object, and write the results to the output table. The policy for
 757 scanning the table and writing the results is totally transparent, as it is handled
 758 at system level by packetizing the table and processing each packet in parallel.

759 With reference to Code Fragment 3, it is possible to rewrite the *execute*
 760 method by exploiting the same row processor and a *TaskExecutionSupport* ob-
 761 ject. In practice, lines 29-38 of the method can be replaced by the lines

```
762 taskExecutionSupport.scan( firstTable ).with(new MyRowProcessorBuilder()).writingTo(outputTable);
763
764 taskExecutionSupport.scan(secondTable).with(new MyRowProcessorBuilder()).writingTo(outputTable);
765
```

766 thus demanding the management of the of the operations to the kernel (and
 767 thus exploiting the native parallelism).

768 Note that argument of *scan* can be any table which is “readable”: for example
 769 a table view generated through the bridge APIs would fit as well. The following
 770 code snippet is a refinement of the previous one, where tables are randomized
 771 prior to start the scan.

```
772 taskExecutionSupport.scan( firstTable .randomize()).with(new
773     ↪ MyRowProcessorBuilder()).writingTo(outputTable);
774
775 taskExecutionSupport.scan(secondTable.randomize()).with(new
776     ↪ MyRowProcessorBuilder()).writingTo(outputTable);
777
```

778 Pipelining can be seen as an enhancement of the above strategy. Within
 779 pipelining, each row of a table is forwarded to the next task as soon as it has been
 780 processed by the current task. Pipes can be assigned to multiple processors.

781 Pipelining is relative to tasks, and not all tasks can support it. It is up to
 782 the user to specify (and guarantee) whether the task can be piped, by overriding
 783 the method *canBePiped*. With reference to Code Fragment 4, it is possible to
 784 enable pipelining simply by adding the following:

```

785
786  @Override
787  public boolean canBePiped() {
788      return true;
789  }
790

```

791 When Rialto detects that two or more task instances, having workflow connec-
792 tions, can work in pipe, the latter are automatically replaced by a single virtual
793 task joining their behavior, but avoiding the burden of maintaining temporary
794 tables working as “buffers” between tasks. This can be especially effective when
795 such temporary, intermediate tables are very large.

Code Fragment 4: A simple task which seamlessly processes rows
in parallel by exploiting Rialto’s API

```

796
797
798  ...
799  @Override
800  // Core task method: Rialto expects the whole task logic is called from here
801  public void execute( final TaskExecutionSupport taskExecutionSupport) {
802      // The following line asks for copying the whole input table
803      // to output table via MyRowProcessorBuilder; Rialto automatically
804      // handles input/output buffering, task parallelization, and so on.
805      taskExecutionSupport.scan(getInputTable()).with(new
806          ↪ MyRowProcessorBuilder()).writingTo(getOutputTable());
807  }
808
809  private static final class MyRowProcessorBuilder implements RowProcessorBuilder {
810      @Override
811      public RowProcessor buildNew() {
812          return new AbstractRowTransformer() {
813              @Override
814              public void process( final Row inputRow, final Row outputRow) {
815                  // In this simple example output row is copied from input row, so it has the same content,
816                  // but Rialto API can be used to manipulate output accordingly
817                  outputRow.copyFrom(inputRow);
818              }
819          }
820      }
821  }
822

```

823 5. Expressiveness by examples: Rialto in Action

824 The purpose of this section is to show how Rialto works on different classes of
825 applications. To this end, we provide one case study on customer segmentation,
826 and another on text classification.

827 5.1. Case Study 1: Marketing Campaigns

828 ACME Household Cleaning Supplies manufactures and sells a wide array of
829 household cleaning supplies over the Internet. They have a very large customer
830 base for which they have gathered information about customers through prod-
831 uct registration cards and customer satisfaction surveys. Annually, ACME mails
832 holiday cards with a promotional voucher for a free sample product. In order
833 to reduce mailing costs and estimate the total cost of the promotion, ACME
834 wants to know which customers in their database are most likely to redeem
835 the vouchers. To this end, the workflow depicted in Figure 2 was developed.
836 *Data Acquisition.* The first step in almost every data mining workflow is the
837 importing of historical data. ACME’s data is contained in a comma-delimited
838 file, which can be imported using the CSV Acquisition task.

839 *Data Preparation.* Rialto comes prepackaged with many of the most common
840 preprocessing tasks. For the purpose of this case study, a first data preparation
841 task is concerned with missing values. To do so, we use the Missing Values
842 statistical tool in Rialto. Once executed, Rialto determines that there are 431
843 missing values in the WORKCLASS attribute and 433 missing values in the
844 OCCUPATION attributes. Now that we have identified that the data is miss-
845 ing values, we apply the Fill Missing Values task to automatically populate
846 the missing data. To this end, Rialto uses the most common values found in
847 WORKCLASS and OCCUPATION to replace the missing values. The second
848 data preparation task is that of putting data in buckets. Indeed, ACME tracks
849 the ages of their customers. ACME’s sales department thinks of customers in
850 the following age groups: under 25, 26 to 35, 36 to 60, and 61+. There are
851 several ways to place this data into buckets, one of which is utilizing Rialto’s

852 Custom Discretize task.

853 *Data Mining.* Now that data is prepared, we want to perform data mining on
854 it to see if an accurate predictive classifier for ACMEs marketing campaign can
855 be created. By using the Hold Out task, we separate data into a training set
856 and a test set, the latter allowing us to validate the accuracy of results. The
857 EC4.5 algorithm (an extension to the commonly used C4.5 algorithm) is then
858 run to produce a predictive classifier for ACME customers.

859 *Model application.* To test the quality of the learned model, the Prediction Join
860 task is applied to the test set. Then, the Confusion Matrix is computed, along
861 with the statistical accuracy of the model. Also, the resulting table *Table5*,
862 where the predicted labels are compared with those of the ideal classification,
863 is *exported* as a CSV file.

864 Once the trained classifier is available, it can be applied to determine which
865 new customers should receive the promotional mailer. To this end, a new work-
866 flow, to be deployed over Rialto server, is constructed. Of course, the predictive
867 classifier expects the data to be prepared in the same manner in which it was
868 trained. Therefore, data representing new customers must again be discretized
869 using the same ranges as the original data. ACME can now use classification
870 results to assess the cost of running a promotion with new customers. Now,
871 instead of sending out the promotions to all the new customers, ACME can
872 confidently send out the promotion to 169 of the total 1500 new customers. A
873 savings of eighty-nine percent!

874 5.2. Case Study 2: Text Classification

875 The application scenario is as follows. We want to induce a model for classify-
876 ing a set of medical abstracts from the OHSUMED data collection (Hersh et al.,
877 1994), a subset of the bibliographic database MEDLINE consisting of peer-
878 reviewed medical literature maintained by the National Library of Medicine.
879 To this end, we have available a subset of pre-classified documents, along with
880 two learning algorithms, namely, Complement Nave Bayes (Rennie et al., 2003)
881 and MaxEntropy (Nigam et al., 1999). Our goal is to compare how well they

perform, in order to select the best one (over the given data).

The OHSUMED subset we use for training is the collection consisting of the first 13,929 documents, from the 50,216 medical abstracts of the year 1991, labelled by the 23 diseases categories of the C codes of the Medical Subject Headings (MeSH) thesaurus (REF).

The KD process consists in the acquisition of the training documents, their pre-processing, models learning by CNB and MaxEnt and, finally, comparison of the induced models. The collapsed learning workflow (i.e., where table nodes are hidden) implementing this process is shown in Figure 10.

Data Acquisition. The input documents are represented as lines of a CSV file. Each line contains the document text (string) and the class labels. Data are acquired through CSV Acquisition task. Since features spaces in text categorization are typically very large (the high dimensionality problem) and, in addition, each document contains only a small subsets of features, it turns out that a document vector is normally largely sparse (i.e., too many 0s occur). Thus, a standard feature vector representation of texts would be highly inefficient. To overcome this drawback, the terms of each document are stored by an event collection (see Section 4.2.1)

Data Preparation. For the purpose of this study, we consider the following pre-processing steps:

- *N-gram Extraction* includes both stop-words and stop-tokens removal, and n-grams of up to 3 stem words are generated. This choice originates from the observation of the presence, in the medical vocabulary, of long phrases such as “immunologic deficiency syndromes”. Of course, the meaning of “immunologic deficiency syndromes” is very different from that of “deficiency” alone.
- *Concept Extraction* (CE) is performed by using the subset of the MeSH thesaurus made of the concept hierarchy having the 23 diseases categories (C codes) as roots. The advantage of CE is that, for example, terms like “arrhythmia” and “heart aneurysm” can merge into the most general

concept “heart diseases”, thus getting a stronger feature, able to better discriminate between the MeSH “cardiovascular diseases” category and the others.

- *Feature Selection* (FS) is performed by selecting the 500 highest scoring terms (a term is either a ngram or a concept) according to the Information Gain function (Forman, 2003). The advantage of applying FS to both grams and concepts is that of selecting discriminating terms with both statistical and semantic qualities (Lewis, 1992).

Preliminarily to the above steps, the *Remove Attributes* filter is applied to eliminate the attribute *Document_Name* (it has no effect on classification performance). After *Ngram Extraction* and *Concept Extraction* have been executed in parallel, the ngram-based and the concept-based document representations are merged by the *Tables Join* node, so as each document is represented in terms of both ngrams and concepts (both stored by event collections). The *Ngram Extraction* filter is set through a configuration mask providing stop-words and stop-tokens lists, maximum ngrams length, etc.. The *Concept Extraction* filter is set through a configuration panel whereby the MeSH thesaurus is provided (as an XML file).

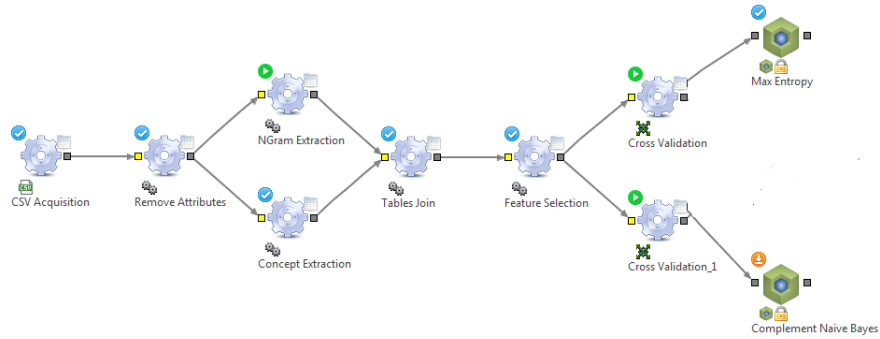


Figure 10: Compact representation of a workflow comparing two text classification algorithms

Data Mining. This step is performed by using both the Complement Nave

931 Bayes and MaxEntropy text classification algorithms. To validate and com-
932 pare the predictive capabilities of the induced models, five-fold cross validation
933 is performed. Overall classification performance are determined by calculating
934 Precision, Recall, and F1-measure metrics.

935 5.3. Discussion

936 The two case studies above were aimed at giving a flavor of how KD processes
937 can be accomplished by using Rialto. To this end, we used some of the pre-
938 packaged plug-ins for data acquisition, pre-processing, mining, etc., and showed
939 how they were connected to create KD workflows (one working on relational
940 data and one on texts). As we have seen, both workflows were characterized by
941 quite standard tasks: acquisition from CSV files, classification models induction
942 performed by classical algorithms (i.e., C4.5, Complement Naive Bayes and
943 MaxEntropy, the latter two for text classification), and data preparation, which
944 confirmed to be the most time consuming phase of KD - in particular, text pre-
945 processing in Case Study 2 consisted of five elementary tasks (from Attribute
946 Removal to Feature Selection - see Figure 10). Though not reported in the above
947 description for space reasons, the results of the two analyses can be inspected
948 by several view plug-ins, displaying data and models in diverse ways. For an
949 instance, one could display the table (not shown in Figure 10) obtained by
950 joining the output tables of the Ngram Extraction and Concept Extraction
951 tasks to inspect the bag of n-grams and the bag of concepts representing each
952 processed document (both bags stored as attributes of type Even Collection).
953 Different variants of the above workflows could be created in order to improve
954 the quality of the respective analyses. For instance, frequency analysis could
955 be performed in Case Study 2, to keep the highest frequent terms, simply by
956 adding a suitable node to the workflow of Figure 10 (Rialto provides plug-ins for
957 the computation of the most important frequency measures). Or, for another
958 example, more advanced feature extraction techniques, such as those based on
959 dependency graphs (Abdulsahib and Kamaruddin, 2015), (Vilares et al., 2015),
960 could be carried out to make machine learning methods better generalize. This

961 goal could be achieved by replacing in the workflow of Figure 10 the n-grams
 962 extractor node by a node implementing dependency-graph-based techniques.
 963 Both these examples show how the plug-in based architecture of Rialto makes
 964 the construction of a KD analysis a simple process, where tasks (nodes) can be
 965 assembled, removed, replaced, so allowing for quick and interactive changes to
 966 the analysis.

967 6. Discussion and Evaluation

968 We refer to (Chen et al., 2007), which identifies the following dimensions
 969 along which a data mining system should be evaluated (see also Section 1):

- 970 • *Managing large data volumes*: is the tool capable of manipulating large
 971 data volumes? Does this require changes in the architecture?
- 972 • *Acquisition from data sources*: can the tool gather data from diverse data
 973 sources, such as databases, Excel, CSV files, etc.?
- 974 • *Mining models*: is the choice of mining algorithms available in the tool
 975 rich enough to support real-world data mining applications?
- 976 • *Data preprocessing*: does the tool offer support to pre-process data before
 977 applying mining algorithms? Is there some guidance in the process?
- 978 • *Data visualization*: has the tool some efficient data visualizations, in order
 979 to support decision making and informed data manipulation?
- 980 • *Open Architecture*: is the tool extendable? Is the architecture open? Is
 981 there a development environment?
- 982 • *Open Community*: is there a thriving user community? Does it have the
 983 ability to develop new features and/or fix bugs?

984 Figure 11 provides a picture of the positioning of Rialto relative the above-
 985 mentioned dimensions.

	Description
Managing large amounts of data	Bridges paradigm allows Rialto to transparently use mass memory or databases, thus allowing the processing of large data volumes without changing the system architecture.
Acquisition from data sources	Rialto acquires dataset in CSV, Excel, ARFF, and others. Rialto also exploits the JDBC API, allowing import of data from database management systems (DBMSs) such as Oracle, MySQL, SQL Server, PostgreSQL, and many others.
Mining models	Rialto provides several algorithms for classification and clustering, like SSEM, Naive Bayes, Maximum Entropy, EC4.5, Rule Learner, CBOD, Travel CBOD, Dolphin, Linear Regression, EM, K-Means, Association Miner.
Data preprocessing	Rialto provides a wide variety of functionalities for data manipulation through a number of predefined plug-ins implementing filter tasks. In addition, it offers a number of facilities at Data Perspective level, where the choice of the available tools is guided through contextual menus.
Data visualization	The system offers a number of graphical tools for data and model visualization; new ad-hoc visualizers can be built and integrated, if necessary.
Open Architecture	Rialto provides a powerful and intuitive API for developing new plugins. Thus, new types of statistics, tasks, models, visualizers, etc., are easily implemented. The new plug-ins transparently take advantage of all the mechanisms of parallelism, pipelining, and memory management. The development environment of Rialto is based on opensource Eclipse IDE.
Open Community	Rialto does not have an open community, as it is not distributed in open source. However, it is widely used in academic environments, and this involves a continuous tool evolution, along with a growing availability of plugins.

Figure 11: Rialto Evaluation

986 *6.1. Experimental Analysis*

987 In this section we report the experimentation performed with the aim of
988 assessing a number of architectural choices affecting the capability of Rialto to
989 scale to large data volumes. To this end, there are some specific features we
990 have analyzed:

- 991 • How bridges affect the performance. Since the architecture strongly relies
992 on the idea of semantic abstraction, the coupling between the semantic
993 part and the data engine has to be carefully analyzed. In particular, with
994 data which are disk-resident, it is crucial to quantify the overhead with
995 respect to memory-resident data, and to make sure that such an overhead
996 is not affected by the data size.
- 997 • Pipelining. Since pipes can be assigned to multiple processors, it is worth
998 measuring how enabling this feature within the tasks improves the overall
999 performance.
- 1000 • Scalability of the parallel table scan. The *TaskExecutionSupport* compo-
1001 nent provides mechanism for parallel processing of elementary user-defined
1002 operations on tables. However, it introduces some overheads that need to
1003 be quantified.
- 1004 • Event Collections. These components are particularly suited for managing
1005 high-dimensional data. We perform some comparative measurements on
1006 memory usage and efficient support to native filtering operations.

1007 It is worth noticing that these choices are a direct consequence of the underlying
1008 model, which enables modularity in the execution of data mining queries. Thus,
1009 the experiments are aimed at demonstrating that those architectural choices
1010 actually optimize the executions, and that they can be profitably exploited in
1011 suitable query evaluation plans based on the underlying algebra.

1012 *Bridges.* In a first set of experiments, we exploit a simple workflow which
1013 loads a table and builds a model on such a table. We use a simple Naive Bayes
1014 Model, which scans the input table to collect statistics. Figure 12 reports the

1015 performance (time in ms) of three different bridges for increasing size of the input
 1016 data. We compare the Rialto's in-memory bridge, and two SQL-based bridges.
 1017 One is based on the H2 open source java database engine¹. We instantiate the
 1018 bridge using the in-memory table storing of H2. Thus, the H2 bridge is just
 1019 an alternative memory bridge which supports SQL queries natively. The other
 1020 SQL-based bridge relies on the widely known commercial MySQL engine².

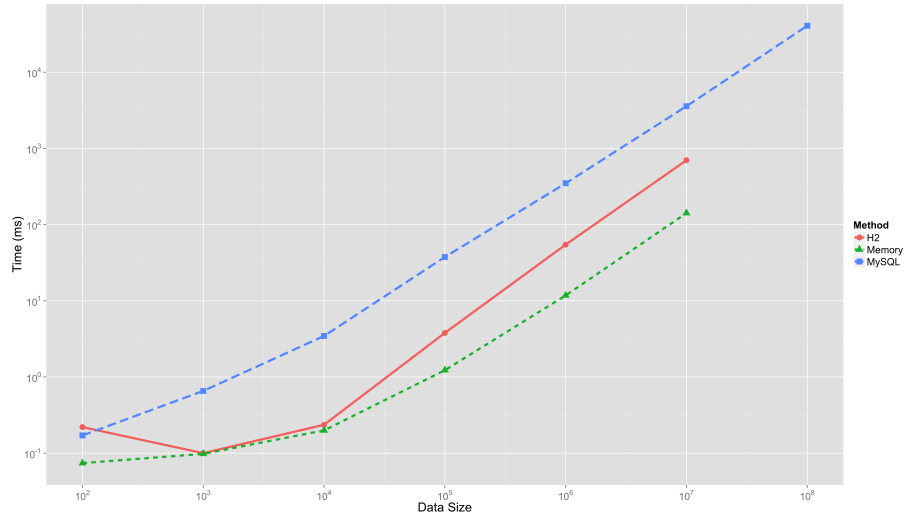


Figure 12: Bridges comparison - Rialto memory bridge (lower dotted line) vs H2 bridge (solid line) and MySQL bridge (upper dotted line)

1021 How we can see from the graph, there is a constant factor which differentiates
 1022 the time performances of the different bridges (so, the overhead is not affected
 1023 by the data size). Also, notice that the in-memory bridges suffer of memory
 1024 limitations, as with the current experimental configuration they are only capable
 1025 of handling tables with at most 10⁷ tuples, whereas the same limitation does
 1026 not hold for the MySQL bridge.

1027 *Pipelining.* In a second set of experiments, we measure the effectiveness of

¹See <http://www.h2database.com> for details.

²See <http://www.mysql.com>.

1028 pipelining. To this purpose, we setup a simple workflow which exploits the Naive
 1029 Bayes model built in the previous set of experiments, and applies it to a new
 1030 table. Thus, the piped tasks in this scenario are *data acquisition* followed by
predictive join (Naive Bayes is applied to the acquired table). Figure 13 shows

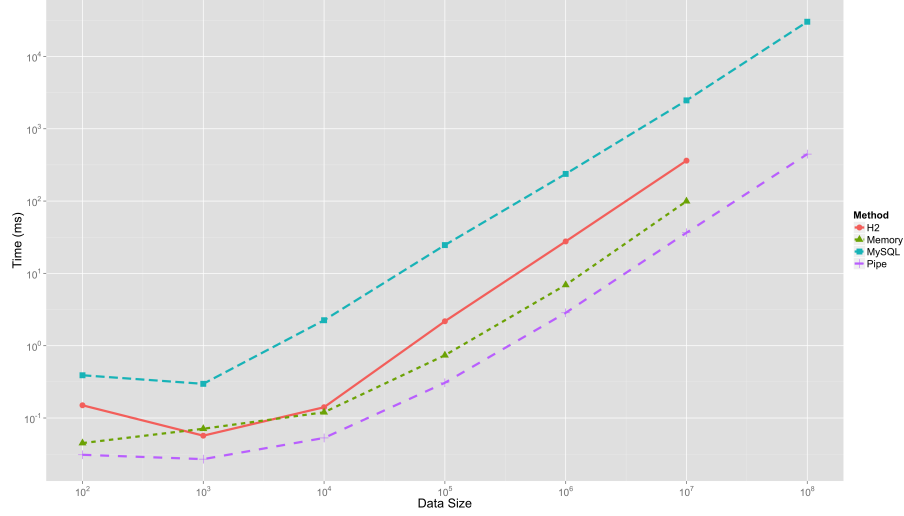


Figure 13: Time performance of piped execution (lower dotted line) vs sequential execution (the three upper lines are those reported in Figure 12).

1031 the performance of the piped version of the tasks (bottom line) compared to the
 1032 sequential versions (where the predictive join is run after the input table has
 1033 been *entirely* created) obtained by using the same bridges of the previous ex-
 1034 periments. There are a few aspects that are worth being mentioned. First, the
 1035 pipelining clearly outperforms all other methods in terms of efficiency. Further,
 1036 since intermediate tables can be deemed as volatile within a workflow, the frag-
 1037 ment of the workflow which makes use of pipelining is essentially independent
 1038 from the size of the input tables. This can be appreciated within the graph by
 1039 noticing that the pipe (bottom line) does not suffer from memory limitations of
 1040 the other in-memory methods (lines in the middle): tuples are passed away as
 1041 long as they are processed, and hence the only memory requirement is the need
 1042 to store a single tuple.

1044 *Parallelism.* The next set of experiments measures the speed-up of the par-
1045 allel table scan. To this purpose, we implemented a filtering task which takes
1046 a table in input, and produces the same table as output. Within the *Row-*
1047 *ProcessBuilder* we inserted a delay mechanism which simulates time consuming
1048 operations on a single row. The length of the delay is parameterized and ranges
1049 from $0\mu s$ to $300\mu s$ in our experiments.

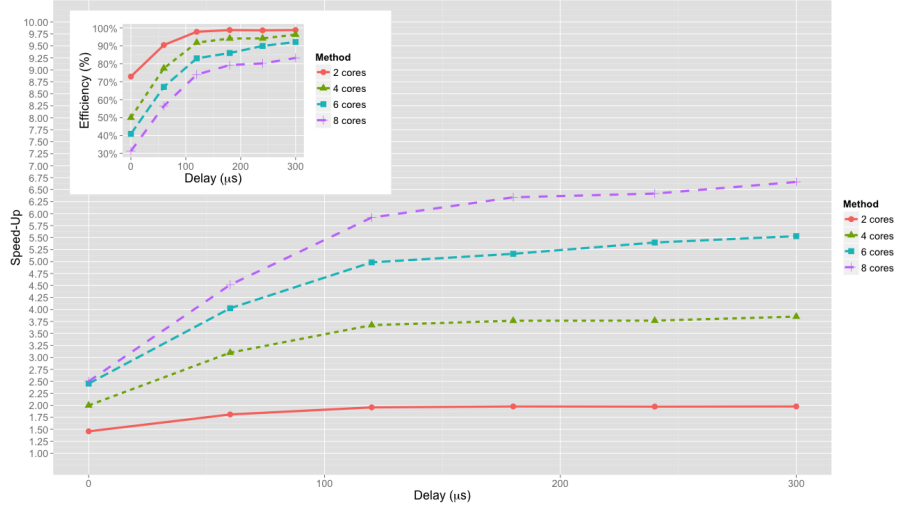


Figure 14: Parallel table scan speed-ups in main memory with 2-, 4-, 6- and 8-core processors

1050 Figure 14 shows the speed-up achieved by running the task on different
1051 architectures. The baseline is a processor with a single core, and the figure shows
1052 the speed-ups with processors having number of cores equal to 2 (bottom line),
1053 4, 6 and 8 (upper line). We can see that, in general, (1) the higher the delay,
1054 the higher the speed-up with a given core configuration, and (2) the higher the
1055 number of cores, the lower the speed-up. In particular, we can observe that the
1056 highest speed-up is obtained with a 2-core processor (varying from 1.50 to nearly
1057 2), while the maximum speed-up of the 8-core configuration is 6.75. In the left
1058 upper corner, the figure shows the efficiency in the exploitation of the different
1059 core configurations. The efficiency represents the percentage of proximity to
1060 the ideal situation: intuitively, a 2-core configuration is 100% efficient if it is

1061 twice faster than the single-core configuration (speed-up=2). Similarly, 4-cores
 1062 should be 4 times faster, and so on. We can notice that efficiency increases as
 1063 long as delay increases and, in general, efficiency values are quite satisfactory -
 1064 ranging from 80% (8 cores) to 100% (2 cores).

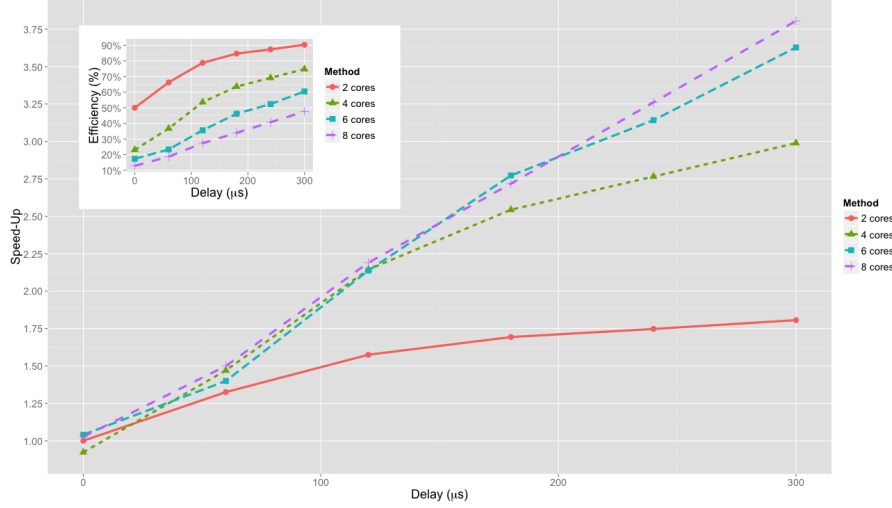


Figure 15: Parallel table scan speed-ups in external memory with 2-, 4-, 6- and 8-core processors

1065 Figure 15 shows the speed-up when an out-of-core bridge is exploited. The
 1066 situation in this case is more problematic, since the retrieval and saving of the
 1067 tuples clearly requires access to external memory. The speed-ups in this case
 1068 are halved and for increasing cores the efficiency cuts significantly.

1069 *Event Collections.* There are essentially two aspects we need to face, namely,
 1070 memory consumption and efficiency of native filtering operations. In order to
 1071 stress the system on the above aspects and to measure its response, we exploited
 1072 some synthetic datasets, containing 10,000 transactions each. A transaction
 1073 consists of a set of items in the form $\{i_1, \dots, i_n\}$ where i_j uniquely identifies
 1074 an item picked from an initial set $\mathcal{I}$. These datasets simulate several real-life
 1075 situations, like e.g. purchase transactions, where $\mathcal{I}$ represents the set of pos-
 1076 sible items which can be purchased, and each transaction represents a set of

1077 purchases relative to an individual; or textual data, where $\mathcal{I}$ represents a vocab-
 1078 ulary and a transaction represents a text document, irrespective of the order of
 1079 words. In the generation process, we vary the size of $\mathcal{I}$ from 10^2 to 10^6 . Then,
 1080 each dataset is generated by repeatedly performing, for each transaction, the
 1081 following steps: (1) sample the size of the transaction through a Poisson distri-
 1082 bution with fixed λ parameter, (2) sample a Multinomial distribution on $\mathcal{I}$ from
 1083 a Dirichlet distribution with fixed prior α , and (3) populate the transaction by
 1084 repeatedly sampling from the multinomial distribution sampled in the previous
 1085 step.

1086 The datasets were generated in a relational format as text files where each
 1087 line represents the pair `(TransactionID, {ItemID})`. We then built a workflow
 1088 which acquires the text file and builds a table T , where each row represents
 1089 a transaction. There are two alternative ways of representing the items of a
 1090 single transaction in T . Either as a binary tuple with fixed length $|\mathcal{I}|$, or as
 1091 an event collection. We built two alternative versions of the workflow, where the
 1092 acquisition task was customized to produce either of the two representations.
 1093 Next, the workflow proceeds with a filtering task which aims at removing the
 1094 top 10% more frequent items.

1095 The rationale of this workflow is twofold. First of all, by monitoring the
 1096 acquisition task, we can measure the memory usage in both representations.
 1097 The binary representation (referred to as *Full Table* in the following) acts as
 1098 a full matrix. By contrast, the representation of a transaction as an event
 1099 collection (referred to as EC) allows us to measure the compression ratio and
 1100 the possible overhead due to internal indices required for the event collection.
 1101 Figure 16 plots the memory usage after the loading of the transactions for both
 1102 representations. To this end, both the in-memory and the H2 bridges are used.
 1103 We can observe that the memory usage of the EC version of the table is low
 1104 and substantially independent from the size of the dataset $\mathcal{I}$, whereas the Full
 1105 Table representation grows very fast with it (upper line). That is, ECs enable
 1106 an effective and compact representation of very large sets of transactions.

1107 Now, let us consider the filtering operation removing frequent items from

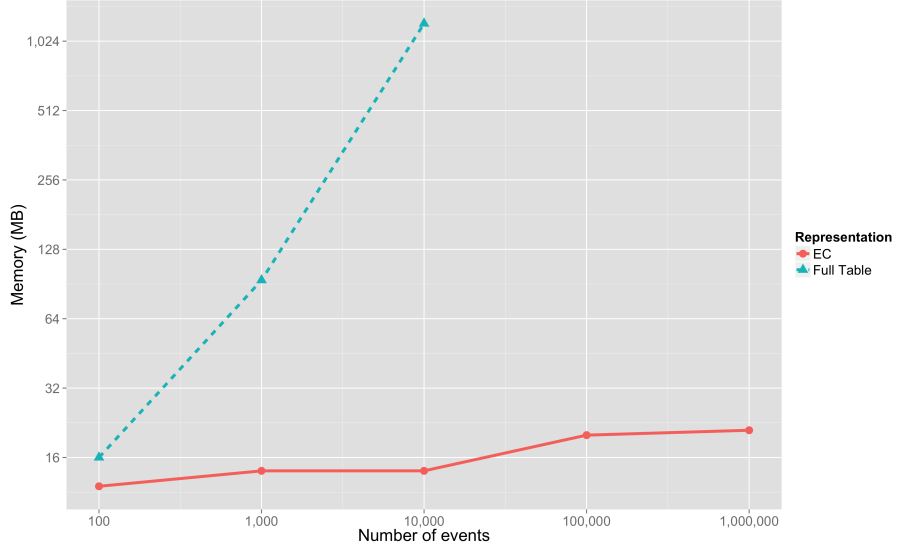


Figure 16: Memory usage when transactions are represented as (1) binary tuples with fixed length (dotted line) and (2) event collections (solid line).

1108 $\mathcal{I}$. In principle, event collections should be more sensitive to such an operation:
 1109 removing a column from a table should be straightforward, whereas removing
 1110 an event from a collection requires searching for such an event. Thus, this
 1111 simple filtering operation is a good test for assessing the performance of ECs
 1112 manipulations. We can see the performance of the removal operation in Figure
 1113 17. Again, the scan of the EC representation is quite constant. By contrast,
 1114 removing a column on the Full Table representation requires to restructure the
 1115 metadata associated with the table, which is a costly operation when the number
 1116 of attributes is overwhelming high.

1117 **7. Ongoing Work**

1118 Besides the above described features, there are also a number of features
 1119 that we are currently working on in order to enhance the Rialto capabilities,
 1120 notably:

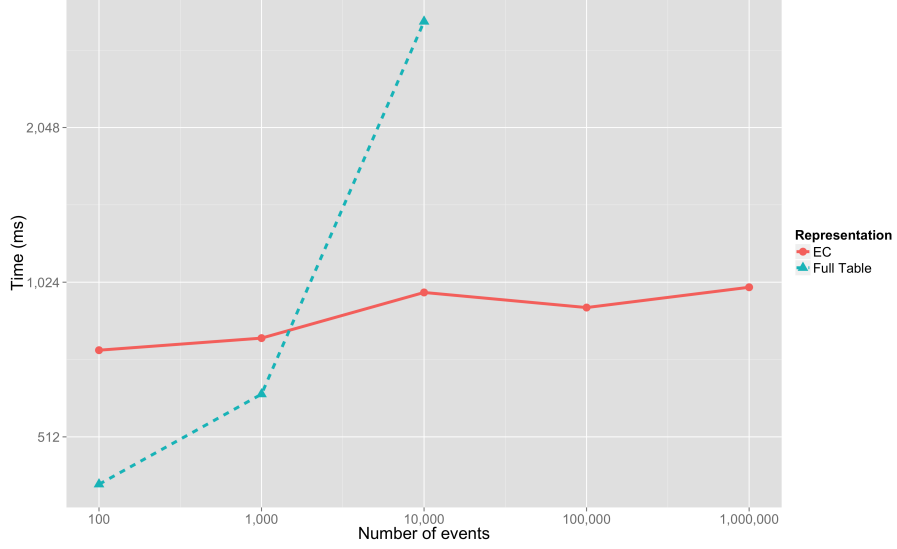


Figure 17: Filtering operation costs in case transactions are represented as (1) binary tuple with fixed length (dotted line) and (2) event collections (solid line)

- 1121 • *Data Streams.* In the current implementation, we assume that data within

1122 Rialto is made of either tables or models. However, in the same paradigm,

1123 we can extend such basic types to comprise data streams, i.e., memoryless

1124 tables with a timestamp associated. In practice, in our model a stream has

1125 the same features of a table (metadata and rows). However, the content

1126 of the table changes over time and it is relative to a specific time window.

1127 Under this perspective, the paradigm of tasks has to be reconfigured as

1128 well, to enable transformation from streams to tables and models, and

1129 viceversa.
- 1130 • *Spatio-temporal data.* Besides the current types, we are working on ex-

1131 tending attributes types to support spatial and spatio-temporal data. This

1132 is especially important for complex applications that require the analysis

1133 of geo-reference data or moving objects.
- 1134 • *Support to Hadoop.* The simple form of parallelism which are provided

1135 to the user, namely parallel table scan and pipelining, can be further en-
1136 hanced by considering the current technologies. We are currently working
1137 on extending the Rialto architecture to support the MapReduce program-
1138 ming model, by developing bridges under the Hadoop framework.

1139 8. Conclusions

1140 In this paper we have first proposed a general framework for modeling KD
1141 processes, named 2W Model, a variant of the 3W Model proposed in (Johnson
1142 et al., 2000). According to this model, the essence of a KD process can be
1143 regarded as the interaction between two separated worlds: the data and the
1144 model world. This way, any knowledge discovery process can be formalized as an
1145 algebraic expression, that is essentially a composition of operators representing
1146 (pseudo)elementary operations on the two worlds.

1147 Then, we have described a KD platform, called Rialto, which relies on the 2W
1148 Model. To this end, it provides a visual interface whereby a KD process can be
1149 modeled as a workflow, where each node is either a table (data world), or a model
1150 (model world) or a task representing an operator allowing a transition between
1151 the two worlds. Using this metaphor, Rialto provides support to all steps of
1152 a KD process, from data acquisition to model exploration and deployment. In
1153 addition, in order to cope with real-world analytical problems, where new *ad*
1154 *hoc*, specific tasks may be needed, Rialto allows incorporation of new (Java)
1155 plug-ins within an open architecture.

1156 Efficiency and scalability are achieved thanks to suitable design choices in-
1157 cluding bridges, event collections and parallelism. We have conducted a number
1158 of experiments showing the effectiveness of the proposed architecture to scale
1159 to large volumes of data.

1160 Download of Rialto can be found at [http://www.exeura.eu/products/](http://www.exeura.eu/products/rialto)
1161 **rialto**.

1162 **References**

- 1163 Abdulsahib, A. K., Kamaruddin, S. S., June 2015. Graph based text representa-
1164 tion for document clustering. *Journal of Theoretical and Applied Information*
1165 *Technology* 76 (1).
- 1166 Berthold, M., et al., 2009. Knime: The konstanz information miner. *SIGKDD*
1167 *Explorations* 11 (1).
- 1168 Blei, D., Ng, A., Jordan, M., March 2003. Latent dirichlet allocation. *Journal*
1169 *of Machine Learning Research* 3, 993–1022.
- 1170 Borgelt, C., 2009. Graph mining: An overview. In: *Proc. 19th Workshop on*
1171 *Computational Intelligence*. KIT Scientific Publishing.
- 1172 Calders, T., Lakshmanan, L. V. S., Ng, R. T., , Paredaens, J., 2006. Expressive
1173 power of an algebra for data mining. *ACM Transactions on Database Systems*
1174 31 (4), 1169–1214.
- 1175 Chaudhuri, S., Narasayya, V., Sarawagi, S., 2002. Efficient evaluation of queries
1176 with mining predicates. In: *Proc. 18th International Conference on Data*
1177 *Engineering (ICDE'03)*. pp. 529–540.
- 1178 Chen, X., Ye, Y., Williams, G., Xu, X., 2007. A survey of open source data
1179 mining systems. *Lecture Notes in Computer Science* 4819, 3–14.
- 1180 chung Fu, T., 2011. A review on time series data mining. *Engineering Applica-*
1181 *tions of Artificial Intelligence* 24 (1), 164–181.
- 1182 Forman, G., 2003. An extensive empirical study of feature selection metrics for
1183 text classification. *Journal of Machine Learning Research* 3, 1289–1305.
- 1184 Gaber, M. M., Zaslavsky, A., Krishnaswamy, S., 2005. Mining data streams: A
1185 review. *ACM SIGMOD Record* 34 (2), 18–26.
- 1186 Giannotti, F., Manco, G., Turini, F., 2004. Specifying mining algorithms with
1187 iterative user-defined aggregates. *IEEE Transactions on Knowledge and Data*
1188 *Engineering* 16 (10), 1232–1246.

- 1189 Hall, M., et al., 2009. The weka data mining software: An update. SIGKDD
1190 Explorations 11 (1).
- 1191 Han, J., Kamber, M., 2001. Data mining: concepts and techniques. Morgan
1192 Kaufmann.
- 1193 Hersh, W., Buckley, C., Leone, T., Hickman, D., 1994. Ohsumed: An interac-
1194 tive retrieval evaluation and new large text collection for research. In: Proc.
1195 17th ACM Int. Conf. on Research and Development in Information Retrieval
1196 (SIGIR'94). pp. 192–201.
- 1197 Hristovski, D., Friedman, C., Rindflesch, T. C., Peterlin, B., 2008. Literature-
1198 based knowledge discovery using natural language processing. Information
1199 Science and Knowledge Management 15, 133–152.
- 1200 Imielinski, T., Mannila, H., 1996. A database perspective on knowledge discov-
1201 ery. Communications of the ACM 39 (11), 58–64.
- 1202 Imielinski, T., Virmani, A., 1999. MSQL: A query language for database mining.
1203 Data Mining and Knowledge Discovery 3 (4), 373–408.
- 1204 Johnson, T., Lakshmanan, L. S., Ng, R. T., 2000. The 3w model and algebra for
1205 unified data mining. In: Proc Int. Conf. on Very Large Databases (VLDB'00).
1206 pp. 21–32.
- 1207 Lewis, D. D., 1992. An evaluation of phrasal and clustered representations on
1208 a text categorization task. In: Proc. 15th ACM Int. Conf. on Research and
1209 Development in Information Retrieval (SIGIR'92). pp. 37–50.
- 1210 Manco, G., et al., 2008. Querying and reasoning for spatiotemporal data mining.
1211 In: Mobility, Data Mining and Privacy - Geographic Knowledge Discovery.
1212 Springer, pp. 335–374.
- 1213 Mikut, R., Reischl, M., 2011. Data mining tools. Wiley Interdisciplinary Re-
1214 views: Data Mining and Knowledge Discovery 1 (5), 431–443.

1215 Nigam, K., Lafferty, J., McCallum, A., 1999. Using maximum entropy for text
1216 classification. In: Proc. IJCAI-99 Workshop on Machine Learning for Infor-
1217 mation Filtering. pp. 61–67.

1218 Ortale, R., et al., 2008. The DAEDALUS framework: progressive querying and
1219 mining of movement data. In: Proc. 16th ACM SIGSPATIAL Int. Symp. on
1220 Advances in Geographic Information System. p. 52.

1221 Pietramala, A., Policicchio, V. L., Rullo, P., Sidhu, I., 2008. A genetic algorithm
1222 for text classification rule induction. In: Proc. European Conf. on Machine
1223 Learning and Knowledge Discovery in Databases (ECMLPKDD’08). Vol. 5212
1224 of Lecture Notes in Computer Science. pp. 188–203.

1225 Quinlan, J. R., 1987. Generating production rules from decision trees. In: Proc.
1226 10th International Joint Conference on Artificial Intelligence (IJCAI’87). pp.
1227 304–307.

1228 Rennie, J. D., Shih, L., Teevan, J., Karger, D. R., 2003. Tackling the poor
1229 assumptions of naive bayes text classifiers. In: Proc. Int. conf. on Machine
1230 Learning (ICML’03). pp. 616–623.

1231 Rullo, P., Policicchio, V., Cumbo, C., Iiritano, S., August 2009. Olex: effective
1232 rule learning for text categorization. IEEE Transactions on Knowledge and
1233 Data Engineering 21 (8), 1118–1132.

1234 S., M., 2005. Data Streams: Algorithms and Applications. Now Publishers Inc.

1235 Sebastiani, F., 2002. Machine learning in automated text categorization. ACM
1236 Computing Surveys 34 (1), 1–47.

1237 Shearer, C., 2000. The crisp-dm model: The new blueprint for data mining.
1238 Journal of Data Warehousing 5, 13–22.

1239 Vilares, M., Fernandez, M., Blanco, A., November 2015. Supporting knowledge
1240 discovery for biodiversity. Data and Knowledge Engineering 100, 34–53.

- 1241 Wang, H., Zaniolo, C., Luo, C., 2003. ATLAS: A small but complete SQL
1242 extension for data mining and data streams. In: Proc. Int. Conf. on Very
1243 Large Databases (VLDB'03). pp. 1113–1116.
- 1244 Ziv, J., Lempel, A., 1977. A universal algorithm for sequential data compression.
1245 IEEE Transactions on Information Theory IT-23, 337–343.