



7-1-2004

The Piazza peer data management system

Alon Y. Halevy
University of Washington

Zachary G. Ives
University of Pennsylvania, zives@cis.upenn.edu

Jayant Madhavan
University of Washington

Peter Mork
University of Washington

Dan Suciu
University of Washington

See next page for additional authors

Copyright 2004 IEEE. Reprinted from *IEEE Transactions on Knowledge and Data Engineering*, Volume 16, Issue 7, July 2004, pages 787-798.
Publisher URL: <http://ieeexplore.ieee.org/xpl/tocresult.jsp?isNumber=29215&puNumber=69>

This material is posted here with permission of the IEEE. Such permission of the IEEE does not in any way imply IEEE endorsement of any of the University of Pennsylvania's products or services. Internal or personal use of this material is permitted. However, permission to reprint/republish this material for advertising or promotional purposes or for creating new collective works for resale or redistribution must be obtained from the IEEE by writing to pubs-permissions@ieee.org. By choosing to view this document, you agree to all provisions of the copyright laws protecting it.

This paper is posted at ScholarlyCommons. http://repository.upenn.edu/cis_papers/47
For more information, please contact repository@pobox.upenn.edu.

The Piazza peer data management system

Abstract

Intuitively, data management and data integration tools should be well-suited for exchanging information in a semantically meaningful way. Unfortunately, they suffer from two significant problems: They typically require a comprehensive schema design before they can be used to store or share information and they are difficult to extend because schema evolution is heavyweight and may break backward compatibility. As a result, many small-scale data sharing tasks are more easily facilitated by non-database-oriented tools that have little support for semantics. The goal of the peer data management system (PDMS) is to address this need: We propose the use of a decentralized, easily extensible data management architecture in which any user can contribute new data, schema information, or even mappings between other peers' schemas. PDMSs represent a natural step beyond data integration systems, replacing their single logical schema with an interlinked collection of semantic mappings between peers' individual schemas. This paper describes several aspects of the Piazza PDMS, including the schema mediation formalism, query answering and optimization algorithms, and the relevance of PDMSs to the Semantic Web.

Keywords

peer data management, data integration, schema mediation, web, databases

Comments

Copyright 2004 IEEE. Reprinted from *IEEE Transactions on Knowledge and Data Engineering*, Volume 16, Issue 7, July 2004, pages 787-798.

Publisher URL: <http://ieeexplore.ieee.org/xpl/tocresult.jsp?isNumber=29215&puNumber=69>

This material is posted here with permission of the IEEE. Such permission of the IEEE does not in any way imply IEEE endorsement of any of the University of Pennsylvania's products or services. Internal or personal use of this material is permitted. However, permission to reprint/republish this material for advertising or promotional purposes or for creating new collective works for resale or redistribution must be obtained from the IEEE by writing to pubs-permissions@ieee.org. By choosing to view this document, you agree to all provisions of the copyright laws protecting it.

Author(s)

Alon Y. Halevy, Zachary G. Ives, Jayant Madhavan, Peter Mork, Dan Suciu, and Igor Tatarinov

The Piazza Peer Data Management System

Alon Y. Halevy, Zachary G. Ives, Jayant Madhavan, Peter Mork, Dan Suciu, and Igor Tatarinov

Abstract—Intuitively, data management and data integration tools should be well-suited for exchanging information in a semantically meaningful way. Unfortunately, they suffer from two significant problems: They typically require a comprehensive schema design before they can be used to store or share information and they are difficult to extend because schema evolution is heavyweight and may break backward compatibility. As a result, many small-scale data sharing tasks are more easily facilitated by non-database-oriented tools that have little support for semantics. The goal of the peer data management system (PDMS) is to address this need: We propose the use of a decentralized, easily extensible data management architecture in which any user can contribute new data, schema information, or even mappings between other peers' schemas. PDMSs represent a natural step beyond data integration systems, replacing their single logical schema with an interlinked collection of semantic mappings between peers' individual schemas. This paper describes several aspects of the Piazza PDMS, including the schema mediation formalism, query answering and optimization algorithms, and the relevance of PDMSs to the Semantic Web.

Index Terms—Peer data management, data integration, schema mediation, Web, databases.

1 INTRODUCTION

WHILE databases and data management tools excel at providing semantically rich data representations and expressive query languages, they have historically been hindered by a need for significant investment in design, administration, and schema evolution. Schemas must generally be predefined in comprehensive fashion rather than evolving incrementally as new concepts are encountered; schema evolution is typically heavyweight and may “break” existing queries. As a result, many people find that database techniques are obstacles to lightweight data storage and large-scale data sharing tasks, rather than facilitators. They resort to simpler and less expressive tools, ranging from spreadsheets to text files, to store and exchange their data. This provides a simpler administrative environment (although some standardization of terminology and description is always necessary), but with a significant cost in functionality. Worse, when a lightweight repository grows larger and more complex in scale, there is no easy migration path to a semantically richer tool.

Conversely, the strength of HTML and the World Wide Web has been easy and intuitive support for ad hoc extensibility—new pages can be authored, uploaded, and quickly linked to existing pages. However, as with flat files, the Web environment lacks rich semantics. Initially, that shortcoming spurred a movement toward XML, which allows data to be semantically tagged. Unfortunately, XML carries many of the same requirements and shortcomings as data management tools: For rich data to be shared among

different groups, all concepts need to be placed into a common frame of reference. XML schemas must be completely standardized across groups or mappings must be created between all pairs of related data sources.

More recently, the desire to complement the Web with more semantics spurred the vision of the Semantic Web [1], which calls for sharing of structured data at Web scale, with queries spanning large numbers of Web sites. Much of the research focused on the Semantic Web is based on treating the Web as a knowledge base defining concepts and relationships. In particular, researchers have developed knowledge representation languages for representing *meanings*—relating them within custom ontologies for different domains—and *reasoning* about the concepts. The best-known example is RDF and the languages that build upon it: RDF Schema and OWL [2]. While there has been much investigation of how to define the meaning of data locally, the issues of large-scale data sharing have yet to be addressed.

Data integration systems have been proposed as a partial solution to the problem of large-scale data sharing [3], [4], [5], [6], [7], [8], [9], [10]. These systems support rich queries over large numbers of autonomous, heterogeneous data sources by exploiting the semantic relationships between the different sources' schemas. An administrator defines a global *mediated schema* for the application domain and specifies semantic mappings between the sources and the mediated schema. We get the strong semantics needed by many applications, and data sources can evolve independently—and, it would appear, relatively flexibly. Yet, in reality, the mediated schema, the integrated part of the system that actually facilitates all information sharing, becomes a bottleneck in the process. Mediated schema design must be done carefully and globally; data sources cannot change significantly or they might violate the mappings to the mediated schema; concepts can only be added to the mediated schema by the central administrator. The ad hoc extensibility of the Web is missing and, as a

• A.Y. Halevy, J. Madhavan, P. Mork, D. Suciu, and I. Tatarinov are with the Department of Computer Science and Engineering, University of Washington, Box 352350, Seattle, WA 98195-2350.

E-mail: {alon, jayant, pmork, suciu, igor}@cs.washington.edu.

• Z.G. Ives is with the Department of Computer and Information Science, University of Pennsylvania, 3330 Walnut Street, Philadelphia, PA 19104-6389. E-mail: zives@cis.upenn.edu.

Manuscript received 24 June 2003; revised 5 Sept. 2003; accepted 10 Sept. 2003.

For information on obtaining reprints of this article, please send e-mail to: tkde@computer.org, and reference IEEECS Log Number TKDESI-0104-0603.

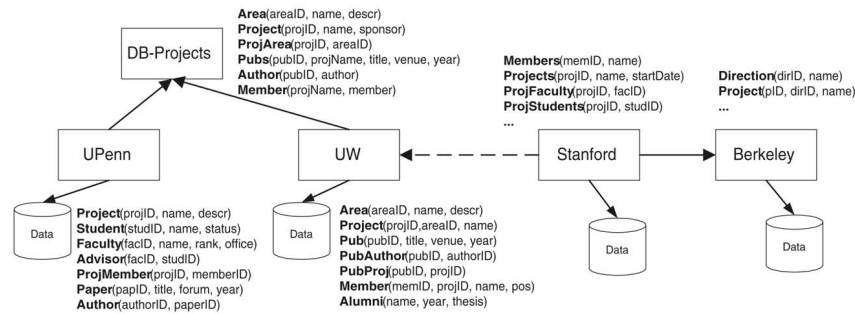


Fig. 1. A PDMS for the database research domain. Arrows indicate that there is (at least a partial) mapping between the relations of the peer schemas. The figure illustrates how two semantic networks can be joined by establishing a single mapping between a pair of peers (UW and Stanford in this case).

result, data integration systems provide limited support for large-scale data sharing.

We believe that there is a clear need for a new class of data sharing tools that preserves semantics and rich query languages, but which facilitates ad hoc, decentralized sharing, and administration of data and defining of semantic relationships. Every participant in such an environment should be able to contribute new data and relate the data to existing concepts and schemas, define new schemas that others can use as frames of reference for their queries, or define new relationships between existing schemas or data providers. We believe that a natural implementation of such a system will be based on a peer-to-peer architecture and, hence, call such a system a *peer data management system* (PDMS). The vision of a PDMS is to blend the extensibility of the HTML Web with the semantics of data management applications.

Example 1.1. The extensibility of a PDMS can best be illustrated with a simple example. Fig. 1 illustrates a peer data management system for supporting a Web of database research-related data. This will be a running example throughout the paper, so we only describe the functionality here. Unlike a hierarchy of data integration systems or mediators, a PDMS supports any arbitrary network of relationships between peers. The true novelty lies in the PDMS's ability to exploit transitive relationships among peers' schemas. The figure shows that two semantic networks can be fully joined together with only a few mappings between similar members of each semantic network (in our example, we only required a single mapping). The new mapping from Stanford to UW enables any query at any of the five peers to access data at *all* other peers through transitive evaluation of semantic mappings. Importantly, the mappings can be defined between the most similar nodes in the two semantic networks; this is typically much easier than attempting to map a large number of highly dissimilar schemas into a single mediated schema (as in conventional data integration).

This paper describes the main contributions of the Piazza PDMS that we have been building at the University of Washington. Most of the data integration literature is based on the relational data model, largely because it is the simplest and cleanest model in which to define properties

and analyze complexity; accordingly, we begin our description of the Piazza system using the relational context. Section 2 presents the logical model underlying and it defines the problem of semantic mediation in a PDMS. Section 3 outlines some of the theoretical results concerning query answering and mediation in a PDMS. Section 4 describes a query answering algorithm for Piazza and explains some of the current research directions we are pursuing for query optimization. One of these directions is the development of techniques for *mapping composition*, which we explain in the Appendix (which can be found on the Computer Society Digital Library at <http://computer.org/tkde/archives.htm>).

Although the relational model is ideal for defining properties of the PDMS, in the real world, XML is the most useful representation for sharing semantically rich data: Most relational and semistructured data sources export to XML and XML is also used for the RDF data format that has been the focus of efforts in the Semantic Web. Thus, our actual Piazza system uses XML as the data model. In Section 5, we describe Piazza's XML support; the Appendix (which can be found on the Computer Society Digital Library at <http://computer.org/tkde/archives.htm>) explains how Piazza can provide an infrastructure for supporting both XML and RDF-based Semantic Web applications.

2 LOGICAL MODEL OF THE PDMS

We begin with a description of the logical model underlying a PDMS (defined using the relational data model; Section 5 details how this definition changes for XML data). Informally, a PDMS consists of a set of data sources (also known as peers) and they are related through *semantic mappings*. A PDMS can be viewed as a strict generalization of data integration systems. In our discussion, we assume a relational data model, we focus on select-project-join queries with a set semantics, and we use the notation of conjunctive queries. In this notation, joins are specified by multiple occurrences of the same variable. Unless explicitly specified, we assume queries do not contain comparison predicates (e.g., $\neq$, $<$). Views refer to named queries.

We assume that each peer defines its own relational *peer schema* whose relations are called *peer relations*; a query in a PDMS will be posed over the relations from a specific peer

schema. Without loss of generality, we assume that relation and attribute names are unique to each peer.

Peers may also contribute data to the system in the form of *stored relations*. Stored relations are analogous to data sources in a data integration system: All queries in a PDMS will be reformulated strictly in terms of stored relations that may be stored locally or at other peers. (Note that not every peer needs to contribute stored relations to the system as some peers may strictly serve as logical mediators to other peers.) We assume that the names of stored relations are distinct from those of peer relations.

Example 2.1. In our example PDMS in Fig. 1, only peer relations are shown. The lines between peers indicate that there is a mapping (described later) between the two peers.

Stored relations containing actual data are provided by the universities: the UPenn, UW, Stanford, and Berkeley peers. DB-Projects is a virtual peer that provides a uniform view over the domain. Stanford and Berkeley, as neighboring universities, came to an agreement to map their schemas directly. The flexibility of the PDMS (due to its ability to evaluate transitive relationships between schemas) becomes evident when two PDMSs are joined. In our example, once a mapping between the Stanford-Berkeley PDMS and the UPenn-UW-DBProjects PDMS is established, queries over any of the five peers will be able to access *all* of the stored relations.

Note that our approach can support evolving schemas very naturally. A new schema version can be treated as an additional peer schema. In general, the new version is likely to be very similar to the previous version, making the problem of specifying a mapping between the versions rather easy. In addition, the resulting mapping is likely to be very accurate.

2.1 Mapping Language for PDMS

Obviously, the power of the PDMS lies in its ability to exploit semantic mappings between peer and stored relations. In particular, there are two types of mappings that must be considered: mappings describing the data within the stored relations (generally, with respect to one or more peer relations) and mappings between peer schemas. At this point, it is instructive to recall the formalisms used in the context of data integration systems since we build upon them in defining our mapping description language.

2.1.1 Mappings in Data Integration

Data integration systems provide a uniform interface to a multitude of data sources through a logical, *mediated* schema. Mappings are established between the mediated schema and the relations at the data sources, forming a two-tier architecture in which queries are posed over the mediated schema and evaluated over the underlying source relations. A data integration system can be viewed as a special case of a PDMS.

Two main formalisms have been proposed for schema mediation in data integration systems. In the first, called *global-as-view* (GAV) [11], [3], [4], [5], the relations in the

mediated schema are defined as views over the relations in the sources. In the second, called *local-as-view* (LAV) [6], [7], [8], the relations in the sources are specified as views over the mediated schema. In fact, in many cases, the source relations are said to be *contained* in a view over the mediated schema as opposed to being exactly equal to it. We illustrate both below.

Example 2.2. The DB-Projects' *Member* peer relation, which mediates UPenn and UW peers, may be expressed using a GAV-like definition. The definition specifies that *Member* in DB-Projects is obtained by a union over the UPenn and UW schemas. Note in our examples, that peer relations are named using a *peer-name:relation-name* syntax:

```
DBProjects : Member(projName, member) : –
    UPenn : Student(sid, member, _),
    UPenn : ProjMember(pid, sid),
    UPenn : Project(pid, projName, _)
DBProjects : Member(projName, member) : –
    UPenn : Faculty(sid, member, _),
    UPenn : ProjMember(pid, sid),
    UPenn : Project(pid, projName, _)
DBProjects : Member(projName, member) : –
    UW : Member(pid, member, _),
    UW : Project(pid, _, projName).
```

We may use the LAV formalism to specify the UW peer relations as views over mediated DB-Projects relations. This formalism is especially useful when there are many data sources that are related to a particular mediated schema. In such cases, it is more convenient to describe the data sources as views over the mediated schema rather than the other way around. In our scenario, DB-Projects may eventually mediate between many universities and, hence, LAV is appropriate for future extensibility. The following illustrates an LAV mapping for UW:

```
UW : Project(projID, areaID, projName) ⊆
    DBProjects : Project(projID, projName),
    DBProjects : ProjArea(projID, areaID).
```

The fundamental difference between the two formalisms is that GAV specifies how to extract tuples for the mediated schema relations from the sources and, hence, query answering amounts to view unfolding. In contrast, LAV is source-centric, describing the contents of the data sources. Query answering requires algorithms for answering queries using views [12], but, in exchange, LAV provides greater extensibility: The addition of new sources is less likely to require a change to the mediated schema.

Before proceeding, we note that all the languages we discuss here are for mapping *schemas*, rather than *data values*. For example, it is very common that a person or company name may appear differently in two data sources. The topic of object identification is currently a very active area of research and beyond the scope of this paper. In commercial systems (e.g., [13]), the problem has usually been addressed

by *concordance tables*, which are binary tables relating the different ways of referring to the same object.

2.1.2 Mappings for PDMS

We now present the relational-model version of *PPL*, the Piazza Peer Language. Our goal in *PPL* is to preserve the features of both the GAV and LAV formalisms, but to extend them from a two-tiered architecture to our more general network of interrelated peer and source relations. Semantic relationships in a PDMS will be specified between pairs (or small sets) of peer (and, optionally, source) relations. Ultimately, a query over a given peer relation may be reformulated over source relations on any peer in the transitive closure of peer mappings.

First, we formally define our two types of mappings, which we refer to as storage descriptions and peer mappings.

Storage descriptions. Each peer contains a (possibly empty) set of *storage descriptions* that specify the data stored at a peer by relating its stored relations to its peer relations. Formally, a storage description is of the form $A : R = Q$, where Q is a conjunctive query over the schema of peer A and R is a stored relation at the peer. The description specifies that A stores in relation R the result of the query Q over its schema.

In many cases, the data that is stored is not *exactly* the definition of the view, but only a subset of it. As in the context of data integration, this situation arises often when the data at the peer may be incomplete (this is often called the *open-world assumption* [14]).¹ Hence, we also allow storage descriptions of the form $A : R \subseteq Q$. We call the latter descriptions *containment* (or inclusion) descriptions versus *equality* descriptions.

Example 2.3. A storage description might relate the stored students relation at peer UPenn to the peer relations:

$$\begin{aligned} \text{UPenn} : \text{students}(\text{sid}, \text{name}, \text{advisor}) &\subseteq \\ \text{UPenn} : \text{Student}(\text{sid}, \text{name}, -), \\ \text{UPenn} : \text{Advisor}(\text{sid}, \text{fid}), \\ \text{UPenn} : \text{Faculty}(\text{fid}, \text{advisor}, -, -). \end{aligned}$$

This storage description says that **UPenn:students** stores a subset of the join of **Student**, **Advisor**, and **Faculty**, which reflects the fact that **UPenn:students** is unlikely to contain information about *all* students in the world; it will probably contain data on “local” students only. Hence, if a **UPenn:Affiliation** peer relation with the corresponding semantics was available, the above storage description could be specified more precisely as follows:

$$\begin{aligned} \text{UPenn} : \text{students}(\text{sid}, \text{name}, \text{advisor}) &= \\ \text{UPenn} : \text{Student}(\text{sid}, \text{name}, -), \\ \text{UPenn} : \text{Advisor}(\text{sid}, \text{fid}), \\ \text{UPenn} : \text{Faculty}(\text{fid}, \text{advisor}, -, -), \\ \text{UPenn} : \text{Affiliation}(\text{sid}, ' \text{UPenn}'). \end{aligned}$$

1. Sometimes, it may be possible to describe the exact contents of a data source with a more refined query, but very often this cannot be done.

Peer mappings. *Peer mappings* provide semantic glue between the schemas of different peers. We have two types of peer mappings in *PPL*. The first are *inclusion* and *equality* mappings (similar to the concepts for storage descriptions). In the most general case, these mappings are of the form $Q_1(\bar{A}_1) = Q_2(\bar{A}_2)$ (or $Q_1(\bar{A}_1) \subseteq Q_2(\bar{A}_2)$ for inclusions), where Q_1 and Q_2 are conjunctive queries with the same arity and $\bar{A}_1$ and $\bar{A}_2$ are *sets* of peers. Query Q_1 (Q_2) can refer to any of the peer relations in $\bar{A}_1$ ($\bar{A}_2$, respectively). Intuitively, such a statement specifies a semantic mapping by stating that evaluating Q_1 over the peers $\bar{A}_1$ will always produce the same answer (or a subset in the case of inclusions) as evaluating Q_2 over $\bar{A}_2$. Note that, since *PPL* allows queries on both sides of the equation, we can accommodate both GAV and LAV-style mappings (and, thus, we can express any of the types of mappings used in data integration. This approach is also known as GLAV [12].

The second kind of peer mappings are called *definitional mappings*. They are datalog rules whose relations (both head and body) are peer relations, i.e., the body cannot contain a query. Formally, as long as a peer relation appears only once in the head of a definitional description, such mappings can be written as equalities. We include definitional mappings in order to obtain the full power of GAV mappings. We distinguish definitional mappings for the following reasons:

- As we show in Section 3, the complexity of answering queries when equality mappings are restricted to being definitional is more attractive than the general case.
- Definitional mappings can easily express disjunction, e.g., $P(x) : \neg P_1(x)$ and $P(x) : \neg P_2(x)$ means that P is the union of P_1 and P_2 (while the pair of mappings $P(x) = P_1(x)$ and $P(x) = P_2(x)$ means that P , P_1 , and P_2 are equal).

In summary, a PDMS N is specified by a set of peers $\{P_1, \dots, P_n\}$, a set of peer schemas $\{S_1, \dots, S_m\}$, and a mapping function from peers to schemas, a set of stored relations $\mathcal{R}_i$ at each peer P_i , a set of peer mappings $\mathcal{L}_N$, and a set of storage descriptions $\mathcal{D}_N$. The storage descriptions and peer mappings provided by a peer P_i may reference stored or peer relations defined by other peers; thus, any peer can extend another peer's relations or use its data.

2.2 Semantics of PPL

Given the peer and stored relations, their mappings, and a query over some peer schema, the PDMS needs to answer the query using the data from the stored relations. To formally specify the problem of query answering, we need to define the semantics of queries. We show below how the notion of *certain answers* [14] from the data integration context can be generalized to our context. Our goal is to formally define what the set of answers to a conjunctive query Q posed over the relations of a peer A is. The challenge arises because the peer schemas are virtual; in fact, some data may only exist partially, if at all, in the system.

Formally, we assume that we are given a PDMS N and an instance for the stored relations, D , i.e., a set of tuples $D(R)$ for each stored relation $R \in (\mathcal{R}_1 \cup \dots \cup \mathcal{R}_n)$. A *data*

instance I for a PDMS N is an assignment of a set of tuples to each relation in each peer (both the peer and stored relations). We denote by $I(R)$ the set of tuples assigned to the relation R by I and we denote by $Q(I)$ the result of computing the query Q over the extensional data in I . To define certain answers, we will consider only the data instances that are consistent with the specification of N :

Definition 2.1 (Consistent data instance). A data instance I is said to be consistent with a PDMS N and an instance D for N 's stored relations if:

1. For every storage description in $\mathcal{D}_N$, if it is of the form $A : R = Q_1$ ($A : R \subseteq Q_1$), then $D(R) = Q_1(I)$ ($D(R) \subseteq Q_1(I)$).
2. For every peer description in $\mathcal{L}_N$:
 - If it is of the form $Q_1(\mathcal{A}_1) = Q_2(\mathcal{A}_2)$, then $Q_1(I) = Q_2(I)$,
 - If it is of the form $Q_1(\mathcal{A}_1) \subseteq Q_2(\mathcal{A}_2)$, then $Q_1(I) \subseteq Q_2(I)$, and
 - If it is a definitional description whose head predicate is p , then let $r_1, \dots, r_m$ be all the definitional mappings with p in the head and let $I(r_i)$ be the result of evaluating the body of r_i on the instance I . Then, $I(p) = I(r_1) \cup \dots \cup I(r_m)$.

Intuitively, a data instance I is consistent with N and D if it describes one possible state of the world (i.e., extension for each of the peer relations) that is allowable given the data and peer mappings and D . We define the certain answers to be those that hold in every possible consistent data instance:

Definition 2.2 (Certain answers). Let Q be a query over the schema of a peer A in a PDMS N , and let D be an instance of the stored relations of N . A tuple $\bar{a}$ is a certain answer to Q if $\bar{a}$ is in $Q(I)$ for every data instance that is consistent with N and D .

Note that, in item 2 of Definition 2.1, we did not require that the extension of p be the least-fixed point model of the datalog rules. However, since we defined certain answers to be those that hold for every consistent data instance, we actually do get the intuitive semantics of datalog for these mappings.

Query answering. Now, we can define the query answering problem for PDMS: Given a PDMS N , an instance of the stored relations D , and a query Q , find all certain answers of Q .

3 COMPLEXITY OF QUERY ANSWERING

Before we present an algorithm for answering queries in Piazza (the focus of Section 4), it is important to understand how tractable query answering is in the PDMS model. In this section, we briefly review some basic results relating to the complexity of finding certain answers with our PDMS mapping language (full details are given in [15]).

The computational complexity of finding all certain answers is well understood for the data integration context with a two-tiered architecture of a mediator and a set of data sources [14]. Here, we show the complexity of query answering in the global context of a PDMS when

the data integration formalisms are used locally. The complexity will depend on the restrictions we impose on peer mappings in PPL .

The focus of our analysis is on data complexity—the complexity of query answering in terms of the total size of the data stored in the peers. Typically, the complexity of query answering is either polynomial, Co-NP-hard but decidable, or undecidable. In the polynomial case, it is possible to find a reformulation of the query into a query that refers only to the stored relations and this reformulation is then optimized and executed. In the latter two cases, it is not possible to find all certain answers efficiently, but it is possible to develop an efficient reformulation algorithm that does not provide all certain answers, but that only returns certain answers.

A key result. Cyclicity of peer mappings plays a very significant role in the complexity of answering queries.

Definition 3.1 (Acyclic peer mappings). A set $\mathcal{L}$ of inclusion peer mappings in PPL is said to be acyclic if the following directed graph is acyclic. The graph contains a node for every peer relation mentioned in $\mathcal{L}$. There is an arc from the node corresponding to R to the node corresponding to S if there is a peer description in $\mathcal{L}$ of the form $Q_1(\mathcal{A}_1) \subseteq Q_2(\mathcal{A}_2)$, where R appears in Q_1 and S appears in Q_2 .

The following theorem characterizes two extreme cases of query answering in PDMS:

Theorem 3.1. Let N be a PDMS specified in PPL .

1. The problem of finding all certain answers to a conjunctive query Q , for a given PDMS N , is undecidable.
2. If N includes only inclusion peer mappings and storage descriptions and the peer mappings are acyclic, then a conjunctive query can be answered in polynomial time data complexity.

The proof of this theorem appears in [15]. The difference in complexity between the first and second bullets shows that the presence of cycles is the culprit for achieving query answerability in a PDMS (note that equalities automatically create cycles). In a sense, the theorem also establishes a limit on the arbitrary combination of the formalisms of LAV and GAV.

The second bullet points out a powerful schema mediation language for PDMS for which query answering can be done efficiently. It shows that LAV and GAV style reformulations can be chained together arbitrarily and extends the results of [16], which combined one level of LAV followed by one level of GAV.

In [15], we also show how the complexity is affected by allowing comparison predicates, identify a few cases where query answering in cyclic PDMS is tractable, and establish the complexity of the PDMS consistency problem. In summary, while the arbitrary use of the data integration formalisms in a PDMS, query answering is undecidable, we have identified a large and useful subset of PPL for which query answering is tractable. Our results are based on an open-world assumption [14] in which peers have incomplete rather than full information. The closed-world assumption, which is necessary for supporting negation in mappings and queries, is known to make the problem of

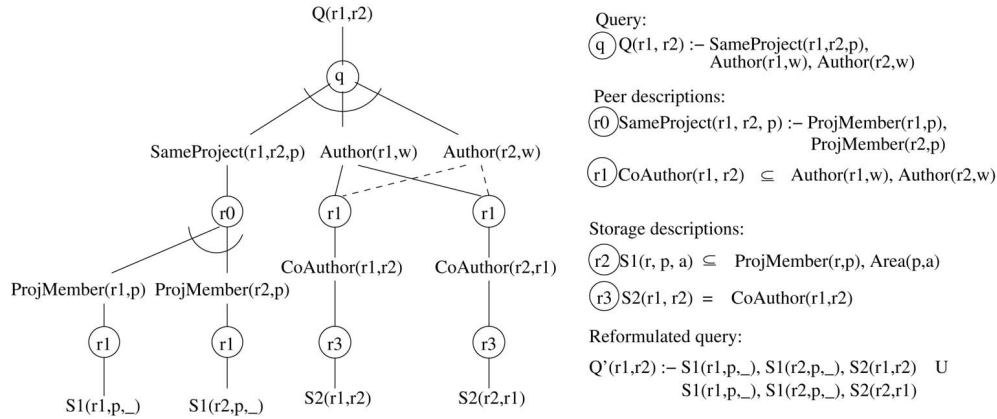


Fig. 2. Reformulation rule-goal tree for Emergency Services domain. Dashed lines represent nodes that are included in the *unc* label (see text).

finding all certain answers much harder (co-NP hard in the size of the data [14])—even for two peers (and, in fact, even without negation).

4 QUERY REFORMULATION ALGORITHM

In this section, we describe an algorithm for query reformulation for PDMSs. The input of the algorithm is a set of peer mappings and storage descriptions and a query Q . The output of the algorithm is a query expression Q' that refers to stored relations *only*. To answer Q , we need to evaluate Q' over the stored relations. The precise method of evaluating Q' is beyond the scope of this paper, but we note that recent techniques for adaptive query processing [17] are well-suited for our context. Furthermore, in this discussion, we assume that all the peer mappings are available at a single location and, hence, all the reformulation is done in a single place. We are currently investigating methods for distributed query reformulation.

The algorithm is sound and complete in the following sense: Evaluating Q' will always *only* produce certain answers to Q . When all the certain answers can be found in polynomial time (according to Section 3), Q' will produce all certain answers. Q' is called the *maximally contained rewriting* of Q [6]: It is a query over the sources that produces all the answers to Q that are possible from any such query.

Before we describe the algorithm, we first provide some intuition on its working and the challenges it faces. Consider a PDMS in which all peer mappings are definitional (similar to GAV mappings in data integration). In this case, the algorithm is a simple construction of a rule-goal tree: Goal nodes are labeled with atoms of the peer relations, and rule nodes are labeled with peer mappings. We begin by expanding each query subgoal according to the relevant definitional peer mappings in the PDMS. When none of the leaves of the tree can be expanded any further, we use the storage descriptions for the final step of reformulation in terms of the stored relations.

At the other extreme, suppose all peer mappings in the PDMS are inclusions in which the left-hand side has a single atom (similar to LAV mappings in data integration). In this case, we begin with the query subgoals and apply an algorithm for answering queries using views (e.g., [12]). We

apply the algorithm to the result until we cannot proceed further and, as in the previous case, we use the storage descriptions for the last step of reformulation.

The first challenge of the complete algorithm is to combine and interleave the two types of reformulation techniques. One type of reformulation replaces a subgoal with a set of subgoals, while the other replaces a set of subgoals with a single subgoal. The algorithm will achieve this by building a rule-goal tree, while it simultaneously marks certain nodes as covering not only their parent node, but also their uncle nodes. We illustrate the algorithm by an example below.

Example 4.1. To illustrate the rule-goal tree,² Fig. 2 shows an example for a simple query. We begin with the query, Q , which asks for researchers who have worked on the same project and also coauthored a paper. Q is expanded into its three subgoals, each of which appears as a goal node. The **SameProject** peer relation (indicating which researchers work on the same project) is involved in a single definitional peer description (r_0), hence we expand the **SameProject** goal node with the rule r_0 and its children are two goal nodes of the **ProjMember** peer relation (each specifying the projects an individual researcher is involved in).

The **Author** relation is involved in an inclusion peer description (r_1). We expand **Author**($r1, w$) with the rule node r_1 and its child becomes a goal node of the relation **CoAuthor**. This “expansion” is of a different nature because of the LAV-style reformulation. Intuitively, we are reformulating the **Author**($r1, w$) subgoal to use the left-hand side of r_1 . The right-hand side of r_1 includes two subgoals of **Author** (with the appropriate variable patterns), so we also mark r_1 as covering its *uncle* node. (In the figure, this annotation is indicated by a dashed line.) Since the peer relation **Author** is involved in a single peer description, we do not need to expand the subgoal **Author**($r2, w$) any further. Note, however, that we must apply description r_1 a second time with the head variables reversed since **CoAuthor** may not be symmetric (because it is $\subseteq$ rather than $=$).

2. More precisely, we actually build a rule-goal DAG, as illustrated in the example.

At this point, since we cannot reformulate the peer mappings any further, we consider the storage descriptions. We find stored relations for each of the peer relations in the tree (S_1 and S_2), and produce the final reformulation. Reformulations of peer relations into stored relations can also be either in GAV or LAV style. In this simple example, our reformulation involves only one level of peer mappings, but, in general, the tree may be arbitrarily deep.

4.1 Optimizations

The second challenge we face is that the rule-goal tree may be huge. First, the tree may be very deep because it may need to follow any path through semantically related peers. Second, the branching factor of the tree may be large because data is replicated at many peers. Hence, it is crucial that we develop effective methods for pruning the tree and for generating first solutions quickly. It is important to emphasize that, while many sophisticated methods have been developed for constructing rule-goal trees in the context of datalog analysis (e.g., [18], [19]), the focus in these works has been developing termination criteria that provide certain guarantees, rather than optimizing the construction of the tree itself.

Several optimization methods can immediately be borrowed from the techniques developed for evaluation of datalog and logic programs, but *lifted* from the data level to the expression level: 1) memoization of nodes and 2) detection of dead ends and useless paths [15].

A more subtle case in which useless paths can be detected is as follows: Suppose we have two sibling goal nodes with labels $p_1(\bar{X})$ and $p_2(\bar{Y})$ and suppose that p_1 appears in a *single* inclusion peer description of the form $V(\bar{Z}) \subseteq p_1(\bar{X}), p_2(\bar{Y})$ and that predicate p_2 appears on the right-hand side of numerous inclusion peer mappings. In this case, the only way to reformulate p_1 will be through V , and V already satisfies the subgoal $p_2(\bar{Y})$. Hence, there is no need to explore any of the other ways of reformulating p_2 : They are all redundant.

Finally, it is likely that many paths in the PDMS will be traversed frequently and, therefore, we would like to develop a set of techniques that may judiciously *precompose* a select set of mapping chains in the network. Composition, in this context, means the following: Given semantic mappings between data sources A and B , and between B and C , is it possible to generate a direct mapping between A and C that is equivalent to the original mappings? Here, equivalence means that, for any query in a given class of queries Q and for any instance of the data sources, using the composed mapping yields exactly the same answer that would be obtained by the two original mappings.

The composition problem is also relevant to a number of static analysis questions that arise in a PDMS. First, by precomputing the composition of mappings, we can also remove redundancies from the resulting reformulation, leading to significant runtime savings. Second, we would like to find redundant paths in the network: Two paths between a pair of nodes A and B are equivalent if, given any query on A , reformulating the query along both paths will result in equivalent queries on B . Third, we note that data from source A can be used in source B only if the necessary

concepts are modeled in each of the nodes on the path between A and B . As a result, when paths in the network get longer, we may witness *information loss*. Hence, we would like to determine whether a path between A and B can possibly be useful for some query and if not, find the weak links and try to improve the mappings there. The problem of mapping composition is discussed in detail in the Appendix (which can be found on the Computer Society Digital Library at <http://computer.org/tkde/archives.htm>).

5 THE PIAZZA PDMS

Early in this paper, our focus has been understanding query answering in the PDMS from a formal perspective. Our goal has been to establish the semantics of our system implementation, which we call Piazza. Our prototype Piazza system is designed to be a scalable foundation for data sharing applications.

In this section, we discuss two aspects of our Piazza implementation that are of special interest. First, we overview our architecture for query answering in the PDMS. Second, in any practical system designed to integrate data, XML makes a much better interchange format than relational data. We provide an overview of the XML version of Piazza's $\mathcal{PPL}$ language.

5.1 System Architecture and Implementation

A Piazza PDMS consists of a set of nodes (physical peers) connected in an overlay network on the Internet. Following the logical PDMS model, every peer may have a peer (mediated) schema and it may optionally provide source data and query processing capabilities. Finally, a peer may specify schema mappings. We note that, in contrast to P2P file sharing systems, we assume the PDMS to be a relative stable environment: Joining a PDMS is a heavyweight operation and data contributors are unlikely to be modem users, so we expect that peers will seldom leave (and if they do, they will notify the system).

Query reformulation is performed at the node that receives the query—this allows enumeration of all possible rewritings and detection of redundant rewritings. The reformulator assumes a global *system catalog* that provides access to all of the mappings that involve a particular peer relation. At each step in the rule-goal tree expansion, the catalog is consulted to expand the frontier. We currently use a centralized catalog and cache the mappings at each peer for performance and robustness. The rewritings from the reformulator are ultimately pipelined to the query processor at the originating node and it is responsible for delegating portions of the query plan to other nodes.

We believe that this area has many opportunities for future work. We plan to reimplement the catalog using distributed hash table techniques (e.g., [22], [23]), which will increase scalability and robustness. We hope to investigate improvements to the rewriting algorithm, both in terms of optimizations and in terms of distributing the work. Finally, we hope to investigate adaptive approaches to distributing the query processing itself across the PDMS.

5.2 Mapping XML Data

Earlier in this paper, we presented the relational version of $\mathcal{PPL}$, our peer-mapping language. We now briefly describe the language we use for mapping between XML nodes in a

Source1.xml schema:	Source2.xml schema:	Source3.rdf OWL class definition:
pubs	authors	Class id = "book"
book*	author*	DataTypeProperty id = "bookTitle"
title	full-name	domain = "#book"
author*	publication*	range = "&xsd:string"
name	title	DataTypeProperty id = "bookAuthor"
publisher*	pub-type	domain = "#book"
name		range = "#author"
		Class id = "author"
		DataTypeProperty id = "authorName"
		domain = "#author"
		range = "&xsd:string"

Fig. 3. An example of three peer (data source) schemas (two are XML sources and one is an RDF source with an OWL ontology). Source1.xml contains books with nested authors; Source2.xml contains authors with nested publications. Indentation illustrates nesting and a * suffix indicates "0 or more occurrences of...", as in a BNF grammar. Source3.rdf is a set of OWL class and property definitions with a slightly simplified notation.

Piazza overlay network, which is more complex due to the richer XML data model. The algorithm for using these mappings for query answering is described in [24] and that paper also discusses issues relating to soundness and completeness of answers.

Each Piazza node has an XML Schema that defines the terminology and the structural constraints of the node. We make a clear distinction between the intended domain of the terms defined by the schema at a node and the actual data that may be stored there. Clearly, the stored data conforms to the terms and constraints of the schema, but the intended domain of the terms may be much broader than the particular data stored at the node. For example, the terminology for publications applies to data instances beyond the particular ones stored at the node. As in the relational case, mappings play two roles. The first is as *storage descriptions* that specify which data is actually stored at a node. This allows us to separate between the intended domain and the actual data stored at the node. For example, we may specify that a particular node contains publications whose topic is computer science that have at least one author from the University of Washington. The second role is as *peer mappings*, which describe how the terminology and structure of one node correspond to those in a second node. The language for storage mappings is a subset of the language for schema mappings, hence our discussion focuses on the latter.

Following the data integration literature, which uses a standard relational query language for both queries and mappings, we might elect to use XQuery [25] for both our query language and our language for specifying mappings. However, we found XQuery inappropriate as a mapping language. First, an XQuery user thinks in terms of the input documents and the transformations to be performed. The mental connection to a required schema for the output is tenuous, whereas our setting requires thinking about relationships between the input and output schemas. Second, the user must define a mapping in its entirety before it can be used. There is no simple way to define mappings incrementally for different parts of the schemas, to collaborate with other experts on developing subregions

of the mapping, etc. Finally, XQuery is an extremely powerful (in fact, Turing-complete) query language and, as a result, some aspects are difficult or even impossible to reason about.

Our approach is to define a mapping language that borrows elements of XQuery, but is more tractable to reason about and can be expressed in *piecewise* form. Mappings in the language are defined as one or more *mapping definitions* and they are *directional* from a source to a target: We take a fragment of the target schema and annotate it with restricted XQuery expressions that define what source data should be mapped into that fragment. The mapping language is designed to make it easy for the mapping designer to visualize the target schema while describing where its data originates.

Conceptually, the results of the different mapping definitions are combined to form a complete mapping from the source document to the target, according to certain rules. The results of different mapping definitions can often be concatenated together to form the document, but, in some cases, different definitions may create content that should all be combined into a single element; Piazza "fuses" these results together based on the output element's unique identifiers (similar to the use of Skolem functions in languages such as XML-QL [26]). A complete formal description of the language is given in [24]. Here, we describe the main ideas of the language and illustrate them through examples.

Each mapping definition begins with an XML template that matches some path or a subtree of a legal instance of the target schema, i.e., a prefix of a legal string in the target schema's grammar. Elements in the template may be annotated with restricted XQuery expressions that bind variables to XML nodes in the source; for each combination of bindings, an instance of the target element will be created. Once a variable is bound, it can be referenced anywhere within its scope, which is defined to be the enclosing tags of the template. Variable bindings can be output as new target data or they can be referenced by other

<pre> <pubs> <book> {: \$a IN document("Source2.xml") \ /authors/author, \$t IN \$a/publication/title/text(), \$typ IN \$a/publication/pub-type/text() WHERE \$typ = "book" :} <title>{ \$t }</title> <author> <name>{: \$a/full-name/text() :}</name> </author> </book> </pubs> </pre>	<pre> <pubs> <book piazza:id={\$t}> {: \$a IN document("Source2.xml") \ /authors/author, \$t IN \$a/publication/title/text(), \$typ IN \$a/publication/pub-type/text() WHERE \$typ = "book" :} <title piazza:id={\$t}>{ \$t }</title> <author piazza:id={\$t}> <name>{: \$a/full-name/text() :}</name> </author> </book> </pubs> </pre>
(a)	(b)

Fig. 4. Simple examples of mappings from the schema of Source2.xml in Fig. 4 to Source1.xml's schema. (a) Initial mapping. (b) Refined mapping that coalesces entries.

query expressions to *correlate* data in different areas of the mapping definition.

Fig. 3 shows an example of two peer XML schemas, *Source1.xml* and *Source2.xml*. Fig. 4a defines a simple mapping from the schema of Source2.xml of Fig. 3 to Source1.xml. We make variable references within {} braces and delimit query expression annotations by {: :}. This mapping definition will instantiate a new book element in the target for every occurrence of variables \$a, \$t, and \$typ, which are bound to the author, title, and publication-type elements in the source, respectively. We construct a title and author element for each occurrence of book. The author name contains a new query expression annotation (\$a/full-name), so this element will be created for each match to the XPath expression (for this schema, there should only be one match).

The example mapping will create a new book element for each author-publication combination. This is probably not the desired behavior since a book with multiple authors will appear as multiple book entries, rather than as a single book with multiple author subelements. To enable the desired behavior in situations like this, Piazza reserves a special *piazza:id* attribute in the target schema for mapping multiple binding instances to the same output: If two elements created in the target have the same tag name and ID attribute, then they will be *coalesced*—all of their attributes and element content will be combined. This coalescing process is repeated recursively over the combined elements.

Example 5.1. See Fig. 4b for an improved mapping that does coalescing of book elements. The sole difference from the previous example is the use of the *piazza:id* attribute. We have determined that book titles in our collection are unique, so every occurrence of a title in the data source refers to the same book. Identical books will be given the same *piazza:id* and coalesced; likewise for their title and

author subelements (but not author names). Hence, in the target, we will see all authors nested under each book entry. This example shows how we can *invert* hierarchies in going from source to target schemas.

Sometimes, we may have detailed information about the values of the data being mapped from the source to the target—perhaps, in the above example, we know that the mapping definition only yields book titles starting with the letter “A.” Perhaps more interestingly, we may know something about the possible values of an attribute present in the target but *absent* in the source—such as the publisher. In Piazza, we refer to this sort of metainformation as *properties*. This information can be used to help the query answering system determine whether a mapping is relevant to a particular query, so it is very useful for efficiency purposes.

Example 5.2. Continuing with the previous schema, consider the partial mapping:

```

<pubs>
  <book piazza:id={$t}>
    {: $a IN document("Source2.xml") /authors/
      author,
      $t IN $a/publication/title/text(),
      $typ IN $a/publication/pub-type/text()
    WHERE $typ = "book"
    PROPERTY $t >= 'A' AND $t < 'B'
    :}
    <title piazza:id={$t}>{$t }</title>
    <author piazza:id={$t}>
      <name>{: $a/full-name/text() :}</name>
    </author>
    [: <publisher>
      <name>{: PROPERTY $this IN
        {"PrintersInc", "PubsInc"} :}</name>

```

```

</publisher>:]
</book>
</pubs>

```

The first PROPERTY definition specifies that we know this mapping includes only titles starting with "A." The second defines a "virtual subtree" (delimited by [: :]) in the target. There is insufficient data at the source to insert a value for the publisher name, but we can define a PROPERTY restriction on the values it *might* have. The special variable \$this allows us to establish a known invariant about the value at the current location within the virtual subtree: Here, it is known that the publisher name must be one of the two values specified. In general, a query over the target looking for books will make use of this mapping; a query looking for books published by BooksInc will not. Moreover, a query looking for books published by PubsInc cannot use this mapping since Piazza cannot tell whether a book was published by PubsInc or by PrintersInc.

6 RELATED WORK

The idea of mediating between different databases using local semantic relationships was explored in federated databases and cooperative databases (e.g., [30], [31], [32], [33]). There, it was assumed that each database in the federation stored data and the focus was on mapping between the stored relations in the federation. Our work differs in several ways. First, the scale of a PDMS is assumed to be much larger and its structure more ad hoc. Joining and leaving a PDMS should be much easier than in a federated database. As a consequence, the relationships between the peers are much looser. Second, peers can play different roles—some provide data, others provide integration services between other peers, and some provide both. As a result, we need to be able to map both relationships among stored relations and among conceptual relations (i.e., extensional versus intentional relations). Third, our focus is on algorithms for chaining through multiple peer mappings in order to locate data relevant to a query.

In [34], we described some of the challenges involved in building a PDMS, focusing on *intelligent data placement*, a technique for materializing views at nodes in the network in order to improve performance and availability. In [35], the authors study a variant of the data placement problem and focus on intelligently caching and reusing queries in an OLAP environment. Recently, [36] described *local relational models* as a formalism for mediating between different peers in a PDMS and a sound and complete algorithm for answering queries using the formalism, but do not describe the expressive power of the formalism compared to previous ones in the data integration literature.

Eduella [37] represents an interesting design point in the XML-RDF interoperability spectrum. Like Piazza, it is built on a peer-to-peer architecture (based on JXTA) and it mediates between different data representations—but it provides query and storage services for RDF, using a variety of underlying stores. Thus, an important focus of the project is on translating the RDF data and queries to the underlying storage format and query language.

Several other projects in the database community are developing peer-to-peer architectures, with slightly different emphases. The Chatty Web [38] focuses on gossip protocols for exchanging semantic mapping information, where mappings are selection-projection queries that they evaluate for information loss. Hyperion [39] focuses on problems relating to mappings between *objects* in different relations, which is another important aspect of mapping between sources. PeerDB [40] takes another approach to mapping between peers: Instead of schema mappings, PeerDB employs an Information Retrieval-based approach to query reformulation. A peer relation (and each of its columns) is associated with a set of keywords. PeerDB reformulates a query over one schema into other peers' schemas by matching the keywords associated with the two schemas. Keywords can be matched directly between any pair of schemas, so chaining of reformulation steps is not required; however, keyword matching may give irrelevant query reformulations, so the user must decide which queries are to be executed.

In the KR community, work on the OBSERVER [41] and Kraft [42] systems have explored a number of issues in distributed ontologies, including mappings from structured sources and approximate mappings between concepts in ontologies.

7 CONCLUSIONS

The concept of the peer data management system emphasizes not only an ad hoc, scalable, distributed peer-to-peer computing environment (which is compelling from a distributed systems perspective), but it provides an easily extensible, decentralized environment for sharing data with rich semantics. This is in contrast to data integration systems, which have a centralized mediated schema and administrator and which, in our experience, impede small, point-to-point collaborations. It also complements the knowledge representation work of the Semantic Web by providing a mechanism for translating between different ontologies' data representations.

We described some of the main highlights of the Piazza PDMS, including

1. a solution to schema mediation in peer data based on a language that uses previous mediation formalisms at the local level to form a network of semantically related peers,
2. a characterization of the theoretical limitations on answering queries in a PDMS,
3. an algorithm for answering queries in such a system, and
4. results concerning the composition of semantic mappings.

We also argued that a PDMS can provide a basis for building applications for the Semantic Web and we showed how to extend Piazza to the XML data model.

Though we have not described these in this paper, we have implemented a prototype of Piazza and built a small Web of database-research related Web sites. We are currently focusing on developing and testing effective methods for query optimization in Piazza and the management and propagation of updates in a PDMS. In addition,

we believe that the management of large collections of semantic mappings raises interesting challenges. Our work on query composition lays the basis for studying several fundamental properties of such networks. We are also interested in studying how one can *boost* a collection of such mappings to improve the ability of nodes to obtain relevant data from other distant nodes in the network. Finally, peer data management is a very rich domain that creates a wealth of new problems, such as how to replicate data and how to reconcile inconsistent data.

The Appendix to this paper, which can be found on the Computer Society Digital Library at <http://computer.org/tkde/archives.htm>, describes the problem of mapping composition and discusses using a PDMS as an infrastructure for the Semantic Web.

ACKNOWLEDGMENTS

This work was funded in part by US National Science Foundation ITR grants IIS-0205635 and IIS-9985114 and a gift from Microsoft Research.

REFERENCES

- [1] T. Berners-Lee, J. Hendler, and O. Lassila, "The Semantic Web," *Scientific Am.*, May 2001.
- [2] M. Dean, D. Connolly, F. van Harmelen, J. Hendler, I. Horrocks, D.L. McGuinness, P.F. Patel-Schneider, and L.A. Stein, "OWL Web Ontology Language 1.0 Reference," <http://www.w3c.org/TR/2002-WD-owl-ref-20020729/>, w3C working draft, 29 July 2002.
- [3] H. Garcia-Molina, Y. Papakonstantinou, D. Quass, A. Rajaraman, Y. Sagiv, J. Ullman, and J. Widom, "The TSIMMIS Project: Integration of Heterogeneous Information Sources," *J. Intelligent Information Systems*, vol. 8, no. 2, pp. 117-132, Mar. 1997.
- [4] L. Haas, D. Kossmann, E. Wimmers, and J. Yang, "Optimizing Queries across Diverse Data Sources," *Proc. 23 Int'l Conf. Very Large Databases*, 1997.
- [5] S. Adali, K. Candan, Y. Papakonstantinou, and V. Subrahmanian, "Query Caching and Optimization in Distributed Mediator Systems," *Proc. SIGMOD*, pp. 137-148, 1996.
- [6] A.Y. Levy, A. Rajaraman, and J.J. Ordille, "Querying Heterogeneous Information Sources Using Source Descriptions," *Proc. 22nd Int'l Conf. Very Large Databases*, pp. 251-262, 1996.
- [7] O.M. Duschka and M.R. Genesereth, "Answering Recursive Queries Using Views," *Proc. 16th ACM SIGACT-SIGMOD-SIGART Symp. Principles of Database Systems*, pp. 109-116, 1997.
- [8] I. Manolescu, D. Florescu, and D. Kossmann, "Answering XML Queries on Heterogeneous Data Sources," *Proc. 27th Int'l Conf. Very Large Data Bases*, pp. 241-250, 2001.
- [9] J.L. Ambite, N. Ashish, G. Barish, C.A. Knoblock, S. Minton, P.J. Modi, I. Muslea, A. Philpot, and S. Tejada, "ARIADNE: A System for Constructing Mediators for Internet Sources (System Demonstration)," *Proc. SIGMOD*, 1998.
- [10] E. Lambrecht, S. Kambhampati, and S. Gnanaprakasam, "Optimizing Recursive Information Gathering Plans," *Proc. 16th Int'l Joint Conf. Artificial Intelligence*, pp. 1204-1211, 1999.
- [11] J.M. Smith, P.A. Bernstein, U. Dayal, N. Goodman, T. Landers, K. Lin, and E. Wong, "Multibase—Integrating Heterogeneous Distributed Database Systems," *Proc. Nat'l Computer Conf.*, 1981.
- [12] A.Y. Halevy, "Answering Queries Using Views: A Survey," *VLDB J.*, vol. 10, no. 4, 2001.
- [13] D. Draper, A.Y. Halevy, and D.S. Weld, "The Nimble Integration System," *Proc. SIGMOD*, 2001.
- [14] S. Abiteboul and O. Duschka, "Complexity of Answering Queries Using Materialized Views," *Proc. 17th ACM SIGACT-SIGMOD-SIGART Symp. Principles of Database Systems*, pp. 254-263, 1998.
- [15] A. Halevy, Z. Ives, D. Suciu, and I. Tatarinov, "Schema Mediation for Large-Scale Semantic Data Sharing," *VLDB J.*, 2003.
- [16] M. Friedman, A. Levy, and T. Millstein, "Navigational Plans for Data Integration," *Proc. 16th Nat'l Conf. Artificial Intelligence*, 1999.
- [17] Z.G. Ives, A.Y. Halevy, and D.S. Weld, "Integrating Network-bound XML Data," *IEEE Data Eng. Bull.*, special issue on XML, vol. 24, no. 2, June 2001.
- [18] A.Y. Halevy, I. Mumick, Y. Sagiv, and O. Shmueli, "Static Analysis in Datalog Extensions," *J. ACM*, vol. 48, no. 5, pp. 971-1012, Sept. 2001.
- [19] D. Srivastava and R. Ramakrishnan, "Pushing Constraint Selections," *Proc. 11th ACM SIGACT-SIGMOD-SIGART Symp. Principles of Database Systems*, pp. 301-315, 1992.
- [20] J. Madhavan and A. Halevy, "Composing Mappings among Data Sources," *Proc. 20th Int'l Conf. Very Large Data Bases*, 2003.
- [21] M.Y. Vardi, "On the Complexity of Bounded-Variable Queries," *Proc. 14th ACM SIGACT-SIGMOD-SIGART Symp. Principles of Database Systems*, pp. 266-276, 1995.
- [22] I. Stoica, R. Morris, D. Karger, M.F. Kaashoek, and H. Balakrishnan, "Chord: A Scalable Peer-to-Peer Lookup Service for Internet Applications," *Proc. ACM SIGCOMM*, 2001.
- [23] S. Ratnasamy, P. Francis, M. Handley, R. Karp, and S. Shenker, "A Scalable Content-Addressable Network," *Proc. ACM SIGCOMM*, 2001.
- [24] A. Halevy, Z. Ives, P. Mork, and I. Tatarinov, "Piazza: Data Management Infrastructure for Semantic Web Applications," *Proc. 12th Int'l World Wide Web Conf.*, 2003.
- [25] D. Chamberlin, D. Florescu, J. Robie, J. Simeon, and M. Stefanescu, "XQuery: A Query Language for XML," World Wide Web Consortium, technical report, Feb. 2001, available from <http://www.w3.org/TR/xquery/>.
- [26] A. Deutsch, M. Fernandez, D. Florescu, and A. Levy, D. Suciu, "A Query Language for XML," <http://www.research.att.com/mff/xml/w3c-note.html>, 1998.
- [27] P. Patel-Schneider and J. Simeon, "Building the Semantic Web on XML," *Proc. Int'l Semantic Web Conf.*, June 2002.
- [28] I. Horrocks, F. van Harmelen, and P. Patel-Schneider, "DAML+OIL," <http://www.daml.org/2001/03/daml+oil-in-dex.html>, Mar. 2001.
- [29] A. Levy and M.-C. Rousset, "Combining Horn Rules and Description Logics in Carin," *Artificial Intelligence*, vol. 104, pp. 165-209, Sept. 1998.
- [30] W. Litwin, L. Mark, and N. Roussopoulos, "Interoperability of Multiple Autonomous Databases," *ACM Computing Surveys*, vol. 22, no. 3, pp. 267-293, 1990.
- [31] R. Krishnamurthy, W. Litwin, and W. Kent, "Language Features for Interoperability of Databases with Schematic Discrepancies," *Proc. SIGMOD*, pp. 40-49, 1991.
- [32] M. Rusinkiewicz, A. Sheth, and G. Karabatis, "Specifying Interdatabase Dependencies in a Multidatabase Environment," *Computer*, vol. 24, no. 12, Dec. 1991.
- [33] T. Catarci and M. Lenzerini, "Representing and Using Inter-schema Knowledge in Cooperative Information Systems," *J. Intelligent and Cooperative Information Systems*, pp. 55-62, 1993.
- [34] S. Gribble, A. Halevy, Z. Ives, M. Rodrig, and D. Suciu, "What Can Databases Do for Peer-to-Peer?" *Proc. ACM SIGMOD WebDB Workshop*, 2001.
- [35] P. Kalnis, W. Ng, B. Ooi, D. Papadias, and K. Tan, "An Adaptive Peer-to-Peer Network for Distributed Caching of Olap Results," *Proc. SIGMOD*, 2002.
- [36] P. Bernstein, F. Giunchiglia, A. Kementsietsidis, J. Mylopoulos, L. Serafini, and I. Zaihayeu, "Data Management for Peer-to-Peer Computing: A Vision," *Proc. WebDB Workshop*, 2002.
- [37] W. Nejdl, B. Wolf, C. Qu, S. Decker, M. Sintek, A. Naeve, M. Nilsson, M. Palmer, and T. Risch, "EDUTELLA: A P2P Networking Infrastructure Based on RDF," *Proc. Int'l WWW Conf.*, 2002.
- [38] K. Aberer, P. Cudre-Mauroux, and M. Hauswirth, "The Chatty Web: Emergent Semantics through Gossiping," *Proc. 12th Int'l World Wide Web Conf.*, 2003.
- [39] M. Arenas, V. Kantere, A. Kementsietsidis, I. Kiringa, R.J. Miller, and J. Mylopoulos, "The Hyperion Project: From Data Integration to Data Coordination," *SIGMOD Record*, Sept. 2003.
- [40] W.S. Ng, B.C. Ooi, K.-L. Tan, and A. Zhou, "PeerDB: A P2P-Based System for Distributed Data Sharing," *Proc. SIGMOD*, 2003.
- [41] E. Mena, V. Kashyap, A.P. Sheth, and A. Illarramendi, "OB-SERVER: An Approach for Query Processing in Global Information Systems Based on Interoperation across Pre-Existing Ontologies," *Distributed and Parallel Databases*, vol. 8, no. 2, pp. 223-271, 2000.

- [42] A. Preece, K. Hui, and P. Gray, "Kraft: An Agent Architecture for Knowledge Fusion," *Int'l J. Cooperative Information Systems*, vol. 10, nos. 1-2, pp. 171-195, 1999.



Alon Y. Halevy received the PhD degree in computer science from Stanford University in 1993. He is an associate professor at the University of Washington. Dr. Halevy's research interests are in data integration, management of XML data, Web site management, knowledge representation and, more generally, the intersection between database and AI technologies. He was a Sloan Fellow (1999-2000), and received the Presidential Early Career Award for Scientists and Engineers (PECASE) in 2000. He serves on the editorial boards of the *VLDB Journal*, the *Journal of Artificial Intelligence Research* and served as program chair the ACM SIGMOD 2003 Conference.



Zachary G. Ives is an assistant professor at the University of Pennsylvania. His primary interests are in adaptive query processing for distributed applications, XML query processing, and developing new techniques for peer-to-peer and collaborative data sharing.



Jayant Madhavan received the BTech degree from the Indian Institute of Technology, Bombay, and the MS degree from the University of Washington. He is a PhD candidate at the University of Washington. His main interests are in the development and application of machine learning techniques for robust solutions in data management and processing. He is currently investigating the use of learning techniques to bridge semantic heterogeneity between disparate data representations.



Peter Mork received the MS degree from Stanford University. He is a PhD candidate at the University of Washington. He is interested in collaborative knowledge-bases for biology and medicine. Currently, he is investigating the role of peer architectures in knowledge sharing.



Dan Suciu is an associate professor of computer science at the University of Washington. He is conducting research in data management, with an emphasis on topics that arise from sharing data on the Internet, such as management of semistructured and heterogeneous data, large-scale message stream processing, security of data. He is a coauthor of the book *Data on the Web: From Relations to Semistructured Data and XML*, holds six US patents, received the 2000 ACM SIGMOD Best Paper Award, and conducted three projects that led to popular software tools in the public domain: XMill, XMLTK, and SilkRoute.



Igor Tatarinov is currently a PhD candidate in the Department of Computer Science and Engineering at the University of Washington. His research interests include peer database systems and XML query processing.

► **For more information on this or any other computing topic, please visit our Digital Library at www.computer.org/publications/dlib.**