

RDF-3X: a RISC-style Engine for RDF

Thomas Neumann
Max-Planck-Institut für Informatik
Saarbrücken, Germany
neumann@mpi-inf.mpg.de

Gerhard Weikum
Max-Planck-Institut für Informatik
Saarbrücken, Germany
weikum@mpi-inf.mpg.de

ABSTRACT

RDF is a data representation format for schema-free structured information that is gaining momentum in the context of Semantic-Web corpora, life sciences, and also Web 2.0 platforms. The “pay-as-you-go” nature of RDF and the flexible pattern-matching capabilities of its query language SPARQL entail efficiency and scalability challenges for complex queries including long join paths. This paper presents the RDF-3X engine, an implementation of SPARQL that achieves excellent performance by pursuing a RISC-style architecture with a streamlined architecture and carefully designed, puristic data structures and operations. The salient points of RDF-3X are: 1) a generic solution for storing and indexing RDF triples that completely eliminates the need for physical-design tuning, 2) a powerful yet simple query processor that leverages fast merge joins to the largest possible extent, and 3) a query optimizer for choosing optimal join orders using a cost model based on statistical synopses for entire join paths. The performance of RDF-3X, in comparison to the previously best state-of-the-art systems, has been measured on several large-scale datasets with more than 50 million RDF triples and benchmark queries that include pattern matching and long join paths in the underlying data graphs.

1. INTRODUCTION

1.1 Motivation and Problem

The RDF (Resource Description Framework) data model has been around for a decade. It has been designed as a flexible representation of schema-relaxable or even schema-free information for the Semantic Web [30]. In the commercial IT world, RDF has not received much attention until recently, but now it seems that RDF is building up a strong momentum. Semantic-Web-style ontologies and knowledge bases with millions of facts from Wikipedia and other sources have been created and are available online [4, 40, 47]. E-science data repositories support RDF as an import/export format

Permission to make digital or hard copies of portions of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyright for components of this work owned by others than VLDB Endowment must be honored.

Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists requires prior specific permission and/or a fee. Request permission to republish from: Publications Dept., ACM, Inc. Fax +1 (212) 869-0481 or permissions@acm.org.

PVLDB '08, August 23-28, 2008, Auckland, New Zealand
Copyright 2008 VLDB Endowment, ACM 978-1-60558-305-1/08/08

and also for selective (thus, query-driven) data extractions, most notably, in the area of life sciences (e.g., [7, 43]). Finally, Web 2.0 platforms for online communities are considering RDF as a non-proprietary exchange format and as an instrument for the construction of information mashups [21, 22, 31].

In RDF, all data items are represented in the form of *(subject, predicate, object)* triples, also known as *(subject, property, value)* triples. For example, information about the movie “Sweeney Todd” could include the following triples:

```
(id1, hasTitle, "Sweeney Todd"),  
(id1, producedInYear, "2007"),  
(id1, directedBy, "Tim Burton"),  
(id1, hasCasting, id2),  
(id2, RoleName, "Sweeney Todd"), (id2, Actor, id11),  
(id1, hasCasting, id3),  
(id3, RoleName, "Mrs. Lovett"), (id3, Actor, id12),  
(id11, hasName, "Johnny Depp"),  
(id12, hasName, "Helena Bonham Carter"), and so on.
```

Note that, although predicate names such as “producedInYear” resemble attributes, there is no database schema; the same database may contain triples about movies with a predicate “productionYear”. A schema may emerge in the long run (and can then be described by the RDF Vocabulary Description Language). In this sense, the notion of RDF triples fits well with the modern notion of data spaces and its “pay as you go” philosophy [18]. Compared to an entity-relationship model, both an entity’s attributes and its relationships to other entities are represented by predicates. All RDF triples together can be viewed as a large (instance-level) graph.

The SPARQL query language is the official standard for searching over RDF repositories. It supports conjunctions (and also disjunctions) of triple patterns, the counterpart to select-project-join queries in a relational engine. For example, we can retrieve the titles of all movies with Johnny Depp by the SPARQL query:

```
Select ?title Where {  
  ?m <hasTitle> ?title. ?m <hasCasting> ?c.  
  ?c <Actor> ?a. ?a <hasName> "Johnny Depp" }
```

Here each of the conjunctions, denoted by a dot, corresponds to a join. The whole query can also be seen as graph pattern that needs to be matched in the RDF data graph. In SPARQL, predicates can also be variables or wildcards, thus allowing schema-agnostic queries.

RDF engines for storing, indexing, and querying have been around for quite a few years; especially, the Jena frame-

work by HP Labs has gained significant popularity [46], and Oracle also provides RDF support for semantic data integration in life sciences and enterprises [11, 29]. However, with the exception of the VLDB 2007 paper by Abadi et al. [1], none of the prior implementations could demonstrate convincing efficiency, failing to scale up towards large datasets and high load. [1] achieves good performance by grouping triples with the same property name into property tables, mapping these onto a column store, and creating materialized views for frequent joins.

Managing large-scale RDF data includes technical challenges for the storage layout, indexing, and query processing:

1. The absence of a global schema and the diversity of predicate names pose major problems for the *physical database design*. In principle, one could rely on an auto-tuning “wizard” to materialize frequent join paths; however, in practice, the evolving structure of the data and the variance and dynamics of the workload turn this problem into a complex sisypus task.
2. By the fine-grained modeling of RDF data – triples instead of entire records or entities – queries with a large number of joins will inherently form a large part of the workload, but the join attributes are much less predictable than in a relational setting. This calls for specific choices of *query processing algorithms*, and for careful *optimization of complex join queries*; but RDF is meant for on-the-fly applications over data spaces, so the optimization takes place at query run-time.
3. As join-order and other execution-plan optimizations require data statistics for selectivity estimation, an RDF engine faces the problem that a suitable granularity of *statistics gathering* is all but obvious in the absence of a schema. For example, single-dimensional histograms on all attributes that occur in the workload’s where clauses – the state-of-the-art approach in relational systems – is unsuitable for RDF, as it misses the effects of long join chains or large join stars over many-to-many relationships.
4. Although RDF uses XML syntax and SPARQL involves search patterns that resemble XML path expressions, the fact that RDF triples form a graph rather than a collection of trees is a major difference to the more intensively researched settings for XML.

1.2 Contribution and Outline

This paper gives a comprehensive, scalable solution to the above problems. It presents a complete system, coined *RDF-3X* (for RDF Triple eXpress), designed and implemented from scratch specifically for the management and querying of RDF data. RDF-3X follows the rationale advocated in [24, 39] that data-management systems that are designed for and customized to specific application domains can outperform generic mainstream systems by two orders of magnitude. The factors in this argument include 1) tailored choices of data structures and algorithms rather than supporting a wide variety of methods, 2) much lighter software footprint and overhead, and as a result, 3) simplified optimization of system internals and easier configuration and self-adaptation to changing environments (e.g., data and workload characteristics).

RDF-3X follows such a RISC-style design philosophy [10], with “reduced instruction set” designed to support RDF. RDF-3X is based on a three key principles:

- *Physical design* is *workload-independent* by creating appropriate indexes over a single “giant triples table”. RDF-3X does not rely on the success (or limitation) of an auto-tuning wizard, but has effectively eliminated the need for physical-design tuning. It does so by building indexes over all 6 permutations of the three dimensions that constitute an RDF triple, and additionally, indexes over count-aggregated variants for all three two-dimensional and all three one-dimensional projections. Each of these indexes can be compressed very well; the total storage space for all indexes together is less than the size of the primary data.
- The *query processor* is *RISC-style* by relying mostly on merge joins over sorted index lists. This is made possible by the “exhaustive” indexing of the triples table. In fact, all processing is index-only, and the triples table exists merely virtually. Operator trees are constructed so as to preserve *interesting orders* [17] for subsequent joins to the largest possible extent; only when this is no longer possible, RDF-3X switches to hash-based join processing. This approach can be highly optimized at the code level, and has much lower overhead than traditional query processors. At the same time, it is sufficiently versatile to support also the various duplicate-elimination options of SPARQL, disjunctive patterns in queries, and all other features that SPARQL requires.
- The *query optimizer* mostly focuses on *join order* in its generation of execution plans. It employs dynamic programming for plan enumeration, with a cost model based on RDF-specific statistical synopses. These statistics include counters of frequent predicate-sequences in paths of the data graph; such paths are potential join patterns. Compared to the query optimizer in a universal database system, the RDF-3X optimizer is simpler but much more accurate in its selectivity estimations and decisions about execution plans.

The scientific contributions of this work are: i) a novel architecture for RDF indexing and querying, eliminating the need for physical database design, ii) an optimizer for large join queries over non-schematic RDF triples, driven by a new kind of selectivity estimation for RDF paths, iii) a comprehensive performance evaluation, based on real measurements with three large datasets, demonstrating large gains over the previously best engine [1] (by a typical factor of 5 and up to 20 for some queries). The source code of RDF-3X and the experimental data are available for non-commercial purposes upon request.

2. BACKGROUND AND STATE OF THE ART

2.1 SPARQL

SPARQL queries [36] are pattern matching queries on triples that constitute an RDF data graph. Syntactic sugar aside, a basic SPARQL query has the form

```
select ?variable1 ?variable2 ...
where { pattern1. pattern2. ... }
```

where each *pattern* consists of *subject*, *predicate*, *object*, and each of these is either a variable or a literal. The query model is query-by-example style: the query specifies the known literals and leaves the unknowns as variables. Variables can occur in multiple patterns and thus imply joins. The query processor needs to find all possible variable bindings that satisfy the given patterns and return the bindings from the projection clause to the application. Note that not all variables are necessarily bound (e.g., if a variable only occurs in the projection and not in a pattern), which results in NULL values.

This pattern matching approach restricts the freedom of the query engine with regard to possible execution plans, as shown in the following example:

```
select ?a ?c
where { ?a label1 ?b. ?b label2 ?c }
```

The user is interested in all *?a* and *?c* that are reachable with certain edges via *?b*. The value of *?b* itself is not part of the projection clause. Unfortunately the pattern matching semantics requires that nevertheless all bindings of *?b* need to be computed. There might be multiple ways from *?a* to *?c*, resulting in duplicates in the output. As this is usually not what the user/application intends, SPARQL introduced two query modifiers: the *distinct* keyword specifies that duplicates must be eliminated, and the *reduced* keyword specifies that duplicates may but need not be eliminated. The goal of the *reduced* keyword was obviously to help RDF query engines by allowing optimizations, but with the *reduced* option the query output has a nondeterministic nature.

Nevertheless, even the default mode of creating all duplicates allows some optimizations. The query processor must not ignore variables that are projected away due to their effect on duplicates, but it does not have to create the explicit bindings. As long as we can guarantee that the correct number of duplicates is produced, the bindings themselves are not relevant. We will use this observation later by counting the number of duplicates rather than producing the duplicates themselves.

2.2 Related Work

Most publicly accessible RDF systems have mapped RDF triples onto relational tables (e.g., RDFSuite [2, 32], Sesame [8, 28], Jena [23, 46], the C-Store-based RDF engine of [1], and also Oracle's RDF_MATCH implementation [11]). There are two extreme ways of doing this: 1) All triples are stored in a single, giant *triples table* with generic attributes *subject*, *predicate*, *object*. 2) Triples are grouped by their predicate name, and all triples with the same predicate name are stored in the same *property table*. The extreme form of property tables with a separate table for each predicate name can be made more flexible, leading to a hybrid approach: 3) Triples are clustered by predicate names, based on predicates for the same entity class or co-occurrence in the workload; each *cluster-property table* contains the values for a small number of correlated predicates, and there may additionally be a "left-over" table for triples with infrequent predicates. A cluster-property table has a class-specific schema with attributes named after the corresponding RDF predicates, and its width can range from a single predicate (attribute) to all predicates of the same entity type.

Early open-source systems like Jena [23, 46] and Sesame

[8, 28] use clustered-property tables, but left the physical design to an application tuning expert. Neither of these systems has reported any performance benchmarks with large-scale data in the Gigabytes range with more than 10 million triples. Oracle [29] has reported very good performance results in [11], but seems to heavily rely on good tuning by making the right choice of materialized join views (coined subject-property matrix) in addition to its basic triples table. The currently fastest RDF engine by [1] uses minimum-width property tables (i.e., binary relations), but maps them onto a column-store system. [41] gives a nice taxonomy of different storage layouts and presents systematic performance comparisons for medium-sized synthetic data and synthetic workload. In contrast to the arguments that [1] gives against the "giant-triples-table" approach, our RDF-3X system shows how to successfully employ a triples table with excellent performance.

The best performing systems, Oracle and the C-Store-based engine, rely on materialized join paths and indexes on these views. The indexes themselves are standard indexes as supported by the underlying RDBMS and column store, respectively. The native YARS2 system [19] proposes exhaustive indexes of triples and all their embedded sub-triples (pairs and single values) in 6 separate B⁺-tree or hash indexes. This resembles our approach, but YARS2 misses the need for indexing triples in collation orders other than the canonical order by subject, predicate, object (as primary, secondary, and tertiary sort criterion). A very similar approach is presented for the HPRD system [6] and available as an option (coined "triple-indexes") in the Sesame system [28]. Both YARS2 and HPRD seem to be primarily geared for simple lookup operations with limited support for joins; they lack DBMS-style query optimization (e.g., do not consider any join-order optimizations, although [6] recognizes the issue). [5] proposes to index entire join paths using suffix arrays, but does not discuss optimizing queries over this physical design. [42] introduces a new kind of path indexing based on judiciously chosen "center nodes"; this index, coined GRIN, shows good performance on small- to medium-sized data and for hand-specified execution plans. Physical design for schema-agnostic "wide and sparse tables" is also discussed in [12], without specific consideration to RDF. All these methods for RDF indexing and materialized views incur some form of physical design problem, and none of them addresses the resulting query optimization issues over these physical-design features.

As for query optimization, [11, 29] and [1] utilize the state-of-the-art techniques that come with the SQL engines on which these solutions are layered. To our knowledge, none of them employs any RDF-native optimizations. [38] outlines a framework for algebraic rewriting, but it seems that the main rule for performance gains is pushing selections below joins; there is no consideration of join ordering. [20] has a similar flavor, and likewise disregards the key problem of finding good join orderings.

Recently, selectivity estimation for SPARQL patterns over graphs have been addressed by [38] and [25]. The method by [38] gathers separate frequency statistics for each subject, each predicate, and each object (label or value); the frequency of an entire triple pattern is estimated by assuming that subject, predicate, and object distributions are probabilistically independent. The method by [25] is much more sophisticated by building statistics over a selected set

of arbitrarily shaped graph patterns. It casts the selection of patterns into an optimization problem and uses greedy heuristics. The cardinality estimation of a query pattern identifies maximal subpatterns for which statistics exist, and combines them with uniformity assumptions about super-patterns without statistics. While [38] seems to be too simple for producing accurate estimates, the method by [25] is based on a complex optimization problem and relies on simple heuristics to select a good set of patterns for the summary construction. The method that we employ in RDF-3X captures path-label frequencies, thus going beyond [38] but avoiding the computational complexity of [25].

3. STORAGE AND INDEXING

3.1 Triples Store and Dictionary

Although most of the prior, and especially the recent, literature favors a storage schema with property tables, we decided to pursue the conceptually simpler approach with a single, potentially huge *triples table*, with our own storage implementation underneath (as opposed to using an RDBMS). This reflects our RISC-style and “no-knobs” design rationale. We overcome the previous criticism that a triples table incurs too many expensive self-joins by creating the “right” set of indexes (see below) and by very fast processing of merge joins (see Section 4).

We store all triples in a (compressed) clustered B⁺-tree. The triples are sorted lexicographically in the B⁺-tree, which allows the conversion of SPARQL patterns into range scans. In the pattern (literal1, literal2, ?x) the literals specify the common prefix and thus effectively a range scan. Each possible binding of ?x is found during a single scan over a moderate number of leaf pages.

As triples may contain long string literals, we adopt the natural approach (see, e.g., [11]) of replacing all literals by ids using a *mapping dictionary*. This has two benefits: 1) it compresses the triple store, now containing only id triples, and 2) it is a great simplification for the query processor, allowing for fast, simple, RISC-style operators (see Section 4.5). The small cost for these gains is two additional *dictionary indexes*. During query translation, the literals occurring in the query are translated into their dictionary ids, which can be done with a standard B⁺-tree from strings to ids. After processing the query the resulting ids have to be transformed back into literals as output to the application/user. We could have used a B⁺-tree for this direction, too, but instead we implemented a direct mapping index [14]. Direct mapping is tuned for id lookups and results in a better cache hit ratio. Note that this is only an issue when the query produces many results. Usually the prior steps (joins etc.) dominate the costs, but for simple queries with many results dictionary lookups are non-negligible.

3.2 Compressed Indexes

In the index-range-scan example given above we rely on the fact that the variables are a suffix (i.e., the *object* or the *predicate* and *object*). To guarantee that we can answer every possible pattern with variables in any position of the pattern triple by merely performing a single index scan, we maintain all six possible permutations of *subject* (*S*), *predicate* (*P*) and *object* (*O*) in six separate indexes. We can afford this level of redundancy because we compress the id triples (discussed below). On all our experimental datasets, the total size for

Gap	Payload	Delta	Delta	Delta
1 Bit	7 Bits	0-4 Bytes	0-4 Bytes	0-4 Bytes
Header		$value_1$	$value_2$	$value_3$

Figure 1: Structure of a compressed triple

```
compress((v1, v2, v3), (prev1, prev2, prev3))
// writes (v1, v2, v3) relative to (prev1, prev2, prev3)
if v1 = prev1 ∧ v2 = prev2
  if v3 - prev3 < 128
    write v3 - prev3
  else encode(0, 0, v3 - prev3 - 128)
else if v1 = prev1
  encode(0, v2 - prev2, v3)
else
  encode(v1 - prev1, v2, v3)

encode(δ1, δ2, δ3)
// writes the compressed tuple corresponding to the deltas
write 128 + bytes(δ1)*25 + bytes(δ2)*5 + bytes(δ3)
write the non-zero tail bytes of δ1
write the non-zero tail bytes of δ2
write the non-zero tail bytes of δ3
```

Figure 2: Pseudo-Code of triple compression

all indexes together is less than the original data.

As the collation order in each of the six indexes is different (SPO, SOP, OSP, OPS, PSO, POS), we use the generic terminology $value_1$, $value_2$, $value_3$ instead of *subject*, *predicate*, *object* for referring to the different columns. The triples in an index are sorted lexicographically by ($value_1$, $value_2$, $value_3$) (for each of the six different permutations) and are directly stored in the leaf pages of the clustered B⁺-tree.

The collation order causes neighboring triples to be very similar: most neighboring triples have the same values in $value_1$ and $value_2$, and the increases in $value_3$ tend to be very small. This observation naturally leads to a compression scheme for triples. Instead of storing full triples we only store the changes between triples. This compression scheme is inspired by methods for inverted lists in text retrieval systems [50], but we generalize it to id triples rather than simple ids. For reasons discussed below, we apply the compression only within individual leaf pages and never across pages.

For the compression scheme itself, there is a clear trade-off between space savings and CPU consumption for decompression or interpretation of compressed items [45]. We noticed that CPU time starts to become an issue when compressing too aggressively, and therefore settled for a byte-level (rather than bit-level) compression scheme. We compute the delta for each value, and then use the minimum number of bytes to encode just the delta. A header byte denotes the number of bytes used by the following values (Figure 1). Each value consumes between 0 bytes (unchanged) and 4 bytes (delta needs the full 4 bytes), which means that we have 5 possible sizes per value. For three values these are $5 * 5 * 5 = 125$ different size combinations, which fits into the payload of the header byte. The remaining *gap* bit is used to indicate a small gap: When only $value_3$ changes, and the delta is less than 128, it can be directly included in the payload of the header byte. This kind of small delta is very common, and can be encoded by a single byte in our scheme.

The details of the compression are shown in Figure 2. The

	byte-wise	bit-wise		
Barton data	RDF-3X	Gamma	Delta	Golomb
6 indexes [GBytes]	1.10	1.21	1.06	2.03
Decompression [s]	3.2	44.7	42.5	82.6

Figure 3: Comparison of byte-wise compression vs. bit-wise compression for the Barton dataset

algorithm computes the delta to the previous tuple. If it is small it is directly encoded in the header byte, otherwise it computes the δ_i values for each of the tree values and calls *encode*. *encode* writes the header byte with the size information and then writes the non-zero tail of the δ_i (i.e., it writes δ_i byte-wise but skips leading zero bytes). This results in compressed tuples with varying sizes, but during decompression the sizes can be reconstructed easily from the header byte. As all operations are byte-wise, decompression involves only a few cheap operations and is very fast.

We tested the compression rate and the decompression time (in seconds) of our byte-wise compression against a number of bit-wise compression schemes proposed in the literature [33]. The results for the Barton dataset (see Section 6) are shown in Figure 3. Our byte-wise scheme compresses nearly as good as the best bit-wise compression scheme, while providing much better decompression speed. The Gamma and Golomb compression methods, which are popular for inverted lists in IR systems, performed worse because, in our setting, gaps can be large whenever there is a change in the triple prefix.

We also experimented with the more powerful LZ77 compression on top of our compression scheme. Interestingly our compression scheme compresses better with LZ77 than the original data, as the delta encoding exhibits common patterns in the triples. The additional LZ77 compression decreases the index size roughly by a factor of two, but increases CPU time significantly, which would become critical for complex queries. Thus, the RDF-3X engine does not employ LZ77.

An important consideration for the compressed index is that each leaf page is compressed individually. Compressing larger chunks of data leads to better compression (in particular in combination with the LZ77 compression), but page-wise compression has several advantages. First, it allows us to seek (via B⁺-tree traversal) to any leaf page and directly start reading triples. If we compressed larger chunks we would often have to decompress preceding pages. Second, the compressed index behaves just like a normal B⁺-tree (with a special leaf encoding). Thus, updates can be easily done like in a standard B⁺-tree. This greatly simplifies the integration of the compressed index into the rest of the engine and preserves its RISC nature. In particular, we can adopt advanced concurrency control and recovery methods for index management without any changes.

3.3 Aggregated Indices

For many SPARQL patterns, indexing partial triples rather than full triples would be sufficient, as demonstrated by the following SPARQL query:

```
select ?a ?c
where { ?a ?b ?c }
```

It computes all *?a* and *?c* that are connected through any predicate, the actual bindings of *?b* are not relevant.

```
compressAggregated((v1, v2), count, (prev1, prev2))
// writes (v1, v2) * count relative to (prev1, prev2)
if v1 = prev1
  if v2 - prev2 < 32 ∧ count < 5
    write (count - 1) * 32 + (v2 - prev2)
  else
    encode(0, v2 - prev2, count)
else
  encode(v1 - prev1, v2, count)
```

Figure 4: Aggregated triple compression

We therefore additionally build *aggregated indexes*, each of which stores only two out of the three columns of a triple. More precisely, they store two entries (e.g., *subject* and *object*), and an aggregated *count*, namely, the number of occurrences of this pair in the full set of triples. This is done for each of the three possible pairs out of a triple and in each collation order (SP, PS, SO, OS, PO, OP), thus adding another six indexes.

The count is necessary because of the SPARQL semantics. The bindings of *?b* do not occur in the output, but the right number of duplicates needs to be produced. Note that the aggregated indexes are much smaller than the full-triple indexes; the increase of the total database size caused by the six additional indexes is negligible.

Instead of (*value*₁, *value*₂, *value*₃), the aggregated indexes store (*value*₁, *value*₂, *count*), but otherwise they are organized in B⁺-trees just like the full-triple compressed indexes. The leaf encoding is slightly different, as now most changes involve a gap in *value*₂ and a low *count* value. The pseudo-code is shown in Figure 4.

Finally, in addition to these indexes for pairs in triples, we also build all three one-value indexes containing just (*value*₁, *count*) entries (the encoding is analogous). While triple patterns using only one variable are probably rare, the one-value indexes are very small, and having them available simplifies query translation.

4. QUERY PROCESSING AND OPTIMIZATION

4.1 Translating SPARQL Queries

The first step in compiling a SPARQL query is to transform it into a calculus representation suitable for later optimizations. We construct a query graph representation that can be interpreted as relational tuple calculus. It would be simpler to derive domain calculus from SPARQL, but tuple calculus is easier to optimize later on.

While SPARQL allows many syntax shortcuts to simplify query formulation, each (conjunctive) query can be parsed and expanded into a set of triple patterns. Each component of a triple is either a literal or a variable. The parser already performs dictionary lookups, i.e., the literals are mapped into ids. Similar to domain calculus, SPARQL specifies that variable bindings must be produced for every occurrence of the resulting literals-only triple in the data. When a query consists of a single triple pattern, we can use our index structures from Section 3 and answer the query with a single range scan. When a query consist of multiple triple patterns, we must join the results of the individual patterns. We thus employ join ordering algorithms on the query graph representation, as discussed further below.

Each triple pattern corresponds to one node in the query graph. Conceptually each node entails a scan of the whole database with appropriate variable bindings and selections induced by the literals. While each of these scans could be implemented as a single index range scan, the optimizer might choose a different strategy (see below). The edges in the query graph reflect the variable occurrences. Two nodes are connected if and only if they have a (query) variable in common.

When the projection clause of a query includes the *distinct* option, we add an aggregation operator that eliminates duplicates in the result. Finally we add a dictionary lookup operator that converts the resulting ids back into strings.

4.2 Optimizing Join Ordering

The key issue for optimizing SPARQL execution plans is join ordering. There is a rich body of literature on this problem, with solutions typically based on various forms of dynamic programming (DP) or randomization (e.g., [13, 15, 26, 34]). However, the intrinsic characteristics of RDF and SPARQL create join queries with particularly demanding properties, which are not directly addressed by prior work:

1. SPARQL queries tend to contain star-shaped subqueries, for combining several attribute-like properties of the same entity. Thus, it is essential to use a strategy that can create bushy join trees (rather than focusing on left-deep or right-deep trees).
2. These star joins occur at various nodes of long join paths, often at the start and end of a path. SPARQL queries can easily lead to 10 or more joins between triples (see, for example, our benchmark queries in Section 6). So exact optimization either requires very fast plan enumeration and cost estimation or needs to resort to heuristic approximations.
3. We would like to leverage the particular strengths of our triple indexes, which encourage extensive use of merge joins (rather than hash or nested-loop joins), but this entails being very careful about preserving interesting orders in the generation of join plans.

The first requirement rules out methods that cannot generate all possible star-chain combinations. The second requirement strongly suggests a fast bottom-up method rather than transformation-based top-down enumeration. The third requirement rules out sampling-based plan enumeration (or randomized optimization methods), as these are unlikely to generate all order-preserving plans for queries with more than 10 joins. In fact, we expect that the most competitive execution plans have a particular form: they would use order-preserving merge-joins as long as possible and switch to hash-joins only for the last few operators.

Our solution is based on the bottom-up DP framework of [26]. It seeds its DP table with scans for the base relations, in our case the triple patterns. The seeding is a two step process. First, the optimizer analyzes the query to check which variable bindings are used in other parts of the query. If a variable is unused, it can be projected away by using an aggregated index (see Section 3.3). Note that this projection conceptually preserves the cardinality through the count information in the index; we discuss this in Subsection 4.4 below. In the second step the optimizer decides which of the applicable indexes to use. There are two factors that

	6-triples star patterns	8-triples star patterns	10-triples star patterns
path length 5	0.7 ms	2.4 ms	6.4 ms
path length 10	2.0 ms	4.6 ms	17.3 ms
path length 20	3.5 ms	13.0 ms	56.6 ms

Figure 5: Optimization times for queries with x patterns in two stars connected by a path of length y

affect the index selection. When the literals in the triple pattern form the prefix of the index key, they are automatically handled by the range scan. Otherwise too many tuples are read and additional selections are required. On the other hand, a different index might produce the triples in an order suitable for a subsequent merge-join later on, which may have lower overall costs than reading too many tuples. The optimizer therefore generates plans for all indexes and uses the plan pruning mechanism to decide if some of them can be discarded early on.

Pruning is done primarily based upon estimated execution costs. That is, the optimizer calls the cost model for each generated plan and prunes equivalent plans that are dominated by cheaper alternatives. This pruning mechanism relies on order optimization [35] to decide if a plan can be dominated by another plan. As the optimizer can use indexes on all triple permutations, it can produce tuples in an arbitrary order, which makes merge joins very attractive. Thus plans are kept if they are more expensive but produce an interesting ordering that can be used later on. Note that orderings are not only created by index scans but also by functional dependencies induced by selections, therefore the order optimization component is non-trivial [35].

Starting with the seeds, larger plans are created by joining optimal solutions of smaller problems. During this process all attribute-handling logic is implemented as reasoning over the available equivalence classes instead of individual variable bindings. Each plan will produce at most one binding for each equivalence class. This both simplifies implicit projections in front of pipeline breakers and allows for automatic detection of transitive join conditions (i.e., $a = b \wedge b = c \Rightarrow a = c$).

Starting with the seeds, larger plans are created by joining optimal solutions of smaller problems that are adjacent in the query graph [26]. When the query contains additional selections due to *FILTER* predicates, they are placed greedily as soon as possible, as they are usually inexpensive to evaluate. If a selection is really expensive, it could be better to integrate it into the DP operator placement as proposed in [9], but we did not investigate this further. The DP method that we implemented along these lines is very fast and is able to compute the exact cost-optimal solution for join queries with up to 20 triple patterns. We measured optimization times (in milliseconds) for a typical SPARQL scenario with two entities selected by star-shaped subqueries and connected by a chain of join patterns. The results are shown in Figure 5.

4.3 Handling Disjunctive Queries

While conjunctive queries are more commonly used, SPARQL also allows certain forms of disjunctions. The *UNION* expression returns the union of the bindings produced by two or more pattern groups. The *OPTIONAL* expressions returns the bindings of a pattern group if there are any re-

sults, and NULL values otherwise. In this context, pattern groups are sets of triple patterns, potentially containing *UNION* and *OPTIONAL* expressions themselves.

During optimization we treat pattern groups in *UNION* and *OPTIONAL* as nested subqueries. That is, we optimize the nested pattern groups first (potentially recursively), and then treat them as base relations with special costs/cardinalities during the optimization of the outer query. For *UNION* we add the union of the results of the pattern groups as if it were a base relation, for *OPTIONAL* we add the result as a base relation using an outer join.

In principle it would be possible to optimize these queries more aggressively, but most interesting optimizations require the usage of bypass plans [37] or other non tree-structured execution plans, which is beyond the scope of this work. And these optimizations would only pay off for complex queries; when the disjunctive elements are simple, our nested optimization scheme produces the optimal solution.

4.4 Preserving Result Cardinality

The standard SPARQL semantic requires that the right number of variable bindings are produced, even if many of them are duplicates. However, from a processing point of view, one should avoid the additional work for producing and keeping duplicates.

We solve this issue by tracking the multiplicity of each tuple during query processing. Scans over unaggregated indexes always produce a multiplicity of 1, while aggregated indexes report the number of duplicates as multiplicity. Join operators multiply the multiplicities to get the number of duplicates of each output tuple. Note that we can optionally switch off the multiplicity tracking if we can statically derive that it has to be 1 in a subplan. When the result is presented to the application/user, the output operator interprets the multiplicities according to the specified query semantics (*distinct*, *reduced*, or *standard*).

4.5 Implementation Issues

Our run-time system includes the typical set of algebraic operators (merge-join, hash-join, filter, aggregation, etc.). One notable difference to other systems is that our run-time system is very RISC-style: most operators merely process integer-encoded ids, consume and produce streams of id tuples, compare ids, etc. Besides simplifying the code, this reduced complexity allows neat implementation tricks.

For example, consider an index-scan operator that uses a B^+ -tree iterator to access the physical data, comparing a triple pattern against the data. Each entry in the triple is either an id attribute that must be produced or bound to a literal, which affects start/stop condition if it is in the prefix or implies a selection if unbound entries come first. Instead of checking for these different conditions at run-time, we can handle them at query compilation time. Each entry is either an id attribute or a literal. There are three entries in a triple, which means there are eight possible combinations. With a single method interface that has eight different implementations of the index scan operator, we can greatly reduce the number of conditional branches in the system code. Besides being faster, each specialized operator is much simpler as it now implements just the logic required for its setting. Note that we only need to specialize the step logic, which is less than 10 lines of code for each specialization.

This RISC-style combination of simplified type system

start (s,p,o)		end (s,p,o)	
number of triples			
number of distinct 2-prefixes			
number of distinct 1-prefixes			
join partners on subject			
	s=s	s=p	s=o
join partners on predicate			
	p=s	p=p	p=o
join partners on object			
	o=s	o=p	o=o

Figure 6: Structure of a histogram bucket

and simple, fast operators leads to very good CPU performance. In our evaluation in Section 6 we include warm-cache times to demonstrate these effects. We realize that these kinds of code-tuning issues are often underappreciated, but are crucial for high performance on modern hardware.

5. SELECTIVITY ESTIMATES

The query optimizer relies upon its cost model in finding the lowest-cost execution plan. In particular, estimated cardinalities (and thus selectivities) have a huge impact on plan generation. While this is a standard problem in database systems, the schema-free nature of RDF data complicates statistics generation. We propose two kinds of statistics. The first one, specialized histograms, is generic and can handle any kind of triple patterns and joins. Its disadvantage is that it assumes independence between predicates, which frequently does not hold in tightly coupled triple patterns. The second statistics therefore computes frequent join paths in the data, and gives more accurate predictions on these paths for large joins. During query optimization, we use the join-path cardinalities when available and otherwise assume independence and use the histograms.

5.1 Selectivity Histograms

While triples conceptually form a single table with three columns, histograms over the individual columns are not very useful as most query patterns touch at least two attributes of a triple. Instead we harness our aggregated indexes, which are perfectly suited for the calculation of triple-pattern selectivities: for each literal or literal pair, we can get the exact number of matching triples with one index lookup. Unfortunately this is not sufficient for estimating join selectivities. Also, we would like to keep all auxiliary structures for the cost model in main memory. Therefore, we aggregate the indexes even further such that each index fits into a single database page and includes information about join selectivity.

Just like the aggregated indexes we build six different statistics, one for each order of the entries in the triples. Starting from the aggregated indexes, we place all triples with the same prefix of length two in one bucket and then merge the smallest two neighboring buckets until the total histogram is small enough. This approximates an equi-depth histogram, but avoids placing bucket boundaries between triples with the same prefix (which are expected to be similar).

For each bucket we then compute the statistics shown in Figure 6. The first three values – the number of triples, num-

ber of distinct 2-prefixes, and number of distinct 1-prefixes – are used to estimate the cardinality of a single triple pattern. Note that this only gives the scan cardinality, i.e., the number of scanned triples, which determines the costs of an index scan. The true result cardinality, which affects subsequent operators, could actually be lower when literals are not part of the index prefix and are tested by selections later on. In this case we derive the result cardinality (and obtain exact predictions) by reordering the literals such that all literals are in the prefix.

The next values are the numbers of join partners (i.e., the result cardinality) if the triples in the bucket were joined to all other triples in the database according to the specified join condition. As there are nine ways to combine attributes from two triples, we precompute nine cardinalities. For example the entry $o = s$ is effectively

$$|\{b|b \in \text{current bucket}\} \bowtie_{b.\text{object}=t.\text{subject}} \{t|t \in \text{all triples}\}|.$$

These values give a perfect join-size prediction when joining a pattern that exactly matches the bucket with a pattern without literals. Usually this is not the case, we therefore assume independence between query conditions and multiply the selectivities of the involved predicates. (Such independence assumptions are standard in state-of-the-art query optimizers for tractability.)

5.2 Frequent Paths

The histograms discussed above have decent accuracy, and are applicable for all kinds of predicates. Their main weakness is that they assume independence between predicates. Two kinds of correlated predicates commonly occur in SPARQL queries. First, “stars” of triple patterns, where a number of triple patterns with different predicates share the same subject. These are used to select specific subjects (i.e., entities based on different attributes of the same entities). Second, “chains” of triple patterns, where the object of the first pattern is the subject of the next pattern, again with given predicates. These chains correspond to long join paths (across different entities). As both of these two cases are common, we additionally build specialized statistics to have more accurate estimators for such queries.

To this end, we precompute the frequent paths in the data graph and keep exact join statistics for them. Frequency here refers to paths with the same label sequence. Note that we use the term path both for chains and stars, the constructions are similar in the two cases. We characterize a path P by the sequence of predicates $p_1, \dots, p_n$ seen in its traversal. Using SPARQL syntax, we define a (chain) path $P_{p_1, \dots, p_n}$ as

$$P_{p_1, \dots, p_n} := \text{select } r_1 \ r_{n+1} \text{ where } \{ (r_1 \ p_1 \ r_2). \\ (r_2 \ p_2 \ r_3). \dots (r_n \ p_n \ r_{n+1}) \}$$

Star paths are defined analogous, the $p_1, \dots, p_n$ are unsorted in this case. We compute the most frequent paths, i.e., the paths with the largest cardinalities, and materialize their result cardinalities and path descriptions $p_1, \dots, p_n$. Using this information we can exactly predict the join cardinality for the frequent paths that occur in a query. Again, we want to keep these statistics in main memory and therefore compute the most frequent paths such that they still fit on a single database page. In our experiments we could store about 1000 paths on one 16KB page.

Finding the most frequent paths requires some care. While it may seem that this is a standard graph-mining issue, the

FrequentPath(k)

// Computes the k most frequent paths

$C_1 = \{P_p | p \text{ is a predicate in the database}\}$

sort C_1 , keep the k most frequent

$C = C_1, i = 1$

do

$C_{i+1} = \emptyset$

for each $p' \in C_i, p$ predicate in the database

if top k of $C \cup C_{i+1} \cup \{P_{p'p}\}$ includes all subpaths of $p'p$

$C_{i+1} = C_{i+1} \cup \{P_{p'p}\}$

if top k of $C \cup C_{i+1} \cup \{P_{pp'}\}$ includes all subpaths of pp'

$C_{i+1} = C_{i+1} \cup \{P_{pp'}\}$

$C = C \cup C_{i+1}$, sort C , keep the k most frequent

$C_{i+1} = C_{i+1} \cap C, i = i + 1$

while $C_i \neq \emptyset$

return C

Figure 7: Frequent Path Mining Algorithm

prevalent methods in that line of research [16, 44, 49], e.g., based on the well-known Apriori frequent-itemset mining algorithm, are not directly applicable.

Unlike the Apriori setting, a frequent path in our RDF-path sense does not necessarily consist of frequent subpaths. Consider a graph with two star-shaped link clusters where all end-nodes are connected to their respective star centers by predicates (edge labels) p_1 and p_2 , respectively. Now consider a single edge with predicate p_3 between the two star centers. In this scenario, the path P_{p_3} will be infrequent, while the path P_{p_1, p_3, p_2} will be frequent. Therefore we cannot simply use the Apriori algorithm.

Another problem in our RDF setting are cycles, which could lead to seemingly infinitely long, infinitely frequent paths. We solve this problem by two means. First, we require that if a frequent path P is to be kept, all of its subpaths have to be kept, too. This is required for query optimization purposes anyway, as we may have to break a long join-path into smaller joins, and it simplifies the frequent-path computation. Second, we rank the frequent paths not by their result cardinalities but by their number of distinct nodes. In a tree these two are identical, but in the presence of cycles we do not count nodes twice.

The pseudo-code of the path mining algorithm is shown in Figure 7. It starts from frequent paths of length one and enlarges them by appending or prepending predicates. When a new path is itself frequent and all of its subpaths are still kept, we add it. We stop when no new paths can be added. Note that, although the pseudo-code shows a nested loop for ease of presentation, we actually use a join and a group-by operator in the implementation. For the datasets we considered in our experiments the 1000 most frequent paths could be determined in a few minutes.

5.3 Estimates for Composite Queries

For estimating the overall selectivity of an entire composite query, we combine the histograms with the frequent paths statistics. A long join chain with intermediate nodes that have triple patterns with object literals is decomposed into subchains of maximal lengths such that only their end nodes have triple patterns with literals. For example, a query like

$?x_1 \ a_1 \ v_1. ?x_1 \ p_1 \ ?x_2. ?x_2 \ p_2 \ ?x_3. ?x_3 \ p_3 \ ?x_4. \\ ?x_4 \ a_4 \ v_4. ?x_4 \ p_4 \ ?x_5. ?x_5 \ p_5 \ ?x_6. ?x_6 \ a_6 \ v_6$

with attribute-flavored predicates a_1, a_4, a_6 , literals v_1, v_4, v_6 , and relationship-flavored predicates p_1 through p_5 will be broken down into the subchains for $p_1-p_2-p_3$ and for p_4-p_5 and the per-subject selections a_1-v_1, a_4-v_4 , and a_6-v_6 . We use the frequent paths statistics to estimate the selectivity of the two join subchains, and the histograms for selections. Then, in the absence of any other statistics, we assume that the different estimators are probabilistically independent, leading to a product formula with the per-subchain and per-selection estimates as factors. If instead of a simple attribute-value selection like $?x_6 a_6 v_6$ we had a star pattern such as $?x_6 a_6 u_6, ?x_6 b_6 v_6, ?x_6 c_6 w_6$ with properties a_6, b_6, c_6 and corresponding object literals u_6, v_6, w_6 , we would first invoke the estimator for the star pattern, using the frequent paths statistics for stars, and then combine them with the other estimates in the product form.

6. EVALUATION

6.1 General Setup

For evaluating the performance of RDF-3X, we used three large datasets with different characteristics and compared the query run-times to other approaches (discussed below). All experiments were conducted on a Dell D620 PC with a 2 Ghz Core 2 Duo processor, 2 GBytes of memory, and running a 64-bit Linux 2.6.24 kernel. For the cold-cache experiments we used the `/proc/sys/vm/drop_caches` kernel interface to drop all filesystem caches before restarting the various systems under test. We repeated all queries five times (including the dropping of caches and the system restart) and took the best result to avoid artifacts caused by OS activity. For warm caches we ran the queries five times without dropping caches, again taking the best run-time.

Our primary comparison is against the *column-store-based approach* presented in [1], which has already been shown to be highly superior to all other DBMS-based approaches in that paper. We implemented the approach as described in [1], but used MonetDB 5.2.0 [27] as a backend instead of C-Store because C-Store is no longer maintained and does not run on our hardware/OS platform. The C-Store web page <http://db.csail.mit.edu/projects/cstore/> suggests using MonetDB instead, and MonetDB worked fine. Note that our setup uses substantially weaker hardware than [1]; in particular the hard disk is about a factor of 6 slower than the very fast RAID used in [1], transferring ca. 32 MB/s in sequential reads. Taking this factor of 6 into account, the performance numbers we got for MonetDB are comparable to the C-Store numbers from [1]. For one query (Q6) MonetDB was significantly faster than a factor of 6 (14s vs. 10s), while for another (Q7) significantly slower (61s vs. 1.4s), but overall MonetDB performed as expected given the slower hard disk.

As a second opponent to RDF-3X, we used *PostgreSQL 8.3 as a triple store* with indexes on the string dictionary and on (subject, predicate, object), (predicate, subject, object), and (predicate, object, subject). This emulates a Sesame-style [8] storage system. We also tried out the current release of a leading commercial database system with built-in RDF support, but could not obtain acceptable performance anywhere near the run-times of the other systems. When using its own RDF query language and despite trying several of its auto-tuning options, it performed significantly slower than the PostgreSQL triple store even for simple queries, and

failed to execute more complex queries in reasonable time. We therefore omitted it from the presentation.

In addition to these DBMS-based opponents, we tried several systems from the semantic web community that are available as open-source code. Unfortunately none of them scaled to the dataset sizes that we used. We first tried the popular Jena2 system [46] which came out of the HP Labs Semantic Web Programme. We used *Jena version 2.5.5* with the SDB 1.0 wrapper and Apache Derby 10.3.2.1, but were unable to import any of our three datasets in 24 hours. Judging from the file growth, the system became continuously slower and did not seem to terminate in a reasonable time. We also tried *Yars2* [19, 48], but again were unable to import any of our datasets in 24 hours. Finally, we tried *Sesame 2.0* [8, 28], which is supposed to be one of the fastest semantic web systems. Sesame 2.0 was able to import the Barton dataset in 13 hours, but then needed ca. 15 minutes for each of the first two queries and crashed due to excessive memory usage for the more complex queries.

Note that both MonetDB and RDF-3X could import the data sets in less than half an hour, and could run the queries in the order of seconds. Other semantic web approaches usually assume that the RDF data fits into main memory, which is not the case here. All experiments below therefore only consider RDF-3X, the column-stored-based approach on top of MonetDB, and the PostgreSQL-based triples store.

Independently of the database system, each of the datasets discussed below is first brought into a factorized form: one file with RDF triples represented as integer triples and one dictionary file mapping from integers to literals. All three systems use the same files as inputs, loading them into fact table(s) and dictionary. The load times of this second phase and the database sizes are shown in Figure 8. The MonetDB sizes are the initial sizes after loading. After running the benchmark the sizes were 2.0 / 2.4 / 6.9 GB. Apparently MonetDB builds some index structures on demand.

6.2 Barton Dataset

For the first experiment we used the Barton Library dataset and the queries proposed as a benchmark in [1]. We processed the data as described in [1], converting it into triple form using the Redland parser, and then imported the triples into our RDF-3X system. In [1] the authors pruned the data due to C-Store limitations (they dropped all triples containing strings longer than 127 bytes and some triples with a huge number of join partners). We left the complete data as it was and imported it directly into all three systems. Overall the data consists of 51,598,328 distinct triples, and 19,344,638 distinct strings. The original data was 4.1 GB in RDF (XML) format, 7.7 GB in triple form, and 2.8 GB in our RDF-3X store including all indexes and the string dictionary.

We used the queries from [1] for our experiment, but as they were given in SQL we had to reformulate them in SPARQL for RDF-3X. Appendix A shows all queries. The results of our measurements are shown in Figure 9. We include also the geometric mean of the query set, which is often used as a workload-average measure in benchmarks (e.g., TPC) and is more resilient to extreme outliers than the arithmetic average.

The first observation is that RDF-3X performs much better than MonetDB for all queries, and MonetDB itself performs much better than PostgreSQL (as reported in [1]). We

	Barton		Yago		Librarything	
	Load time	DB size	Load time	DB size	Load time	DB size
RDF-3X	13 min	2.8 GB	25 min	2.7 GB	20 min	1.6 GB
MonetDB	11 min	1.6 GB	21 min	1.1 GB	4 min	0.7 GB
PostgreSQL	30 min	8.7 GB	25 min	7.5 GB	20 min	5.7 GB

Figure 8: Database load after triple construction

	Q1	Q2	Q3	Q4	Q5	Q6	Q7	geom. mean
cold caches								
RDF-3X	0.14	3.10	31.49	11.57	18.82	2.75	32.61	5.9
MonetDB	5.66	11.70	54.66	34.77	80.96	14.14	61.52	26.4
PostgreSQL	28.32	181.00	291.04	224.61	199.07	207.72	271.20	167.8
warm caches								
RDF-3X	0.001	1.17	2.22	1.58	0.49	1.20	1.26	0.4
MonetDB	0.65	1.41	3.63	9.59	77.53	1.86	2.48	3.8
PostgreSQL	8.15	174.41	286.76	26.80	8.77	206.46	231.79	64.3

Figure 9: Query run-times in seconds for the Barton dataset

first discuss the results for RDF-3X vs. MonetDB. When comparing the cold-cache times and the warm-cache times, it becomes clear that disk I/O has a large impact on the overall run-times. RDF-3X simply reads less data due to its highly compressed index structures, therefore outperforming MonetDB in the cold-cache case by a typical factor of 2 to 5, and sometimes by more than 10. In the warm-cache case the differences are typically smaller but still substantial (factor of 2, sometimes much higher). An interesting data point is query Q4, which is relatively expensive in terms of constructed join pairs, and where RDF-3X performs very well even in CPU-dominated warm-cache case. Furthermore, we observe that a third critical aspect besides I/O and CPU usage is memory consumption. Query Q5 has a very large intermediate result. MonetDB apparently materializes parts of these intermediate results in main memory. As a consequence only few database pages can be buffered, which significantly hurts warm-cache behavior.

PostgreSQL has problems with this dataset, due to the nature of the queries for this benchmark. Nearly all queries are aggregations queries (usually aggregating by predicate), and the result cardinality is large which entails expensive dictionary lookups. For other, more natural, kinds of RDF queries, PostgreSQL performs much better, as we will see in the next two subsections.

To get an idea how a Yars2-style system could scale we experimentally disabled all aggregated indices. This increased the geometric means to 9.52s (cold) and 1.04s (warm), which is significantly slower than RDF-3X. This is still much faster than the other systems, though, in particular due to our run-time system and query optimizer.

6.3 Yago Dataset

The Barton dataset is relatively homogeneous, as it describes library data. As a second dataset we therefore used Yago [40] which consists of facts extracted from Wikipedia (exploiting the infoboxes and category system of Wikipedia) and integrated with the WordNet thesaurus. The Yago dataset contains 40,114,899 distinct triples and 33,951,636 distinct strings, consuming 3.1 GB as (factorized) triple dump. RDF-3X needs 2.7 GB for all indexes and the string

dictionary. As queries we considered three different application scenarios – entity-oriented, relationship-oriented, and queries with unknown predicates – and derived eight benchmark queries, shown in Appendix A. These queries are more “natural” than the Barton queries, as they are standard SPARQL without any aggregations and with explicitly given predicates. On the other hand, the queries are much larger (requiring more many-to-many joins) and thus more difficult to optimize and execute.

The results are shown in Figure 10. Again, RDF-3X clearly outperforms the other two systems for both cold and warm caches, by a typical factor of 5 to 10. Here PostgreSQL performed much better than MonetDB. This is most likely caused by the poor join orderings in MonetDB. The warm-cache run-times are nearly as high as the cold-cache times, which indicates that MonetDB creates large intermediate results.

In general this dataset is much more challenging for the query optimizer, as queries are more complex and selectivity estimates are important. While testing our system, we noticed that selectivity mis-estimations can easily cause slowdown by a factor of 10-100 on this dataset. RDF-3X shows excellent performance regarding both the run-time execution and the choice of execution plans by the optimizer.

6.4 LibraryThing Dataset

As a third dataset we used a partial crawl of the LibraryThing book-cataloging social network www.librarything.com. It consists of 9989 users, 5,973,703 distinct books (personally owned by these users), and the tags that the users have assigned to these books. Overall the dataset consists of 36,203,751 triples and 9,352,954 distinct strings, consuming 1.8 GB in its original form and 1.6 GB in RDF-3X. One particularity of this dataset is that it has a heterogeneous link structure. In our RDF representation, each tag is mapped to a predicate, linking the user to the book she tagged. As the number of different tags is very large, the dataset contains 338,824 distinct predicates, whereas the other two datasets contained only 285 and 93 distinct predicates, respectively. While other mappings onto RDF may be possible, we used this extremely non-schematic approach as a stress test for

	A1	A2	A3	B1	B2	B3	C1	C2	geom. mean
cold caches									
RDF-3X	0.29	0.28	1.20	0.28	0.99	0.33	2.23	4.23	0.73
MonetDB	43.55	44.13	55.49	62.94	182.39	72.22	101.66	157.11	78.29
PostgreSQL	1.62	6.31	5.46	3.04	117.51	4.71	29.84	59.64	10.66
warm caches									
RDF-3X	0.02	0.02	0.02	0.01	0.05	0.01	0.61	1.44	0.04
MonetDB	36.92	32.96	34.72	49.95	64.84	52.22	84.41	131.35	54.62
PostgreSQL	0.08	0.43	0.20	0.11	7.33	0.12	0.31	50.37	0.56

Figure 10: Query run-times in seconds for the Yago dataset

all competing systems.

This data makes compression more difficult for RDF-3X, and causes serious problems for MonetDB. MonetDB was unable to handle 338,824 tables, creating millions of files in the file system and swapping all the time. We therefore used a hybrid storage scheme for MonetDB for this dataset. We partitioned the 1000 most commonly used predicates as described in [1], and placed the remaining triples (ca. 12%) in one big triples table. We again constructed three kinds of queries: book-oriented, user-oriented, and navigating book and user chains (see Appendix A). In contrast to the Yago dataset, there were few predicates that occurred in millions of triples, which lowered the impact of join-ordering decisions. On the other hand, the data itself is very inhomogeneous so that selectivities are more difficult to predict.

The results are shown in Figure 11. RDF-3X performs very well, outperforming the opponents by a typical factor of at least 5 and more than 30 in some cases. Between MonetDB and PostgreSQL there is no clear winner. Overall MonetDB seems to perform better, but it crashed two times. It refused to execute query B3 ("too many variables"), probably because it included three patterns with variable predicates (and thus at least 3000 scans). In query C2 it crashed after 15 minutes due to lack of disk space, as it had materialized a 20 GB intermediate result (which is more than 10 times the size of the whole database).

The query A3 stands out by its high run-times. It performs many joins with relatively unselective predicates (book authors, etc.), which are expensive. The other "difficult" queries (B3, C1, C2) are not that difficult per se, they just require the right choice of execution plans. B3, for example, finds all users with books tagged as English, French, and German. PostgreSQL starts this query by collecting all pairs of books a user has, which is prohibitively expensive. The optimizer of RDF-3X, on the other hand, chooses a plan that collects for each tag the users with such books and then joins the results, which is much more efficient.

7. CONCLUSION

This paper has presented the RDF-3X engine, a RISC-style architecture for executing SPARQL queries over large repositories of RDF triples. As our experiments have shown, RDF-3X outperforms the previously best systems by a large margin. In particular, it addresses the challenge of schema-free data and, unlike its opponents, copes very well with data that exhibits large diversity of property names. The salient features of RDF-3X that lead to these performance gains are 1) exhaustive but very space-efficient triple indexes that eliminate the need for physical-design tuning, 2) a streamlined execution engine centered around very fast merge joins,

3) a smart query optimizer that chooses cost-optimal join orderings and can do this efficiently even for long join paths (involving 10 to 20 joins), 4) a selectivity estimator based on statistics for frequent paths that feeds into the optimizer's cost model.

Our future work includes further improvements of the query optimizer (e.g., based on magic sets) and support for RDF search features that go beyond the current SPARQL standard. Along the latter lines, one direction is to allow more powerful wild-card patterns for entire paths, in the spirit of the XPath descendants axis but for graphs rather than trees. Proposals for extending SPARQL have been made [3], but there is no implementation yet. A second direction is to provide ranking of query results, based on application-specific scoring models. This calls for top-k query processing and poses challenging issues for algorithms and query optimization. Finally, full SPARQL support requires some additional information, in particular typing information. We feel that this can be included in the dictionary, but determining the best encoding relative to runtime performance and compression rate needs more work.

8. REFERENCES

- [1] D. J. Abadi, A. Marcus, S. Madden, and K. J. Hollenbach. Scalable semantic web data management using vertical partitioning. In *VLDB*, 2007.
- [2] S. Alexaki et al. The ics-forth rdfsuite: Managing voluminous rdf description bases. In *SemWeb*, 2001.
- [3] K. Anyanwu, A. Maduko, and A. P. Sheth. Sparq2l: towards support for subgraph extraction queries in rdf databases. In *WWW*, 2007.
- [4] S. Auer et al. Dbpedia: A nucleus for a web of open data. In *ISWC/ASWC*, 2007.
- [5] L. Baolin and H. Bo. Path queries based rdf index. In *SKG*, 2005.
- [6] L. Baolin and H. Bo. Hprd: A high performance rdf database. In *NPC*, 2007.
- [7] BioPAX: Biological Pathways Exchange. <http://www.biopax.org/>.
- [8] J. Broekstra, A. Kampman, and F. van Harmelen. Sesame: An architecture for storing and querying rdf data and schema information. In *Spinning the Semantic Web*, 2003.
- [9] S. Chaudhuri and K. Shim. Optimization of queries with user-defined predicates. *TODS*, 24(2), 1999.
- [10] S. Chaudhuri and G. Weikum. Rethinking database system architecture: Towards a self-tuning risc-style database system. In *VLDB*, 2000.
- [11] E. I. Chong et al. An efficient sql-based rdf querying

	A1	A2	A3	B1	B2	B3	C1	C2	geom. mean
cold caches									
RDF-3X	0.28	1.01	21.85	0.14	0.34	4.17	0.28	1.21	0.89
MonetDB	2.14	1.41	1220.09	1.63	2.20	*	1.66	>15min/*	>8.16
PostgreSQL	20.78	1.43	715.64	0.88	2.13	> 8h	5108.01	1031.63	>93.91
warm caches									
RDF-3X	0.05	0.15	0.95	0.01	0.12	1.61	0.03	0.26	0.13
MonetDB	0.82	0.77	1171.82	0.56	0.63	*	0.59	>15min/*	>4.39
PostgreSQL	12.31	0.05	611.41	0.02	0.66	> 8h	5082.34	1013.01	>30.43

* system crashed, see description

Figure 11: Query run-times in seconds for the LibraryThing dataset

- scheme. In *VLDB*, 2005.
- [12] E. Chu, J. L. Beckmann, and J. F. Naughton. The case for a wide-table approach to manage sparse relational data sets. In *SIGMOD*, 2007.
- [13] D. DeHaan and F. W. Tompa. Optimal top-down join enumeration. In *SIGMOD*, 2007.
- [14] A. Eickler, C. A. Gerlhof, and D. Kossmann. A performance evaluation of oid mapping techniques. In *VLDB*, 1995.
- [15] C. A. Galindo-Legaria, A. Pellenkoff, and M. L. Kersten. Fast, randomized join-order selection - why use transformations? In *VLDB*, 1994.
- [16] L. Getoor and C. P. Diehl. Link mining: a survey. *SIGKDD Explorations*, 7(2), 2005.
- [17] G. Graefe. Query evaluation techniques for large databases. *ACM Comput. Surv.*, 25(2), 1993.
- [18] A. Y. Halevy, M. J. Franklin, and D. Maier. Principles of dataspace systems. In *PODS*, 2006.
- [19] A. Harth, J. Umbrich, A. Hogan, and S. Decker. Yars2: A federated repository for querying graph structured data from the web. In *ISWC/ASWC*, 2007.
- [20] O. Hartig and R. Heese. The sparql query graph model for query optimization. In *ESWC*, 2007.
- [21] A. Hogan and A. Harth. The expertfinder corpus 2007 for the benchmarking development of expert-finding systems. In *International ExpertFinder Workshop*, 2007.
- [22] D. Huynh, S. Mazzocchi, and D. R. Karger. Piggy bank: Experience the semantic web inside your web browser. *J. Web Sem.*, 5(1), 2007.
- [23] Jena: a Semantic Web Framework for Java. <http://jena.sourceforge.net/>.
- [24] M. Kersten and A. P. Siebes. An organic database system. Technical report, CWI, 1999.
- [25] A. Maduko et al. Estimating the cardinality of rdf graph patterns. In *WWW*, 2007.
- [26] G. Moerkotte and T. Neumann. Analysis of two existing and one new dynamic programming algorithm for the generation of optimal bushy join trees without cross products. In *VLDB*, 2006.
- [27] Monetdb. <http://monetdb.cwi.nl/>.
- [28] Openrdf. <http://www.openrdf.org/index.jsp>.
- [29] Oracle technical network, semantic technologies center. http://www.oracle.com/technology/tech/semantic_technologies/index.html.
- [30] W3C: Resource Description Framework (RDF). <http://www.w3.org/RDF/>.
- [31] RDFizers. <http://simile.mit.edu/wiki/RDFizers>.
- [32] ICS-FORTH RDF suite. <http://athena.ics.forth.gr:9090/RDF/>.
- [33] F. Scholer et al. Compression of inverted indexes for fast query evaluation. In *SIGIR*, 2002.
- [34] P. G. Selinger et al. Access path selection in a relational database management system. In *SIGMOD*, 1979.
- [35] D. E. Simmen, E. J. Shekita, and T. Malkemus. Fundamental techniques for order optimization. In *SIGMOD*, 1996.
- [36] W3C: SPARQL Query Language for RDF. <http://www.w3.org/TR/rdf-sparql-query/>.
- [37] M. Steinbrunn et al. Bypassing joins in disjunctive queries. In *VLDB*, 1995.
- [38] M. Stocker et al. Sparql basic graph pattern optimization using selectivity estimation. In *WWW*, 2008.
- [39] M. Stonebraker et al. One size fits all? part 2: Benchmarking studies. In *CIDR*, 2007.
- [40] F. M. Suchanek, G. Kasneci, and G. Weikum. Yago: a core of semantic knowledge. In *WWW*, 2007.
- [41] Y. Theoharis, V. Christophides, and G. Karvounarakis. Benchmarking database representations of rdf/s stores. In *International Semantic Web Conference*, 2005.
- [42] O. Udrea, A. Pugliese, and V. S. Subrahmanian. Grin: A graph based rdf index. In *AAAI*, 2007.
- [43] uniprot RDF. <http://dev.isb-sib.ch/projects/uniprot-rdf/>.
- [44] N. Vanetik and E. Gudes. Mining frequent labeled and partially labeled graph patterns. In *ICDE*, 2004.
- [45] T. Westmann et al. The implementation and performance of compressed databases. *SIGMOD Record*, 29(3), 2000.
- [46] K. Wilkinson et al. Efficient rdf storage and retrieval in jena2. In *SWDB*, 2003.
- [47] W3C: RDF/OWL representation of WordNet. <http://www.w3.org/TR/wordnet-rdf/>.
- [48] Yars2. <http://sw.deri.org/svn/sw/2004/06/yars>.
- [49] F. Zhu, X. Yan, J. Han, and P. S. Yu. gprune: A constraint pushing framework for graph pattern mining. In *PAKDD*, 2007.
- [50] J. Zobel and A. Moffat. Inverted files for text search engines. *ACM Comput. Surv.*, 38(2), 2006.

APPENDIX

A. SPARQL QUERIES

For completeness we include the SPARQL queries used in our evaluation.

Barton Dataset. As the queries in [1] were given in SQL, we had to reformulate them in SPARQL. We abbreviate some constants here, the queries are discussed in [1]. We had to extend the SPARQL projection clause a bit to get equivalent queries. *count* is like *distinct* but includes the number of occurrences. *duplicates* is like *count* but only returns bindings that are produced at least twice.

Q1: select count ?c where { ?a a ?c }

Q2: select count ?bp where { ?as a <Text>; ?bp ?bo. filter (?bp in <predicate list>)}

Q3: select duplicates ?bp ?bo where { ?as a <Text>; ?bp ?bo. filter (?bp in <predicate list>)}

Q4: select duplicates ?bp ?bo where { ?as a <Text>; ?bp ?bo; <language> <iso639-2b/fre>. filter (?bp in <predicate list>)}

Q5: select ?as ?co where { ?as <origin> <marcorg/DLC>; <records> ?bo. ?bo a ?co. filter (?co != <Text>)}

Q6: select count ?ap where { { ?as a <Text> } union { ?as <records> []; a <Text> } ?as ?ap [] . filter (?ap in <predicate list>)}

Q7: select ?as ?bo ?co where { ?as <point> "end"; <encoding> ?bo; a ?co }

Yago Dataset. We grouped the queries thematically into three groups. The first group consists of oriented facts, e.g.: "scientists from Switzerland with a doctoral advisor from Germany" (A1). The second group is relationship oriented, e.g. "two actors from England playing together in the same movie" (B1). The third group examines relationships with unknown predicates, e.g. "two scientists related to the same city" (C1).

A1: select ?gn ?fn where { ?gn <givenNameOf> ?p. ?fn <familyNameOf> ?p. ?p <type> "scientist"; <bornInLocation> ?city; <hasDoctoralAdvisor> ?a. ?a <bornInLocation> ?city2. ?city <locatedIn> "Switzerland". ?city2 <locatedIn> "Germany". }

A2: select ?n where { ?a <isCalled> ?n; <type> "actor"; <livesIn> ?city; <actedIn> ?m1; <directed> ?m2. ?city <locatedIn> ?s. ?s <locatedIn> "United.States". ?m1 <type> "movie"; <producedInCountry> "Germany". ?m2 <type> "movie"; <producedInCountry> "Canada". }

A3: select distinct ?n ?co where { ?p <isCalled> ?n. { ?p <type> "actor" } union { ?p <type> "athlete" } ?p <bornInLocation> ?c. ?c <locatedIn> ?s. ?s <locatedIn> ?co. ?p <type> ?t. filter(?t reaches "politician" via <subClassOf>)}

B1: select distinct ?n1 ?n2 where { ?a1 <isCalled> ?n1; <livesIn> ?c1; <actedIn> ?movie. ?a2 <isCalled> ?n2; <livesIn> ?c2; <actedIn> ?movie. ?c1 <locatedIn> "England". ?c2 <locatedIn> "England". filter (?a1 != ?a2)}

B2: select ?n1 ?n2 where { ?p1 <isCalled> ?n1; <bornInLocation> ?city; <isMarriedTo> ?p2. ?p2 <isCalled> ?n2; <bornInLocation> ?city. }

B3: select distinct ?n1 ?n2 where { ?n1 <familyNameOf> ?p1. ?n2 <familyNameOf> ?p2. ?p1 <type> "scientist"; <hasWonPrize> ?award; <bornInLocation> ?city. ?p2 <type> "scientist"; <hasWonPrize> ?award; <bornInLo-

cation> ?city. filter (?p1 != ?p2)}

C1: select distinct ?n1 ?n2 where { ?n1 <familyNameOf> ?p1. ?n2 <familyNameOf> ?p2. ?p1 <type> "scientist"; [] ?city. ?p2 <type> "scientist"; [] ?city. ?city <type> <site> filter (?p1 != ?p2)}

C2: select distinct ?n where { ?p <isCalled> ?n; [] ?c1. [] ?c2. ?c1 <type> <village>; <isCalled> "London". ?c2 <type> <site>; <isCalled> "Paris". }

LibraryThing Dataset. Similar to the Yago setting we used three query groups. First queries on books (e.g., A1 "books tagged with romance, love, suspense, mystery"), second queries on users (e.g., B1 "users who like crime novels and Arthur Conan Doyle and have friends who like romances and Jane Austen"), and third queries with chains over books and users (e.g., C1 "books tagged with romance by users who have friends or friends of friends who have tagged books with documentary which have also been tagged with thriller").

A1: select distinct ?title where { ?b <hasTitle> ?title. [] <romance> ?b. [] <love> ?b. [] <suspense> ?b. [] <mystery> ?b. }

A2: select distinct ?title where { ?b <hasTitle> ?title. ?u <romance> ?b; <love> ?b; <suspense> ?b. }

A3: select distinct ?title where { ?b <hasTitle> ?title; <hasAuthor> ?a. ?u <mystery> ?b; <romance> []. ?b2 <hasAuthor> ?a. [] <children> ?b2. }

B1: select distinct ?u where { ?u <crime> []; <hasFavoriteAuthor> "Arthur Conan Doyle"; <hasFriend> ?f. ?f <romance> []; <hasFavoriteAuthor> "Jane Austen" . }

B2: select distinct ?u where { { ?u <documentary> ?b1; <suspense> ?b1 } union { ?u <biography> ?b2; <suspense> ?b2 } union { ?u <documentary> ?b3; <mystery> ?b3 } union { ?u <biography> ?b4; <mystery> ?b4 } }

B3: select distinct ?u where { ?u [] ?b1; [] ?b2; [] ?b3. [] <english> ?b1. [] <german> ?b2. [] <french> ?b3. }

C1: select distinct ?u where { { ?u <romance> ?b1; <hasFriend> ?f1. ?f1 <biography> ?b2. [] <thriller> ?b2. } union { ?u <romance> ?b1. <hasFriend> ?f1. ?f1 <hasFriend> ?f2. ?f2 <biography> ?b2. [] <thriller> ?b2. } }

C2: select distinct ?u ?u2 where { ?u <hasFavoriteAuthor> ?a1; <america> []; <hasInterestingLibrary> ?u2. ?b1 <hasAuthor> ?a1. [] <europe> ?b1. ?u2 <hasFavoriteAuthor> ?a2; <europe> []. ?b2 <hasAuthor> ?a2. [] <america> ?b2. }