

Benchmark Datasets for the Semantic Web: A Collection for Systematic Machine Learning Evaluations

¹RACHITA RAUL, Gandhi Institute of Excellent Technocrats, Bhubaneswar, India

²PRIYANKA MOHANTY, Indus College of Engineering, Bhubaneswar, Odisha, India

Abstract. Several methods for machine learning on the Semantic Web have been put out in recent years. However, because to a dearth of publicly accessible, recognised benchmark datasets, no in-depth comparisons of those methodologies have been made. We present a collection of 22 benchmark datasets in this work, ranging in size from small to large. Such a set of datasets can be utilised to carry out systematic comparisons of methods and quantitative performance assessment.

Keywords: Linked Open Data Machine learning Datasets Benchmarking

1 Introduction

In the recent years, applying machine learning to Semantic Web data has drawn a lot of attention. Many approaches have been proposed for different tasks at hand, ranging from reformulating machine learning problems on the Semantic Web as traditional, propositional machine learning tasks to developing entirely novel algorithms. However, systematic comparative evaluations of different approaches are scarce; approaches are rather evaluated on a handful of often project-specific datasets, and compared to a baseline and/or one or two other systems.

In contrast, evaluations in the machine learning area are often more rigorous. Approaches are usually compared using a larger number of standard datasets, most often from the UCI repository¹. With a larger set of datasets used in the evaluation, statements about statistical significance are possible as well [3].

At the same time, collections of benchmark datasets have become quite well accepted in other areas of Semantic Web research. Notable examples include the Ontology Alignment Evaluation Initiative (OAEI) for ontology matching², the *Berlin SPARQL Benchmark*³ for triple store performance, the Lehigh University Benchmark (LUBM)⁴ for reasoning, or the Question Answering over Linked Data (QALD) dataset⁵ for natural language query systems.

In this paper, we introduce a collection of datasets for benchmarking machine learning approaches for the Semantic Web. Those datasets are either existing RDF datasets, or external classification or regression problems, for which the instances have been enriched with links to the Linked Open Data cloud [14]. Furthermore, by varying the number of instances for a dataset, scalability evaluations are also made possible.

2 Related Work

Recent surveys on the use of Semantic Web for machine learning organize the proposed approaches in several categories, i.e., approaches that use Semantic Web data for machine learning [16], approaches that perform machine learning on the Semantic Web [11], and approaches that use machine learning techniques to create and improve Semantic Web data [8,16]. Furthermore, there are some challenges, like the *Linked Data Mining Challenge*⁶ or the *Semantic-Web enabled Recommender Systems Challenge*⁷, which usually focus on only a few datasets and a very specific problem setting.

3 Datasets

Our dataset collection has three categories: (i) existing datasets that are commonly used in machine learning experiments, (ii) datasets that were generated from official observations, and (iii) datasets generated from existing RDF datasets. Each of the datasets in the first two categories are initially linked to DBpedia⁸. This has two main reasons, (1) DBpedia being a cross-domain knowledge base usable in datasets from very different topical domains, and (2) tools like DBpedia Lookup and DBpedia Spotlight making it easy to link external datasets to DBpedia. However, DBpedia can be seen as an entry point to the Web of Linked Data, with many datasets linking to and from DBpedia. In fact, we use the RapidMiner Linked Open Data extension [9], to retrieve external links for each entity to YAGO⁹ and Wikidata¹⁰. Such links could be exploited for systematic evaluation of the relevance of the data of different LOD dataset in different learning tasks.

In the dataset collection, there are four datasets that are commonly used for machine learning. For these datasets, we first enrich the instances with links to LOD datasets, and reuse the already defined target variable to perform machine learning experiments:

- The *Auto MPG* dataset¹¹ captures different characteristics of cars, and the target is to predict the fuel consumption (MPG) as a regression task.
- The *AAUP* (American Association of University Professors) dataset contains a list of universities, including eight target variables describing the salary of different staff at the universities¹². We use the average salary as a target variable both for regression and classification, discretizing

the target variable into “high”, “medium” and “low”, using equal frequency binning.

- The *Auto 93* dataset¹³ captures different characteristics of cars, and the target is to predict the price of the vehicles as a regression task.
- The *Zoo* dataset captures different characteristics of animals, and the target is to predict the type of the animals as a classification task.

For those datasets, cars, universities, and animals are linked to DBpedia based on their name.

The second category of datasets contains a list of datasets where the target variable is an observation from different real-world domains, as captured by official sources. Again, the instances were enriched with links to LOD datasets. There are thirteen datasets in this category:

- The *Forbes* dataset contains a list of companies including several features of the companies, which was generated from the Forbes list of leading companies 2015¹⁴. The target is to predict the company’s market value as a classification and regression task. To use it for the task of classification we discretize the target variable into “high”, “medium”, and “low”, using equal frequency binning.
 - The *Cities* dataset contains a list of cities and their quality of living, as captured by Mercer [7]. We use the dataset both for regression and classification.
 - The *Endangered Species* dataset classifies animals into endangered species¹⁵.
- The *Facebook Movies* dataset contains a list of movies and the number of Facebook likes for each movie¹⁶. We first selected 10,000 movies from DBpedia, which were then linked to the corresponding Facebook page, based on the movie’s name and the director. The final dataset contains 1,600 movies, which was created by first ordering the list of movies based on the number of Facebook likes, and then selecting the top 800 movies and the bottom 800 movies. We use the dataset for regression and classification.
- Similarly, the *Facebook Books* dataset contains a list of books and the number of Facebook likes. Each book was linked to the corresponding Facebook page using the book’s title and the book’s author. Again, we selected the top 800 books and the bottom 800 books, based on the number of Facebook likes.
- The *Metacritic Movies* dataset is retrieved from Metacritic.com¹⁷, which contains an average rating of all time reviews for a list of movies [12]. The initial dataset contained around 10,000 movies, from which we selected 1,000 movies from the top of the list, and 1,000 movies from the bottom of the list. We use the dataset both for regression and classification.
 - Similarly, the *Metacritic Albums* dataset is retrieved from Metacritic.com¹⁸, which contains an average rating of all time reviews for a list of albums [13].
- The *HIV Deaths Country* dataset contains a list of countries with the number of deaths caused by HIV, as captured by the World Health Organization¹⁹. We use the dataset both for regression and classification.
- Similarly, the *Traffic Accidents Deaths Country* dataset contains a list of countries with the number of deaths caused by traffic accidents²⁰.
- The *Energy Savings Country* dataset contains a list of countries with the total amount of energy savings of primary energy in 2010²¹, which was downloaded from WorldBank²². We use the dataset both for regression and classification.
- Similarly, the *Inflation Country* dataset contains a list of countries with the inflation rate for 2011²³.
- The *Scientific Journals Country* dataset contains a list of countries with a number of scientific and technical journal articles published in 2011²⁴.
- The *Unemployment French Region* dataset contains a list of regions in France with the unemployment rate, used in the SemStats 2013 challenge [10].

Again, for those datasets, the instances (cities, countries, etc.) are linked to DBpedia. For datasets which are used for classification and regression, the regression target was discretized using equal frequency binning, usually into a *high* and a *low* class.

The third, and final, category contains datasets that were generated from existing RDF datasets, where the value of a certain property is used as a classification target. There are five datasets in this category:

- The *Drug-Food Interaction* dataset contains a list of drug-recipe pairs and their interaction, i.e., “negative” and “neutral” [6]. The dataset was retrieved from FinkilOD²⁵. Furthermore, each drug is linked to DrugBank²⁶. We drew a stratified random sample of 2,000 instances from the complete dataset. When generating the features, we ignore the foodInteraction property in DrugBank, since it highly correlates with the target variable.
- The *AIFB* dataset describes the AIFB research institute in terms of its staff, research group, and publications. In [1] the dataset was first used to predict the affiliation (i.e., research group) for people in the dataset. The dataset contains 178 members of a research group, however the smallest group contains only 4 people, which is removed from the dataset, leaving 4 classes. Also, we remove the employs relation, which is the inverse of the *affiliation* relation.
- The *AM* dataset contains information about artifacts in the Amsterdam Museum [2]. Each artifact in the

dataset is linked to other artifacts and details about its production, material, and content. It also has an artifact category, which serves as a prediction target. We have drawn a stratified random sample of 1,000 instances from the complete dataset. We also removed the material relation, since it highly correlates with the artifact category.

– The *MUTAG* dataset is distributed as an example dataset for the DL-Learner toolkit²⁷. It contains information about complex molecules that are potentially carcinogenic, which is given by the isMutagenic property.

– The *BGS* dataset was created by the British Geological Survey and describes geological measurements in Great Britain²⁸. It was used in [17] to predict the lithogenesis property of named rock units. The dataset contains 146 named rock units with a lithogenesis, from which we use the two largest classes.

An overview of the datasets is given in Tables 1, 2, and 3. For each dataset, we depict the number of instances, the machine learning tasks in which the dataset is used (C stands for classification and R stands for regression), the source of the dataset, and the LOD datasets to which the dataset is linked. For each dataset, we depict basic statistics of the properties of the LOD datasets, i.e., average, median, maximum and minimum number of types, categories, outgoing relations (rel out), incoming relations (rel in), outgoing relations including values (rel-vals out) and incoming relations including values (rel-vals in). The datasets, as well as a detailed description, a link quality evaluation, and licensing information, can be found online²⁹.

From the given statistics, we can infer the following observations: (i) DBpedia contains significantly less owl:sameAs links to YAGO, compared to Wikidata;

(ii) DBpedia provides the highest number of types and categories on average per entity; (iii) Wikidata contains the highest number of outgoing and incoming relations for most of the datasets; (iv) YAGO contains the highest number of outgoing and incoming relations values for most of the datasets.

Table 1. Datasets statistics

Dataset					types				categories				rel out				rel in				rel-vals out				rel-vals in			
Name	Source	Task	LOD	#links	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min
Auto MPG	UCI ML	R	DBpedia	371	29.70	31	46	5	11.20	10	25	2	13.48	13	27	3	5.62	5	25	1	16.50	15	70	0	36.65	23	509	0
			YAGO	331	13.99	16	21	0	9.26	9	23	0	8.76	9	18	0	16.96	2	138	1	77.08	70	278	0	3,236.24	60	28,418	0
			Wikidata	371	1.05	1	3	0	0.29	0	3	0	20.20	18	61	9	5.32	5	31	1	13.92	12	54	4	59.33	21	755	3
AAUP	JSE	R/C (c=3)	DBpedia	960	24.40	28	41	0	9.38	9	20	0	12.68	15	28	0	8.20	7	36	0	11.74	11	66	0	62.18	23	2,488	0
			YAGO	889	10.49	11	17	0	3.31	3	11	0	11.37	12	15	0	13.61	3	138	1	85.83	68	446	0	2,455.27	110	28,418	1
			Wikidata	959	2.13	2	5	0	0.88	1	2	0	30.71	29	83	0	8.51	7	44	0	22.38	21	97	0	296.92	20	31,777	0
Auto 93	JSE	R	DBpedia	93	28.76	31	43	5	11.13	10	25	3	12.69	12	22	8	4.92	5	72	1	14.35	11	64	4	22.60	18	642	0
			YAGO	80	13.80	16	19	0	9.09	10	18	0	8.37	10	11	0	21.09	2	138	2	59.33	59	129	0	4,025.90	46	28,418	4
			Wikidata	93	1.00	1	2	0	0.12	0	1	0	17.31	17	26	9	3.56	3	81	1	11.23	11	25	4	19.91	19	57	3
Zoo	UCI ML	C (c=3)	DBpedia	101	8.61	11	26	0	4.67	3	34	0	8.22	9	15	3	3.54	3	81	1	13.26	11	87	1	146.28	24	3,686	2
			YAGO	8	0.74	0	13	0	0.15	0	6	0	0.63	1	8	0	127.23	138	138	2	5.39	1	156	0	26,173.23	28,418	28,418	3
			Wikidata	101	1.00	1	2	0	0.67	1	2	0	29.69	35	57	3	8.28	7	27	0	18.20	21	45	1	125.82	92	785	0
Forbes	Forbes	R/C (c=2)	DBpedia	1,585	14.77	19	62	0	4.87	4	52	0	10.15	11	27	0	2.76	2	27	0	10.44	10	136	0	14.30	4	1,925	0
			YAGO	1,003	7.28	10	33	0	2.35	2	42	0	7.57	11	21	0	52.07	2	138	1	34.42	27	510	0	10,531.37	107	28,418	1
			Wikidata	1,189	0.82	1	4	0	0.22	0	3	0	16.59	16	137	0	5.00	5	52	0	12.69	10	207	0	30.14	107	28,418	1
Cities	Mercer	R/C (c=3)	DBpedia	212	31.28	35	53	0	6.98	7	26	0	18.08	19	38	0	25.66	25	68	0	16.26	13	131	0	1,474.57	678	19,810	0
			YAGO	187	16.66	19	30	0	4.46	4	15	0	13.75	15	32	0	23.56	9	138	2	222.54	214	681	0	8,087.34	3,555	72,320	5
			Wikidata	212	2.11	2	9	0	1.34	4	6	0	69.08	67	153	6	39.99	37	108	1	105.29	89	390	2	5,298.23	1,599	99,865	1
FB Books	Facebook	R/C (c=2)	DBpedia	1,600	19.08	20	42	0	5.15	5	23	0	11.15	11	20	0	1.64	2	7	0	7.04	7	60	0	2.80	2	420	0
			YAGO	1,334	8.37	10	24	0	2.03	2	15	0	8.41	10	13	0	25.32	3	138	1	24.37	22	149	0	4,735.50	8	28,418	1
			Wikidata	1,578	1.00	1	3	0	0.01	0	1	0	21.19	22	55	0	3.15	3	17	0	16.41	16	69	0	7.47	4	165	0
FB Movies	Facebook	R/C (c=2)	DBpedia	1,600	24.90	27	55	0	12.50	11	60	0	12.43	13	21	0	1.46	1	12	0	11.65	12	51	0	4.96	2	110	0
			YAGO	1,339	12.08	14	32	0	6.51	6	27	0	8.39	10	17	0	26.89	6	138	1	55.01	47	280	0	4,682.42	43	28,418	1
			Wikidata	1,585	1.01	1	4	0	0.04	0	1	0	48.75	48	107	0	2.22	1	22	0	56.37	53	372	0	20.75	12	230	0
Metacritic Albums	Metacritic	R/C (c=2)	DBpedia	1,600	17.92	19	36	0	4.27	4	26	0	10.85	12	17	2	2.63	3	7	0	8.92	9	63	0	5.28	3	50	0
			YAGO	1,444	7.22	8	19	0	3.22	3	20	0	8.05	9	10	0	16.02	3	138	1	40.27	32	361	0	2,749.90	10	28,418	1
			Wikidata	1,576	0.99	1	2	0	0.00	0	1	0	17.64	18	45	0	4.00	5	9	1	11.73	12	49	0	8.77	7	54	1

Table 2. Datasets statistics

Dataset					types				categories				rel out				rel in				rel-vals out				rel-vals in			
Name	Source	Task	LOD	#links	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min
Metacritic Movies	Metacritic	R/C (c=2)	DBpedia	2,000	24.38	27	45	0	11.87	11	42	0	12.54	14	19	3	1.35	1	7	0	11.42	12	30	0	3.56	2	31	0
			YAGO	1,588	11.79	14	19	0	6.43	6	28	0	8.34	10	11	0	0.28	6	138	1	148.84	43	216	0	4,960.84	37	28,418	1
			Wikidata	1,981	0.98	1	1	0	0.03	0	1	0	47.86	49	99	0	0.198	1	13	0	52.70	53	237	0	15.77	11	117	0
HIV Deaths Country	WHO	R/C (c=2)	DBpedia	114	35.69	37	52	0	12.61	13	23	0	23.59	24	28	3	34.26	31	89	6	27.75	25	162	10	4,828.36	1,065	70,426	24
			YAGO	108	13.90	15	24	0	9.28	9	18	0	0.28	31	35	0	0.15	9	138	1	5,302.34	244	1,267	0	12,464.42	4,879	112,032	550
			Wikidata	114	4.12	4	8	1	1.483	5	6	0	120.87	119	173	7	55.68	51	148	2	229.46	210	595	2	45,671.15	4,971	669,273	66
Traffic Accidents Country	WHO	R/C (c=2)	DBpedia	146	36.40	38	53	0	13.12	13	23	0	23.40	24	28	1	137.87	36	94	8	27.44	24	162	0	7,528.18	1,587	218,957	77
			YAGO	139	14.29	15	27	0	9.62	10	16	0	0.28	31	35	0	0.14	9	138	1	5,345.03	290	2,104	0	17,882.47	6,126	423,559	693
			Wikidata	146	4.42	4	10	1	4.94	5	6	0	124.31	121	191	7	61.68	55	148	2	242.38	213	713	2	85,575.10	7,369	1,557,157	66
Energy Savings Country	WorldBank	R/C (c=2)	DBpedia	162	36.07	38	53	0	13.12	13	23	0	23.46	24	28	1	136.74	33	94	8	26.72	23	162	0	6,876.80	1,440	218,957	77
			YAGO	152	14.09	15	27	0	9.52	10	16	0	0.27	31	35	0	0.16	9	138	1	5,329.28	279	2,104	0	16,969.96	5,821	423,559	757
			Wikidata	162	4.41	4	10	1	4.92	5	6	0	123.36	119	191	7	60.02	55	148	2	238.69	210	713	2	77,485.01	5,810	1,557,157	66
Inflation Country	WorldBank	R/C (c=2)	DBpedia	160	36.00	38	53	0	13.11	13	23	0	23.46	24	28	1	136.74	33	94	8	26.85	24	162	0	6,947.59	1,440	218,957	77
			YAGO	150	14.09	15	27	0	9.44	10	16	0	0.27	31	35	0	0.16	9	138	1	5,331.16	279	2,104	0	17,114.88	5,821	423,559	693
			Wikidata	160	4.39	4	10	1	4.88	5	6	0	123.23	119	191	7	60.12	55	148	2	237.94	210	713	2	78,453.16	5,810	1,557,157	66
Scientific Journals Country	WorldBank	R/C (c=2)	DBpedia	160	36.00	38	53	0	13.11	13	23	0	23.46	24	28	1	136.74	33	94	8	26.85	24	162	0	6,947.59	1,440	218,957	77
			YAGO	150	14.09	15	27	0	9.44	10	16	0	0.27	31	35	0	0.16	9	138	1	5,331.16	279	2,104	0	17,114.88	5,821	423,559	693
			Wikidata	160	4.39	4	10	1	4.88	5	6	0	123.23	119	191	7	60.12	55	148	2	237.94	210	713	2	78,453.16	5,810	1,557,157	66

Unemployment French Region	SemStats	R/C (c=2)	DBpedia YAGO Wikidata	26 26 26	16.38 8.92 1.35	21 8 1	32 14 3	0.73 8.277 1.258	3 2 3	15 8 4	0.781 12.42 186.23	9 12 84	10 14 119	3 14.19 74.34.00	13 4 33	24 6 51	7.7.19 3.81.12 21.83.12	7.19 60.299 79.157	1 28 58	975.88 1,793.19 332.69	969 1,424 193	2,292 4,527 1,464	37 88 137
Endangered Species	a-z-animals	R/C (c=2)	DBpedia YAGO Wikidata	301 65 301	11.84 2.48 1.05	12 0 1	33 16 4	0.6.32 0.0.76 0.0.44	5 0 0	34 12 6	0.10.77 0.1.78 0.34.32	11 0 37	25 9 137	0 0 3	2.96 108.62 6.94	3 138 678	55 1.9.53 0.22.04	11.87 0.136 22.400	0 0 1	566.25 22,286.36 21,909.94	15 28,418 70	114,742 28,418 6,460,930	1 1 0
Drug-Food Interaction	FinkiLO D	C (c=2)	DBpedia YAGO Wikidata DrugBan k FinkiLO D	1,989 588 1,908 2,000 2,000	8.83 4.46 1.96 2.00 1.00	4 0 2 2 1	38 31 3 2 1	0.5.46 0.0.68 0.0.01 2 1	5 0 0 0 1	18 6 1 1 1	0.12.65 0.2.15 0.45.92 61.68 3.00	14 0 47 64 3	15 8 79 71 3	0.1.40 0.99.69 0.2.78 41.1.70 3.0.00	1 138 2 2 0	5 138 17 2 0	0.8.63 1.7.28 1.34.99 0.41.96 0.1.00	3 0 27 41 1	12.0 61.0 159.0 132.14 1.1	34.71 20,427.08 32.25 62.49 0.00	24 28,418 26 50 0	158.0 28,418 487.4 211.0 0.0	0 1 0 0 0

Dataset			types				rel out				rel in				rel-val out				rel-val in			
Name	Task	#links	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min	avg	med	max	min
AIFB	C (c=4)	176	1.4	1	2	1	7.1	7	9	5	2.0	2	5	0	18.2	7	219	2	19.8	9	246	0
AM	C (c=11)	1,000	1.0	1	1	1	19.8	20	29	9	0.6	1	3	0	21.9	20	283	7	3.2	1	273	0
MUTAG	C (c=2)	340	1.0	1	1	1	9.8	10	14	5	1	1	1	0	65.8	56	465	4	1	1	1	1
BGS	C (c=2)	146	1.0	1	1	1	29.7	31	36	21	1.4	2	4	0	25.2	24	54	15	2.7	2	12	0

4 Conclusion and Outlook

In this paper, we have introduced a collection of 22 benchmark datasets for machine learning on the Semantic Web. So far, we have concentrated on classification and regression tasks. There are methods to derive clustering and outlier detection benchmarks from classification and regression datasets [4,5], so that extending the dataset collection for such unsupervised tasks is possible as well. Furthermore, as many datasets on the Semantic Web use extensive hierarchies in the form of ontologies, building benchmark datasets for tasks like *hierarchical multi-label classification* [15] would also be an interesting extension.

Acknowledgments. The work presented in this paper has been partly funded by the German Research Foundation (DFG) under grant number PA 2373/1-1 (Mine@LOD), and the Dutch national program COMMIT.

References

1. Bloehdorn, S., Sure, Y.: Kernel methods for mining instance data in ontologies. In: Aberer, K., et al. (eds.) ASWC/ISWC -2007. LNCS, vol. 4825, pp. 58–71. Springer, Heidelberg (2007). doi:10.1007/978-3-540-76298-0 5
2. Boer, V., Wielemaker, J., Gent, J., Hildebrand, M., Isaac, A., Ossenbruggen, J., Schreiber, G.: Supporting linked data production for cultural heritage institutes: the Amsterdam museum case study. In: Simperl, E., Cimiano, P., Polleres, A., Corcho, O., Presutti, V. (eds.) ESWC 2012. LNCS, vol. 7295, pp. 733–747. Springer, Heidelberg (2012). doi:10.1007/978-3-642-30284-8 56
3. Demšar, J.: Statistical comparisons of classifiers over multiple data sets. J. Mach. Learn. Res. 7, 1–30 (2006)
4. Emmott, A.F., Das, S., Dietterich, T., Fern, A., Wong, W.K.: Systematic construction of anomaly detection benchmarks from real data. In: Proceedings of the ACM SIGKDD Workshop on Outlier Detection and Description, pp. 16–21. ACM (2013)
5. Farber, I., Günnemann, S., Kriegel, H.P., Kröger, P., Müller, E., Schubert, E., Seidl, T., Zimek, A.: On using class-labels in evaluation of clusterings. In: MultiClust: Workshop on Discovering, Summarizing and Using Multiple Clusterings (2010)
6. Jovanovik, M., Bogojeska, A., Trajanov, D., Kocarev, L.: Inferring cuisine-drug interactions using the linked data approach. Scientific reports 5 (2015)
7. Paulheim, H.: Generating possible interpretations for statistics from linked open data. In: Simperl, E., Cimiano, P., Polleres, A., Corcho, O., Presutti, V. (eds.) ESWC 2012. LNCS, vol. 7295, pp. 560–574. Springer, Heidelberg (2012). doi:10.1007/978-3-642-30284-8 44
8. Rettinger, A., Lösch, U., Tresp, V., d’Amato, C., Fanizzi, N.: Mining the semantic web. Data Min. Knowl. Disc. 24(3), 613–662 (2012)
9. Ristoski, P., Bizer, C., Paulheim, H.: Mining the web of linked data with rapid-miner. Web Semant. Sci. Serv. Agents WWW 35, 142–151 (2015)
10. Ristoski, P., Paulheim, H.: Analyzing statistics with background knowledge from linked open data. In: Workshop on Semantic Statistics (2013)
11. Ristoski, P., Paulheim, H.: Semantic web in data mining and knowledge discovery: a comprehensive survey. Web Semant. 36, 1–22 (2016)
12. Ristoski, P., Paulheim, H., Svátek, V., Zeman, V.: The linked data mining challenge 2015. In: KNOW@ LOD (2015)
13. Ristoski, P., Paulheim, H., Svátek, V., Zeman, V.: The linked data mining challenge 2016. In: KNOW@ LOD (2016)
14. Schmachtenberg, M., Bizer, C., Paulheim, H.: Adoption of the linked data best practices in different topical domains. In: Mika, P., et al. (eds.) ISWC 2014. LNCS, vol. 8796, pp. 245–260. Springer, Heidelberg (2014). doi:10.1007/978-3-319-11964-9 16
15. Silla Jr., C.N., Freitas, A.A.: A survey of hierarchical classification across different application domains.

Data Min. Knowl. Disc. **22**(1), 31–72 (2011)

16. Tresp, V., Bundschuh, M., Rettinger, A., Huang, Y.: Towards machine learning on the semantic web. In: da Costa, P.C.G., et al. (eds.) URSW 2005-2007. LNCS, vol. 5327, pp. 282–314. Springer, Heidelberg (2008)