

Veracity assessment of online data

Marianela García Lozano^{a,b,*}, Joel Brynielsson^{a,b}, Ulrik Franke^c, Magnus Rosell^a,
Edward Tjörnhamar^{a,b}, Stefan Varga^{b,d}, Vladimir Vlassov^b

^a FOI Swedish Defence Research Agency, SE-164 90 Stockholm, Sweden

^b KTH Royal Institute of Technology, SE-100 44 Stockholm, Sweden

^c RISE Research Institutes of Sweden, P.O. Box 1263, SE-164 29 Kista, Sweden

^d Swedish Armed Forces Headquarters, SE-107 85 Stockholm, Sweden

ARTICLE INFO

Keywords:

Veracity assessment
Credibility
Data quality
Online data
Social media
Fake news

ABSTRACT

Fake news, malicious rumors, fabricated reviews, generated images and videos, are today spread at an unprecedented rate, making the task of manually assessing data veracity for decision-making purposes a daunting task. Hence, it is urgent to explore possibilities to perform automatic veracity assessment. In this work we review the literature in search for methods and techniques representing state of the art with regard to computerized veracity assessment. We study what others have done within the area of veracity assessment, especially targeted towards social media and open source data, to understand research trends and determine needs for future research.

The most common veracity assessment method among the studied set of papers is to perform text analysis using supervised learning. Regarding methods for machine learning much has happened in the last couple of years related to the advancements made in deep learning. However, very few papers make use of these advancements. Also, the papers in general tend to have a narrow scope, as they focus on solving a small task with only one type of data from one main source. The overall veracity assessment problem is complex, requiring a combination of data sources, data types, indicators, and methods. Only a few papers take on such a broad scope, thus, demonstrating the relative immaturity of the veracity assessment domain.

1. Introduction

As the internet has become a significant source of information for many, the need to assess the veracity of statements to, e.g., identify the spreading of false information, is apparent. Since individuals, companies, organizations, etc.—i.e., almost anyone—can write and post anything on the web, the information is often incomplete, ambiguous, contradicting, biased, or wrong. Further, due to the large amounts of heterogeneous information and the velocity with which it is created, it quickly becomes unfeasible to manually assess its veracity. A decision support system is only as good as its underlying data. The question of data veracity especially comes to mind whenever data retrieved from social media and other open sources is utilized. Hence, automatic, i.e., computerized, methods and tools capable of processing and assessing large amounts of data are needed.

The terms veracity and veracity assessment deserve a few words of introduction. The concept of veracity was introduced and became widely used among computer scientists after it, in 2012, was proposed

as the fourth “V” [18, 91, 96] (the other ones being Volume, Variety, and Velocity) of big data [57]. In a blog post Snow [96] argues that trusted data needed to be defined separately in the era of big data, with its generally easy access to large volumes of heterogeneous data. Snow [96] states that “I believe that the definition of trusted data depends on the way you are using the data and applying it to your business.” Furthermore, veracity is presented as a concept that “deals with uncertain or imprecise data” which is an important property to take into account when data is analyzed and ultimately used for decision-making.

In a cursory overview of different veracity definitions in dictionaries, see Fig. 1, one can observe proposals in which aspects of accuracy, credibility, truthfulness and quality can be used to delimit the term. These aspects represent several different but equally valid views of veracity that are interrelated. Hence, we note that it is hard to define such a broad term in a succinct manner.

In the big data domain, data scientists and researchers have tried to give more precise descriptions and/or definitions of the veracity

* Corresponding author.

E-mail addresses: mgl@kth.se (M. García Lozano), joel@kth.se (J. Brynielsson), ulrik.franke@ri.se (U. Franke), magnus.rosell@foi.se (M. Rosell), edwardt@kth.se (E. Tjörnhamar), svarga@kth.se (S. Varga), vladv@kth.se (V. Vlassov).

<https://doi.org/10.1016/j.dss.2019.113132>

Received 6 May 2019; Received in revised form 16 July 2019; Accepted 16 August 2019

Available online 18 October 2019

0167-9236/ © 2019 The Authors. Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

veracity	
- <i>the quality of being true, honest, or accurate</i>	https://dictionary.cambridge.org
- <i>i) conformity with truth or fact; accuracy, ii) devotion to the truth; truthfulness, iii) power of conveying or perceiving truth, iv) something true.</i>	https://www.merriam-webster.com
- <i>i) habitual observance of truth in speech or statement; truthfulness, ii) conformity to truth or fact; accuracy, iii) correctness or accuracy, as of the senses or of a scientific instrument.</i>	https://www.dictionary.com
- <i>unwillingness to tell lies</i>	https://www.vocabulary.com

Fig. 1. English online dictionary definitions of veracity as of November 2018.

concept. Some proposals are in line with the dictionary definitions of Fig. 1, while others take an approach of using corresponding negated terms, or both. An IBM report from 2012 describes veracity as “data uncertainty,” referring to the ability of “managing the reliability and predictability of inherently imprecise data types” [91]. Another IBM report from the same year, states that veracity has to do with managing “data in doubt” and relates it to “uncertainty due to data inconsistency, incompleteness, ambiguities, and deception” [18]. In the corresponding presentation the author also gives what could be interpreted as a definition of veracity, i.e., “truthfulness, accuracy or precision, correctness.” There are many more examples; Lukoianova and Rubin [68] propose a veracity framework with three main veracity dimensions outlined by “objectivity, truthfulness, credibility and their opposites,” and Ramachandramurthy et al. [83] state that veracity “focuses on Information Quality (IQ).” It is worth pointing out that many of the veracity aspects such as data quality, truth, credibility, and trustfulness assessment, were not new and had been addressed by researchers in related settings, e.g., decision support and information systems, before the big data inclusion in 2012 [2, 28, 70, 77, 104].

Another related veracity assessment concept is the indicator; an indicator is a predefined phenomenon of interest that may, or may not, be present in the data. The occurrences of one or several indicators can be used to facilitate the veracity assessment process. The indicators may also affect the confidence in an assessment positively or negatively. Indicators can also in themselves be assessed with regard to veracity. Whether a specific approach targets a single indicator or solves the whole veracity assessment problem is in many cases context-dependent: determining user credibility can for example both be a purpose in itself and be thought of as a veracity assessment indicator.

In sum, despite many researchers’ efforts, we assert that there is no prevalent generally agreed upon definition of veracity in academia. In this work we refrain from adding yet another definition, but simply use a list of terms that are often mentioned in conjunction with veracity related to big data. They include truth, trust, uncertainty, credibility, reliability, noisy, anomalous, imprecise, and quality. As will be discussed in Section 2, such terms were used in the search strings employed in the literature study presented herein. We have also chosen to include studies of approaches, methods and algorithms related to indicators that may help with veracity assessment of data.

1.1. Purpose and problem statement

The purpose of this work is to review the approaches, methods, algorithms, and tools which are used or proposed by the research community for automatic *veracity assessment* (VA) of open source data¹, thereby obtaining a view of the state of the art for this domain. By open source data we refer to information published online such as social media posts, blog posts, forum entries, newspaper articles, whether on the shallow or the deep web. Hence, the research question to be studied is the following:

- Which approaches, methods, algorithms, and tools are used or proposed for automatic veracity assessment of open source data?

1.2. Outline

The remainder of this paper is structured as follows. Section 2 describes the chosen methodology and our choices in the execution of it. Section 3 contains the results of the undertaken systematic literature survey. This is followed by a synthesis and gap analysis discussion of the obtained results in Section 4, while the last section sums up the work and presents conclusions.

2. Review methodology

We address the research question in Section 1 by conducting a *systematic literature review* (SLR). An SLR aims to study scientific literature in an unbiased and reproducible way, aiming to find all existing works that fit the set criteria. Reasons for including or excluding studies are explicitly stated and agreed upon before searching for relevant studies. In this SLR the guidelines proposed by Kitchenham and Charters [52] were followed, which are briefly outlined in the following section.

2.1. SLR methodology

The first step in a systematic literature review is to formulate a main

¹ In this paper the concept of open source data is used as a type of analogy for the concept of open source intelligence (OSINT).

research question together with inclusion and exclusion criteria. The research question should embrace the purpose of the review and the inclusion and exclusion criteria help focus the scope of the research that is included in the survey. To reliably assess papers in a consistent manner, a review protocol along with instructions to the reviewers is created.

The next step is to design and plan the search strategy in the form of key terms combined into suitable search strings which are applied to relevant databases. This is preferably done with the aid of a professional librarian.

Once results have been gathered, the inclusion and exclusion criteria are applied, filtering and narrowing down the final set of papers to review. This is an iterative process, starting by looking only at the title, keywords and authors, then reading the abstract, and in the final iteration reading the full text to decide on whether to include the paper.

In Sections 2.2–2.5 the application of the SLR methodology in the present study is described.

2.2. Search strategy

Based on the purpose and objectives of the survey and our previous knowledge of research within the domain, the following set of keywords was used as a basis for a search conducted by a professional librarian: veracity, credibility, assessment, social media, open source data, rumors, and fake news.

We expanded the set of keywords with related terms and synonyms which were primarily gathered by analysis of the dictionary definitions (see Fig. 1) and Google searches. With the aid of expertise within library search methodology and online libraries we finally obtained a set of search questions, see Table 1, which were then applied to the list of online databases, see Fig. 2. The list of search strings is not an exhaustive list of keyword combinations since that would only inflate the amount of results without adding much to the findings. The search string list is rather based on a trial and error process looking for coverage and relevance.

The number of hits each database generated using the search strings can be seen in Fig. 2. Note that not all of the inclusion and exclusion criteria have been applied at this stage. The total number of hits was 5047. Since i) some papers are indexed by multiple databases, and ii) different search strings sometimes triggered the same papers, this number contains duplicates.

Table 1
Search strings.

Nr	Search string
1	"assessment*" AND ("credibility" OR "veracity") AND "social media"
2	"assessment*" AND "credibility" AND ("fake news" OR "misinformation")
3	"assessment*" AND "fake news" AND "open source data"
4	"assessment*" AND ("lie*" OR "truth*") AND "social media"
5	"assessment*" AND ("rumor*" OR "rumour*")
6	"assessment*" AND "open source data" AND "social media"
7	"credibility" AND ("facebook" OR "twitter") AND ("fake news" OR "misinformation")
8	"credibility" AND "social media" AND ("fake news" OR "misinformation")
9	("credibility" OR "veracity") AND "open source data"
10	"fake news" AND ("misinformation*" OR "reliability" OR "truth*") AND "social media"
11	"misinformation" AND ("instagram" OR "snapchat")
12	"open source data" AND "social media"
13	"validate*" AND ("facebook" OR "twitter") AND "fake news"
14	"validate*" AND ("facebook" OR "twitter") AND "misinformation"
15	"veracity" AND ("facebook" OR "fake news" OR "twitter" OR "misinformation")
16	"veracity" AND ("rumor*" OR "rumour*") AND "social media"
17	"veracity" AND ("lie*" OR "truth*")

2.3. Inclusion and exclusion criteria

The list of inclusion and exclusion criteria used to filter the search results are:

1. Only papers related to automatic/computerized approaches, methods, algorithms and tools are included.
2. Only papers using or discussing open source data are included.
3. Only assessment studies with the purpose of assessing veracity or some related aspect are included. That is, research related to, e.g., the veracity phenomenon as such is excluded.
4. Only research published between 2013 and 2017 is included, which provides a cut-off criterion and at the same time provides a recent view of the methodologies in use.
5. Research published in any other language than English (British and American spelling) is excluded, i.e., only work available to the wider research audience is included.

Due to the automatic filtering possibilities inherent in the databases, criteria number four and five were used in the searches. Some non-English results were still obtained, but eliminated later in the process.

2.4. Study selection process

The database searches were conducted in February 2018, and in the subsequent months the filtering and reviewing process took place. The literature selection consisted of an iterative funnel-like process, see Fig. 2, where the search results were screened based on the SLR method according to Section 2.1. With the inclusion and exclusion criteria at hand, filtering of the data base results was done based on title, keywords, and authors, resulting in a total of 346 papers left. In the following iteration we read the abstracts and were able to remove 159 more papers, leaving us with 187 papers. The last filtering iteration was based on a cursory glance at the full papers, resulting in a set of 112 papers. This final set of papers were read and reviewed in full. On closer inspection, however, five more papers were removed from the final set due to non-complacency with the inclusion-exclusion criteria.

2.5. Review protocol and objectives

Based on the main research question and purpose of the study, a review protocol which was used by all reviewers to analyze the chosen papers was developed (see Appendix A for the full protocol). The research question part of the review protocol is further divided into six groups: approaches, methods, algorithms, tools, data, and miscellaneous questions about issues indirectly related to the main research question. This division was done to provide a good basis for analysis of the papers and synthesis of the results.

3. Results

Out of the 107 papers that were reviewed there is a clear trend in the publication year. The majority of the papers, i.e., 65%, are published in the last two years of the explored time range, i.e., between 2016 and 2017, as can be seen in Fig. 3a. The majority of the papers are also published by authors affiliated with a university or institute, and about a quarter of the papers have a mixture of affiliations, e.g., university with company or university and institute, see Fig. 3b. The one country which most authors have as affiliation is USA with participation in 37 of the publications, see Fig. 3c. China, which was the runner up country, has representation in 15 papers. Grouping the country affiliation in geographical regions, the most productive region is Europe with participation in 50 of the publications, i.e., almost half of the examined papers, see Fig. 3d.

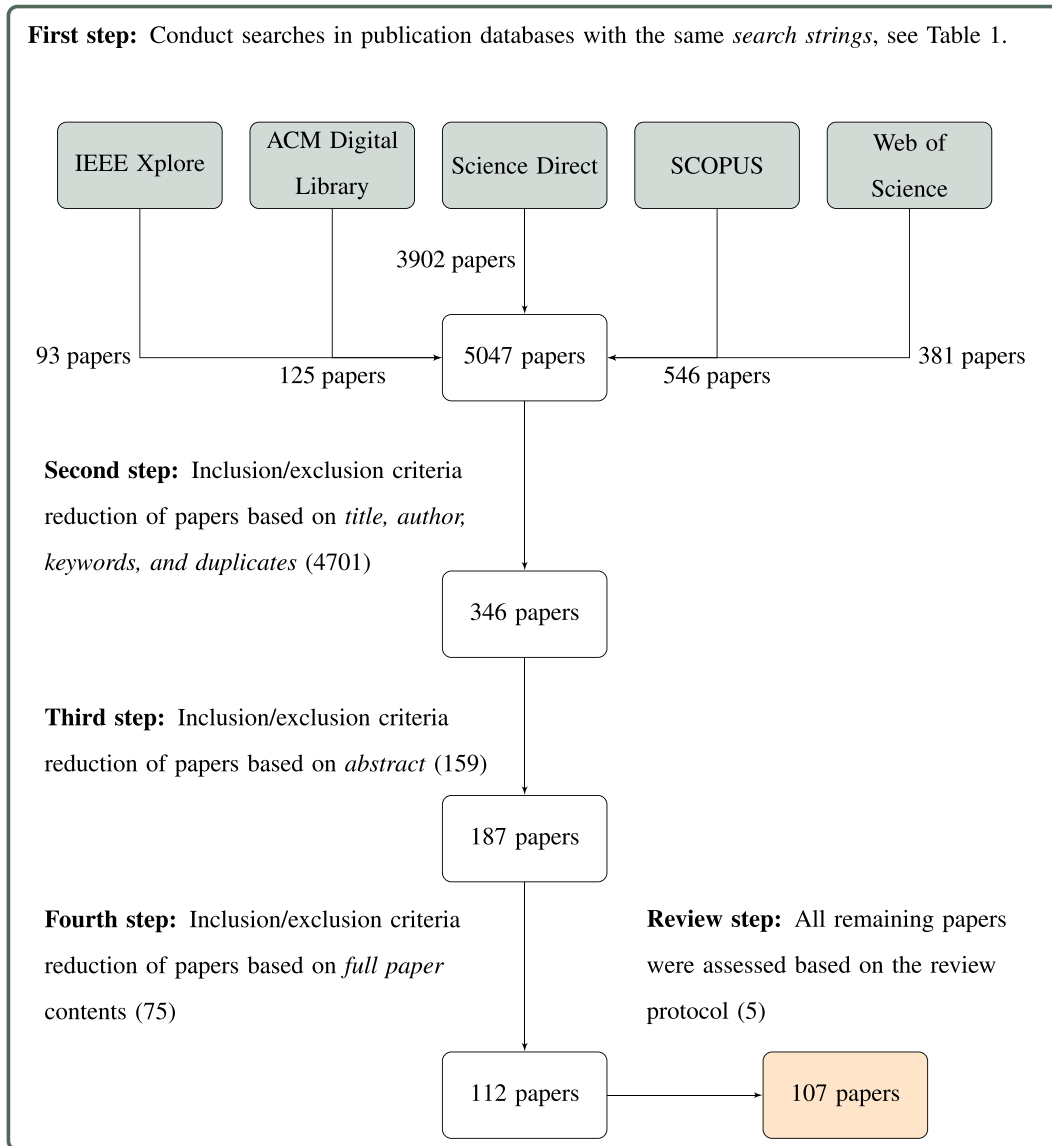


Fig. 2. Paper selection and review process. Numbers in parenthesis are number of papers removed from previous step.

3.1. Approaches

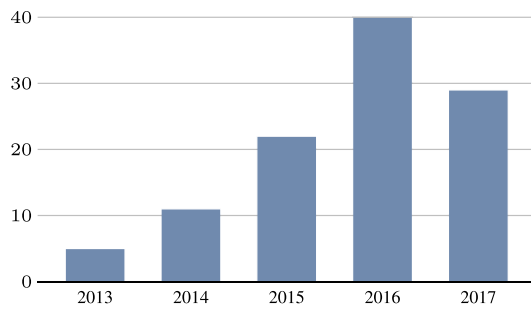
As discussed in Section 1, veracity assessment often, but not always, makes use of indicators as a basis for making the end assessment, and in many cases it is context-dependent whether a specific approach targets a single indicator or solves the whole veracity assessment problem: determining user credibility can for example both be a purpose in itself and be thought of as an indicator [1]. With this precaution in mind it is still interesting to note that the investigated papers can be roughly divided into two broad equally sized categories dependent on their focus: about half of the papers set out to perform the veracity assessment directly [5, 37, 48, 62, 84], while the other half of the papers have a clear focus on the indicators (in themselves or as a means of performing the overarching veracity assessment) [3, 6, 56, 61, 88].

Two main indicator “dimensions” can be discerned. The first indicator dimension is related to the data origin, with indicators derived from, i) message content, ii) meta data, and iii) external sources. Focusing on approaches used for solely looking at the actual message content itself, indicators for sentiment/affect and opinion/stance are the predominant ones [25, 36]. The ingenuity when it comes to the aforementioned external sources is large, including, e.g., crowd

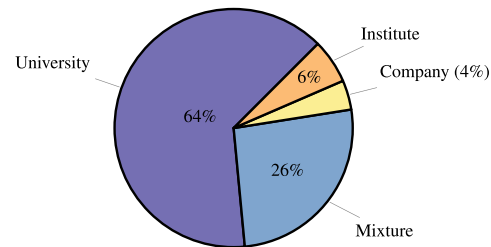
sourcing (letting own users tag the tweet) [89], and could give rise to further division into several dimensions of external sources.

The other indicator dimension can be related to some underlying modeling aspect where the algorithm developer starts with an idea of some aspect that can be used for veracity assessment and tries to model this aspect to confirm or disprove the veracity. One example of this modeling dimension is coordinated behavior where, e.g., many users exhibiting similar behavior could be used as an indicator [1]. The indicators used typically relate to the assumptions being made regarding the intended end user application, e.g., the availability of databases for verifying claims [82, 85, 90, 125], whether additional messages can be used for comparison, etc. As a consequence, many interesting examples of special cases that can only be used in a specific context exist. For example, facts related to soccer games can be used to make good assessments specifically related to soccer claims [44], and meta data concerning geographical positioning can be combined with knowledge regarding traffic patterns [23] and points of interest [6] to improve the veracity assessment in infrastructure contexts.

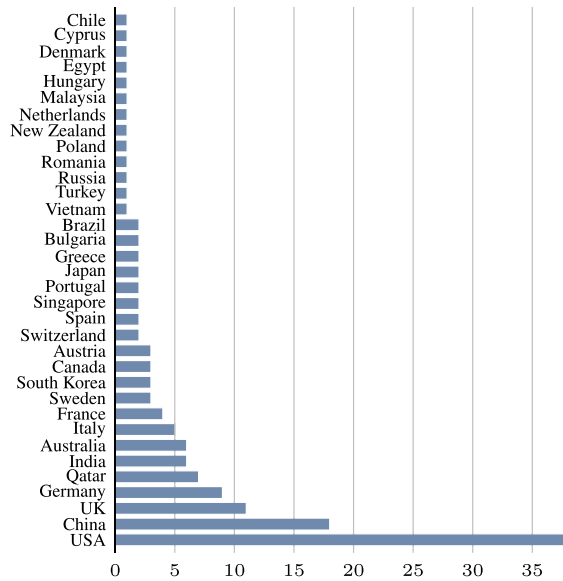
In the papers/approaches making use of indicators, the motivations for the choice of indicators can be divided into three about equally sized categories: i) related work is used to motivate the indicator(s), e.g., [16,



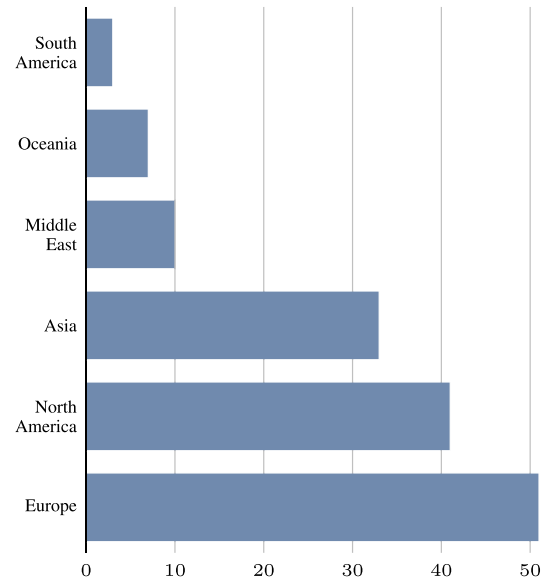
(a) Publication year.



(b) Affiliation.



(c) Countries.



(d) Geographical region.

Fig. 3. Reviewed paper statistics.

58, 122], ii) a convincing argument based on intuition is provided, e.g., [1, 14, 85], or iii) the paper serves in itself to investigate and motivate the indicator(s) used, e.g., [6, 64, 78].

Concerning quantification and presentation of the veracity assessment result, the typical veracity assessment approach calculates a probability to be used for presenting some kind of discrete result depending on the application at hand. In most cases a binary yes/no answer is calculated, e.g., [42, 118, 126], but in some cases the probability measure is used for more fine-grained quantification on a scale [35] and sometimes there are more than two classes to be distinguished between [37, 62]. The few exceptions that stand out include cases where the algorithm design necessitates alternative quantification methods where, e.g., a relative score is calculated and used for ranking different alternatives [78], and cases with alternative means of presentation using, e.g., heat maps [127].

All but a few papers present some kind of a more or less scientific evaluation of the result. Depending on the foreseen application and focus, these evaluation efforts typically target i) the invented method and/or algorithm, ii) the assessment itself, iii) the data, and iv) the end user application. Although much related to the application and focus, it is still interesting to note that roughly two thirds of the evaluation efforts relate to the presented method/algorithm [42, 46, 93], while the rest of the evaluations, i.e., one third, are directed towards the veracity assessment itself [45, 95, 125]. Some papers include evaluations of several aspects, a handful of the papers evaluate the data [17, 25, 40,

62, 92, 126], and yet a few papers include evaluations related to the envisioned end user application where things such as tool usability is included [35].

3.2. Methods

This section seeks to epitomize the methods used for automated veracity assessments. The results obtained in the literature review revealed that not all articles actually describe a complete process for this. There are examples of vague or imprecise research questions, and articles where only parts of the process are addressed, e.g., the algorithms or the process of calculating some score. Others describe inventions or methods to create training data. There are also other literature reviews. In sum, it is difficult to present general characteristics describing the most common methods due to this diversity. Out of the articles that describe semi-automatic and automatic veracity assessment procedures in some detail, however, which constitute the majority of the articles, there are some general steps that can be discerned. First, there is the data acquisition phase, in which a data set is typically downloaded according to some criterion. Second, there is a pre-processing stage in which the data is arranged, and possibly classified. Third, extraction of features that are needed for the following calculations commences. Fourth, some algorithm is used to calculate workable numeric values. Finally, some classifier that determines the final assessment is invoked.

Most papers present one or more explicitly stated research

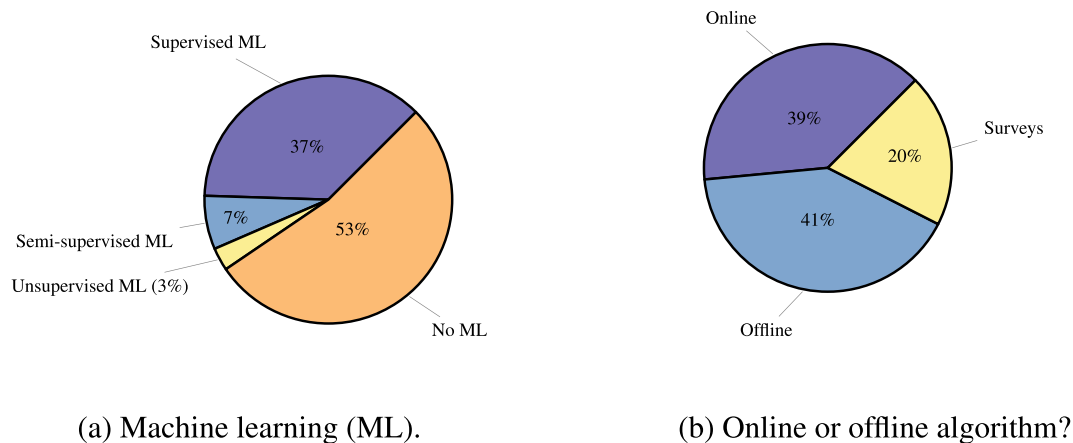


Fig. 4. Algorithm statistics.

questions, e.g., differences between rumors and counter-rumors [16], exploiting topology properties to assess whether a Weibo post is a rumor or not [118], automatic detection of relevance in social networks [25]. Somewhat surprising, quite a few papers lack explicitly stated traditional research questions. However, for many of these papers plausible ulterior research questions and purposes can often be inferred by analyzing the used methods and approaches. Yet in some papers, original research questions are hard to discern or absent. As previously discussed in Section 1, the selected sample of research papers demonstrate a lack of consensus regarding the definition of crucial terms, such as credibility, truthfulness, and veracity. This, what can be called semantic inexactitude, in the context of the research challenges presented here, contributes to muddle the clarity and precision of the posed research questions.

The main aim of our structured literature review is to evaluate articles that seek to determine veracity or a comparable property of a statement in an automated fashion (see Section 1). Hence, we have contributions that seek to determine credibility [61, 73, 75], truthfulness [40, 49, 121], rumors [86, 97, 116], and geolocation [30, 71, 87, 122].

A range of papers do not examine veracity *per se*, but rather develop methods for doing so. On this meta-level authors have developed algorithms [1, 24], novel inventions [35, 125], or created training data (sets) suitable for further research [71, 123], and for example, an approach to combine relevance and credibility scores into a single value [63]. Another category contains the secondary research articles constituting of literature reviews, though with slightly different scopes, [33, 43, 86].

The majority of the papers propose methods that are semi-automatic, that is: some part of the process requires manual intervention, e.g., the downloading of data, labeling, the determination of thresholds, result assessments [16, 38, 60]. In a second category, some articles claim to produce fully automated veracity assessments [27, 85, 98]. A third category do not claim to perform automatic veracity assessments, but the proposed solutions were judged by the authors of this paper to be fully automated with limited additional work, e.g., [90, 101].

The *application fields* in which statements of veracity were to be examined include, most commonly, the news production business. Much interest was shown for potential or established news outlets that produce or distribute news [105]. More specifically, some aim to judge newsworthiness (i.e., newsworthy events) [13], while others try to distinguish actual news items from informal chat [25]. Other application fields include health related information [126], and politics [14]. A few papers include geospatial information [6, 30, 122] as an indicator.

About a dozen papers seek to study phenomena such as *rumors* and *hoaxes* from different perspectives, e.g., [11, 16, 56]. Again, the notion

of semantic inexactitude that we previously mentioned, applies to terms such as rumors and hoaxes as well—neither of which are consistently defined. This means that what is treated like a hoax in one paper, can be labeled as a rumor in another.

The majority of the proposed approaches that were found in this review use Twitter as source data. However, most of the methods are judged to be versatile enough to also use other data source types, e.g., [66, 98, 120].

In general, *detection* and *propagation* methods are studied. Some want to detect and determine whether an item is or is likely to become a rumor, e.g., [42, 97, 118, 124]. Others seek to track how rumors or misinformation spread, e.g., [64, 121, 127]. Some have a more peripheral interest, such as the interplay between rumors and counter-rumors, e.g., [16], as well as the detection of users who spread rumors, e.g., [20, 85]. With regard to hoaxes, some want to examine misinformation in the form of hoaxes [56].

3.3. Algorithms

Almost half of the papers report using machine learning (ML), e.g., [36, 37, 85] (see Fig. 4a). Of these the clear majority use supervised machine learning, e.g., [16, 44, 51], of which a smaller number use some variant of semi-supervised methods, e.g., [10, 34, 60], and only very few use unsupervised methods, e.g., [94, 119, 120]. Of the other half of the papers, some present methods that are not based on machine learning, some are surveys, and some describe data or user behaviors. A large number of different algorithms are used, and some papers try several, or use a combination of several algorithms to achieve their end result, e.g., [3, 8, 48]. Some papers develop specific algorithms for the problem, e.g., [6, 117, 123], while others use well-known algorithms (as part of their method), such as support vector machines, e.g., [66], naïve Bayes, e.g., [118], random forests, e.g., [12], clustering algorithms, e.g., [49], methods for logistic regression, e.g., [5].

About two fifths of the papers claim that their algorithms work online, e.g., [5, 11, 119], see Fig. 4b. Another two fifths describe algorithms that only work offline, e.g., [16, 32, 120]. The last fifth of the papers contain surveys, descriptions of data, or of user behavior, e.g., [29, 100, 127].

For evaluating the methods almost half of the papers use a measure based on the confusion matrix between the result of the algorithm and a known categorization, such as precision, recall, f-measure, e.g., [79, 88, 125]. There are many other measures used for evaluation, and if we count them all almost 70% of the papers make some kind of evaluation, e.g., [94, 106, 124]. Of these, 55% of the papers use machine learning, e.g., [3, 16, 39], which also means that 80% of the machine learning papers make some kind of evaluation.

3.4. Tools

This section gives an overview of the tools employed for veracity assessment by the authors of the studied papers. Around 45% of the papers report the details of all, or parts of, the used tools. Some of the reviewed papers do, however, not implement anything since they are of a visionary, methodological or survey type. The rest of the papers contain no or very sparse information on the used tools. The reported tools and libraries that are in the studied literature belong to a few subfields of data science, information management and artificial intelligence, namely, natural language processing (NLP), machine learning and big data analytics, i.e., large-scale data processing. The use of tools and libraries from different subfields is motivated by the tasks and corresponding steps in veracity assessment, e.g., linguistic analysis of textual data, data collection, and network analysis. Common NLP tools used in the reviewed papers are i) the Natural Language Toolkit (NLTK) [67], ii) Stanford CoreNLP [69], iii) the Stanford dependency parser [21], iv) TweepoParser, a Twitter dependency parser [53], v) the Linguistic Inquiry and Word Count (LIWC) [80], and vi) semantic similarity word vectors like Stanford GloVe [81] and Word2Vec [72].

Rather many of the studied papers, e.g., [30, 32, 55, 103], report the use of Twitter APIs from its developer platform. In particular, the Search API and the Account Activity API, for collecting tweets, finding historic tweets, and obtaining user account statistics, are used. Other examples of tools used in the papers for processing tweets are i) Apache Flume [108], used for streaming tweets from the Twitter API based on a predefined set of keywords [4, 78], ii) Apache Spark [111], a distributed/cluster computing solution used to process tweets [30], iii) networkx [112], a Python library for creating and manipulating complex networks, e.g., used to construct tweet propagation graphs [103], iv) NeuroLab [113], a neural network library for Python used by, e.g., [32], v) scikit-learn [114], ML tools/library in Python used by, e.g., [27, 39], vi) Apache Hadoop [109], a framework for distributed processing of large data sets across clusters of computers, e.g., [13, 22], and vii) Apache Hive [110], a data warehouse software project built on top of Apache Hadoop for providing data query and analysis using SQL used by, e.g., [4].

Almost 18% of the papers state that they use open source tools, but the real number is probably much higher since a majority of the papers provide no or little information of used tools and implementation details, see Fig. 5.

Only around one tenth of the papers have made their tooling publicly available, usually through a web-link, e.g., [50, 74, 92]. Thus, the majority do not provide any details.

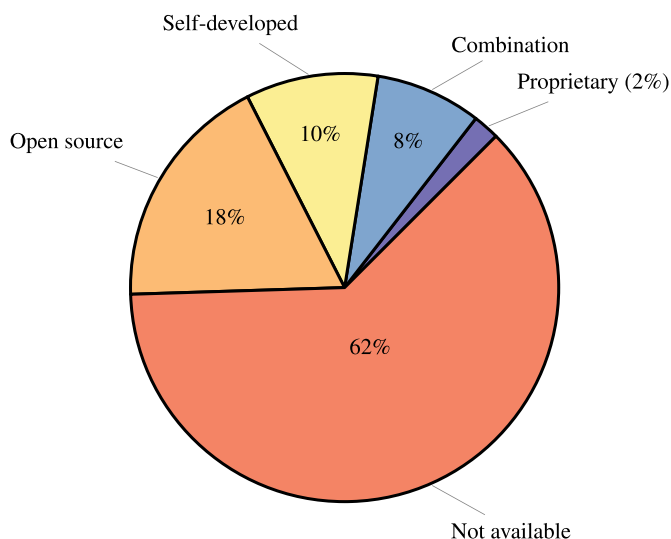


Fig. 5. Nature of used and developed software.

3.5. Data

An overwhelming majority of the research approaches include mining text of some form, see Fig. 6a. Text data types range from i) microblogs, e.g., [6, 13], i.e., short status updates on open social networks, and ii) short texts, e.g., [5, 61, 102], i.e., under 500 words, to iii) long texts, e.g., [60, 82, 92], i.e., more than 500 words. Graph data is the second most common data type. The “not applicable” class consists primarily of exploratory surveys, e.g., [86, 100], books, e.g., [9, 33], and visionary papers, e.g., [29].

Mining veracity assessments themselves, e.g., [38, 125], for veracity assessment is more common than algorithms which process images, e.g., [38, 46], geospatial data, e.g., [31, 87, 122], generic data (any kind of data), e.g., [1, 24], and keyword based approaches, e.g., [74, 95]. Waveform mining, e.g., sound and optical data, in any form, is completely absent in the approaches presented in the studied papers.

As for the data sources, the vast majority of the papers rely on microblogging services such as Twitter [11, 23, 32, 35, 46, 90] and Weibo [34, 49, 105, 118] (see Fig. 6b). Relying on news agencies is more common among papers where the authors are affiliated with China than other countries [47, 105]. Using “fact baseline” sources such as Wikipedia and DBpedia [17, 93, 94] or news agencies [10, 25, 105], is as commonly relied upon as review sites such as Dianping [102], TripAdvisor [5], or Yelp [26, 27]. Geospatial sources, such as Four-Square [122] or GIS services as well as image sources like Instagram [115], are only used by a handful of papers.

As depicted in Fig. 7, almost two thirds of the papers use their own collected authentic data, e.g., [41, 54, 102], whereas one fourth rely on already collected known data sets, e.g., [5, 74]. One fifth of the papers lack details regarding their data set acquisition process, e.g., [43, 99, 117]. A few papers rely on synthetic data, e.g., [8, 94, 106], as part of the data acquisition process, and sometimes also combine the approach by using either authentic or known data sets. Only one paper combined the usage of authentic data and known data [65].

One out of five papers indicate how to access the data sets, mostly as web URLs, e.g., [37]. The most common way among these papers is to share the data sets via GitHub, e.g., [12, 50, 127]. Otherwise we were unable to find any commonalities between papers with regards to data set sharing services: one used Dropbox [47], some used plain web servers, e.g., [37, 56], and one explicitly stated that the data set was available upon request [13].

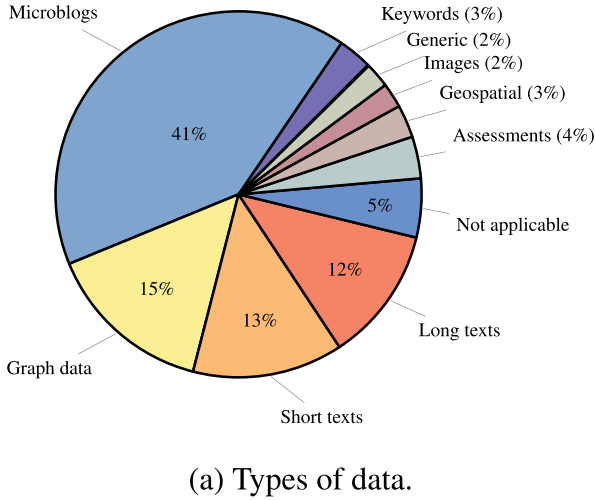
One of the review protocol questions investigated the possibility for data set reuse, and specifically whether the available data can be used for veracity assessment benchmarking. The criteria for whether the data set is benchmarkable is i) if the data is publicly available, and ii) if there is a suitable performance metric target with regards to veracity or an indicator. One in six papers contain enough details regarding their data sets for them to be usable as benchmark data, e.g., [5, 27, 50, 62, 93].

As previously mentioned in Section 3.4, a common data source is Twitter. The data gathering is mostly done using the Twitter streaming API, which is used for collecting microblog posts and annotating them with meta data. The second most popular data gathering method is web crawling and scraping, e.g., [26, 37, 43, 115].

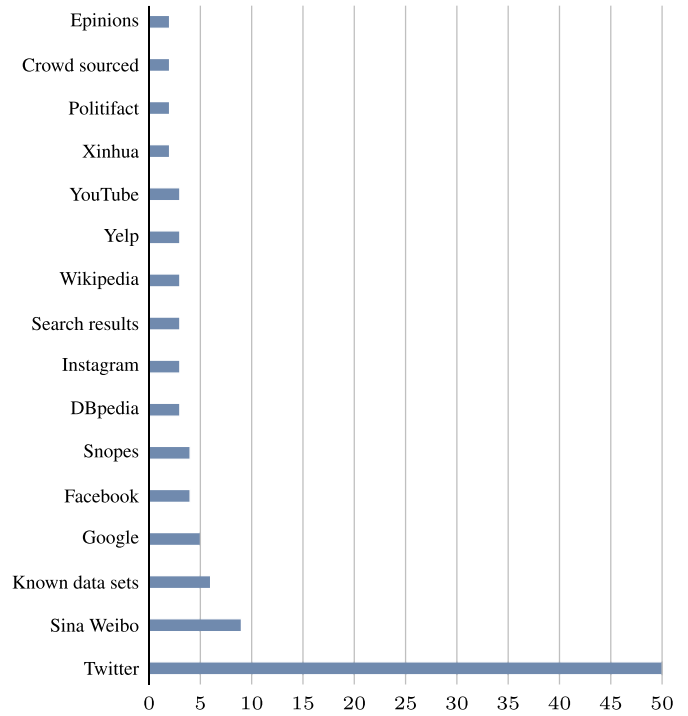
Regarding the data selection criteria, the most common approach is keyword based, e.g., [4, 32, 54, 115], while automatically mining keywords for data acquisition is not that common, e.g., [127]. Many papers select data based on time frame ranges, rather than topics, e.g., [23, 48, 49, 84].

Regarding data use, the papers mainly focus on content and to a lesser extent on meta data and social graph structures. Popular features include (where applicable) number of replies, “retweets,” number of connections, number of positive/negative words, entity frequency, and word class percentages.

The majority of the papers do not include any specific statistical analysis or amendments related to skewed data. The minority cases consist of the papers for which the proposed method could work on



(a) Types of data.



(b) Data sources used in two or more papers.

Fig. 6. Types of data, and data sources.

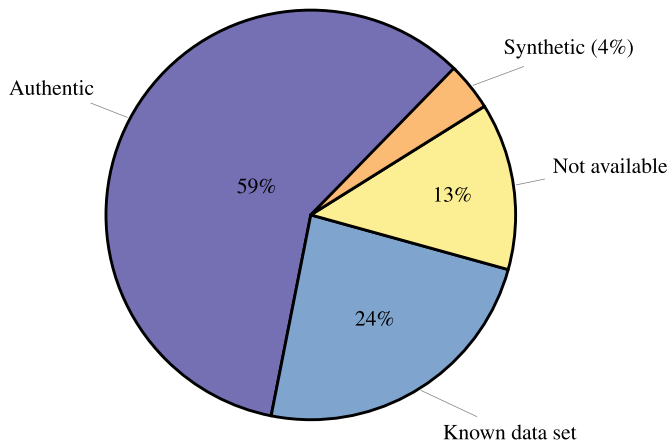


Fig. 7. Data acquisition process.

3.6. Miscellaneous

As previously discussed in Sections 1 and 3.2 there is a semantic inexactitude and terminological breadth present regarding the definitions of central terms such as veracity and its closest concepts, i.e., credibility, truth, quality, etc. The results show that only a handful of the papers analyzed offer explicit definitions of veracity or veracity assessment. Thus, the following definition is given by Jamil et al. [43], in turn based on Bennett-Woods [7]: “Data veracity refers to principles of truth-telling, and it is grounding in respect for persons and the concept of autonomy.” Bodnar et al. [11] give a twofold definition in passing: “veracity referring to the accuracy and truthfulness of the data as well as the ability of the data to predict trends.” Robin et al. [87] equate veracity with truthfulness: “Veracity refers to the degree of truthfulness associated with a data set,” as does Debattista et al. [22]: “conformity with truth or facts.” Wang et al. [105] also remark very briefly: “veracity (trustworthiness of various data).”

Conroy et al. [19] implicitly define veracity by how it is compromised: “Veracity is compromised by the occurrence of intentional deceptions.” A similar approach is taken by Bhattacharjee et al. [10], stating that “[t]he objective of a news veracity detection system is to identify various types of potentially misleading or false information.”

A few other papers give operationalizations intended only for the paper itself, such as “[t]he probability of a tweet to be a counter-rumor is referred as tweet veracity in this paper” [16].

However, most papers do not define veracity or veracity assessment. Instead they introduce, define, or discuss other related terms and concepts. Examples include deception, e.g., [99], misinformation, e.g., [54, 64], credibility, e.g., [25, 39, 50, 73, 101], reliability, e.g., [121], believability, e.g., [25, 32, 100], trust and trustworthiness, e.g., [1, 10, 32, 35], truthfulness, e.g., [94, 121], and truth discovery, e.g., [8, 31]. Again, some give operationalizations intended only for the particular paper, such as a “rumor is defined as any information posted on Twitter, that many people believe to be true, but it contrasts with the

multi-sourced data under non-independent identically distributed assumptions, e.g., [84], makes skewness adjustments, e.g., [93], or performs data exploration as part of the evaluation [54]. The most common assumption concerning the data is that it has the right membership, i.e., that it belongs to a rumor/event, e.g., [3, 15, 86, 105], is a review, e.g., [101], or concerns the topic, e.g., [39, 50, 55]. It should also be noted that papers without any explicitly mentioned assumptions regarding data distribution might still have implicit assumptions. Similarly, the methods proposed in most papers rely on the data exhibiting a specific shape, i.e., that entries contain certain strict features or value ranges, e.g., [10, 36, 55, 87, 106], albeit not being explicit about it. Some papers also rely upon the veracity of features of a data point being inherently correct, e.g., have correct geographical information [6, 71], or have assigned credibility scores [64].

news tweets from the verified news channels” [42].

Turning to legal and ethical issues of automated veracity assessment, these are absent in all but a single paper. Webb et al. [107] alone discuss ethical issues as a prominent part (Section 4) of their effort to define a research agenda on the governance and regulation of social media. However, they do not discuss ethics directly related to veracity assessment. No paper discusses legal issues.

Assessing the relevance of the papers to the main research question, as introduced in Section 1, most of the papers read are of high or medium relevance, as is to be expected given the selection process and search strategy as outlined in Section 2.1. However, some papers are assessed to be of low relevance.

As expected, the breadth of the scientific contributions made in the reviewed papers is significant. Even though the papers all address some aspect of veracity assessment of online data, the ranges of methods, algorithms, tools, data, etc., are substantial. Nevertheless, a “typical” paper i) proposes some kind of method or algorithm that is either entirely novel or more commonly an addition or improvement to an existing one, ii) applies it to some interesting data, and iii) evaluates the results. However, as the papers apply their methods to investigate interesting phenomena, they sometimes also make positive social science contributions, e.g., about the characteristics of Wikipedia hoaxes [56], about differences between true and false health rumors [126], and about the interplay between fake news promoters and grass-roots responses [92]. Another kind of contribution found, as previously described in Section 3.5, is the introduction of data sets subsequently made available to the research community, e.g., [50, 73, 85]. Unfortunately, another not too uncommon deviation from the typical paper structure outlined above is that the evaluation is missing, very narrow, or flawed in some other way.

A large majority of the papers contain primary research contributions. In addition, a handful of secondary research contributions, i.e., literature reviews, were included [19, 22, 33, 43, 86, 99, 100, 107]. Furthermore, some papers are best characterized as positional, i.e., discussing interesting ideas for future research rather than making full contributions in their own right.

Unsurprisingly, the large majority of papers are directed towards the scientific community, mostly that consisting of computer scientists. When a particular application or interest group is mentioned, journalism (including both the supply-side of journalists writing news articles and the demand-side of consumers reading them) is the most common [19, 40, 65, 66, 97]. Other perspectives include marketing [5], e-commerce [117], medicine [75], social network moderation [47, 48], and the military [59].

4. Synthesis and discussion

In this section we synthesize and discuss the results presented in Section 3. The section consists of four subsections containing discussions on i) approaches and methods, ii) algorithms, tools, and data, iii) gaps, and iv) validity and reliability.

4.1. Approaches and methods

Looking at the descriptions of indicators, methods, and definitions (mostly implicit) of veracity used in the papers, three broad categories of veracity operationalizations can be discerned: i) implicit features, ii) explicit fact checking, and iii) appeal to authority.

The *implicit features* approach is by far the most common. Roughly, the idea is that claims that are (in some sense) non-veridical differ from claims that are veridical in other, non-veracity, properties. Such properties include stylistic text features such as length and wordings [37, 56, 82], URL features such as link densities [56] or domain names [97], temporal distributions [92, 103, 117], (social network) distribution patterns [1, 92, 117], and user account features [89, 97].

The *explicit fact checking* approach is rare, but a few examples were

found [59, 94]. The idea is to compare a claim made to an existing body of knowledge so as to determine if it is veridical. Typically this involves representing the claim as a subject-predicate-object triple, and then using graph-methods to compare it to existing knowledge triples.

The *appeal to authority* approach, in its most crude form, is also rare. The idea is that a claim is veridical if it is claimed by an authoritative source. For example, a photo can be trusted if shared by a trusted source 30 min after the event [115], and a claim can be considered veridical if supported by the majority [76] or by verified news channels [42].

It should be noted that the mentioned three approaches are often combined to achieve better results. For example, a moderate appeal to authority is often blended into the implicit features approach by, e.g., including some PageRank-like features among the other implicit features considered [41, 54, 82].

4.2. Algorithms, tools, and data

That most papers that are concerned with machine learning have used supervised methods is not a surprise. Veracity estimation is a very difficult task and the veracity is probably in many cases dependent on factors external to what is available to the algorithm. Therefore, in most situations the results of these algorithms should be subject to manual consideration. In such scenarios unsupervised methods could prove quite useful as a complement, and provide the human with more information.

It is remarkable that two fifths of the papers describe algorithms/methods that work online, considering how complex veracity assessment is. It should probably be understood that these online algorithms, i.e., algorithms that work with streaming data, scale well over processing cores. On the other hand, methods that do not work online (two fifths of the papers) can potentially work in some kind of batch version, although this may require extra resources to update knowledge over the entire data set.

There seems to be a big problem with reproducibility in veracity assessment research. Many papers do not share source code, models, and data. This can be in the form of missing URLs, due to updated web pages and absent servers, or even underspecified details in the paper. Some authors rely on known data sets and software, but fail to disclose versioning constraints or what parts were used. None of the papers provide DOI links to point out the research materials used or to publish trained models.

The main data source in many papers is Twitter. Hence, it is unsurprising that many of the reviewed papers are text oriented. The majority of the tooling is adapted thereafter and is mostly focused on different types of linguistic analysis, supervised machine learning, and big data analytics.

One would expect many proposals to contain intricate computations relying on diverse data sources and data type sets. However, a majority of the papers have a narrow focus using only one data source and/or data type for a specific algorithm, which can be seen as an indication of the immaturity of the field. Only a handful of papers use, or are adapted for, multiple sources, which in many cases would be necessary in a real application, see, e.g., [10, 74, 105].

4.3. Gap analysis

A gap analysis based on the obtained results and synthesis is presented in this section. The identified gaps summarize the main challenges that have been identified through the systematic literature review.

Multiple sources and data types. Of the analyzed papers very few approaches or methods are adapted to handle multiple sources and/or data types. Since one of the pillars of source criticism or information evaluation is the comparison of information from multiple sources and data types, this should also by extension be a criterion for future automatic veracity

assessment systems. One could argue that sources like Twitter and other microblogs are in essence multiple sources since the expressed opinions come from various individuals. However, the format is limited and the expressed opinions/information to a very low degree come from authoritative sources. Also, even though other data types such as links, images, sound, and video, are sometimes embedded, very few of the approaches make full use of these additional data types.

Common definitions of core terminology. As discussed in the introduction, there is no common definition of the core terminology related to veracity or veracity assessment. The analysis of the selected papers showed that the lack of consensus is also present in related terms, e.g., credibility, rumor, and source, making it cumbersome for the research community to compare results and follow the state of the art within the domain.

Reproducibility. Another challenge which was identified in the synthesis is the difficulty of reproducing obtained results. Lack of details or accessibility to data sets, code, and used tools, make reproducing results difficult if not impossible.

Data sets suitable for benchmarking. One of the identified gaps is the limitation of suitable data sets with which the research community can compare results and follow the development of methods.

Deep learning and transfer learning. Machine learning has, with recent years' reemergence of deep learning, made giant leaps and has had unprecedented success in a number of fields. However, the use of deep learning techniques in the evaluated paper set is very low, and a research gap is clearly present. This is also related to the previous point—the lack of suitable data sets—which further limits effective use of machine learning.

Scalable online methods and data. Many of the used approaches and methods are theoretically scalable or applicable in an online setting. However, the majority of the reviewed papers' results come from experiments which have not focused on scalability or streaming data. For a realistic open source data veracity assessment application, these two aspects (scalability and ability to handle streamed data) are probably crucial.

4.4. Validity and reliability

The main strength of validity of the present study is the rigor and transparency of the method employed, adhering to the guidelines of Kitchenham and Charters [52]. In practice, this means that all papers reviewed were selected from databases of renowned peer-reviewed sources, and match explicit inclusion criteria, as listed in Section 2. Thus, the selected papers should comprise a representative selection of the research done in the veracity assessment community.

A moderate threat to validity relates to vocabulary and search strings—the queries listed in Table 1 reflect a Western bias in terms of services (e.g., Facebook, Twitter, Instagram) and language (English). Still, this threat should not be exaggerated—the vast majority of high-impact computer science research is published in English regardless of origin (as is also suggested by the diverse distribution of countries and geographical regions in Fig. 3), and the services mentioned in the queries are truly global, even though there are countries where they are barely used.

A small threat to validity is that there is a bias in the review protocol towards computer science in general and machine learning in particular. Social science terms and methods are not similarly reflected. However, this largely reflects a legitimate delimitation of the research questions, and the residual threat to validity is minor.

A moderate threat to reliability is related to the review protocol, where some questions, notably in part 2 in the protocol (see

Appendix A), can be interpreted in different ways. Though every effort was made to ensure reviewer agreement on these questions, conclusions should be interpreted in light of this risk.

5. Conclusions

The main purpose of this work has been to investigate which approaches, methods, algorithms, and tools that are used or proposed for automatic veracity assessment of open source data. In the use of open source data its veracity is important to consider should the data be used for decision-making in itself or as part of a decision support system. The purpose was also to see how far the research community has progressed since the introduction of veracity (assessment) in big data back in 2012. Using a structured literature review method, papers have been identified, selected and evaluated following a predefined assessment protocol. The protocol was constructed for the purpose of analyzing the research literature targeting veracity assessment of heterogeneous and unstructured open source data, including social media.

One of the things revealed in the results is that in the years that have passed since the inception of veracity in big data, researchers have not reached consensus on a veracity (assessment) definition. Despite this, there is some convergence in the methods used to assess veracity. Three main veracity assessment research approaches were found. The implicit features approach hypothesizes that non-veridical statements differ from veridical statements not only concerning the actual claim but also in other aspects that can be used for assessment. Next, the explicit fact checking approach makes use of external data to evaluate a claim in relation to existing knowledge. Finally, the appeal to authority approach stipulates that a claim can be trusted if it is also claimed or can be verified by an authoritative source. Legal and ethical aspects have unfortunately been discussed to a very low degree. A reproducibility problem can also be seen where many papers are lacking in data gathering details, data sets are not publicly available, and details regarding toolsets and implementation are sparse.

The identified gaps in the current literature mainly consist of i) a general lack of approaches and methods adapted to multiple sources and data types, ii) a lack of consensus in the definitions of core terms, iii) reproducibility challenges, iv) very few available data sets suitable for benchmarking purposes, v) low use of recent advancements made in machine learning, and vi) a lack of research efforts targeting scalable solutions for managing streaming data.

Acknowledgments

We gratefully acknowledge the help obtained from the librarian Alexis Wiklund, for performing the initial database literature searches. This work was supported by the Swedish Armed Forces' research and development program and the European Union Horizon 2020 program (grant agreement no. 832921).

Appendix A. Review protocol

This section contains the review protocol used by the authors to analyze the selected papers.

1. General information

- (a) Internal ID
- (b) Title
- (c) Authors
- (d) Abstract
- (e) Publication year (actual publication date, i.e., not the “online first” date)
- (f) Author background (affiliated to company, university, government institution, a mixture)
- (g) Countries (i.e., author affiliation countries)

- (h) BibTeX reference (including the fields “doi” and/or “url”)
2. Research questions
- (a) Approaches
- What aspect of VA does the approach target, i.e., do the authors try to assess trustworthiness, credibility, formal correctness, explicit lies, bot vs. human, etc.?
 - Do the authors try to extract/mine an indicator related to VA or do they try to determine VA directly?
 - Which indicator(s) do the authors target (e.g., stance, geographical location, social network)?
 - Do the authors motivate the choice of indicator? If so, how?
 - How is the VA or indicator quantified (a scale, a confidence interval, a binary response, a heatmap color, etc.)?
 - What is evaluated, e.g., the assessment itself, the method, the data, etc.?
- (b) Methods
- Are there one or more distinct research questions? If yes, what is it/what are they?
 - Method: describe the procedural VA steps taken in chronological order
 - Is the VA method fully automated or semi-automatic (requiring manual intervention)? If semi-automatic, what intervention(s) are required?
 - Is the method (apparently) applicable for different data sources?
- (c) Algorithms
- Which algorithms do they employ (mention only the algorithms directly involved in the VA or related task)?
 - Is the algorithm(s) based on an ML method? If so, what type (supervised, unsupervised, reinforcement, etc.)?
 - Do the algorithms handle online or offline data (streams)?
 - What quality assessment measure(s) is used (precision, accuracy, entropy, etc.)?
- (d) Tools
- Which tools do the authors employ?
 - Have the tools been developed by the authors themselves or have they used proprietary software, open source, etc.?
 - Has their code/tool been made publicly available (if yes, how)?
- (e) Data
- Which data types (tweet, picture, sound, article, long/short text, etc.)?
 - Which data sources (Twitter, Facebook, Wikipedia, RSS, blog posts, etc.)?
 - Which data sets are used (gathered by themselves, a known data set, synthetic, authentic)?
 - If gathered (produced) by the researchers, has the data been made available? If so, where?
 - If the data was collected or produced by the researchers, how was it done?
 - Is the data usable for benchmarking?
 - What were the data selection criteria (keywords, time frame, accounts, etc.)?
 - Which parts of the data do they use in the VA, e.g., do they use content, meta data, network data, feature types?
 - Are there any particular assumptions made regarding the data or its distribution (if yes, which)?
- (f) Miscellaneous
- Does the paper give a definition of veracity and/or veracity

- assessment? If so, what is the definition? (copy/paste from the paper)
 - Does the paper discuss ethical issues related to automatic veracity assessment? If yes, which?
 - Does the paper discuss legal aspects related to automatic veracity assessment? If yes, which?
 - Relevance to main research question (high, medium, low)?
 - Summary of statements (contributions) made in the article
 - Type of paper (primary, secondary, tertiary, other)
 - Perspective/interest group
3. The paper (qualitative assessment)
- Strengths of the paper
 - Weaknesses of the paper
 - Subjective assessment/reflection (state of the art or not, worth reading or not, etc.)

Appendix B. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.dss.2019.113132>. The supplementary data contains the full list of reviewed papers for this study.

References

- M.-A. Abbasi, H. Liu, Measuring user credibility in social media, in: A.M. Greenberg, W.G. Kennedy, N.D. Bos (Eds.), *Social Computing, Behavioral-Cultural Modeling and Prediction*, Springer Berlin Heidelberg, Berlin, Heidelberg, 2013, pp. 441–448, https://doi.org/10.1007/978-3-642-37210-0_48.
- A. Abdul-Rahman, S. Hailes, Supporting trust in virtual communities, *Proceedings of the 33rd Annual Hawaii International Conference on System Sciences*, IEEE, 2000, pp. 9–pp.
- A. Aker, A. Zubiaga, K. Bontcheva, A. Kolliakou, R. Procter, M. Liakata, Stance classification in out-of-domain rumours: a case study around mental health disorders, in: G.L. Ciampaglia, A. Mashhadi, T. Yasseri (Eds.), *Social Informatics*, Springer International Publishing, Cham, 2017, pp. 53–64.
- K.T. Ashwin, P. Kammarpally, K. George, Veracity of information in twitter data: a case study, 2016 International Conference on Big Data and Smart Computing (BigComp), 2016, pp. 129–136, <https://doi.org/10.1109/BIGCOMP.2016.7425811>.
- L. Ball, J. Elworthy, Fake or real? The computational detection of online deceptive text, *Journal of Marketing Analytics* 2 (2014) 187–201, <https://doi.org/10.1057/jma.2014.15>.
- J. Bendler, S. Wagner, T. Brandt, D. Neumann, Taming uncertainty in big data, business & information systems engineering 6 (2014) 279–288, <https://doi.org/10.1007/s12599-014-0342-4>.
- D. Bennett-Woods, *Ethics at a Glance*, Regis University, 2005.
- L. Berti-Equille, Data veracity estimation with ensembling truth discovery methods, 2015 IEEE International Conference on Big Data (Big Data), 2015, pp. 2628–2636, <https://doi.org/10.1109/BigData.2015.7364062>.
- L. Berti-Equille, J. Borge-Holthoefer, Veracity of data: from truth discovery computation algorithms to models of misinformation dynamics, *Synthesis Lectures on Data Management* 7 (2015) 1–155.
- S.D. Bhattacharjee, A. Talukder, B.V. Balantrapu, Active learning based news veracity detection with feature weighting and deep-shallow fusion, 2017 IEEE International Conference on Big Data (Big Data), 2017, pp. 556–565, <https://doi.org/10.1109/BigData.2017.8257971>.
- T. Bodnar, C. Tucker, K. Hopkinson, S.G. Bilén, Increasing the veracity of event detection on social media networks through user trust modeling, 2014 IEEE International Conference on Big Data (Big Data), 2014, pp. 636–643, <https://doi.org/10.1109/BigData.2014.7004286>.
- C. Buntain, J. Golbeck, Automatically identifying fake news in popular Twitter threads, 2017 IEEE International Conference on Smart Cloud (SmartCloud), 2017, pp. 208–215, <https://doi.org/10.1109/SmartCloud.2017.40>.
- C. Castillo, M. Mendoza, B. Poblete, Predicting information credibility in time-sensitive social media, *Internet Research* 23 (2013) 560–588, <https://doi.org/10.1108/IntR-05-2012-0095>.
- C. Chang, Y. Zhang, C. Szabo, Q.Z. Sheng, Extreme user and political rumor detection on Twitter, in: J. Li, X. Li, S. Wang, J. Li, Q.Z. Sheng (Eds.), *Advanced Data Mining and Applications*, Springer International Publishing, Cham, 2016, pp. 751–763, https://doi.org/10.1007/978-3-319-49586-6_54.
- O.K. Cheng, R.Y. Lau, A multi-perspective methodology for detecting low-quality contents in social media, 2014 International Conference on Advanced ICT (ICAICTE-2014), Citeseer, 2014.
- A.Y. Chua, S. Banerjee, A study of tweet veracity to separate rumours from counter-rumours, *Proceedings of the 8th International Conference on Social Media & Society*, ACM, 2017, p. 4.
- G.L. Ciampaglia, P. Shiralkar, L.M. Rocha, J. Bollen, F. Menczer, A. Flammini,

- Computational fact checking from knowledge networks, *PLoS One* 10 (2015) 1–13, <https://doi.org/10.1371/journal.pone.0128193>.
- [18] I. Claverie-Berge, Solutions Big Data IBM, (2012) presentation slides.
- [19] N.J. Conroy, V.L. Rubin, Y. Chen, Automatic deception detection: methods for finding fake news, Proceedings of the 78th ASIS&T Annual Meeting: Information Science with Impact: Research in and for the Community, ASIST '15, American Society for Information Science, Silver Springs, MD, 2015, pp. 82:1–82:4 <http://dl.acm.org/citation.cfm?id=2857070.2857152>.
- [20] A. Dang, M. Smit, A. Mohammad, R. Minghim, E.E. Milios, Toward understanding how users respond to rumours in social media, 2016 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM), 2016, pp. 777–784, <https://doi.org/10.1109/ASONAM.2016.7752326>.
- [21] M.-C. De Marneffe, B. MacCartney, C.D. Manning, et al., Generating typed dependency parses from phrase structure parses. *Lrec*, vol. 6, 2006, pp. 449–454.
- [22] J. Debatista, C. Lange, S. Scerri, S. Auer, Linked 'Big' Data: Towards a Manifold Increase in Big Data Value and Veracity, 2015 IEEE/ACM 2nd International Symposium on Big Data Computing (BDC), 2015, pp. 92–98, <https://doi.org/10.1109/BDC.2015.34>.
- [23] D. Elias, F. Nadler, I. Cornwell, S. Grant-Muller, T. Heinrich, UNIETD - assessment of third party data as information source for drivers and road operators, *Transportation Research Procedia* 14 (2016) 2035–2043, <https://doi.org/10.1016/j.trpro.2016.05.171>.
- [24] X.S. Fang, Q.Z. Sheng, X. Wang, A.H. Ngu, Value veracity estimation for multi-truth objects via a graph-based approach, Proceedings of the 26th International Conference on World Wide Web Companion, WWW '17 Companion, International World Wide Web Conferences Steering Committee, Republic and Canton of Geneva, Switzerland, 2017, pp. 777–778, <https://doi.org/10.1145/3041021.3054212>.
- [25] A. Figueira, M. Sandim, P. Fortuna, An approach to relevancy detection: contributions to the automatic detection of relevance in social networks, in: Á. Rocha, A.M. Correia, H. Adeli, L.P. Reis, M. Mendonça Teixeira (Eds.), *New Advances in Information Systems and Technologies*, Springer International Publishing, Cham, 2016, pp. 89–99.
- [26] J. Fontanarava, G. Pasi, M. Viviani, An ensemble method for the credibility assessment of user-generated content, Proceedings of the International Conference on Web Intelligence, WI '17, ACM, New York, NY, 2017, pp. 863–868, <https://doi.org/10.1145/3106426.3106464>.
- [27] J. Fontanarava, G. Pasi, M. Viviani, Feature analysis for fake review detection through supervised classification, 2017 IEEE International Conference on Data Science and Advanced Analytics (DSAA), 2017, pp. 658–666, <https://doi.org/10.1109/DSAA.2017.51>.
- [28] D. Gambetta, et al., Can we trust trust, *Trust: making and breaking cooperative relations* 13 (2000) 213–237.
- [29] M. García Lozano, U. Franke, M. Rosell, V. Vlassov, Towards automatic veracity assessment of open source information, 2015 IEEE International Congress on Big Data, 2015, pp. 199–206, <https://doi.org/10.1109/BigDataCongress.2015.36>.
- [30] M. García Lozano, J. Schreiber, J. Brynielson, Tracking geographical locations using a geo-aware topic model for analyzing social media data, *Decision Support Systems* 99 (2017) 18–29, <https://doi.org/10.1016/j.dss.2017.05.006> location Analytics and Decision Support.
- [31] D.A. Garcia-Ulloa, L. Xiong, V. Sunderam, Truth discovery for spatio-temporal events from crowdsourced data, Proceedings of the VLDB Endowment 10 (2017) 1562–1573, <https://doi.org/10.14778/3137628.3137662>.
- [32] G. Giasemidis, C. Singleton, I. Agraftotis, J.R.C. Nurse, A. Pilgrim, C. Willis, D.V. Greetham, Determining the veracity of rumours on Twitter, in: E. Spiro, Y.-Y. Ahn (Eds.), *Social Informatics*, Springer International Publishing, Cham, 2016, pp. 185–205.
- [33] A.L. Ginsca, A. Popescu, M. Lupu, et al., Credibility in information retrieval, foundations and trends[®] in information retrieval 9 (2015) 355–475.
- [34] Q. Guo, W.W. Huang, K. Huang, X. Liu, Information credibility: a probabilistic graphical model for identifying credible influenza posts on social media, in: X. Zheng, D.D. Zeng, H. Chen, S.J. Leischow (Eds.), *Smart Health*, Springer International Publishing, Cham, 2016, pp. 131–142, https://doi.org/10.1007/978-3-319-29175-8_12.
- [35] A. Gupta, P. Kumaraguru, C. Castillo, P. Meier, TweetCred: real-time credibility assessment of content on Twitter, in: L.M. Aiello, D. McFarland (Eds.), *Social Informatics*, Springer International Publishing, Cham, 2014, pp. 228–243.
- [36] M. Hardalov, I. Koychev, P. Nakov, In search of credible news, in: C. Dichev, G. Agre (Eds.), *Artificial Intelligence: Methodology, Systems, and Applications*, Springer International Publishing, Cham, 2016, pp. 172–180.
- [37] R.A. Igawa, S.B. Jr, K.C.S. Paulo, G.S. Kido, R.C. Guido, M.L.P. Júnior, I.N. da Silva, Account classification in online social networks with LBCA and wavelets, *Information Sciences* 332 (2016) 72–83, <https://doi.org/10.1016/j.ins.2015.10.039>.
- [38] B. Ionescu, A. Popescu, M. Lupu, A.L. Gînsca, B. Boteanu, H. Müller, Div150Cred: a social image retrieval result diversification with user tagging credibility dataset, Proceedings of the 6th ACM Multimedia Systems Conference, MMSys '15, ACM, New York, NY, 2015, pp. 207–212, <https://doi.org/10.1145/2713168.2713192>.
- [39] J. Ito, J. Song, H. Toda, Y. Koike, S. Oyama, Assessment of tweet credibility with LDA features, Proceedings of the 24th International Conference on World Wide Web, ACM, 2015, pp. 953–958.
- [40] E. Jaho, E. Tzoannos, A. Papadopoulos, N. Sarris, Alethiometer: a framework for assessing trustworthiness and content validity in social media, Proceedings of the 23rd International Conference on World Wide Web, WWW '14 Companion, ACM, New York, NY, 2014, pp. 749–752, <https://doi.org/10.1145/2567948.2579324>.
- [41] P. Jain, V. Singh, CredRank: evaluating tweet credibility during high impact events, 2016 2nd International Conference on Contemporary Computing and Informatics (IC3I), 2016, pp. 553–557, <https://doi.org/10.1109/IC3I.2016.7918025>.
- [42] S. Jain, V. Sharma, R. Kaushal, Towards automated real-time detection of misinformation on Twitter, 2016 International Conference on Advances in Computing, Communications and Informatics (ICACCI), 2016, pp. 2015–2020, <https://doi.org/10.1109/ICACCI.2016.7732347>.
- [43] N.B.C.E. Jamil, I.B. Ishak, F. Sidi, L.S. Affendey, A. Mamat, A systematic review on the profiling of digital news portal for big data veracity, *Procedia Computer Science* 72 (2015) 390–397, <https://doi.org/10.1016/j.procs.2015.12.154> the Third Information Systems International Conference 2015.
- [44] B. Janssen, M. Habib, M. van Keulen, Truth assessment of objective facts extracted from Tweets: a case study on World Cup 2014 game facts, Proceedings of the 13th International Conference on Web Information Systems and Technologies, vol. 1, 2017, pp. 187–195, <https://doi.org/10.5220/0006185101870195>.
- [45] S. Jeong, G. Noh, H. Oh, C.-k. Kim, Follow spam detection based on cascaded social information, *Information Sciences* 369 (2016) 481–499, <https://doi.org/10.1016/j.ins.2016.07.033>.
- [46] Z. Jin, J. Cao, H. Guo, Y. Zhang, J. Luo, Multimodal fusion with Please supply the year of publication. recurrent neural networks for rumor detection on microblogs, in: Proceedings of the 2017 ACM on Multimedia Conference, MM '17 (pp. 795–816), ACM, doi:10.1145/3123266.3123454.
- [47] Z. Jin, J. Cao, Y.G. Jiang, Y. Zhang, News credibility evaluation on microblog with a hierarchical propagation model, 2014 IEEE International Conference on Data Mining, 2014, pp. 230–239, <https://doi.org/10.1109/ICDM.2014.91>.
- [48] Z. Jin, J. Cao, Y. Zhang, J. Luo, News verification by exploiting conflicting social viewpoints in microblogs, *AAAI*, 2016, pp. 2972–2978.
- [49] Z. Jin, J. Cao, Y. Zhang, J. Zhou, Q. Tian, Novel visual and statistical image features for microblogs news verification, *IEEE Transactions on Multimedia* 19 (2017) 598–608, <https://doi.org/10.1109/TMM.2016.2617078>.
- [50] M. Kakol, R. Nielek, A. Wierzbicki, Understanding and predicting Web content credibility using the Content Credibility Corpus, *Information Processing & Management* 53 (2017) 1043–1061.
- [51] Y.A. Kim, M.A. Ahmad, Trust, distrust and lack of confidence of users in online social media-sharing communities, *Knowledge-Based Systems* 37 (2013) 438–450, <https://doi.org/10.1016/j.knosys.2012.09.002>.
- [52] B.A. Kitchenham, S.M. Charters, Guidelines for Performing Systematic Literature Reviews in Software Engineering, Version 2.3, Technical Report EBSE-2007-01, Keele University and Durham University, United Kingdom, 2007.
- [53] L. Kong, N. Schneider, S. Swayamdipta, A. Bhatia, C. Dyer, N.A. Smith, A dependency parser for tweets, Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP), 2014, pp. 1001–1012.
- [54] K.P.K. Kumar, G. Geethakumari, Detecting misinformation in online social networks using cognitive psychology, *Human-centric Computing and Information Sciences* 4 (2014) 14, <https://doi.org/10.1186/s13673-014-0014-x>.
- [55] K.P.K. Kumar, A. Srivastava, G. Geethakumari, A psychometric analysis of information propagation in online social networks using latent trait theory, *Computing* 98 (2016) 583–607, <https://doi.org/10.1007/s00607-015-0472-7>.
- [56] S. Kumar, R. West, J. Leskovec, Disinformation on the Web: impact, characteristics, and detection of Wikipedia hoaxes, Proceedings of the 25th International Conference on World Wide Web, WWW '16, International World Wide Web Conferences Steering Committee, Republic and Canton of Geneva, Switzerland, 2016, pp. 591–602, <https://doi.org/10.1145/2872427.2883085>.
- [57] D. Laney, 3D data management: controlling data volume, velocity and variety, *META Group Research Note* 6 (2001) 1.
- [58] P. Lendvai, U.D. Reichel, T. Declerck, Factuality drift assessment by lexical markers in resolved rumors, Proceedings of the 1st International Workshop on Semantic Change & Evolving Semantics, 2016.
- [59] G. Levchuk, E. Blasch, Probabilistic graphical models for multi-source fusion from text sources, 2015 IEEE Symposium on Computational Intelligence for Security and Defense Applications (CISDA), 2015, pp. 1–10, <https://doi.org/10.1109/CISDA.2015.7208640>.
- [60] G. Levchuk, C. Shabarekh, Using soft-hard fusion for misinformation detection and pattern of life analysis in OSINT, in: T.P. Hanratty, J. Llinas (Eds.), Proceedings of SPIE Vol. 10207, Next-Generation Analyst V, SPIE, Bellingham, WA, 2017, <https://doi.org/10.1117/12.2263546>.
- [61] R. Li, A. Suh, Factors influencing information credibility on social media platforms: evidence from Facebook Pages, *Procedia Computer Science* 72 (2015) 314–328, <https://doi.org/10.1016/j.procs.2015.12.146>.
- [62] W.Y. Lim, M.L. Lee, W. Hsu, iFACT: an interactive framework to assess claims from Tweets, Proceedings of the 2017 ACM on Conference on Information and Knowledge Management, CIKM '17, ACM, New York, NY, 2017, pp. 787–796, <https://doi.org/10.1145/3132847.3132995>.
- [63] C. Lioma, J.G. Simonsen, B. Larsen, Evaluation measures for relevance and credibility in ranked Lists, Proceedings of the ACM SIGIR International Conference on Theory of Information Retrieval, ICTIR '17, ACM, New York, NY, 2017, pp. 91–98, <https://doi.org/10.1145/3121050.3121072>.
- [64] I. Litou, V. Kalogeraki, I. Katakis, D. Gunopulos, Efficient and timely misinformation blocking under varying cost constraints, *Online Social Networks and Media* 2 (2017) 19–31, <https://doi.org/10.1016/j.osnem.2017.07.001>.
- [65] X. Liu, Q. Li, A. Nourbakhsh, R. Fang, M. Thomas, K. Anderson, R. Kociuba, M. Vedder, S. Pomerville, R. Wudali, R. Martin, J. Duprey, A. Vachher, W. Keenan, S. Shah, Reuters tracer: a large scale system of detecting and verifying real-time news events from Twitter, Proceedings of the 25th ACM International on Conference on Information and Knowledge Management, CIKM '16, ACM, New York, NY, 2016, pp. 207–216, <https://doi.org/10.1145/2983323.2983363>.

- [66] X. Liu, A. Nourbakhsh, Q. Li, S. Shah, R. Martin, J. Duprey, Reuters tracer: toward automated news production using large scale social media data, 2017 IEEE International Conference on Big Data (Big Data), 2017, pp. 1483–1493, <https://doi.org/10.1109/BigData.2017.8258082>.
- [67] E. Loper, S. Bird, NLTK: The Natural Language Toolkit, arXiv preprint cs/0205028 (2002).
- [68] T. Lukoianova, V.L. Rubin, Veracity roadmap: is big data objective, truthful and credible? *Advances In Classification Research Online* 24 (2013) 1.
- [69] C. Manning, M. Surdeanu, J. Bauer, J. Finkel, S. Bethard, D. McClosky, The Stanford CoreNLP natural language processing toolkit, Proceedings of 52nd annual meeting of the association for computational linguistics: system demonstrations, 2014, pp. 55–60.
- [70] S.P. Marsh, Formalising Trust as a Computational Concept, Ph.D. thesis University of Stirling, Stirling, United Kingdom, 1994.
- [71] S.E. Middleton, V. Krivcovs, Geoparsing and geosemantics for social media: spatio-temporal grounding of content propagating rumours to support trust and veracity analysis during breaking news, *ACM Transactions on Information Systems* 34 (2016) 1–27 <https://eprints.soton.ac.uk/390820/>.
- [72] T. Mikolov, I. Sutskever, K. Chen, G.S. Corrado, J. Dean, Distributed representations of words and phrases and their compositionality, *Advances in neural information processing systems*, 2013, pp. 3111–3119.
- [73] T. Mitra, E. Gilbert, CREDBANK: a large-scale social media corpus with associated credibility annotations, Proceedings of the Ninth International Conference on Web and Social Media, ICWSM 2015, University of Oxford, Oxford, UK, May 26–29, 2015, 2015, pp. 258–267 <http://www.aaai.org/ocs/index.php/ICWSM/ICWSM15/paper/view/10582>.
- [74] T. Mitra, G.P. Wright, E. Gilbert, A parsimonious language model of social media credibility across disparate events, Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing, ACM, New York, NY, 2017, pp. 126–145, <https://doi.org/10.1145/2998181.2998351>.
- [75] S. Mukherjee, G. Weikum, C. Danescu-Niculescu-Mizil, People on drugs: credibility of user statements in health communities, Proceedings of the 20th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '14, ACM, New York, New York, 2014, pp. 65–74, <https://doi.org/10.1145/2623330.2623714>.
- [76] Y. Namiyama, N. Segawa, Y. Ikegami, K. Kawai, T. Kawabe, S. Tsuruta, High precision credibility analysis of information on Twitter, 2013 International Conference on Signal-Image Technology Internet-Based Systems, 2013, pp. 909–915, <https://doi.org/10.1109/SITIS.2013.148>.
- [77] D.E. O'Leary, Blog mining-review and extensions: 'from each according to his opinion', *Decision Support Systems* 51 (4) (2011) 821–830.
- [78] J. Paryani, A.K. T.K., K.M. George, Entropy-based model for estimating veracity of topics from Tweets, in: N.T. Nguyen, G.A. Papadopoulos, P. Jędrzejowicz, B. Trawinski, G. Vossen (Eds.), *Computational Collective Intelligence*, Springer International Publishing, Cham, 2017, pp. 417–427, https://doi.org/10.1007/978-3-319-67077-5_40.
- [79] J. Pasternack, D. Roth, Latent credibility analysis, Proceedings of the 22nd International Conference on World Wide Web, WWW '13, ACM, Rio de Janeiro, Brazil, 2013, pp. 1009–1020, <https://doi.org/10.1145/2488388.2488476>.
- [80] J.W. Pennebaker, M.E. Francis, R.J. Booth, Linguistic inquiry and word count: LIWC 2001, Mahway: Lawrence Erlbaum Associates 71 (2001) 2001.
- [81] J. Pennington, R. Socher, C. Manning, Glove: global vectors for word representation, Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP), 2014, pp. 1532–1543.
- [82] K. Popat, S. Mukherjee, J. Strötgen, G. Weikum, Credibility assessment of textual claims on the Web, Proceedings of the 25th ACM International on Conference on Information and Knowledge Management, CIKM '16, ACM, New York, NY, 2016, pp. 2173–2178, <https://doi.org/10.1145/2983323.2983661>.
- [83] S. Ramachandramurthy, S. Subramaniam, C. Ramasamy, Distilling big data: refining quality information in the era of yottabytes, *The Scientific World Journal* 2015 (2015).
- [84] P. Rao, A. Katib, C. Kamhoua, K. Kwiat, L. Njilla, Probabilistic inference on Twitter data to discover suspicious users and malicious content, in: 2016 IEEE International Conference on Computer and Information Technology (CIT) (pp. 407–414). doi:10.1109/CIT.2016.29.
- [85] B. Rath, W. Gao, J. Ma, J. Srivastava, From retweet to believability: utilizing trust to identify rumor spreaders on Twitter, Proceedings of the 2017 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining 2017, ASONAM '17, ACM, New York, NY, 2017, pp. 179–186, <https://doi.org/10.1145/3110025.3110121>.
- [86] C. Reuter, M.-A. Kaufhold, R. Steinfurt, Rumors, Fake news and social bots in conflicts and emergencies: towards a model for believability in social media, Proceedings of the 14th ISCRAM Conference - Albi, France, 2017.
- [87] L. Robin, B. Mark, C. Philip, C. Martin, From big noise to big data: toward the verification of large data sets for understanding regional retail flows, *Geographical Analysis* 48 (2016) 59–81, <https://doi.org/10.1111/gean.12081>.
- [88] S.P. Ros, A.P. Canelles, M.G. Pérez, F.G. Mármol, G.M. Pérez, Chasing offensive conduct in social networks: a reputation-based practical approach for Frisber, *ACM Transactions on Internet Technology* 15 (2015) 1–20, <https://doi.org/10.1145/2797139>.
- [89] D. Saez-Trumper, Fake tweet buster: a webtool to identify users promoting fake news on Twitter, Proceedings of the 25th ACM Conference on Hypertext and Social Media, HT '14, ACM, New York, NY, 2014, pp. 316–317, <https://doi.org/10.1145/2631775.2631786>.
- [90] J. Sampson, F. Morstatter, L. Wu, H. Liu, Leveraging the implicit structure within social media for emergent rumor detection, Proceedings of the 25th ACM International on Conference on Information and Knowledge Management, CIKM '16, ACM, New York, NY, 2016, pp. 2377–2382, <https://doi.org/10.1145/2983323.2983697>.
- [91] M. Schroeck, R. Shockey, J. Smart, D. Romero-Morales, P. Tufano, Analytics: the real-world use of big data, *IBM Global Business Services* 12 (2012) 1–20 <https://public.dhe.ibm.com/common/ssi/ecm/gb/en/gbe03519usen/global-business-services-global-business-services-gb-executive-brief-gbe03519usen-20180209.pdf>.
- [92] C. Shao, G.L. Ciampaglia, A. Flammini, F. Menczer, Hoaxy: A platform for tracking online misinformation, Proceedings of the 25th International Conference Companion on World Wide Web, WWW '16 Companion, International World Wide Web Conferences Steering Committee, Republic and Canton of Geneva, Switzerland, 2016, pp. 745–750, <https://doi.org/10.1145/2872518.2890098>.
- [93] B. Shi, T. Weninger, Discriminative predicate path mining for fact checking in knowledge graphs, *Knowledge-Based Systems* 104 (2016) 123–133, <https://doi.org/10.1016/j.knsys.2016.04.015>.
- [94] P. Shiralkar, A. Flammini, F. Menczer, G.L. Ciampaglia, Finding streams in knowledge graphs to support fact checking, 2017 IEEE International Conference on Data Mining (ICDM), 2017, pp. 859–864, <https://doi.org/10.1109/ICDM.2017.105>.
- [95] M. Sirivianos, K. Kim, J.W. Gan, X. Yang, Leveraging social feedback to verify online identity claims, *ACM Transaction Web* 8 (2014) 9:1–9:38, <https://doi.org/10.1145/2543711>.
- [96] D. Snow, Adding a 4th V to BIG Data - Veracity, <http://dsnowondb2.blogspot.se/2012/07/adding-4th-v-to-big-data-veracity.html>, (2012).
- [97] L. Toloşi, A. Tagarev, G. Georgiev, An analysis of event-agnostic features for rumour classification in Twitter, Proceedings of the Tenth International AAAI Conference on Web and Social Media, 2016, pp. 151–158 <https://www.aaai.org/ocs/index.php/ICWSM/ICWSM16/paper/view/13197/12859>.
- [98] M. Torky, R. Baserre, R. Ibrahim, A.E. Hassanien, G. Schaefer, I. Korovin, S.Y. Zhu, Credibility investigation of newsworthy tweets using a visualising Petri net model, 2016 IEEE International Conference on Systems, Man, and Cybernetics (SMC), 2016, pp. 003894–003898, <https://doi.org/10.1109/SMC.2016.7844842>.
- [99] A. Vartapetian, L. Gillam, Deception detection: dependable or defective? *Social Network Analysis and Mining* 4 (2014) 166, <https://doi.org/10.1007/s13278-014-0166-8>.
- [100] M. Viviani, G. Pasi, Credibility in social media: opinions, news, and health information - a survey, *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 7 (2017).
- [101] M. Viviani, G. Pasi, A multi-criteria decision making approach for the assessment of information credibility in social media, in: A. Petrosino, V. Loia, W. Pedrycz (Eds.), *Fuzzy Logic and Soft Computing Applications*, Springer International Publishing, Cham, 2017, pp. 197–207.
- [102] M. Viviani, G. Pasi, Quantifier guided aggregation for the veracity assessment of online reviews, *International Journal of Intelligent Systems* 32 (2017) 481–501.
- [103] S. Vosoughi, M.N. Mohsenvand, D. Roy, Rumor gauge: predicting the veracity of rumors on Twitter, *ACM Transactions on Knowledge Discovery from Data* 11 (2017) 50:1–50:36, <https://doi.org/10.1145/3070644>.
- [104] R.Y. Wang, D.M. Strong, Beyond accuracy: what data quality means to data consumers, *Journal of Management Information Systems* 12 (1996) 5–33.
- [105] X. Wang, X. Luo, H. Liu, Measuring the veracity of web event via uncertainty, *Journal of Systems and Software* 102 (2015) 226–236, <https://doi.org/10.1016/j.jss.2014.07.023>.
- [106] X. Wang, Q.Z. Sheng, X.S. Fang, X. Li, X. Xu, L. Yao, Approximate truth discovery via problem scale reduction, Proceedings of the 24th ACM International on Conference on Information and Knowledge Management, ACM, 2015, pp. 503–512.
- [107] H. Webb, P. Burnap, R. Procter, O. Rana, B.C. Stahl, M. Williams, W. Housley, A. Edwards, M. Jirotk, Digital wildfires: propagation, verification, regulation, and responsible innovation, *ACM Transactions on Information Systems* 34 (2016) 15:1–15:23, <https://doi.org/10.1145/2893478>.
- [108] Website, Apache Flume, <https://flume.apache.org/>, (Visited 2018).
- [109] Website, Apache Hadoop, <https://hadoop.apache.org/>, (Visited 2018).
- [110] Website, Apache Hive, <https://hive.apache.org/>, (Visited 2018).
- [111] Website, Apache Spark, <https://spark.apache.org/>, (Visited 2018).
- [112] Website, NetworkX, <https://networkx.github.io/>, (Visited 2018).
- [113] Website, NeuroLab, <https://pythonhosted.org/neurolab/>, (Visited 2018).
- [114] Website, scikit-learn, <https://scikit-learn.org/stable/index.html>, (Visited 2018).
- [115] S. Wiegand, S.E. Middleton, Veracity and velocity of social media content during breaking news: analysis of November 2015 Paris Shootings, Proceedings of the 25th International Conference Companion on World Wide Web, WWW '16 Companion, International World Wide Web Conferences Steering Committee, Republic and Canton of Geneva, Switzerland, 2016, pp. 751–756, <https://doi.org/10.1145/2872518.2890095>.
- [116] B. Xie, Y. Wang, C. Chen, Y. Xiang, Gatekeeping behavior analysis for information credibility assessment on Weibo, in: J. Chen, V. Puri, C. Su, M. Yung (Eds.), *Network and System Security*, Springer International Publishing, Cham, 2016, pp. 483–496.
- [117] S.-R. Yan, X.-L. Zheng, Y. Wang, W.W. Song, W.-Y. Zhang, A graph-based comprehensive reputation model: exploiting the social context of opinions to enhance trust in social commerce, *Information Sciences* 318 (2015) 51–72, <https://doi.org/10.1016/j.ins.2014.09.036>.
- [118] Y. Yang, K. Niu, Z. He, Exploiting the topology property of social network for rumor detection, 2015 12th International Joint Conference on Computer Science and Software Engineering (JCSSE), 2015, pp. 41–46, <https://doi.org/10.1109/JCSSE.2015.7219767>.
- [119] S. Yao, M.T. Amin, L. Su, S. Hu, S. Li, S. Wang, Y. Zhao, T. Abdelzaher, L. Kaplan,

- C. Aggarwal, A. Yener, Recursive ground truth estimator for social data streams, 2016 15th ACM/IEEE International Conference on Information Processing in Sensor Networks (IPSN), 2016, pp. 1–12, <https://doi.org/10.1109/IPSN.2016.7460719>.
- [120] S. Yao, S. Hu, S. Li, Y. Zhao, L. Su, L. Kaplan, A. Yener, T. Abdelzaher, On source dependency models for reliable social sensing: algorithms and fundamental error bounds, 2016 IEEE 36th International Conference on Distributed Computing Systems (ICDCS), 2016, pp. 467–476, <https://doi.org/10.1109/ICDCS.2016.75>.
- [121] D.Y. Zhang, R. Han, D. Wang, C. Huang, On robust truth discovery in sparse social media sensing, 2016 IEEE International Conference on Big Data (Big Data), 2016, pp. 1076–1081, <https://doi.org/10.1109/BigData.2016.7840710>.
- [122] B. Zhao, D.Z. Sui, True lies in geospatial big data: detecting location spoofing in social media, *Annals of GIS* 23 (2017) 1–14, <https://doi.org/10.1080/19475683.2017.1280536>.
- [123] L. Zhao, T. Hua, C.-T. Lu, R. Chen, A topic-focused trust model for Twitter, *Computer Communications* 76 (2016) 1–11.
- [124] Z. Zhao, P. Resnick, Q. Mei, Enquiring minds: early detection of rumors in social media from enquiry posts, *Proceedings of the 24th International Conference on World Wide Web, International World Wide Web Conferences Steering Committee*, 2015, pp. 1395–1405.
- [125] S. Zhi, Y. Sun, J. Liu, C. Zhang, J. Han, ClaimVerif: A real-time claim verification system using the web and fact databases, *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management, CIKM '17*, ACM, New York, NY, 2017, pp. 2555–2558, <https://doi.org/10.1145/3132847.3133182>.
- [126] Z. Zili, Z. Ziqiong, L. Hengyun, Predictors of the authenticity of Internet health rumours, *Health Information & Libraries Journal* 32 (2015) 195–205, <https://doi.org/10.1111/hir.12115>.
- [127] A. Zubiaga, M. Liakata, R. Procter, G. Wong Sak Hoi, P. Tolmie, Analysing how people orient to and spread rumours in social media by looking at conversational threads, *PLoS One* 11 (2016) e0150989, <https://doi.org/10.1371/journal.pone.0150989>.

Marianela García Lozano is a senior scientist at the Swedish Defence Research Agency (FOI) since 2001. Her research interests include information and knowledge modeling, veracity assessment, software development in distributed systems, web mining, machine learning, and natural language processing. Marianela received her M.Sc. degree in Computer Science and Engineering in 2003 and her Licentiate degree in Electronic and Computer Systems in 2010 from the Royal Institute of Technology (KTH). Marianela's licentiate thesis is on the topic of distributed systems.

Joel Brynielsson is a research director at the Swedish Defence Research Agency (FOI) and an associate professor at the Royal Institute of Technology (KTH). He previously worked as an assistant professor at the Swedish Defence University. Joel is Docent (Habilitation) in Computer Science (2015), and holds a Ph.D. in Computer Science (2006) and an M.Sc. in Computer Science and Engineering (2000) from KTH. His research interests include uncertainty management, information fusion, probabilistic expert systems,

the theory and practice of decision-making, command and control, operations research, game theory, web mining, privacy-preserving data mining, cyber security, and computer security education. He is the author or co-author of more than 150 papers and reports devoted to these subjects.

Ulrik Franke is a senior researcher at Research Institutes of Sweden (RISE). Prior to joining RISE, he was a senior scientist at the Swedish Defence Research Agency (FOI). His research interests include IT service availability, enterprise architecture, cyber insurance, and cyber situational awareness. He received his M.Sc. and Ph.D. degrees in 2007 and 2012, respectively, both from the Royal Institute of Technology (KTH) in Stockholm, Sweden.

Magnus Rosell is a scientist at the Swedish Defence Research Agency (FOI), where he manages a long-term research project on semi-automatic intelligence analysis. He previously worked at Recorded Future, a web intelligence company, where he designed and implemented essential parts of the core engine for extracting events from free text. Magnus holds a Ph.D. in Computer Science (2009) and an M.Sc. in Engineering Physics (2002) from the Royal Institute of Technology (KTH). His research interests include natural language processing, machine learning, data and web mining, decision support, and crisis management.

Edward Tjörnhammar is a Ph.D. candidate at the Royal Institute of Technology (KTH) since 2015 and a research engineer at the Swedish Defence Research Agency (FOI) since 2006. His interests include distributed systems, data mining, and machine learning. Edward received his M.Sc. degree in Computer Science and Engineering in 2012 from the Royal Institute of Technology. Edward's master's thesis is on the topic of distributed systems.

Stefan Varga Swedish Armed Forces, is a professional Ph.D. student (Computer Science) at the Royal Institute of Technology. Major (air force) Varga has worked in the military specialty fields of air surveillance, communications, and intelligence. He is an armed forces military specialist in command and control systems development. Stefan is a graduate from the Advanced Management Program at the Information Resources Management College of the U.S. National Defense University. He is a NATO cyber security professional trained by the U.S. Naval Post Graduate School and the NATO School Oberammergau, Germany. His research interests include cyber security, cyber situational awareness, and decision support.

Vladimir Vlassov is a professor in Computer Systems at the Royal Institute of Technology (KTH) in Stockholm, Sweden. Prior to coming to KTH in 1993, he was an assistant and associate professor at the Electrotechnical University LETI of Saint Petersburg, Russia (1985/1993). He was a visiting scientist at MIT (1998), and a researcher at the University of Massachusetts Amherst (2004). Vladimir has co-authored more than 150 research papers. His research interests include big data analytics, data-intensive computing, autonomic computing, and distributed and parallel computing.