

A fuzzy Dempster–Shafer classifier for detecting Web spams[☆]

Moitrayee Chatterjee^{a,*}, Akbar Siami Namin^b

^a Computer Science Department, New Jersey City University, 2039 John F Kennedy Boulevard, Jersey City, NJ 07305, United States of America

^b Computer Science Department, Texas Tech University, 1012 Boston Ave, Lubbock, TX 79409, United States of America

ARTICLE INFO

CCS Concepts:

Security and privacy
Software security engineering

Keywords:

Dempster–Shafer Theory
Dempster–Shafer Combination
Basic probability assignment
Mass function
Belief
Plausibility
Fuzzy reasoning
Classification
Web spam

ABSTRACT

The Web spam identification problem can be modeled as an instance of the conventional classification problem. Web spams aim at deceiving web crawlers by advertising certain Web pages through elevation of their page rankings superficially than their actual weights. Web spams are intended to produce fraudulent results of web search queries and degenerate the client's experience by directing users to fake Web pages. We present a fuzzy evidence-based methodology for identifying Web spams by which the *spamcity* of web hosts is formulated as a reasoning problem in the presence of uncertainty. However, any classification task intrinsically suffers from incomplete or vague evidence and ambiguity in the class assignment based on evidence. In this work, we combine fuzzy reasoning as the decision maker for selecting the most suitable evidence in a multi-source Dempster–Shafer (DS) based classification algorithm. The introduced approach has the benefit of providing more reliable solution to detect spams without any prior information. The evidence theory offers flexible support that takes into account the multi-dimensional nature of implementation decisions. The experimental results show that the fuzzy reasoning in combination with DS theory, reduces the conflicts among evidence leading to enhanced classification results. The aim of this paper is to describe the potential of fuzzy reasoning and the Dempster–Shafer Theory (DST) as a decision model for the web spams classification problem.

1. Introduction

The term *Web spam*, or *spamdexing*, was proposed in 1996 by Eric Convey [1] and soon was perceived as one of the key security problems of the Internet search engines [2]. As of today, most prominent Search Engine Organizations (SEO) are focused on dealing with harmful data retrieval as this is the most negative impact caused by spams and thus arising research challenges in this area [3].

In addition to breaking down the attributes of query items, spams may ruin the reputation of the underlying search engine causing loss of its users and thus customers. These purposeful malicious Web pages are also the major tool for propagating phishing activities. For example, Eiron et al. [4] investigated 100 million Web pages utilizing Page Rank calculation [5] and found that 11 out of 20 results were adultery sites, but had higher ranks through maneuvering the contents and connections to the pages.

Web spams also drive a search engine organization to squander a lot of computational and memory resources. In 2005, the aggregate budgetary adversities caused by spams were estimated at \$50 billion [6]. The estimate reached \$130 billion in 2009 [7].

We all live in an era of data explosion and retrieval of precise and reliable data is of prime importance. Due to spamming or spamdexing, retrieved online data can become misleading. Hence, we need unceasing adaptation of the methodologies to identify new spams in systems. Technological advances cause introduction of new spams to the existing systems, every now and then and the prevailing detection methodology might fail due to the lack of prior knowledge about the newly introduced spam generated using some novel technologies.

To effectively classify newly introduced spam pages, typical approaches such as the classical Bayesian theorem would require prior knowledge. However, on some circumstances, the prior knowledge may not be available. As an alternative approach, if we utilize the evidence-based approach of Dempster–Shafer Theory (DST), we could measure the improbability of newly introduced spam pages, using the DST's distinctive capability to handle the uncertainty. DST can combine evidence from multiple resources and can thus help in classifying a web page as spam or not spam.

Evidence fusion is an integral part of any typical classification task. The underlying application domain may exhibit various types of uncertainty that render unreliable classification results, for example: (1) evaluation error for deciding class-association, (2) the associated

[☆] This work is supported by the National Science Foundation, United States of America under Grants No: 1516636, 1723765, 1821560.

* Corresponding author.

E-mail address: mchatterjee@njcu.edu (M. Chatterjee).

variables that are restrictive or deficient for probability based estimates, (3) stochastic noise and so on. Since, the data from individual sources are either incomplete or tainted, they introduce uncertainty and/or ambiguity to the application. To avoid the impact of uncertainty on the classification outcome, information from various data sources can be combined. However, when the evidence generated from such applications are fused for classification, the result may not be optimal. Hence, it is necessary to increase the discriminative proficiency of the classifier. One way to achieve better classifier performance is to select the evidence that are more impactful on the classification task.

In most cases, evidence are presented as a measurable quantity and referred to as vectors. Selecting relevant vectors for classification would lead to reduced computational cost and improved accuracy. One such approach for vector selection could be fuzzy rules. The possibility of adapting fuzzy set-based classification models in managing uncertainty-related issues have been intensively studied, and their ability has been demonstrated empirically in numerous applications such as clustering including [8], pattern recognition Pedrycz et al. [9], and classification [10]. However, the fuzzy *If-Then* rules can also be applied to select appropriate features before fusing them for classification. While the Dempster–Shafer Theory can provide powerful mechanism for evidence fusion, fuzzy logic can be utilized for selecting the highly correlated evidence.

The Dempster–Shafer Theory (DST), first introduced by Dempster [11] and then developed by Shafer [12], on the utilization of probabilities with upper and lower limits, is essentially a generalization of Bayes theorem. DST has been extensively studied and advanced in the field of artificial intelligence (AI) and expert systems [13–15], with specific accentuation set on combining evidence from various sources. Multi-characteristic decision making is a field of research in which different schemes have been introduced to make “preferred choices” such as assessment, prioritization, and determination over the possible outcomes that are portrayed by numerous, and generally contradicting criteria. In this paper, we present the fundamental ideas of the DST of evidence for addressing the problem of identification of Web spams, specifying its roots and correlations with the more customary Bayesian hypothesis. DST models the web spam threats by utilizing the bounds of *belief* and *plausibility* of some evidence and to achieve that the DST gathers qualitative evidential support.

Degrees of belief and plausibility computed by DST might be utilized to demonstrate and measure the abstract credibility by incomplete or faulty evidence. Researchers have studied the combination between fuzzy sets and the belief functions of DST, and proposed powerful methods for incorporating them (e.g., [16]).

Our approach of combining fuzzy modeling with DST aims at eliminating the disadvantages ingrained to either of the approach. While the Fuzzy modeling helps to select the most suitable evidence from different sources at different time, DST contributes to integrate the evidence for final class assignment.

Motivation and Contributions

Inspired by the uncertainty modeling capability provided by fuzzy rules and DST, we build an improved classification framework extending our previous work [17]. In this paper, we introduce an improved approach for Web spam classification using the mathematical model of evidence theory (i.e., the Dempster–Shafer Theory (DST)). DST has the unique capabilities to model uncertainty. More specifically, DST is proficient in representation of knowledge even with lack of prior knowledge about all possible outcomes. DST also provides the computational framework for combining evidence from multiple sources resulting in a numerical score for each of the opinions and beliefs. Leveraging the functionalities provided by DST, the work presented in [17] aimed at classifying content and link spams. This paper extends our work [17] on Web spam detection using Dempster–Shafer Theory in accordance with fuzzy rules.

In [17], we applied DST’s rule of combination to fuse the independent evidence. DST’s rule of combination ascribes the mass of

conflicting evidence to the null hypothesis. DST causes unreasonable mass assignment to null hypothesis. So, our classifier in [17] rejected the null hypothesis with irrational mass (mass assignment outside the range of 0 through 1) and suffers from lack of specificity (true negative rate) for the Web spam identification task. To overcome the shortcomings of our previous DST-based classifier [17], in this paper, we propose to combine a fuzzy logic-based evidence selection approach with DST based classifier. This approach reduces the conflict among evidence before fusion and attain better classification results.

For example, let us consider that there are two different sources of evidence s_1 and s_2 . s_1 and s_2 provide three observations supporting the status of a web host h . s_1 supports one observation, that indicates the host is a spam and it also provides another observation that neither specifies the web host as spam or nonspam. While s_2 provides observation that h is a spam. Intuitively, it might seem that the web host should be classified as spam, but while quantifying and assigning mass values to these three observations and combining them using DST’s combination rule, the theory detects conflict and does not classify it as either spam or nonspam. Rather, h is assigned “undecided” status due to the presence of conflicting evidence.

In this paper, we improved upon our previous approach by utilizing fuzzy **If-Then** rules. So that instead of leaving a host as unclassified, we select the observation that is best representative of s_1 ’s evidence. For example, **IF** s_1 is *undecided* about h at time t **AND** s_1 observes h is a *spam* at $t + 1$, **THEN** s_1 supports h is a *spam*. Where **AND** is a fuzzy logical operator. This approach reduces the conflicts among evidence and improves the classification results when the evidence are obtained from multiple sources to combine. The key contributions to our previous work [17] are as follows:

- Employ fuzzy modeling to select the most appropriate candidate for evidence fusion and reduce imprecision due to redundant data. This hybrid approach ensures higher interpretability and better classification performance.
- Use DST for spam classification, as an alternative to the traditional classification approaches, to model the uncertainty reasoning for multi source unsupervised classification task using the functionalities provided by DST.
- Present a detailed evaluation of the proposed model through case studies.

In the rest of the paper, we provide the background and motivation of our work and relevant definitions (Section 2), followed by discussion on related research (Section 3). Then, we move onto the demonstration of the exploitation of the fuzzy modeling and DST approach to show that fuzzy-Dempster Shafer is a dependable and robust model for making a decision from conflicting proofs (Section 4). The general methodology is presented in Section 5. Following this we discuss how this hypothesis can be utilized to display web spam classification problem (Section 6) as well as model user-reliability (Section 7) by means of numerical results. Section 8 concludes the paper and highlights future research directions.

2. Related work

Identifying novel features to build low cost web spam filters that reduce uncertainty in real time has been an active research area. Ntoulas et al. [18] proposed a decision tree based approach for content-based spam detection. The authors extracted feature set based on the common design practices of search engines and applied C4.5 decision tree algorithm for classifying the web pages as spam or normal. Androutsopoulos et al. [19] implemented a Naïve Bayes classifier that trains only on, or a combination of lemmatized words and phrases and non-textual attributes like attachments of, an email message to indicate if it is a spam.

Amitay et al. [20] extracted structural features of websites and applied decision rules using *See5* and *C5.0* [21] to identify the patterns of the websites and detect their functionalities. Their work could identify spams from the clusters of websites with similar fingerprint or structural patterns.

Support Vector Machine (SVM) has been a prevalent classifier choice for spam detection. Kolari et al. [22] created catalog of local and global features for spam blogs or *splogs* and applied SVM classifier and logistic regression to set the *splogs* apart. Abernethy et al. [23] utilized the page features and the structure of the hyperlinks and trained a linear learning model using SVM to classify spams. Ott et al. [24] work is targeted towards identifying “deceptive opinion spam” on products reviews. In their work, they combined computational linguistics with psychological aspects and trained a Naïve Bayes and an SVM classifier to identify the spam reviews.

The machine learning approaches have shown significant results in classifying web spams. However, handling inconsistent and incomplete information on the web with the changing web standards, makes it imperative for the machine learning models to retrain to integrate newer information so that the newer patterns in web structure and information on the web are leveraged in identifying the spam contents. The role of discriminative patterns in the features are crucial for building an inexpensive and efficient classifier.

Fuzzy logic has been extensively used to identify patterns for classification problems. Keller et al. [25] employed fuzzy integration for solving a multi-sensory data fusion for target recognition and for classifying handwritten characters. Their work establishes the efficiency of fuzzy integration for different application areas such as feature combination, classifier data fusion and sensor evidence fusion.

The machine learning based classifiers exhibit superior performance. Depending on the application areas, they require balanced data set to train. This could be a problem in many real time applications like web spam classification, where the systems are faced with inherent uncertainty, incompleteness and ambiguity of information. Dempster–Shafer Theory has widely been used to address the classification in such problem domains.

In their work, Le Hegarat-Masclé et al. [26] illustrated how Dempster–Shafer Theory can be applied for multisource data fusion for classification of remote sensors. Chatterjee and Namin [17] utilized the uncertainty reasoning capabilities of Dempster–Shafer Theory and combined evidence from multiple sources for web spam classification. However, the classifier sensitivity suffered in the presence of conflicting evidence.

Recent work of Xiao [27,28], and Xiao et al. [29] show DST-base approach can be modified to resolve conflicting evidence and reduce uncertainty for better decision making. Boston et al. [30] compared the classification results from DST based classifier with that of fuzzy detectors. Binaghi et al. [31] explored the hybrid fuzzy Dempster Shafer model (FDS) for classification tasks, that overcomes the drawbacks of DST.

Motivated by the existing literature, this paper overcomes the drawback of the classifier presented in [17] by combining fuzzy logic with Dempster–Shafer Theory. This paper implements an inexpensive, fast, efficient that achieve better classification results in web spam detection in the presence of incomplete, conflicting and ambiguous information.

3. Basic principles of Dempster–Shafer evidence theory

The Dempster–Shafer theory, commonly referred as the Evidence Theory, is a generic interpretation of the Bayesian theory of subjective probability. Despite the fact that the Bayesian concept requires probabilities for every result, evidence theory enables us to base “*degrees of belief*” of the outcome of an event, formulated from the probabilities of related events. These *degrees of belief* could possibly have the logical properties of probabilities. However, the extent *beliefs* differ from probabilities will depend upon how closely the two results are

connected. The Dempster–Shafer theory relies upon two considerations: (1) gaining degrees of belief for one result from subjective probabilities for a related result, and (2) Dempster’s standard for unification of such degrees of belief when they rely upon autonomous evidence.

Traditional Bayesian guideline faces trouble when applied with vague information sources. In such cases, we have the “Bayesian dogma of precision”, by which the information pertaining to unverifiable but real parameters, must be exhibited by conventional, unequivocally suggested, probability distribution. DST is a general expansion of Bayesian rule that can utilize the available information very effectively. Furthermore, DST offers several advantages, including the capabilities to transfer probabilistic measures to focal elements, and assigning the probabilistic values to the frame of discernment.

3.1. A formal elaboration on Dempster Shafer Theory (DST)

In this section, we review the fundamental concepts and terminologies related to DST. Suppose θ is the “*frame of discernment*” and it consists of an exhaustive and exclusive set of postulations as $\{h_1, h_2, \dots, h_n\}$. The *basic probability assignment* (bpa) is defined as $m : 2^\theta \rightarrow [0, 1]$ such that 2^θ is the power set of θ . Any subset x of the frame of discernment θ for which $m(x)$ is non-zero is known as a “*focal element*” and denotes the confidence in x .

DST is driven by the *frame of discernment* (θ), which is characterized by the finite set of propositions and suppositions (i.e., perceptions and conceivable outcomes) for the event space. It is the super set of every single imaginable state. For instance, while determining a patient illness, θ would be the set account for all conceivable illnesses. As another example is rolling a dice. While rolling a dice, θ would be a set accounting for all possible sides of the dice. The power set 2^θ is the set of all admissible sub-sets of θ including the empty-set ϕ . For instance, if

$$\theta = \{a, b\} \quad (1)$$

then

$$2^\theta = \{\phi, \{a\}, \{b\}, \theta\} \quad (2)$$

Every component in the power set adds to the representation of the decision making process. For example, the recommendation of “this host is a spam” allocates more inclinations to parts of θ that are spams and contains all and simply the states where that recommendation is legitimate.

3.2. Dempster’s Rule of Combination

The Dempster Rule of Combination (DRC) is the premise of Dempster–Shafer hypothesis. The proportions of *Belief* and *Plausibility* are consequent of basic likelihood assignments. DRC combines two or more distinct, independent evidence sources through *bpa(m)*, i.e., basic probability assignments over mass function. Mass functions are unique representations of the belief functions. These independent sources of evidence are intrinsically the subset of the frame of discernment. DRC provides the standard methods for combining general, potentially infinite, sets of conceivable outcomes.

DRC exploits the conjunctive operation (AND) to combine evidence. Thus, a decision is made through performing conjunctive operations on the opinions acquired from independent evidence sources. For instance, consider m_1 and m_2 as the masses associated with two different observations, then the combined mass or the joint *bpa*($m_{1,2}$) is expressed and computed from the accrued bpa’s of the masses of the two observations as follows:

$$m_{12}(A) = (m_1 \oplus m_2)(A) \quad (3)$$

$$(m_1 \oplus m_2)(A) = \left(\frac{1}{1-K} \right) \sum_{B \cap C = A \neq \phi} m_1(B)m_2(C) \quad (4)$$

$$m_{12}(\phi) = 0 \quad (5)$$

$$K = \sum_{B \cap C = \phi} m_1(B)m_2(C) \quad (6)$$

Where A , B and C are independent bodies of evidence, $\oplus$ is the notation to signify combination of masses, K is the mass related with opposing beliefs and it is determined as the results of the mass of every null intersections. The count of K is commutative and associative, but not idempotent or persistent.

The denominator $1 - K$ in DRC is the normalization factor. It has the effect of comprehensively ignoring the conflict and attributing any mass related with conflict to the invalid set [32,33]. Consequently, ascribing all the conflicting masses to invalid set would yield unreasonable results in spite of high conflict in certain event space. Without the normalization procedure, $\oplus$ becomes proportionate to the Dempster rule. The denominator $1 - K$ enables quick and clear convergence to an outcome.

4. Background on Fuzzy feature selection

For circumstances, in which the data cannot be measured correctly in a quantitative structure, it might be presented as a subjective manner. For instance, consider the cases when we qualify human judgments. We frequently use words from natural language rather than numerical qualities. In different cases, exact quantitative data cannot be expressed because it is possible that it is inaccessible or the expense for its calculation is too high. The utilization of Fuzzy Sets Theory has given excellent outcomes for subjective data modeling [34] and it has demonstrated to be useful in solving numerous issues, such as decision making [35,36] and information retrieval [37]. In this section, we provide the theoretical background on fuzzy aggregation and the application of it to select the features for classifier.

We applied fuzzy If-Then rule to labeled data to reduce the linguistic set (spam, nonspam, undecided). The selected fuzzy sets are optimal for DST combination for a better classification results. A fuzzy If-Then rule takes the form of:

$$\text{If } x_1 \text{ is } A \text{ Then } y_1 \text{ is } B \quad (7)$$

Expression (7) can mathematically be written as:

$$\text{If } A \text{ Then } B \text{ or } A \rightarrow B \quad (8)$$

We define fuzzy rules in the form:

If feature f_1 is e_1 **AND** feature f_n is e_n **Then** $h \rightarrow$ nonspam

Where $e_1, \dots, e_n$ constitutes fuzzy linguistic set (in our case evidence), h represents host and "AND" is the fuzzy logical operator.

4.1. Aggregation of rule

The knowledge-base of the system is bound to generate more than one rule and to achieve the overall rule-based membership. If we consider our knowledge base as a fuzzy system, and we have two evidence x_1 , x_2 for host h_1 . Input x_1 and x_2 have three linguistic variables supporting host status *spam*, *nonspam*, *borderline*. The output h_1 has two linguistic variables *spam*, *nonspam*. The rule-base comprises the rules like:

$$\text{If } x_1 \text{ is } \textit{spam} \text{ AND } x_2 \text{ is } \textit{spam} \text{ Then } h_1 \text{ is } \textit{spam} \quad (9)$$

$$\text{If } x_1 \text{ is } \textit{nonspam} \text{ AND } x_2 \text{ is } \textit{borderline} \text{ Then } h_1 \text{ is } \textit{nonspam} \quad (10)$$

The rules are aggregated and defined by membership function (μ) as:

$$\mu(h) = \max^i [\min[\mu^i(x_1), \mu^i(x_2)]] \quad (11)$$

Where $i = 1, 2, \dots, k$. k is the number of rules. In this case, since we defined two rules in Eq. (9) and Eq. (10) value of k for this knowledge base is 2.

5. Methodology

Our classification algorithm identifies an object under consideration as represented by attributes vector. The classifier looks for common rules in the vector representation of the object for defining a class label.

In our work, the classification process integrated in the algorithm works in five distinct steps: (1) Data pre-processing, (2) Rule translation using fuzzy modeling, (3) Mass assignment, (4) Belief combination, and (5) Normalization. The overview of the steps is shown in Fig. 1.

The steps (3), (4) and (5) constitute the DST based classifier and represented using the $\oplus$ symbol in Fig. 1. Following are the elaboration on each step:

Steps 1: Data Preprocessing. Data pre-processing enables identification of the class labels, that constitutes the DST's frame of discernment. The dataset has four distinct observation provided for any specific host: "spam", "nonspam", "unknown", and "borderline." Our algorithm selects the distinct observations for all the hosts and forms the frame of discernment (θ) for the classification model. Hence, the Eq. (1) for our model becomes: $\theta = \{\textit{nonspam}, \textit{spam}, \textit{borderline}, \textit{unknown}\}$.

Step 2: Rule Translation using Fuzzy Modeling. Fuzzy feature selection incorporates the human intuition to assign numerical values for features using linguistic rules. It decides which information qualities are vital for accurately defining the relationships between classes. Accordingly, the important characteristics must be maintained and the remaining parts of data can be discarded. This step defines the fuzzy rule for selecting the most suitable assessment per host for better classification results.

Step 3: Mass Assignment. The algorithm iterates over each host. For each host, it calculates mass values or basic probability assignments (BPA), for the distinct list of assessor over θ . Once the BPAs for all assessors are computed, the values are then provided to DRC for belief combination.

Step 4: Belief Combination. Once the algorithm combines all the BPAs, for a particular host, from different assessors, that combined value is used to classify the host. This algorithm is administered by the formulae described in Section 3. However, the algorithm requires pre-processed data before it can calculate the mass function for each host and each assessor and combines them iteratively using DRC and decide the host's status.

Step 5: Normalization of Results. The result from DRC needs re-distribution for conflicting masses and is achieved through normalization. Normalization divides the combined masses by the difference between 1 and the mass of the empty set. That way, the non-zero masses are allotted to the empty intersections.

5.1. Experimental setup

We applied the aforementioned fuzzy-rule for evidence selection and combined the evidence using Dempster-Shafer Combination rule to two different case studies, as described in the following section.

For both case studies, we utilized *UK-WEBSPAM2007* dataset.¹ We build a fuzzy-Dempster Shafer model to characterize a system for creating and executing a conventional web spam classifier, which is aimed at improving performance of our previous DST based classifier for Web spam detection [17]. The dataset has been prepared by collective effort of volunteers to aid the research work on Web spam identification. It includes 105.9 million pages and over 3.7 billion links to about 114,529 pages.

Table 1 captures a snapshot of the data. It shows the set of observations on *admin-to-go.co.uk* (represented by numeric host-id 322), from different assessors, at different timestamp. The *Label* column records the observation of the assessors. *Period* column captures the type of assessment.

This data set employed to present two different case studies:

¹ <http://chato.cl/webspam/datasets/uk2007>.

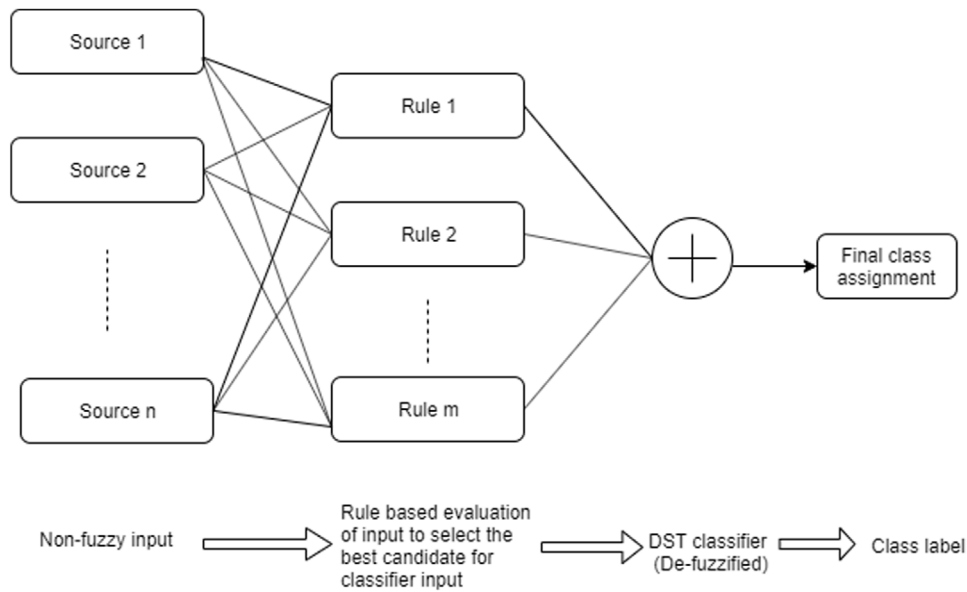


Fig. 1. Schematic diagram of the proposed Fuzzy-DST classifier.

Table 1

Sample data: Assessment on Host *admin-to-go.co.uk*.

Host	Assessor	Label	Timestamp	Period
322	J44	–	1197388571	VIEW
	J44	Spam	1197388733	INITIAL
	J49	–	1197389053	VIEW
	J49	Borderline	1197389155	INITIAL
	J49	–	1198072545	VIEW
	J49	Spam	1198072563	REVISED

(1) We observed that at different timestamp, assessors have changed their assessment of the hosts, leading to conflict, while combining evidence for classification.

(2) We observed that some assessors are more efficient in assessing the hosts. Hence in a second case study, we exploited the capabilities of Dempster Shafer combination framework to quantify the *trustworthiness* of the assessors.

6. Case study I: Time-variant observations

The main goal for the first case study is to leverage the time dependent observations for predicting the host status. We utilized the fuzzy If-Then rules over the linguistic observation set for each of the assessor. This resulted in a reduced number of non-conflicting evidence. The reduced set of evidence is then combined using DRC for final class assignments of the host.

The steps for fuzzy feature selection and DST based classification is shown in Algorithm 1 to elaborate the steps involved in classifying a host as spam or non spam. The algorithm takes the entire data frame containing the list of hosts and their different types of assessment by different assessors at various time period and it produces the host and its associated label as output.

The algorithm iterates for each host, and using the If-Then fuzzy rules, selects the assessment of each host. The assessments of the hosts are further computed to assign qualitative mass values. Finally, all the assessment for a particular host is combined using DRC that results in the final classification label of the host.

To give an insight about the data, Table 2 presents a snapshot of the results for five different hosts. The number of distinct evidence sources are number of assessor providing assessments for each of the web host. The “Actual” label is the actual class assignment of the hosts and the

Table 2

Snapshot of experimental results (Case Study I).

Host	# Distinct evidence sources	Actual		Predicted	
		Score	Label	Score	Label
109belfast.boys-brigade.org.uk	4	0	nonspam	0	nonspam
admin-to-go.co.uk	2	1	spam	1	spam
www.aaxon.org.uk	6	1	spam	1	spam
bradleyhire.co.uk	2	–	undecided	–	undecided
www.manaction.co.uk	7	0	nonspam	0	nonspam

Actual score is a quantitative measure of the level of “spamcity” of the host. Predicted score and Predicted labels are the output produced by the Algorithm 1.

For example, let us consider Algorithm 1 selects the host 322 in Table 1. Then it iterates and identifies the 2 distinct evidence sources (assessor) related to it. Next, using the fuzzy If-Then rule (inside the nested for loop of the algorithm), it selects the most recent observation about this host provided by each assessor. So in this case, the algorithm selects *spam* as an observation by assessor J44 and *spam* by assessor J49 as well. The most recent (or most relevant) observations from each assessor are then translated into mass (bpa) assignment. For basic probability assignment we applied Laplace’s definition of probability [38], we divide the specific observation (spam, nonspam etc.) by the total number of observations of the host. The calculated mass is then provided to the DST based combination rule defined in Eq. (3). Finally, the DST rule based combination generates the combined mass for each class label for the hosts using Eqs. (4), (5) and (6). Our algorithm is designed to select the class label that has a combined mass greater than 50% for a particular host to define its (the host’s) status.

6.1. Confusion matrix

Confusion Matrix is a table, which is usually utilized to portray the performance of a classification model (i.e., “classifier”) on an arrangement of test information for which the results are known. Each column of the matrix represents samples from the predicted class while each row represents samples in the actual class (or vice versa). The name originates from the fact that it makes it easy to see if the system is

ALGORITHM 1: Fuzzy-DST algorithm for Webspam classification.

input : Dataframe containing HostList, AssessorList, LabelsList
output: List of hosts and their respective class labels

$n \leftarrow$ Count of distinct hosts;
 $k \leftarrow$ Count of distinct labels;
for $i \leftarrow \text{HostList}[1]$ **to** $\text{HostList}[n]$ **do**
 $bpa \leftarrow$ empty massList;
 $bca \leftarrow$ empty beliefList // contains k distinct classes and mass functions
 $m \leftarrow$ Count of distinct assessor;
 for $j \leftarrow \text{AssessorList}[1]$ **to** $\text{AssessorList}[m]$ **do**
 if $\exists$ nonspam cases given by Assessor j for Host i **then** $ns_{ij} \leftarrow 1$;
 if $\exists$ spam cases given by Assessor j for Host i **then** $s_{ij} \leftarrow 1$;
 if $\exists$ unknown cases given by Assessor j for Host i **then** $u_{ij} \leftarrow 1$;
 if $\exists$ borderline cases given by Assessor j for Host i **then** $b_{ij} \leftarrow 1$;
 $t_{ij} \leftarrow (ns_{ij} + s_{ij} + u_{ij} + b_{ij})$;
 if $ns_{ij} > 0$ **then**
 append bca_j for nonspam as (ns_{ij}/t_{ij}) computed for assessor j
 else
 bca_j for nonspam $\leftarrow 0$
 end
 if $s_{ij} > 0$ **then**
 append bca_j for spam as (s_{ij}/t_{ij}) computed for assessor j
 else
 bca_j for spam $\leftarrow 0$
 end
 if $b_{ij} > 0$ **then**
 append bca_j for borderline as (b_{ij}/t_{ij}) computed for assessor j
 else
 bca_j for borderline $\leftarrow 0$
 end
 if $u_{ij} > 0$ **then**
 append bca_j for unknown as (u_{ij}/t_{ij}) computed for assessor j
 else
 bca_j for unknown $\leftarrow 0$
 end
 update bpa_i as $(bca_j \oplus bpa_i)$;
 // Combination of mass function using equation (3)
 ;
end
 $bpa_i \leftarrow \text{normalized}(bpa_i)$ // Normalize mass using equation (4), (5), (6)
 $class_i \leftarrow$ Label associated with highest mass for Host i // The label associated with highest mass in bpa is the host's class
end

Table 3
Confusion matrix.

		Prediction		
		Nonspam	Spam	
Actual	Nonspam	TP = 4777	FP = 182	4959
	Spam	FN = 358	TN = 0	358
		5135	182	

“confused” while deciding amongst the two classes (i.e. commonly mis-labeling one as another). This allows more itemized investigation than minor extent of right conjectures (i.e., precision). Table 3 shows the confusion matrix produced by the Fuzzy-DST based model introduced in this paper.

In a typical confusion matrix, the definitive performance is indicated by having the values on non-diagonal cells to be close to zero (i.e., FP and FN to be close to zero. As Table 3 reports, the values of FP and FN are 182 and 358, respectively, which is an indicator of permissible level of performance manifested by the enhanced DST-based model.

Table 4

Relevance measures for UK-WEBSHAM-2007 using Fuzzy-DST Model.

Parameter	Parameter description	Computed value
True Positive (TP)	Correctly Predicted	4777
True Negative (TN)	Incorrectly Predicted	0
False Positive (FP)	Correctly Rejected	182
False Negative (FN)	Incorrectly Rejected	358

Table 5

Relevance measures, their formula and the corresponding values obtained by the proposed model.

Relevance measures	Formula to compute	Computed value
Precision	$TP/(TP + FP)$	0.963
Recall	$TP/(TP + FN)$	0.93
True Negative Rate	$TN/(TN + FP)$	0.9945
Accuracy	$(TP+TN+FP+FN)$	0.9927
F-Score	$2 * \frac{(Precision * Recall)}{(Precision + Recall)}$	0.9962

6.2. Performance evaluation

For classification problems, *precision* and *recall* (or *sensitivity* and *probability of detection*) calculations are performed for relevance measurements. Precision is expressed as a probabilistic value to signify that a randomly selected prediction is correct and it is calculated as the quantity of true positives divided by the whole of true and false positives; whereas, *Recall* is the probabilistic value of correctly predicted instance that has been randomly selected. Recall is characterized as the quantity of true positives divided by the total number of true positives and false negatives. Before further explanation, we present the exposition of TP, TN, FP and FN in Table 4 as well as the number of instances for those parameters achieved through our proposed Fuzzy-DST classification model.

In a typical data classification problem:

- A precision score proximate to 1.0 suggests that predictions made by the proposed model is close to truth; whereas, a recall score proximate to 1.0 suggests that the model could perform all the predictions of the referred classes.
- Specificity (or True Negative Rate) is characterized by probabilistic value of the correctly identified instances that do not belong to the referred class. A prediction with a specificity score proximate to 1.0 implies a prediction was correct reject from a class relationship.
- Accuracy is the statistical measurement of a model's capability to incorporate or separate any instance to a class relationship. The accuracy of a model signifies how perfectly it can identify class memberships.
- F-Score emphasize both precision and recall at the same time by taking their *harmonic mean*. We computed the basic F-score to minimize the impact of either higher number of FP or high FN. F-score values closer to 1.0 suggests a prefect model.

To estimate the model's effectiveness, the authors calculated the precision, recall, true negative rate, accuracy and F-score. The formula for calculating the aforementioned relevance measures along with the values obtained for these relevance measures obtained by this model are presented in Table 5.

The model performance presented in Table 5 is slightly better when compared to Ntoulas et al. [18], which has reported 82.1% recall and 84.2% precision for spam classification, and 97.5% recall and 97.1% precision for non spam identification; whereas, Androutsopoulos et al. [19] reported 92.3% precision and 80.0% recall for spam identification before application of filters.

Table 6
Statistical measures of performance.

Category	χ^2	R^2	Relative error
Spam	0.008	0.998	0.024
Nonspam	0.011	0.915	0.031
Undecided	0.004	0.91	0.022
Total Dataset	0.016	0.97	0.009

6.3. Statistical tests

6.3.1. χ^2 Test

We performed χ^2 , (chi-squared) goodness of fit test, to measure the significance of differences between the actual results and the predictions measured by this method. A very small chi-square statistics signifies a good fit; whereas, higher statistics indicates the prediction and actual values are not related. The χ^2 value is measured by the following formula:

$$\chi^2 = \frac{\sum_i (O_i - E_i)^2}{\sum_i E_i} \quad (12)$$

Where O and E are the predicted and actual values, respectively. The χ^2 results for the dataset and Fuzzy-DST based model are presented in Table 6.

With respect to Table 6, we observe that the proposed fuzzy-DST based model performs very well in detecting spam hosts ($\chi^2 = 0.008$). The approach, also, does a very good job when dealing with detecting non-spam hosts ($\chi^2 = 0.011$) and a good job when it is difficult to decide about a host being spam or not ($\chi^2 = 0.004$). Overall, achieving $\chi^2 = 0.016$ is an indication of the good performance observed by the proposed model.

6.3.2. The R^2 test: The coefficient of determination

The R-squared (R^2) test is also performed on the results. This test represents the proportion of the variance obtained by predictions in comparison to the actual values. The value closer to 1 shows a perfect prediction (i.e., the two variables demonstrate a strong relationship); whereas, a zero value is an indication of a poor prediction. If E_i are the actual values and O_i are the predicted values then R^2 is defined as:

$$R^2 = 1 - \frac{S_{res}}{S_{tot}} \quad (13)$$

Where S_{res} is the *residual sum square* and defined as:

$$S_{res} = \sum_i (E_i - O_i)^2 \quad (14)$$

and S_{tot} is the *total sum square* and defined as:

$$S_{tot} = \sum_i (O_i - \bar{O}_i)^2 \quad (15)$$

where $\bar{O}_i$ is defined as the mean of the n actual data:

$$\bar{O}_i = \frac{1}{n} \sum_{i=1}^n O_i \quad (16)$$

Table 6 also reports the R^2 statistics of the model's results. From the results, it can be seen that this model has the performance close to 1 in the case of spam prediction (0.998 to be exact). Overall, the $R^2 = 0.97$ is an indication of a very good prediction.

6.3.3. Relative error

The relative error or *approximation error* calculates disparity between an actual and the corresponding prediction. Relative error (η) is calculated using the following formula:

$$\eta = 1 - \frac{\vartheta_{approx}}{\vartheta} \quad (17)$$

Table 7
Rank correlations.

	Correlation Γ	Kendall's τ	Spearman's ρ
Spam	0.91	0.91	0.95
Nonspam	0.94	0.92	0.94
Undecided	0.9	0.9	0.9
Average Dataset	0.935	0.9	0.911

Table 8
Illustrative example.

Assessor	Assessment of host		Actual label
	#1281	#10480	
J48	Nonspam	Nonspam	Nonspam
J49	Nonspam	Nonspam	Nonspam

Where ϑ_{approx} is the predicted value and ϑ is the actual. The relative errors between actual and prediction of our work are presented in Table 6. As mentioned earlier, our model overall has only 0.09% errors while predicting spams.

6.3.4. Rank correlations

The rank correlation coefficient computes the level of comparability between actual and predicted outcome. It is computed in order to evaluate the ordinal association between actual and predicted class labels. In this work, the rank correlation between the actual and predicted data is calculated. The highest of Pearson correlation is $\Gamma = 1$, which implies that 100% of the population support the hypothesis, i.e., the actual class and the model predicted class are similarly ranked. On the other hand, a relationship of $\Gamma = 0$ indicates that there is no linear relationship between actual and predicted class. Furthermore, a relationship of $\Gamma = -1$ indicates that the actual and predicted class are inversely correlated. Kendall's τ and Spearman's ρ are two known rank correlations. Table 7 presents the Pearson and rank correlations for the actual and predicted class relationships obtained for the proposed model.

Consistent with the other analyses, the fuzzy-DST model performs satisfactorily in detecting spam hosts (i.e., correlation of 0.95). Overall, the regular and rank correlations for all data are remarkably high (0.935, 0.9, 0.911).

7. Case study II: User reliability modeling

Whenever a system is dependent on a human assessor, then the actions or decision made by them is important to produce a certain output by the system. Therefore, it is important to recognize the reliability of the assessors. The problem is similar to recognizing the correct or faulty behavior of a sensor of a certain system. In this section, we present that the Dempster Shafer combination can be utilized to quantify the reliability of assessors. This quantitative modeling using DST can lead to a potential decision about the assessors and classify them as "trustworthy" or "untrustworthy" from fusing the information.

7.1. Illustrative example

We elaborate the trust computation process using an example. Let us consider two hosts *bristolwest-libdems.org.uk* (host id 1281) and *www.acquireland.co.uk* (host id 10480). 1281 is assessed by J48 and J49, and 10480 is assessed by J20, J40, J48 and J49. If we are to compute the trust for J48 and J49 based on the two hosts status, (as presented in Table 8), we compute the trustworthiness as:

- Capturing opinion's (i.e., evidence) of each assessor.
- Translating evidence into mass assignment.
- Combining the mass using DST for overall class assignment of the host.

- Comparing assessor’s assessment with the host’s label. If they match, then we label the assessor as trustworthy and if they do not match then untrustworthy.

Both J48 and J49 support evidence that 1281 and 10480 are nonspam, which matches their class assignment. So, we can draw conclusion that J48 and J49 are “trustworthy” assessor. Algorithm 2 presents the steps involved in user reliability measure.

ALGORITHM 2: The DRC-based algorithm for user reliability.

input : Dataframe containing HostList, AssessorList, LabelsList

output: List of hosts and their respective class labels

```

n ← Count of distinct hosts;
k ← Count of distinct labels;
for i ← HostList[1] to HostList[n] do
    bpa ← empty massList;
    bca ← empty beliefList // contains k distinct classes and
    mass functions
    m ← Count of distinct assessor;
    for j ← AssessorList[1] to AssessorList[m] do
        if ∃ nonspam cases given by Assessor j for Host i then  $ns_{ij} \leftarrow 1$ ;
        if ∃ spam cases given by Assessor j for Host i then  $s_{ij} \leftarrow 1$ ;
        if ∃ unknown cases given by Assessor j for Host i then  $u_{ij} \leftarrow 1$ ;
        if ∃ borderline cases given by Assessor j for Host i then  $b_{ij} \leftarrow 1$ ;
         $t_{ij} \leftarrow (ns_{ij} + s_{ij} + u_{ij} + b_{ij})$ ;
        if  $ns_{ij} > 0$  then
            append  $bca_j$  for nonspam as  $(ns_{ij}/t_{ij})$  computed for assessor j
        else
             $bca_j$  for nonspam ← 0
        end
        if  $s_{ij} > 0$  then
            append  $bca_j$  for spam as  $(s_{ij}/t_{ij})$  computed for assessor j
        else
             $bca_j$  for spam ← 0
        end
        if  $b_{ij} > 0$  then
            append  $bca_j$  for borderline as  $(b_{ij}/t_{ij})$  computed for assessor j
        else
             $bca_j$  for borderline ← 0
        end
        if  $u_{ij} > 0$  then
            append  $bca_j$  for unknown as  $(u_{ij}/t_{ij})$  computed for assessor j
        else
             $bca_j$  for unknown ← 0
        end
        update  $bpa_i$  as  $(bca_j \oplus bpa_i)$ ;
        // Combination of mass function using equation (3)
    end
     $bpa_i \leftarrow \text{normalized}(bpa_i)$  // Normalize mass using equation (4),
    (5), (6)
     $class_i \leftarrow$  Label associated with highest mass for Host i // The label
    associated with highest mass in bpa is the host’s
    class
    if  $class_i = assessment_j$  then
        Mark assessor_j as “trustworthy”;
    end
end

```

The Algorithm 2, is an extension to the Algorithm 1, with an additional steps to compute the user trust based on their evaluation of the hosts.

7.2. Results

Following the aforementioned steps, and as described in Algorithm 2, we processed reliability of 105 assessor over their assessments and class assignments of 5318 different web hosts. of Table 9 shows the results:

Table 9

The assessors classification.

Assessor reliability	Count
Trustworthy	67
Untrustworthy	15
Undecided	23
Total Dataset	105

Trust plays an important role in deciding which assessor is more trustworthy and which are not. The levels of trust and distrust for independent sources of evidence can efficiently be modeled using DST. The results shown in Table 9 indicates to the fact that DST could set apart the trustworthy and untrustworthy assessors apart while it fail to classify 23% of the assessors due to conflicting evidence.

8. Conclusion and future work

Web spam recognition has experienced a couple of eras: beginning from straightforward content-based strategies to methodologies utilizing complex link mining and client behavior mining systems. Even though web spams are continuously evolving we can summarize [39] the web spam detection methodologies as:

- (1) Recognize examples of spam, i.e., discover pages that contain particular type of spam, and ignore crawling and ranking of such pages.
- (2) Counteract spamming, that is, making particular spamming procedures difficult to utilize. For example, a web crawler could distinguish itself as a normal web program application so as to abstain from cloaking.
- (3) Offset the impact of spamming. Nowadays, web crawlers utilize varieties of the principal ranking techniques that amounts to some level of potency of the spam.

This paper presented a fuzzy Dempster–Shafer Theory for addressing the Web spams identification problem. Both the theories are appealing for their potential for handling uncertainty in the presence of incomplete knowledge. Dempster–Shafer Theory encourages the accumulation of evidence assembled at different degrees of specifics. The fuzzy rule adds another layer of uncertainty handling capabilities to the classifier, eliminating the most irrelevant evidence. In a framework where all the evidence either affirms or dis-affirms singleton speculations, the blend of proof by means of the Dempster–Shafer Theory is computationally straightforward but can suffer from indecision if the evidence are conflicting. Because of DST’s current rule format arrangement, our previous web spam detection model performs satisfactory but failed to actualize the class assignments for large number of instances. Hence, the fuzzy rule based evidence selection is introduced in this work, where the fuzzy rule selects the most fitting evidence for combination with DRC. Moreover, the DRC computation would increment exponentially and become multifaceted in nature if the structure of the evidence changes.

The approach presented in this paper can still become computationally expensive with introduction of multiple evidence sources and multi criteria evidence selections. To overcome that, the proposed integration of fuzzy rule for DST based classifier can be introduced as a multi-layer neural networks. Then, the neural networks can be trained to mitigate the risk of exponential computations. The trained model can then be deployed to make the decision making process tractable. Further work can be coordinated to adjust the hypothesis in the neural networks or by confining the evidence domain.

CRedit authorship contribution statement

Moitrayee Chatterjee: Conceptualization, Methodology, Formal analysis, Software, Validation, Writing - original draft. **Akbar Siami Namin:** Conceptualization, Supervision, Writing - review & editing, Funding acquisition.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

References

- [1] Convey Eric. Porn sneaks way back on web. *The Boston Herald* 1996;28.
- [2] Chen Jingnian, Huang Houkuan, Tian Shengfeng, Qu Youli. Feature selection for text classification with Naive Bayes. *Expert Syst Appl* 2009;36(3):5432–5.
- [3] Fetterly Dennis. Adversarial information retrieval: The manipulation of web content. *ACM Comput Rev* 2007.
- [4] Eiron Nadav, McCurley Kevin S, Tomlin John A. Ranking the web frontier. In: *Proceedings of the 13th international conference on world wide web*; 2004. p. 309–18.
- [5] Page Lawrence, Brin Sergey, Motwani Rajeev, Winograd Terry. The PageRank citation ranking: Bringing order to the web. Technical report, Stanford InfoLab; 1999.
- [6] Jennings Richi. The global economic impact of spam. *Ferris Res* 2005.
- [7] Jennings R. Cost of spam is flattening—our 2009 predictions. *Ferris Res* 2009.
- [8] Gómez-Skarmeta Antonio Fernandez, Delgado Miguel, Vila M Amparo. About the use of fuzzy clustering techniques for fuzzy model identification. *Fuzzy Sets and Systems* 1999;106(2):179–88.
- [9] Pedrycz Witold. Fuzzy sets in pattern recognition: methodology and methods. *Pattern Recognit* 1990;23(1–2):121–46.
- [10] Setnes Magne, Roubos Hans. GA-fuzzy modeling and classification: complexity and performance. *IEEE Trans Fuzzy Syst* 2000;8(5):509–22.
- [11] Dempster Arthur P. A generalization of Bayesian inference. *J R Stat Soc Ser B Stat Methodol* 1968;30(2):205–32.
- [12] Shafer Glenn. A mathematical theory of evidence, vol. 42. Princeton university press; 1976.
- [13] Shafer Glenn, et al. Probability judgment in artificial intelligence and expert systems. *Statist Sci* 1987;2(1):3–16.
- [14] Beynon Malcolm, Cosker Darren, Marshall David. An expert system for multicriteria decision making using Dempster Shafer theory. *Expert Syst Appl* 2001;20(4):357–67.
- [15] Hall DL, Llinas J. An introduction to multisensor data fusion. *Proc IEEE* 1997;85(1):6–23. <http://dx.doi.org/10.1109/5.554205>.
- [16] Yager Ronald R, Filev Dimitar P. Including probabilistic uncertainty in fuzzy logic controller modeling using Dempster-Shafer theory. *IEEE Trans Syst Man Cybern* 1995;25(8):1221–30.
- [17] Chatterjee Moitrayee, Namin Akbar Siami. Detecting web spams using evidence theory. In: *2018 IEEE 42nd annual computer software and applications conference (COMPSAC)*, vol. 2. IEEE; 2018. p. 695–700.
- [18] Ntoulas Alexandros, Najork Marc, Manasse Mark, Fetterly Dennis. Detecting spam web pages through content analysis. In: *Proceedings of the 15th international conference on world wide web*; 2006. p. 83–92.
- [19] Androutsopoulos Ion, Koutsias John, Chandrinou Konstantinos V, Spyropoulos Constantine D. An experimental comparison of naive Bayesian and keyword-based anti-spam filtering with personal e-mail messages. In: *Proceedings of the 23rd annual international ACM SIGIR conference on research and development in information retrieval*; 2000. p. 160–7.
- [20] Amitay Einat, Carmel David, Darlow Adam, Lempel Ronny, Soffer Aya. The connectivity sonar: detecting site functionality by structural patterns. In: *Proceedings of the fourteenth ACM conference on hypertext and hypermedia*; 2003. p. 38–47.
- [21] Quinlan J Ross. Data mining tools See5 and C5. 0. 2004. <http://www.rulequest.com/see5-info.html>.
- [22] Kolari Pranam, Java Akshay, Finin Tim, Oates Tim, Joshi Anupam, et al. Detecting spam blogs: A machine learning approach. In: *Proceedings of the national conference on artificial intelligence*, vol. 21. Menlo Park, CA; Cambridge, MA; London; AAAI Press; MIT Press; 1999; 2006. p. 1351.
- [23] Abernethy Jacob, Chapelle Olivier, Castillo Carlos. Web spam identification through content and hyperlinks. In: *Proceedings of the 4th international workshop on adversarial information retrieval on the web*; 2008. p. 41–4.
- [24] Ott Myle, Choi Yejin, Cardie Claire, Hancock Jeffrey T. Finding deceptive opinion spam by any stretch of the imagination. 2011, arXiv preprint [arXiv:1107.4557](https://arxiv.org/abs/1107.4557).
- [25] Keller James M, Gader Paul, Tahani Hossein, Chiang Jung-Hsien, Mohamed Magdi, et al. Advances in fuzzy integration for pattern recognition. *Fuzzy Sets and Systems* 1994;65(2):273–83.
- [26] Le Hegarat-Masclé Sylvie, Bloch Isabelle, Vidal-Madjar Daniel. Application of Dempster-Shafer evidence theory to unsupervised classification in multisource remote sensing. *IEEE Trans Geosci Remote Sens* 1997;35(4):1018–31.
- [27] Xiao Fuyuan. EFMCDM: Evidential fuzzy multicriteria decision making based on belief entropy. *IEEE Trans Fuzzy Syst* 2019.
- [28] Xiao Fuyuan. CED: A distance for complex mass functions. *IEEE Trans Neural Netw Learn Syst* 2020.
- [29] Xiao Fuyuan, Cao Zehong, Jolfaei Alireza. A novel conflict measurement in decision making and its application in fault diagnosis. *IEEE Trans Fuzzy Syst* 2020.
- [30] Boston J Robert. A signal detection system based on Dempster-Shafer theory and comparison to fuzzy detection. *IEEE Trans Syst Man Cybern C* 2000;30(1):45–51.
- [31] Binaghi Elisabetta, Madella Paolo. Fuzzy Dempster-Shafer reasoning for rule-based classifiers. *Int J Intell Syst* 1999;14(6):559–83.
- [32] Yager Ronald R. On the Dempster-Shafer framework and new combination rules. *Inf Sci* 1987;41(2):93–137.
- [33] Yager Ronald R. Quasi-associative operations in the combination of evidence. *Kybernetes* 1987.
- [34] Zadeh Lotfi A. The concept of a linguistic variable and its application to approximate reasoning—I. *Inf Sci* 1975;8(3):199–249.
- [35] Herrera Francisco, Herrera-Viedma Enrique, Verdegay José L. Choice processes for non-homogeneous group decision making in linguistic setting. *Fuzzy Sets and Systems* 1998;94(3):287–308.
- [36] Ngan Tran Thi, Tuan Tran Manh, Minh Nguyen Hai, Dey Nilanjan, et al. Decision making based on fuzzy aggregation operators for medical diagnosis from dental X-ray images. *J Med Syst* 2016;40(12):280.
- [37] Herrera-Viedma Enrique, López-Herrera Antonio Gabriel, Luque María, Porcel Carlos. A fuzzy linguistic IRS model based on a 2-tuple fuzzy linguistic approach. *Int J Uncertain Fuzziness Knowl-Based Syst* 2007;15(02):225–50.
- [38] Laplace Pierre-Simon. Pierre-Simon Laplace philosophical essay on probabilities: Translated from the fifth french edition of 1825 with notes by the translator, vol. 13. Springer Science & Business Media; 1998.
- [39] Spirin Nikita, Han Jiawei. Survey on web spam detection: principles and algorithms. *ACM SIGKDD Explorat Newslett* 2012;13(2):50–64.