

Article

Effectiveness of Machine Learning Approaches Towards Credibility Assessment of Crowdfunding Projects for Reliable Recommendations

Wafa Shafqat ¹, Yung-Cheol Byun ^{1,*}  and Namje Park ² ¹ Department of Computer Engineering, Jeju National University, Jeju 63243, Korea; wafashafqat@jejunu.ac.kr² Department of Computer Education, Teachers College, Jeju National University, Jeju 63243, Korea; namjepark@jejunu.ac.kr

* Correspondence: ycb@jejunu.ac.kr

Received: 22 October 2020; Accepted: 16 December 2020; Published: 18 December 2020



Abstract: Recommendation systems aim to decipher user interests, preferences, and behavioral patterns automatically. However, it becomes trickier to make the most trustworthy and reliable recommendation to users, especially when their hardest earned money is at risk. The credibility of the recommendation is of magnificent importance in crowdfunding project recommendations. This research work devises a hybrid machine learning-based approach for credible crowdfunding projects' recommendations by wisely incorporating backers' sentiments and other influential features. The proposed model has four modules: a feature extraction module, a hybrid LDA-LSTM (latent Dirichlet allocation and long short-term memory) based latent topics evaluation module, credibility formulation, and recommendation module. The credibility analysis proffers a process of correlating project creator's proficiency, reviewers' sentiments, and their influence to estimate a project's authenticity level that makes our model robust to unauthentic and untrustworthy projects and profiles. The recommendation module selects projects based on the user's interests with the highest credible scores and recommends them. The proposed recommendation method harnesses numeric data and sentiment expressions linked with comments, backers' preferences, profile data, and the creator's credibility for quantitative examination of several alternative projects. The proposed model's evaluation depicts that credibility assessment based on the hybrid machine learning approach contributes efficient results (with 98% accuracy) than existing recommendation models. We have also evaluated our credibility assessment technique on different categories of the projects, i.e., suspended, canceled, delivered, and never delivered projects, and achieved satisfactory outcomes, i.e., 93%, 84%, 58%, and 93%, projects respectively accurately classify into our desired range of credibility.

Keywords: LDA; LSTM; crowdfunding; project recommendation system; optimization; deep learning

1. Introduction

Recommendation systems aim to assist users in daily decision-making processes and are being utilized by perpetually developing online business ventures. Crowdfunding is a platform that plays the role of a venture capitalist for entrepreneurs with creative minds. The recommendation in crowdfunding becomes trickier and complicated than offline businesses due to many challenges, such as information scrutiny and less proficient investors [1]. Moreover, online data is less reliable and inclined to alteration, making it difficult for investors to rely on a new business idea [2]. Therefore, the credibility assessment of crowdfunding projects becomes an absolute necessity to mitigate the risk of fraud. Crowdfunding is undoubtedly becoming popular as a study [3] shows that approximately 6,445,080 fundraising campaigns were hosted in 2019, with gaming companies being the most successful in

generating profit. It is also predicted that the growth of transaction rate annually will reach up to 5.8%, resulting in a total amount of 1180.5 million dollars by 2024 [4]. Despite the remarkable development and inexhaustible possibilities that crowdfunding provides, the challenges and risks of trust, reliability, transparency, etc., are equally daunting, seemingly mounting ones. This study attempts to plug that credibility gap by analyzing and filtering key players' features towards trust-building among investors and the creator. In addition to basic campaign features, we also concentrate on the comments section of the crowdfunding sites, which plays a significant role in fighting against the considerable risks of deceitful online events.

In this paper, we propose a credibility formulation for project recommendations based on a hybrid model. Our proposed architecture has four modules: a text analysis module for project comments, a deep learning module, a credibility estimation module, and a recommendation module. The input data is based on comments-related features and other project-related features. The comment-related features are derived from a comment section of a campaign where users leave their feedback about particular topics; other project-related features include project funding goal, creator's experience, number of images/videos/updates/comments, etc. We perform tokenization, streaming, stop words removal, and data normalization in the data preprocessing layer. In the next step, we perform parameters estimation and topic modeling through latent Dirichlet allocation (LDA). LDA clusters the words with the same meaning in a single topic and is passed to the long short-term memory (LSTM) layer, where the input data consists of the word embeddings, topic embeddings per time-step, and topic distributions.

The LSTM is then trained against new comments to generate sentiments, i.e., positive, negative, and neutral sentiments. As the output from the LSTM layer, we get topic class, accuracy, and project classification. These results are used in the recommendation module to equate and analyze the product's credibility through sentiment score and authenticity calculation. For optimization, we compute the objective function that has both maximization and minimization function. We also formulate the authenticity score and credibility of the project. Through our developed Equation, we get the credibility score and recommended product as our output. We have developed various equations step by step to build the recommendation system considering the critical aspects of positive and negative comments from users and mentioned how authenticity and credibility inter-related for project evaluation are. Our results show that the proposed approach is feasible for all scenarios and achieves high accuracy in recommendation result and authenticity level evaluation and low error rate. The proposed model's evaluation depicts that credibility assessment based on the hybrid machine learning approach contributes efficient results (with 98% accuracy) than existing recommendation models. We have also evaluated our credibility assessment technique on different categories of the projects and achieved satisfactory outcomes. As 95%, 89%, 58%, and 96% of the projects from their respective categories, i.e., suspended, canceled, delivered, and never delivered projects categories were accurately classified into our desired range of credibility.

The rest of this paper includes related works in Section 2, data in Section 3, the proposed method in Section 4, results in Section 5, and conclusion in Section 6.

2. Related Works

In this era of internet and digitalization, an enormous amount of textual data is generated at a high rate. Text data analysis applications are widespread, starting from customer review analysis to extracting and finding a large dataset's hidden meaning. Blei proposes a novel approach to recognize the topics, which ultimately led to sentiments classification, documents classification, and unlocked relatively many assessment prospects for textual data [5]. Topic models are of crucial importance for the illustration of discrete data and are used in different research fields such as medical sciences [6], software engineering [7], geography [8], and political sciences [9], etc. There are many topic modeling techniques; each has its strengths and limitations. The most frequently used approaches include

latent semantic analysis (LSA) [10], probabilistic latent semantic analysis (PLSA) [11], latent Dirichlet allocation (LDA) [12], and correlated topic model (CTM) [13].

LSA's primary focus is to generate different representations of texts based on vectors to create semantic content [10,14]. These vector representations are designed to choose related words by computing the similarity among text data. LSA has many applications such as keyword matching, word quality assessment, power collaborative learning, guidance in career choices, making optimal teams [15], reduction of dimensions [16], and identification of research trends [17]. PLSA was introduced to fix the limitations of LSA [18]. It has many implications, including the differentiation of the words with several meanings and clustering of words that share similar contexts [19]. In [20], PLSA is introduced as an aspect model based on a latent variable responsible for linking observations with unseen class variables. In addition to introducing advancements in LSA, PLSA has many other applications, including recommender systems and computer vision [21–23]. LDA model aims to overcome the limitations of LSA and PLSA in capturing the exchangeability of document words.

LDA being an unsupervised approach for topic modeling, has recently become very popular, mainly for topic discovery in a large corpus. In [24], LDA is used for text mining that is based on Bayesian topic models. LDA is also a generative and probabilistic model that attempts to imitate the writing task. Therefore, it attempts to produce a document if a topic is given. There is a variety of LDA based algorithms used in different domains, including author–topic analysis [25], LDA based bioinformatics [26], temporal text mining [27], supervised topic models, and latent co-clustering, etc. In simple words, LDA's fundamental idea is that each document is represented as a mixture of topics. Each topic represents a discrete probability distribution reflecting each word's likelihood to occur in a specific topic. Therefore, a document is described as probability distributions of words in each topic. Certainly, LDA has many applications such as role discovery [28], emotion topic [29], automatic grading of essays [30], and email filtering [31], etc. Biterm topic modeling (BTM) is a topic modeling approach over short texts. These topic modeling methods are becoming a significant job because of the pervasiveness of the short texts available on the internet. BTM is also used to discover discriminative and comprehensible latent topics from short text [32].

Recommendation system (RS) is an intelligent system that suggests items to users that might interest them. Some of the practical example applications of RSs include movie, book, tourist spot recommendations, etc. It is a point of amusement to discover how, “People you may know” feature on Facebook or LinkedIn. In a personalized RS, users get item suggestions based on their past behaviors and social networks-based interpersonal relationships. There are four categories of personalized recommendation systems based on the approach, content-based filtering, collaborative filtering (CF), knowledge-based filtering, and hybrid. A novel clustering method is proposed in [33] that uses the latent class regression model as a baseline model, which considers both the general ratings and textual reviews. In [34], a system that assesses a user's location as an attribute of a recommendation system is proposed. A recommendation method is suggested in [35], which investigates the difference between user feedback to discover a customer's preferences. It considers user ratings and focuses on the sparsity issue of the data. In [36], a CF method is being suggested that uses ratings of different items and feedbacks on various social networks such as Twitter.

A convolutional neural network (CNN) devised by Krizhevsky et al. is referred to as deep CNN [37] that learned 1000 semantic concepts for training based on ImageNet Large Scale Visual Recognition Challenge (ILSVRC) 2012 dataset. Deep CNN proposed by [38] is not suitable for the clothing domain. Therefore, fully connected layers have been included between the seventh and eighth layers to fill the gap between semantics and mid-level features. In [39], the author built a CNN model for the classification of the music genre. This model comprises two convolutional layers, one fully connected layer, and two max-pooling layers. Further, there are ten softmax units with a logistic regression layer to classify the music genre.

Xin Liu et al. [40] used a fusion of matrix factorization and LDA to build a web content-based recommendation model that recommends to the user's fake credibility information to analyze their

reaction and improve the model. Schwarz et al. [41] considered measuring webpage popularity, page rank metric, and popularity of a web page to assess a user's web credibility. Studies [42,43] have shown that varied linguistic features, writing styles, and project creators' patterns reveal how communication impacts crowdfunding projects' success. Generally, crowdfunding success is predicted by extracting LDA's semantic features and then by feature selection and data mining [44]. Most of the literature studies are focused on simple embeddings and have not considered using words plus topic embeddings for LSTM training. We have incorporated these embeddings to make more meaningful recommendations that are highly authentic and trustworthy. Moreover, our methodology is novel because we focus on crowdfunding comments to analyze and formulate their impact on the crowdfunding project's credibility. In other sections, we have presented how we have overcome the shortcomings of the literature studies to build a system that considers several factors to recommend credible crowdfunding projects.

3. Proposed Credibility Formulation for Project Recommendation Based on Hybrid Model

This section elaborates the credibility assessment formulation based on learned topics from text and other vital features. The proposed approach uses LDA and LSTM as underlying methods for the credibility assessment process. The overall procedure is divided into multiple tasks as shown in Figure 1, which primarily includes data collection, features selection, text data analysis for topics discovery, topic classification, and formulation for credibility estimation and recommendations. Each task is elaborated separately in the following subsections.

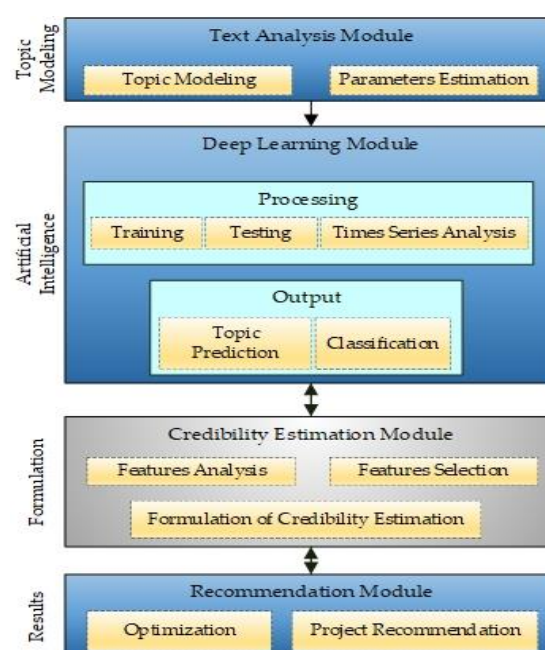


Figure 1. Layered view of the proposed approach.

3.1. Input Data

Each crowdfunding project is rich with the information and data it has in terms of the project's data and user's profile data. The project-based data includes many elements such as project description, duration, number of backers, numbers of comments, and project's success status, etc. Similarly, the user's profile data is related to the creator's information such as name, ID, linked social networks, number of friends, number of created or backed projects, etc. We are mainly focusing on the comments section of a project as comments reveal a lot of information about a project's status and its creator's behavior through backers' experiences. In addition to features extracted from the comment section, we have also focused on the statistical features such as the number of comments, updates,

pledged amount, number of backers, etc. We also recorded time delay between different posts to track the project creator's activities. We have collected data from a famous reward-based crowdfunding platform, i.e., Kickstarter. Its mission is to bring creative projects to life that belong to 15 different categories and eight sections: arts, comics & illustration, design and tech, film, food and craft, games, music, and publishing. Table 1 describes the data in detail. There is no limitation on the length of a comment.

Table 1. Input data characteristics.

Data Characteristics	Specifications
Total number of projects	600
Total number of comments (before cleaning)	645,251
Total number of comments (after cleaning)	504,184
Average comments per project	841
Training data	70%
Test data	30%

The temporal patterns of a review, interaction patterns between project backers and creators, the average timeline required from the proposal stage to the approval state varies for every project. The importance of social link and user description in assessing credibility is described in later sections.

3.2. Data Pre-Processing

This unit is in charge of several jobs. It first tokenizes the comments into multiple words. Then these tokenized words are passed through the cleansing unit. Here, all the punctuations are removed, and words are passed through the stemming unit. This unit lower cases all the words and convert each word to its root. (e.g., working is replaced with work). Then, we filter out all the stop words. Stop words are used in any language for grammatical reasons (e.g., a, an, is, etc.) after this processing comment is passed to LDA for further processing.

Then we label those clusters into meaningful topics. Therefore, after LDA, we have topic distributions representing the probability of a topic in a document and word distributions representing the probability of a word in a topic. These probability distributions are then prepared as an LSTM input. For LSTM, the word embedding and topic embedding are also generated. These embedding against each new input comment are trained in an LSTM network. The topic classes are distributed in three basic types of sentiments, i.e., positive sentiments, negative sentiments, and neutral. Therefore, the percentage of each topic class is calculated and assigned a sentiment class accordingly.

3.3. LDA and LSTM Based Hybrid Model

The preprocessed data is passed to the hybrid module responsible for the data's primary processing. Here, data is first handed over to the topic modeling process, where LDA is applied. The number of topics and Dirichlet parameters is initiated. LDA generates clusters of words that have the highest similarity.

A. Topic Discovery and Classification

We used LDA for topics discovery in the comments data. We used comments to discover topics as the comments left by backers can present their emotions, feelings, thoughts, and experiences related to the project. Therefore, reviews or comments are powerful enough to shape other's decisions. Figure 2 elaborates on the overall process of LDA. Each project's input data is in the form of comments; each comment is treated as one document that results in N documents per project. Data preprocessing is a crucial and vital part of any NLP technique; therefore, we perform essential yet necessary

preprocessing tasks on input data such as removing quotes, stop words, and URLs, tokenization and stemming, etc. Data preprocessing has been influenced by the paper [42], which helps us work with short-texts and proves that LDA works with equivalent efficiency. Once the data is preprocessed, LDA is performed where we set the Dirichlet parameters to calculate desired distributions. We present the output in terms of probability distributions of topics over projects and word distributions over documents. This output is then used as input for the next step, where we use these discovered topics as ground truth and train our LSTM model to predict the topic class of new comments. All the learned topics are divided into different classes, and each class depicts a specific sentiment.

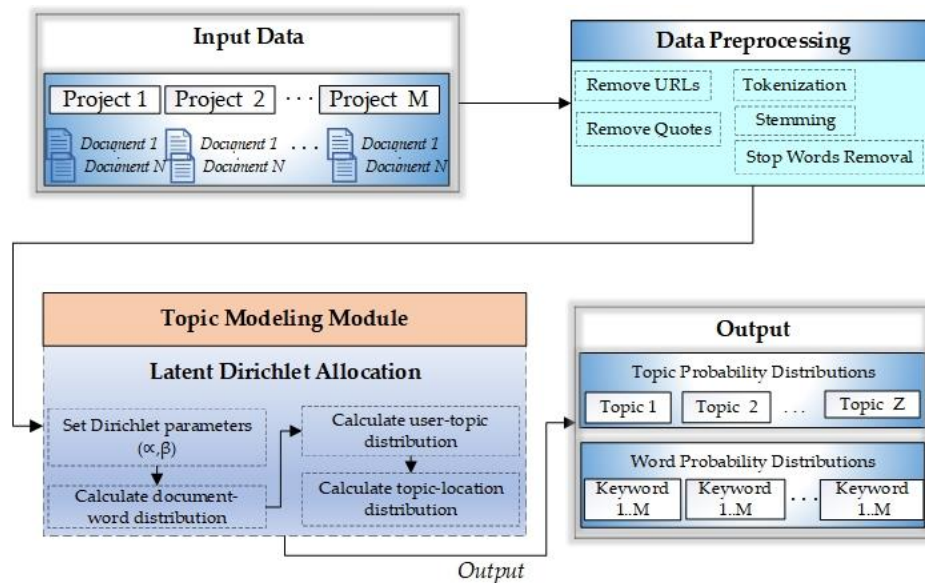


Figure 2. Latent Dirichlet allocation (LDA) process.

B. Deep Learning using LSTM

We are using a bidirectional LSTM for capturing the context dependencies concerning time. A bidirectional LSTM is analyzed in its natural order and inverse order when an input is provided to capture maximum dependencies within the data. We are using a 128-unit LSTM (bidirectional) for this purpose. The preprocessing module's input is passed to an embedding layer that converts the input into a 64-bit vector representation. This representation is then processed by the LSTM layer, which is then connected to a dense layer. This layer helps to consolidate the LSTM results. The output layer gives the probability distribution of the output category. The detailed architecture of the proposed approach is presented in Figure 3.

3.4. Project Credibility Estimation

In this section, we present a detailed explanation of the credibility module. The overall process of deriving formulas steps by steps to estimate the credibility of a project is delivered. Trust is an ultimate significant element in any domain that helps to gain the customer's confidence. It is valid for e-commerce sites and online social networks, as well. Therefore, multiple trust-aware recommender systems are being proposed that adopt user's trust statements and their personal or profile data to improve the quality of recommendations considerably.

As we target crowdfunding projects, we aim to formulate an equation to calculate any project's credibility before recommending it to a user. A highly credible recommendation is a project that most likely reflects the user-defined interests and categories with higher chances of its delivery. It must also reflect the lowest probability of factors that can disturb the project's trustworthiness, such as communication delays and less frequent updates, etc. A credible project can precisely be defined as a

project with the maximum likelihood of completing and delivering to the backers within the promised period. Various factors are associated with a project's credibility; we define and link a documents' credibility with its estimated authenticity score range. A project's authenticity is a multi-fold view of different and latent aspects, such as latent aspects of a creator's profile and all his or her external social links. It also involves the frequency of account usage and updates from creators. In other words, keeping the backers up to date with each development or progress in the project can earn more credibility points. In addition to that, factors such as the most frequent keywords used, promises related to product delivery or rewards delivery, and investors' sentiments are also crucial. These sentiments of backers are discovered during the LDA process to find latent topics in their comments. There can be multiple topics in a document, and each topic represents a particular class of sentiments. As shown in Table 2 [45], we identified 12 topic classes labeled Topic-1 to Topic-12. The number of topics was varied between 2 and 30 during the experiments to find the optimal number of topics.

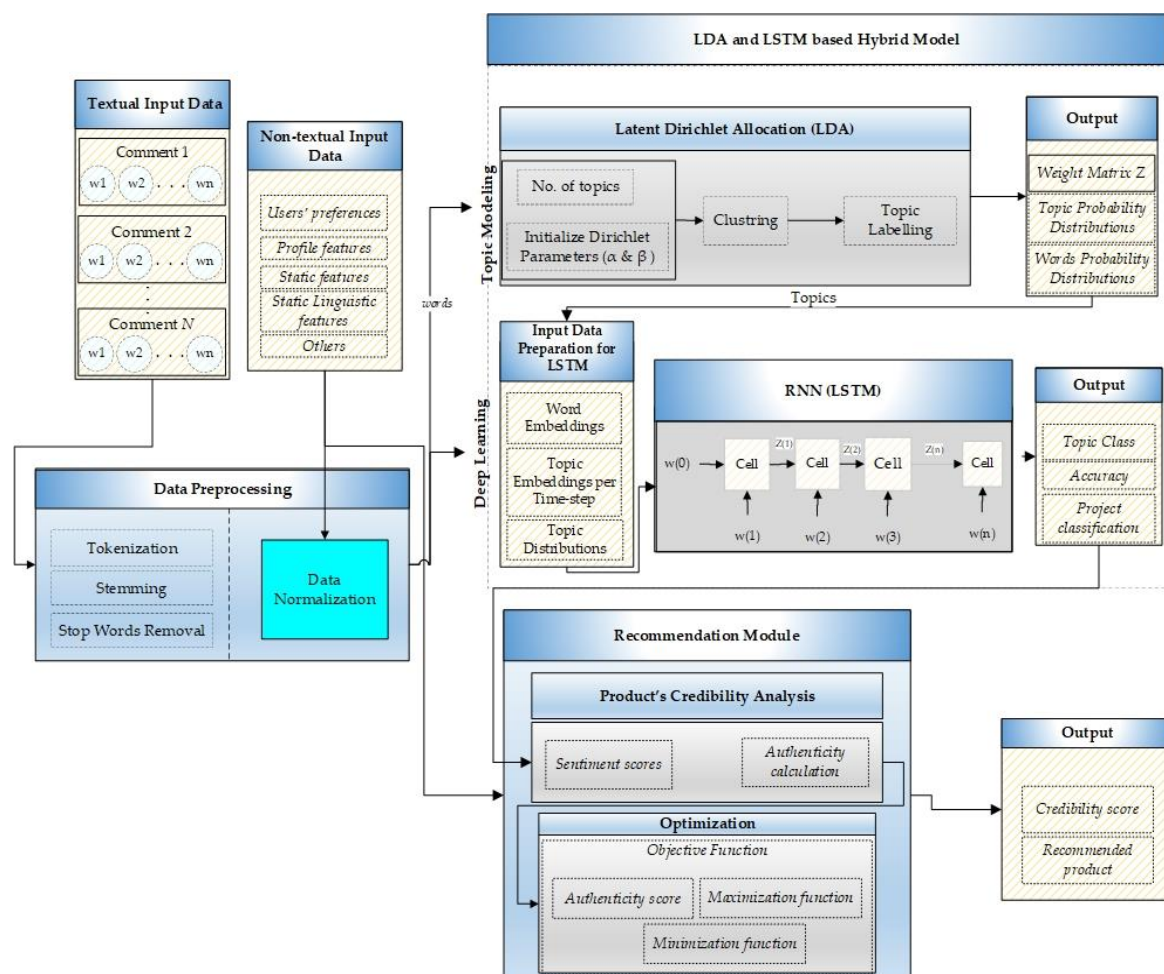


Figure 3. Detailed view of the proposed approach.

The coherence score was increasing as the number of topics was growing. We selected and evaluated the topics based on the coherence score before flattening out, i.e., 12 topics. After training LSTM, the classification of each comment is done into one of these topic classes. We have divided these sentiment classes into three categories, and this division is customized based on the problem, i.e., credibility assessment. These categories are referred to as A, B, and C. Category A is responsible for extremely negative comments, which is represented by Topic-4 to Topic-7; category B means negative reviews, which is characterized by Topic-1 to Topic-3; and category C is representing positive or neutral reviews which are represented by Topic-8 to Topic-12. More emphasis is laid on the negative

comments because the negative comments and reviews significantly impact the viewer's mind and decision-making process than positive comments regarding credibility or trust. Therefore, we divided the negative comments into extremely unfavorable class A and negative class B.

Table 2. Topic classes identified using LDA analysis.

Topic Classes	Labels	Popular Words
Topic-1	Waiting for update	Wait, waiting, update, posted, long, silence
Topic-2	Refund inquiry	Return, refund, back, need, please, amount, invoke, demand, reimbursement, want, request, money
Topic-3	Rewards inquiry	Reward, approval, desired, pledge, invoke, request
Topic-4	Legal actions	Legitimacy, filed, report, case, lied, legally, sue, petition, criminal, complaint, signed
Topic-5	Communication gap	Reply, communication, delay, years, last, please, information, progress, since, silent
Topic-6	Product never received	Never, product, receive, deliver, released
Topic-7	Fraud	Ridiculous, Disappointed, Scamming, Scam, fraud, fraudster, ran, product, delete, fail, immoral, thieves
Topic-8	Shipment Information	Ship, delivery, address, time, date, weight, charges, replacement, cancel, received
Topic-9	Product experience	Working, status, advance, easy, difficult, understand, battery, condition
Topic-10	Product description	Weight, color, length, height, memory, soft, dark, light, colorful, single, quantity
Topic-11	Excited about product	Loved, tremendous, excited, awesome, excellent, super, happy, lovely
Topic-12	Product received	Received, fast, today, time, days, quick, late, ago, early, delivered

To evaluate our selected topics, we measured the agreement between two raters using Cohen's Kappa coefficient [46] and followed the process mentioned in [47] to assess our LDA model. Two students (student A and student B) from different laboratories who were unaware of our proposed methodology and had no prior knowledge about the list of LDA topics were requested to extract topics from 250 sampled reviews.

Student A and student B were not allowed to communicate or discuss their thought process behind labeling each review. Student A and student B could identify 9 and 11 topics, respectively. Student A had seven topics in common with LDA, whereas student B had ten topics common with LDA. Among all the topics, we selected six topics that were most common among the two students' topics to measure our LDA model's reliability, as shown in Table 3. As we can see from Table 3, student A and student B have a high degree of agreement for all six topics. The LDA model and respective students' contract is also relatively high, as indicated by the Kappa coefficient.

Category A is for extremely negative comments and severe nature and typically reflects anger by filing lawsuits or complaints. Category B is for relatively simple and generic negative comments that reflect emotions of sadness or disappointment. The classification is based on the nature of malicious content. All other comments belong to category C. The purpose behind this arrangement with more emphasis upon negative comments is the underlying prominence or impact of the malicious content on a product's credibility. Table 4 summarizes the parameters used for authenticity measures with their definitions and notations. In addition to sentiments, we have also included other relevant and impactful features such as readability of content referred to as readScore, the existence of a profile picture, etc.

Table 3. Reliability Assessment of LDA model using Cohen’s Kappa coefficient.

Topic Classes	Student A-LDA		Student B-LDA		Student A–Student B	
	Kappa	Overlap	Kappa	Overlap	Kappa	Overlap
Topic-1	0.67 (Moderate)	207 (11)	0.65 (Moderate)	203 (10)	0.69 (Moderate)	223 (22)
Topic-2	0.69 (Moderate)	219 (17)	0.63 (Moderate)	211 (18)	0.79 (Moderate)	234 (24)
Topic-5	0.81 (High)	223 (19)	0.78 (Moderate)	219 (17)	0.85 (High)	239 (25)
Topic-6	0.83 (High)	227 (21)	0.81 (High)	226 (19)	0.89 (High)	238 (21)
Topic-7	0.82 (High)	226 (19)	0.80 (High)	227 (15)	0.87 (High)	230 (20)
Topic-10	0.76 (Moderate)	220 (20)	0.83 (High)	227 (21)	0.85 (High)	236 (19)

Table 4. Definitions of the parameters of authenticity.

No.	Authenticity Parameters	Description	Notations
Content related Features			
1	Sentiments (-ve)	Percentage of comments that belong to class A	$eNegA$
2	Sentiments (-ve)	Percentage of comments that belong to class B	$NegB$
3	Sentiments (+ve)	Percentage of comments that belong to class C	$PosC$
4	Readability score	This score reflects the clarity of content, i.e., how easy or difficult it is to understand.	$read_{Score}$
Profile related Features			
5	Profile picture	To check profile picture exists or not	$picY$
6	Social Links	The total number of links (external or social network links, e.g., Facebook, Twitter, etc.),	$Links_{Ext}$
7	Communication delay	The time delay between two consecutive posts (update or comment) by the creator	$delay_{comm}$

Hence, by incorporating all the factors mentioned above, we have formulated an equation that helps calculate a given project’s authenticity. To figure the authenticity of a project, it must first fulfill the eligibility criteria given in Equation (1). Once a project passes the eligibility criteria, Equation (2) is used to calculate the authenticity of it. The eligibility criteria are based on a project’s content and partially on the profile associated features in Equation (1).

$$Eligibility_{criteria} = -(eNegA + \alpha * picY) \quad (1)$$

Here, α represents the weight associated with the existence of a profile picture. The weightage assigned to $picY$ is lower than the weightage of $eNegA$ because of the level of impact asserted by each parameter. The value of α is set to 0.4. From the above Equation, we define the ranges for both the parameters.

$$eNegA = \begin{cases} 0 & \text{if } A \leq 0 \\ 1 & \text{if } A > 0 \end{cases} \quad (2)$$

Similarly,

$$picY = \begin{cases} 0 & \text{if profile picture} = \text{Exists} \\ 1 & \text{if profile picture} = \text{Does not exists} \end{cases} \quad (3)$$

Hence from above Equations (2) and (3), we have

$$Eligibility_{project} = \begin{cases} 0 & \text{favorable} \\ < 0 \geq -0.4 & \text{can be considered} \\ < -0.4 & \text{unfavorable} \end{cases} \quad (4)$$

Therefore, based on Equation (1) and following the conditions in Equation (4), we can list all possible scenarios of eligibility in Table 5. The content in *eNegA* is extremely unfavorable as one can sense fears, suspicion, and frustrations in it. Therefore, this category is handled independently to alleviate the probability of any unreliable recommendation. For a reliable project, it must be free from any of the comments in *eNegA* category. Thus, we used this to set our eligibility criteria. The objective function targets getting the maximum percentage of positive comments, i.e., category C. It also targets to get the maximum number of social links of the project's creator.

Table 5. All possible cases for a project's eligibility criteria.

<i>eNegA</i>	<i>picY</i>	Explanation	Eligibility $Eligibility_{project} = -(eNegA + \alpha \cdot picY)$ ($\alpha = 0.4$)
0	0	The author's profile picture exists and no comment belongs to category A of comments	0 (favorable)
0	1	The author's profile picture does not exist and no comment belongs to category A of comments	−0.4 (can be considered)
1	0	The author's profile picture exists and Some comments belong to category A of comments	−1 (unfavorable)
1	1	The author's profile picture does not exist and Some comments belong to category A of comments	−1.4 (unfavorable)

In crowdfunding, a backer's faith and confidence rely on the content authenticity and creator's limpidity. Therefore, these aspects are fundamental to a project's success. Table 4 shows that the factor $delay_{comm}$ is one prime feature of the project, representing a creator's communication styles such as his updates and comments. This feature, $delay_{comm}$ can be defined as the average time gap between any consecutive posts by the project creator in an update or a comment. It shows the communication rate of a project creator towards the development of a project. Due to the impact of $delay_{comm}$, the project's authenticity will be damaged if the communication delay upsurges.

After observing and estimating all the relevant features, all the values are normalized between 0 to 1. Here, 0 represents the least authentic feature, and 1 illustrates the highly authentic feature. In other words, these values depict the trustworthiness of a project. Equation (5) below describes this relationship, i.e., the higher the authenticity is, the higher the reliability of a project turns out.

$$Authenticity_{project} \propto Credibility_{project} \quad (5)$$

As a result, a project has different credibility levels, i.e., extremely low, low, normal, high, and extremely high credibility. Each credibility level falls into another degree of authenticity range. The extremely low and low credible projects have higher chances of getting forged. It means the projects with lower credibility levels have the utmost possibilities of fighting with non-payments, no communication or communication delays, delays in posts by the creator in the form of updates or comments, and late or no deliveries. Therefore, such projects are not favorable to be recommended to backers to invest in. Instead, a project with a higher credibility level (high or extremely high credibility) is undoubtedly a profitable project recommended to backers. It has the maximum probability of on-time delivery with more consistent patterns of communication throughout its duration.

For any recommendation system, the percentage of positive and negative reviews is pre-eminent as it reflects a user's attitude towards a product. Therefore, we assess the following points wisely:

1. A fundamental requirement for a product to be reliable is to have a maximum percentage of positive comments and a minimum negative comments rate. A product with a relatively high number of negative reviews becomes less favorable. Therefore, Equations (6) and (7) represent the relationship of comments with authenticity.

$$\text{Authenticity} \propto [\text{PosCi}] \quad (6)$$

and

$$\text{Authenticity} \propto [1/\text{NegBi}] \quad (7)$$

where the percentage of positive and negative comments is referred to as PosCi and NegBi , respectively.

2. The accessibility of social and profile information such as profile links, display pictures, number of friends or followers, etc., are persuasive and compelling elements for a profile's credibility. Thus, the more a project creator shares personal and relevant information, the easier it gets to earn trust. Therefore, we can say,

$$\text{Authenticity} \propto [\text{Links}_{\text{Ext}}] \quad (8)$$

In the above Equation (8), $\text{Links}_{\text{Ext}}$ is the number of links a person provides for his/her external social media networks, such as Facebook, Twitter, etc.

3. The clarity of speech also plays a vital role in trust development. If the content is easy to follow and understand, a user will easily connect and comprehend it. It helps diminish the misunderstandings, and the confidence level of the reader increases. Therefore,

$$\text{Authenticity} \propto [1/\text{read}_{\text{Score}}] \quad (9)$$

In Equation (9), $\text{read}_{\text{Score}}$ is the readability score of a document. If $\text{read}_{\text{Score}}$ is high, the document is difficult to follow or to understand. The lower the readability score is, the higher probability is to understand it fast.

4. The communication patterns are the key to trust maintenance. A smoother and consistent communication can help people to put their trust in it. If there is no communication from the product creator, it will cause frustration and anger in backers and lose their interests. Therefore, the communication delay should be minimized between the creator's posts.

$$\text{Authenticity} \propto [1/\text{delay}_{\text{comm}}] \quad (10)$$

In Equation (10), $\text{delay}_{\text{comm}}$ is the average delay between any successive posts, i.e., comments or updates by the project creator. The higher delays will negatively affect project authenticity.

5. Hence, we can summarize the factors mentioned above as

$$\text{Authenticity} \propto [\text{PosCi}, \text{Links}_{\text{Ext}}] \quad (11)$$

also,

$$\text{Authenticity} \propto [1/\text{NegBi}, \text{delay}_{\text{comm}}, \text{read}_{\text{Score}}] \quad (12)$$

By combining Equations (11) and (12), Equation (13) is formulated as below,

$$\text{Authenticity} \propto [\text{PosCi}, \text{Links}_{\text{Ext}}/\text{NegBi}, \text{delay}_{\text{comm}}, \text{read}_{\text{Score}}] \quad (13)$$

6. We divide the Equation into two parts; the similar factors based on their priority are combined. Hence, Equation (14) combines sentiment-based factors.

$$\text{Authenticity} = [\text{PosCi}/\text{NegBi}] \quad (14)$$

This factor is only associated with product comments. For higher authenticity, *PosCi* has to be greater than *NegBi*. We have combined other features related to the product or creator into one Equation as,

$$\text{Authenticity} = [\text{Links}_{\text{Ext}}/\text{read}_{\text{Score}} + \text{delay}_{\text{comm}}] \quad (15)$$

7. Then combine all the factors in one place results into Equation (16) as below,

$$\text{Authenticity}_{\text{project}} = \left[\sum_{i=1}^n \frac{\text{PosCi}}{\text{NegBi}} + \left(\frac{\text{Links}_{\text{Ext}}}{\text{read}_{\text{Score}} + \text{delay}_{\text{comm}}} \right) \right] \quad (16)$$

8. At the final step, we apply optimizations and formulate our objective functions. We have both maximization and minimization functions. The maximization function maximizes the values for favorable factors, and the minimization function underrates the cost of the least desirable parameters. Hence, we can now formulate the credibility estimation in terms of maximization and minimization functions in Equation (17).

$$\text{Credibility}_{\text{project}} = \left[\sum_{i=1}^n \frac{\max(\text{PosCi})}{\min(\text{NegBi})} + \left(\frac{\max(\text{Links}_{\text{Ext}})}{\min(\text{read}_{\text{Score}}) + \min(\text{delay}_{\text{comm}})} \right) \right] \quad (17)$$

For the above Equation, we can define the ranges of all the parameters as below in Equations (18)–(22).

$$\text{NegBi} = \begin{cases} 0 & \text{if \%age of negative comments} = 0 \\ 1 & \text{if \%age of negative comments} = 100\% \end{cases} \quad (18)$$

$$\text{PosCi} = \begin{cases} 0 & \text{if \%age of positive comments} = 0 \\ 1 & \text{if \%age of positive comments} = 100\% \end{cases} \quad (19)$$

$$\text{Links}_{\text{Ext}} = \begin{cases} 0 & \text{if Number of links} = 0 \\ > 0 \leq 9 & \text{if number of links} > 0 \end{cases} \quad (20)$$

The value of $\text{Links}_{\text{Ext}}$ was decided based on the maximum number of external links provided by the project creator. In our case, the maximum number of links a person can provide is considered to be 9. Therefore, $\text{Links}_{\text{Ext}}$ can have any value between 0 and 9.

$$\text{read}_{\text{Score}} = \begin{cases} \text{near } 1 & \text{Comprehensible (easy to understand)} \\ \geq 50 \leq 100 & \text{Incomprehensible or vague (difficult to understand)} \end{cases} \quad (21)$$

$$\text{delay}_{\text{comm}} = [0 - 365 \text{ days}] \quad (22)$$

Following Table 6, we can define the maximum and minimum ranges of each parameter.

Table 6. The value ranges for each credibility parameters.

No.	Parameters for Credibility	Maximum Range	Minimum Range
Content-based			
1	<i>NegBi</i>	100	1
2	<i>PosCi</i>	100	1
3	<i>read_{score}</i>	100	1
Profile-based			
5	<i>Links_{Ext}</i>	9	0
6	<i>delay_{comm}</i>	365	0

4. Implementation and Experimental Setup

In this section, we present our implementation environment, along with the experimental setup, in detail. This section also explains the evaluation metrics used for results assessment.

4.1. Experimental Setup

The core system components include Ubuntu 18.04.1 as an operating system (LTS version), 32 Gb memory, and Nvidia GeForce 1080 as a graphics processing unit (GPU). In addition to the core system component, we used python language for development along with Tensorflow API.

4.2. Evaluation Metrics

The performance of our system is measured by using the following evaluation metrics.

1. Accuracy: The accuracy of the model is calculated by using the following formula as shown in Equation (23)

$$Accuracy = 1 - \frac{\|Y - \hat{Y}\|_F}{\|Y\|_F} \quad (23)$$

where Y & $\hat{Y}$ and represent the actual data and predicted data, respectively.

2. Root mean square error (RMSE): The RMSE is calculated using Equation (24).

$$RMSE = \sqrt{\frac{1}{MN} \sum_{j=1}^M \sum_{i=1}^N (y_{ij} - \hat{y}_{ij})^2} \quad (24)$$

where y_{ij} and $\hat{y}_{ij}$ are subsets of Y & $\hat{Y}$ and represent the actual data and predicted data at the j th time sample in the i th session, respectively. M is the total time samples, and N is the number of projects. RMSE is precisely used to evaluate the prediction error. The smaller the value of RMSE is, the better is prediction rate or score according to Equation (25).

$$Prediction\ rate \propto \frac{1}{RMSE} \quad (25)$$

While accuracy is used to detect predictions' precision, it has an opposite effect than RMSE on the prediction rate, as shown in Equation (26). The higher the value of accuracy is, the better is the prediction rate.

$$Prediction\ rate \propto Accuracy \quad (26)$$

5. Results

This section presents the results and analysis for crowdfunding project recommendations based on the user's previous interests and credibility. In Section 5.1, we offer a study of the recommendation

results for crowdfunding projects. In Section 5.2, we report the accuracy of the proposed model results compared with other models. Table 7 shows the selection criteria for credible projects with different levels of credibility.

Table 7. Selection Criteria for Credibility.

Decision	Credibility Level of a Project	Range of Authenticity (Score)
Highly undesired	Extremely Low	$(>0 \leq 0.2)$
Not desired	Low	$(>0.2 \leq 0.4)$
Can be considered	Normal	$(>0.4 \leq 0.6)$
Desired	High	$(>0.6 \leq 0.8)$
Highly Desired	Extremely High	$(>0.8 \leq 1.0)$

5.1. Statistical Analysis of Recommendation Results

To observe the results of our recommendation module, we used ground truth data. This data includes 100 projects, 55 non-scams, 20 suspended projects, ten canceled projects, and 15 successfully funded projects. Our proposed model works efficiently with high accuracy in both scenarios, i.e., when projects are from the same category or different categories. This dataset exemplifies all possible use case scenarios. We can evaluate how well our recommendation system performs on each type of project in terms of its funding status. The above Figure 4 shows the percentage for each category of crowdfunding projects.

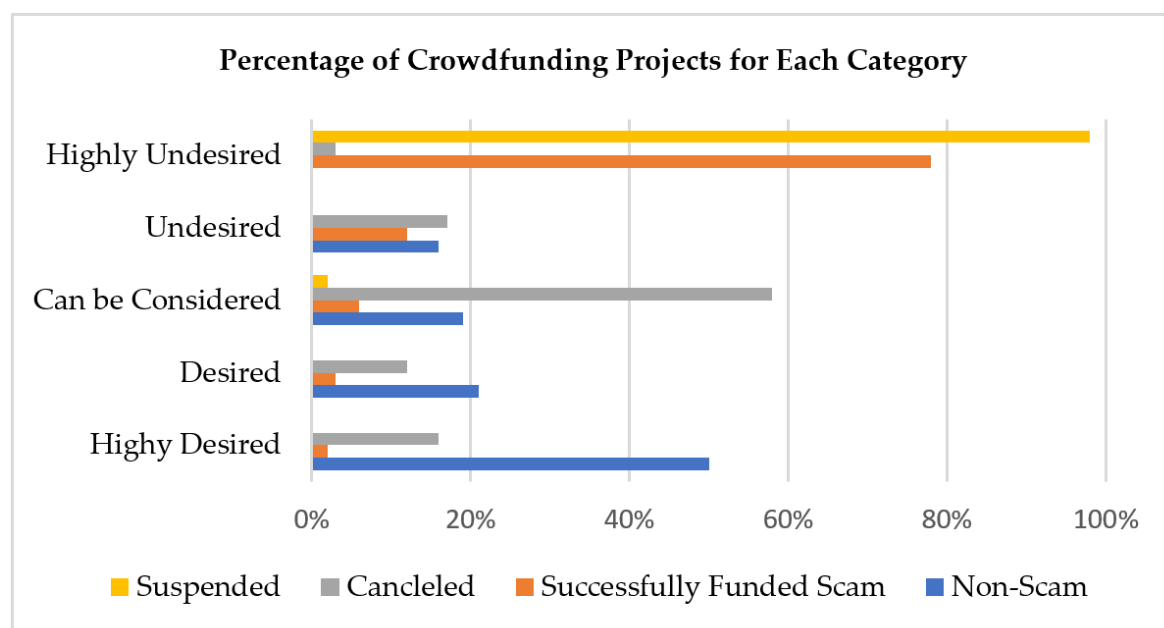


Figure 4. Percentage of crowdfunding projects for each category.

These types are categorized based on the funding status of a project, i.e., “Non-Scam” are projects that are successfully delivered after successful funding; “successfully funded scam” are projects that successfully raised the required funds but failed to deliver; “canceled” category represents projects that have been withdrawn by the project creator before its funding period expires; and “suspended” type means those projects which have been discontinued by the platform in case they figure out any suspicious activity or content.

We have tested our model on all the categories mentioned above of projects to find their authenticity. In Figure 5, we have estimated the authenticity levels for the suspended projects. The x-axis shows

the authenticity levels between 0 and 1, and the y -axis presents the percentage of suspended projects. The estimated authenticity level for 93% of the suspended projects falls in the range (0–0.2).

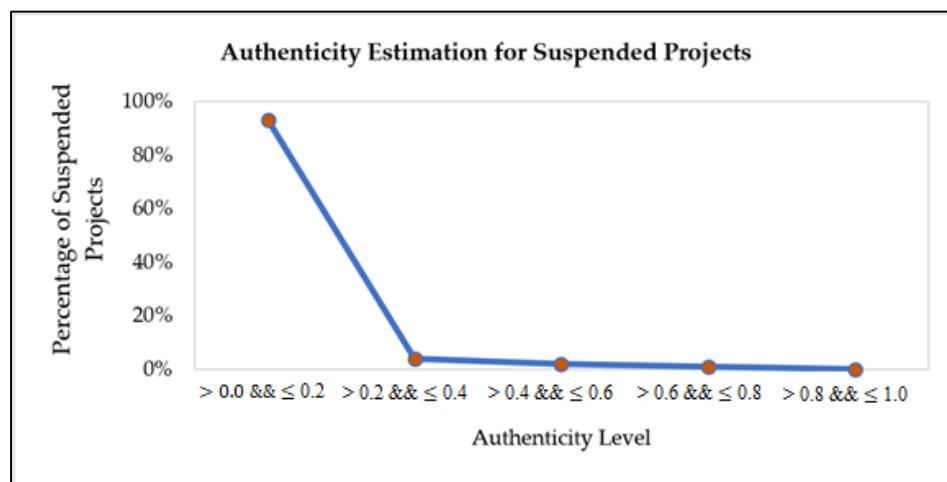


Figure 5. Authenticity estimation for suspended projects.

The data used for this experiment included 80 suspended projects, and 74 projects were falling into the highly undesired range of credibility. This means that these projects are highly undesirable for backers. That is true because these projects are being suspended for some suspicious activities. From this range, we can interpret that most of the projects that have been suspended fail to fulfill the selection criteria of the credibility assessment tool for the recommendation. The statistical analysis of the results is presented in Table 8.

Table 8. Statistical analysis of credibility assessment of suspended projects (Total projects = 80).

Credibility Level	Number of Projects
Highly undesired	74
Not desired	3
Can be considered	2
Desired	1
Highly Desired	0

In Figure 6, we have evaluated the authenticity levels for 70 canceled projects. The estimated authenticity level for 84% of the canceled projects falls in the range (0.4–0.6). Hence, keeping the risk factor in mind, these projects can be considered for investments. The statistical analysis is presented in Table 9. These projects are canceled for multiple reasons, such as lack of funding and budget issues during development phases.

The results show that for most cases, the predicted authenticity level range is between 0 and 0.2. We used 120 undelivered projects, i.e., successfully funded but never delivered projects, and out of these selected projects, 112 projects did not meet the credibility criteria and are highly undesired projects. It represents that regardless of successfully raising funds, backers are disappointed with the progress and development. For such projects, comments play a vital role in understanding a creator's behavior towards his investors after successfully collecting the desired funds. False promises, long delays in communication, or disappearance from the platform are the essential characteristics found in such cases.

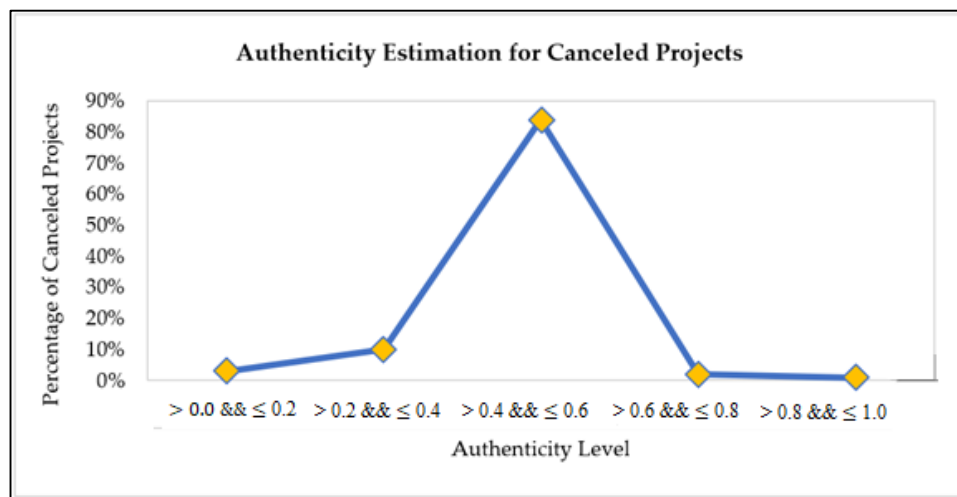


Figure 6. Authenticity estimation for canceled projects.

Table 9. Statistical analysis of credibility assessment of canceled projects (Total projects = 70).

Credibility Level	Number of Projects
Highly undesired	2
Not desired	7
Can be considered	59
Desired	1
Highly Desired	1

Figure 7 presents the accuracy of the recommendation results on the successfully funded projects that didn't deliver, i.e., scam projects.

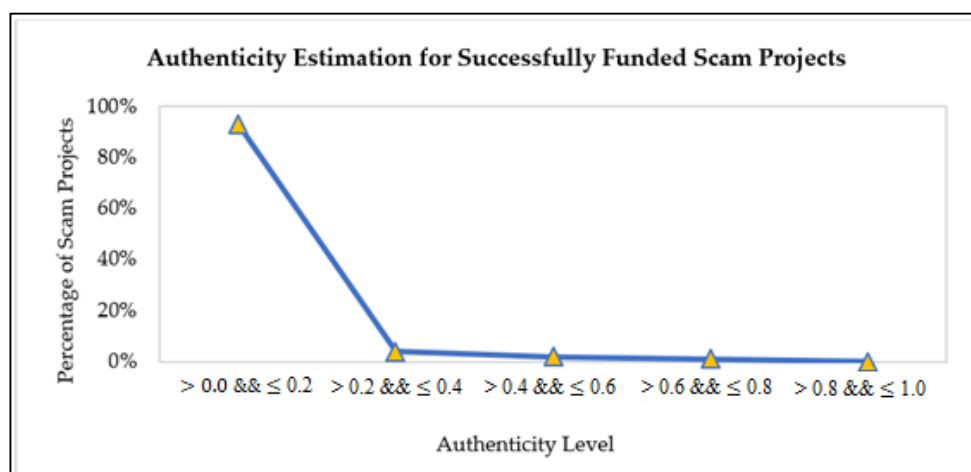


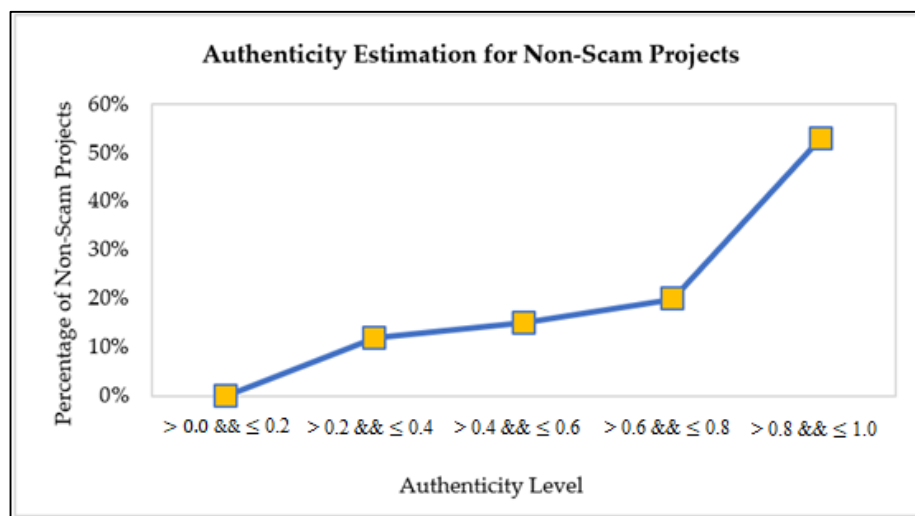
Figure 7. Authenticity estimation for successfully funded scam projects.

The statistical analysis is presented in Table 10. These projects are undelivered and counted as scam projects because the creators didn't fulfill the promises and lacked transparency during the project's development phase.

Table 10. Statistical analysis of credibility assessment of undelivered projects (Total projects = 120).

Credibility Level	Number of Projects
Highly undesired	112
Not desired	5
Can be considered	1
Desired	2
Highly Desired	0

Figure 8 presents an exciting trend. It shows the authenticity level estimation accuracy for non-scam projects.

**Figure 8.** Authenticity estimation for Non-Scam projects.

The inclination depicts rare chances for an authentic and genuine project to have parameters that can lower authenticity scores. Frequently projects are falling in the range of 0.4 to 0.9 and reflecting that comments are going in the gray range, usually for less risky projects.

The statistical analysis is in Table 11. These projects are successfully delivered to the backers. We used 120 successfully delivered projects for this experiment, and 64 of them were highly credible.

Table 11. Statistical analysis of credibility assessment of delivered projects (Total projects = 120).

Credibility Level	Number of Projects
Highly undesired	0
Not desired	14
Can be considered	18
Desired	24
Highly Desired	64

Different learning rates are used for experiments, i.e., 0.1, 0.01, and 0.001 referred to as LR_0.1, LR_0.01, and LR_0.001, respectively, in Figure 9 that evaluate RMSE for a different number of iterations. It can be detected that the testing errors start to get decreased if the learning rate gets smaller. For example, it represents that with a shorter learning rate value, the system's performance improves.

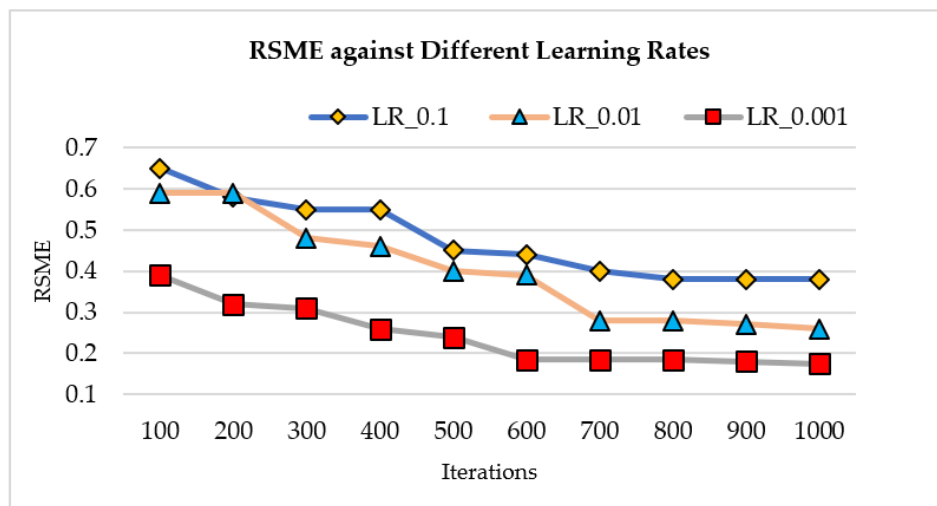


Figure 9. Testing error against different learning rates.

5.2. Comparison with Other Approaches

Here, we have used RMSE as an evaluation metric to evaluate our technique with different ML approaches such as basic NN, bidirectional LSTM, an integrated model of recurrent neural network (RNN) and LDA referred as RNN-LDA, etc. Table 12 presents the RMSE value as a comparison with other models.

Table 12. Evaluation Metrics for Applied Machine Learning Approaches.

ML Approaches	RMSE
Neural Network (NN)	3.37
Bidirectional-LSTM	1.43
RNN-LDA	2.01
NN-LDA	0.82
LDA-LSTM	0.30

Figure 10 presents the accuracy percentage of different models in comparison with our proposed model. It shows that topic models, combined with deep learning models, can achieve better performance than other models.

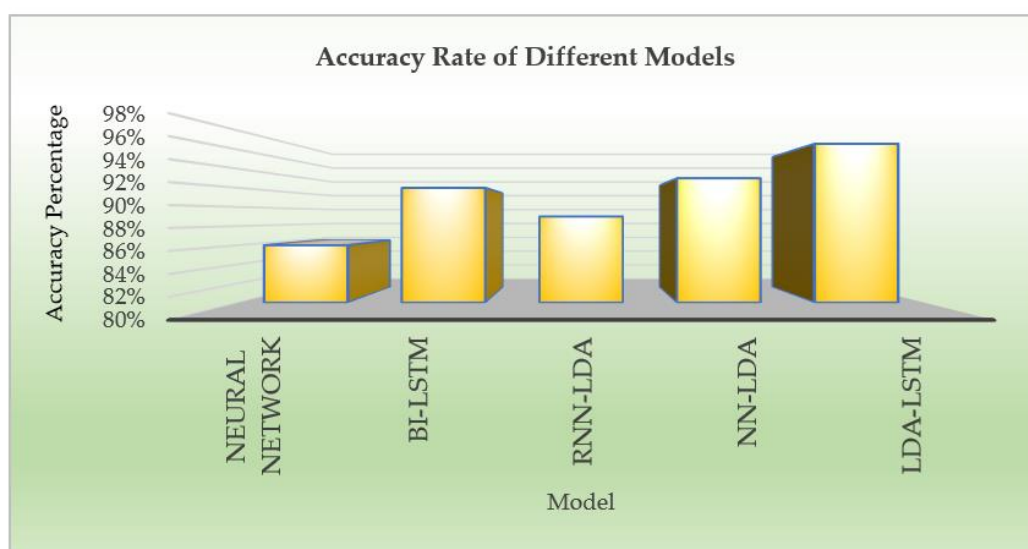


Figure 10. Accuracy Rate of Different Models.

6. Conclusions

We have proposed the methodology for measuring a project's credibility to build a recommendation system. The proposed method uses textual and non-textual data. This system is developed to help the users in selecting reliable and trustworthy options in their preferred categories. The proposed method is a hybrid model of LDA-LSTM and topic modeling that joins the benefits of both (1) LSTMs that captures time dependencies for class and topic prediction and (2) topic modeling that extracts topics that nicely summarize the content. A case study on crowdfunding is performed to analyze and test the proposed system's behavior. We have also embedded an optimized recommendation strategy based on a project's credibility.

This study aims to overcome the limitations of topic models and deep learning and get the most out of both approaches. The main objectives include:

- Finding ways to preserve the contextual dependencies as traditional topic models are based on the bag-of-words approach, so there is a high probability of missing contextual and temporal dependencies.
- Recommendation tools are in use for a long time now; finding the recommended project's credibility is a potential target of this research.

This joint model of LDA-LSTM exploits words and topic embedding, and the temporal data attain 96% accuracy in predicting the topic categories accurately. The topics classes discovered were also evaluated in the context of helping investors identify suspicious campaigns. The prediction quality can be improved if we find out different configurations of comments concerning a project's timeline. We experimented with this by dividing the comments into five various batches of comments. We have not considered projects that have less than 50 comments to maintain the quality of the results.

Many developed applications for recommendation systems in different fields have been proposed. Our proposed approach is a novel approach to recommend a credible crowdfunding project to the best of our knowledge. Moreover, none of the works have focused on crowdfunding comments to find discussion trends and their impact on project credibility. Hence, in crowdfunding, this approach can be used to recommend safe or secure projects to investors. In Table 13, we show a comparison analysis of our proposed model with existing models of recommender systems. In [47–50], the authors use Kickstarter for predictions of the project's success. Others are related to taking comments and updates for estimating the completion of projects and then recommend them to the user. We summarize this research work's contributions: (1) a hybrid method is proposed for reliable and promising recommendations. This approach can model user preferences and word representations in a typical and dynamic style to empower the active measurement of the semantic similarity among the user's preferences and the words. (2) The proposed algorithm is to infer the dynamic embeddings of both the documents and words. We offer a credibility measurement approach for reliable recommendations. The results show that our proposed method outperforms similar state-of-the-art methods significantly.

Table 13. Comparison of our proposed approach with existing recommender systems.

Recommendation Systems	LDA	LDA-LSTM	Language Assessment	Optimization	Comments	Credibility Assessment
[49]	X	X	✓	X	X	X
[50]	✓	X	✓	X	X	X
[51]	✓	X	✓	X	X	X
[52]	✓	X	✓	X	X	X
[53]	X	X	✓	X	X	X
[54]	✓	X	✓	X	✓	X
Proposed Model	✓	✓	✓	✓	✓	✓

7. Discussion

Recommendation systems help users in their decision-making process. Many applications of these systems nowadays in every domain, e.g., location recommendation to tourists, product recommendation to online buyers, restaurant recommendation, route recommendation for travelers, etc. In other words, the need and importance of recommendation systems are not limited to just one platform; crowdfunding is also taking advantage of such applications to make this platform trustworthy for their investors. In this paper, we propose a hybrid model for crowdfunding project recommendations to backers.

The main contribution of this study is we evaluate different features of a campaign to assess its credibility. A credibility assessment is required to build the trust of backers in a campaign. If a backer is partially aware of a campaign's outcome, he can easily decide on investing in it or not. It is essential to build trustworthy recommendation systems, especially when users' hard-earned money is at risk.

We have tried to delve into the details of a campaign and analyze the outcomes of different campaigns based on their funding status. The hybrid model based on topic modeling and deep learning can (1) learn latent topics in comments, (2) to predict the outcome of a project based on the topics discovered so far, and (3) the credibility formulation process carefully evaluates the impact of each feature on the result of a project.

Author Contributions: W.S. conceived the idea for this paper, designed the experiments, wrote the article, assisted in algorithms implementation, and assisted with design and simulation; Y.-C.B. finalized, evaluated proof-read the manuscript, and supervised the work; N.P. did investigation, proof reading and evaluation. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Acknowledgments: This work was supported by the Ministry of Education of the Republic of Korea and the National Research Foundation of Korea (NRF-2019S1A5C2A04083374), and also following are the results of a study on the "Leaders in INdustry-university Cooperation+" Project, supported by the Ministry of Education and National Research Foundation of Korea.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Alexander, J.E.; Tate, M.A. *Web Wisdom: How to Evaluate and Create Web Page Quality*; L. Erlbaum Associates, Inc.: Hillsdale, NJ, USA, 1999.
2. Grabner-Kräuter, S.; Kaluscha, E.A. Empirical research in online trust: A review and critical assessment. *Int. J. Hum. -Comput. Stud.* **2003**, *58*, 783–812. [CrossRef]
3. Available online: <https://www.mordorintelligence.com/industry-reports/crowdfunding-market> (accessed on 20 October 2020).
4. Available online: <https://www.statista.com/outlook/335/100/crowdfunding/worldwide> (accessed on 20 October 2020).
5. Pergola, G.; Gui, L.; He, Y. TDAM: A topic-dependent attention model for sentiment analysis. *Inf. Process. Manag.* **2019**, *56*, 102084. [CrossRef]
6. Karami, A. *Fuzzy Topic Modeling for Medical Corpora*; University of Maryland: Baltimore County, MD, USA, 2015.
7. Asuncion, H.U.; Asuncion, A.U.; Taylor, R.N. Software traceability with topic modeling. In Proceedings of the 2010 ACM/IEEE 32nd International Conference on Software Engineering, Cape Town, South Africa, 1–8 May 2010; Volume 1, pp. 95–104.
8. Ghosh, D.; Guha, R. What are we 'tweeting' about obesity? Mapping tweets with topic modeling and Geographic Information System. *Cartogr. Geogr. Inf. Sci.* **2013**, *40*, 90–102. [CrossRef]
9. DiMaggio, P.; Nag, M.; Blei, D. Exploiting affinities between topic modeling and the sociological perspective on culture: Application to newspaper coverage of US government arts funding. *Poetics* **2013**, *41*, 570–606. [CrossRef]
10. Dumais, S.T. Latent semantic analysis. *Annu. Rev. Inf. Sci. Technol.* **2004**, *38*, 188–230. [CrossRef]
11. Brants, T.; Chen, F.; Tsochantaridis, I. Topic-based document segmentation with probabilistic latent semantic analysis. In Proceedings of the Eleventh International Conference on Information and Knowledge Management, McLean, VA, USA, 4–9 November 2002; pp. 211–218.

12. Blei, D.M.; Ng, A.Y.; Jordan, M.I. Latent Dirichlet Allocation. *J. Mach. Learn. Res.* **2003**, *3*, 993–1022.
13. Blei, D.; Lafferty, J. Correlated topic models. *Adv. Neural Inf. Process. Syst.* **2006**, *18*, 147.
14. Landauer, T.K.; Foltz, P.W.; Laham, D. An introduction to latent semantic analysis. *Discourse Process.* **1988**, *25*, 259–284. [[CrossRef](#)]
15. Landauer, T.K. Applications of latent semantic analysis. In Proceedings of the Annual Meeting of the Cognitive Science Society, Fairfax, VA, USA, 7–10 August 2002; Volume 24.
16. Sidorov, G. Latent Semantic Analysis (LSA): Reduction of Dimensions. In *Syntactic n-grams in Computational Linguistics*; Springer: Cham, Switzerland, 2019; pp. 17–19.
17. Sehra, S.; Singh, J.; Rai, H. Using latent semantic analysis to identify research trends in openstreetmap. *ISPRS Int. J. Geo-Inf.* **2017**, *6*, 195. [[CrossRef](#)]
18. Hofmann, T. Probabilistic latent semantic analysis. *Uncertain. Artif. Intell.* **1999**, 289–296. [[CrossRef](#)]
19. Hofmann, T. Unsupervised learning by probabilistic latent semantic analysis. *Mach. Learn.* **2001**, *42*, 177–196. [[CrossRef](#)]
20. Kakkonen, T.; Myller, N.; Sutinen, E.; Timonen, J. Comparison of Dimension Reduction Methods for Automated Essay Grading. *Educ. Technol. Soc.* **2008**, *11*, 275–288.
21. Liu, S.; Xia, C.; Jiang, X. Efficient Probabilistic Latent Semantic Analysis with Sparsity Control. In Proceedings of the IEEE International Conference on Data Mining, Sydney, NSW, Australia, 13–17 December 2010; pp. 905–910.
22. Romberg, S.; Hörster, E.; Lienhart, R. *Multimodal pLSA on Visual Features and Tags*; The Institute of Electrical and Electronics Engineers Inc.: Cancun, Mexico, 2009; pp. 414–417.
23. Wu, H.; Wang, Y.; Cheng, X. *Incremental Probabilistic Latent Semantic Analysis for Automatic Question Recommendation*; ACM: New York, NY, USA, 2008; pp. 99–106.
24. Shen, Z.Y.; Sun, J.; Shen, Y.D. Collective Latent Dirichlet Allocation. In Proceedings of the Eighth IEEE International Conference on Data Mining, Pisa, Italy, 15–19 December 2008; pp. 1019–1025.
25. Rosen-Zvi, M.; Griffiths, T.; Steyvers, M.; Smyth, P. The author-topic model for authors and documents. In Proceedings of the 20th Conference on Uncertainty in Artificial Intelligence, Banff, AB, Canada, 7–11 July 2004; pp. 487–494.
26. Liu, L.; Tang, L.; Dong, W.; Yao, S.; Zhou, W. An overview of topic modeling and its current applications in bioinformatics. *SpringerPlus* **2016**, *5*, 1608. [[CrossRef](#)]
27. Wang, X.; McCallum, A. Topics over time: A non-markov continuous-time model of topical trends. In Proceedings of the International Conference on Knowledge Discovery and Data Mining, Philadelphia, PA, USA, 20–23 August 2006; pp. 424–433.
28. McCallum, A.; Wang, X.; Corrada-Emmanuel, A. Topic and role discovery in social networks with experiments on enron and academic email. *J. Artif. Intell. Res.* **2007**, *30*, 249–272. [[CrossRef](#)]
29. Bao, S.; Xu, S.; Zhang, L.; Yan, R.; Su, Z.; Han, D.; Yu, Y. Joint Emotion-Topic Modeling for Social Affective Text Mining. In Proceedings of the 2009 Ninth IEEE International Conference on Data Mining, Miami, FL, USA, 6–9 December 2009; pp. 699–704.
30. Kakkonen, T.; Myller, N.; Sutinen, E. Applying latent Dirichlet allocation to automatic essay grading. *Lect. Notes Comput. Sci.* **2006**, *4139*, 110–120.
31. Bergholz, A.; Chang, J.; Paaß, G.; Reichartz, F.; Strobel, S. Improved phishing detection using model-based features. In Proceedings of the CEAS 2008—The Fifth Conference on Email and Anti-Spam, Mountain View, CA, USA, 21–22 August 2008.
32. Cheng, X.; Yan, X.; Lan, Y.; Guo, J. Btm: Topic modeling over short texts. *IEEE Trans. Knowl. Data Eng.* **2014**, *26*, 2928–2941. [[CrossRef](#)]
33. Wang, H.; Lu, Y.; Zhai, C. Latent aspect rating analysis on review text data: A rating regression approach. In Proceedings of the 16th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Washington, DC, USA, 25–28 July 2010; pp. 783–792.
34. Park, M.H.; Hong, J.H.; Cho, S.B. Location-based recommendation system using bayesian user’s preference model in mobile devices. In *International Conference on Ubiquitous Intelligence and Computing*; Springer: Berlin/Heidelberg, Germany, 2007; pp. 1130–1139.
35. Pazzani, M.J.; Billsus, D. Content-based recommendation systems. In *the Adaptive Web*; Springer: Berlin/Heidelberg, Germany, 2007; pp. 325–341.

36. Greer, C.F.; Ferguson, D.A. Using Twitter for promotion and branding: A content analysis of local television Twitter sites. *J. Broadcasting Electron. Media* **2011**, *55*, 198–214. [CrossRef]
37. Krizhevsky, A.; Sutskever, I.; Hinton, G.E. Imagenet classification with deep convolutional neural networks. *Adv. Neural Inf. Process. Syst.* **2012**, *60*, 84–90. [CrossRef]
38. Lin, K.; Yang, H.F.; Liu, K.H.; Hsiao, J.H.; Chen, C.S. Rapid clothing retrieval via deep learning of binary codes and hierarchical search. In Proceedings of the 5th ACM on International Conference on Multimedia Retrieval, Shanghai, China, 23–26 June 2015; pp. 499–502.
39. Chiliguano, P.; Fazekas, G. Hybrid music recommender using content-based and social information. In Proceedings of the 2016 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), Shanghai, China, 20–25 March 2016.
40. Liu, X.; Nielek, R.; Adamska, P.; Wierzbicki, A.; Aberer, K. Towards a highly effective and robust Web credibility evaluation system. *Decis. Support Syst.* **2015**, *79*, 99–108. [CrossRef]
41. Schwarz, J.; Meredith, M. Augmenting web pages and search results to support credibility assessment. In Proceedings of the SIGCHI Conference on Human Factors in Computing Systems, Vancouver, BC, Canada, 7–12 May 2011; pp. 1245–1254.
42. Mitra, T.; Gilbert, E. The language that gets people to give: Phrases that predict success on kickstarter. In Proceedings of the 17th ACM Conference on Computer Supported Cooperative Work & Social Computing, Baltimore, MD, USA, 15–19 February 2014; pp. 49–61.
43. Parhankangas, A.; Renko, M. Linguistic style and crowdfunding success among social and commercial entrepreneurs. *J. Bus. Ventur.* **2017**, *32*, 215–236. [CrossRef]
44. Yuan, H.; Lau, R.Y.; Xu, W. The determinants of crowdfunding success: A semantic text analytics approach. *Decis. Support Syst.* **2016**, *91*, 67–76. [CrossRef]
45. Shafqat, W.; Byun, Y.C. Topic Predictions and Optimized Recommendation Mechanism Based on Integrated Topic Modeling and Deep Neural Networks in Crowdfunding Platforms. *Appl. Sci.* **2019**, *9*, 5496. [CrossRef]
46. Cohen, J. A coefficient of agreement for nominal scales. *Educ. Psychol. Meas.* **1960**, *20*, 37–46. [CrossRef]
47. Jung, Y.; Suh, Y. Mining the voice of employees: A text mining approach to identifying and analyzing job satisfaction factors from online employee reviews. *Decis. Support Syst.* **2019**, *123*, 113074. [CrossRef]
48. Albalawi, R.; Yeap, T.H.; Benyoucef, M. Using Topic Modeling Methods for Short-Text Data: A Comparative Analysis. *Front. Artif. Intell.* **2020**, *3*, 42. [CrossRef]
49. Desai, N.; Gupta, R.; Truong, K. *Plead or Pitch? The Role of Language in Kickstarter Project Success*; Stanford University: Stanford, CA, USA, 2015.
50. Sawhney, K.; Tran, C.; Tuason, R. *Using Language to Predict Kickstarter Success*; Stanford University: Stanford, CA, USA, 2016.
51. Westerlund, M.; Singh, I.; Rajahonka, M.; Leminen, S. Can short-text project summaries predict funding success on crowdfunding platforms? In *Proceedings of the ISPIM Conference Proceedings*; The International Society for Professional Innovation Management (ISPIM): Manchester, UK, 2019; pp. 1–15.
52. Siering, M.; Koch, J.A.; Deokar, A.V. Detecting fraudulent behavior on crowdfunding platforms: The role of linguistic and content-based cues in static and dynamic contexts. *J. Manag. Inf. Syst.* **2016**, *33*, 421–455. [CrossRef]
53. Cumming, D.J.; Hornuf, L.; Karami, M.; Schweizer, D. *Disentangling Crowdfunding from Fraudfunding*; SSRN, 2016; Available online: <https://ssrn.com/abstract=2828919> (accessed on 20 October 2020).
54. Tran, T.; Lee, K.; Vo, N.; Choi, H. Identifying on-time reward delivery projects with estimating delivery duration on kickstarter. In Proceedings of the 2017 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining, Sydney, Australia, 31 July–3 August 2017; ACM: New York, NY, USA, 2017; pp. 250–257.

Publisher’s Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



© 2020 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).