

Fake News Detection Using Machine Learning

Shankar H B*, Mrs. Geetha N B**

*(Computer Science & Engineering, University BDT College, Davanagere

Email: shankarhb18@gmail.com)

** (Computer Science & Engineering, University BDT College, Davanagere

Email: geethanb291979@gmail.com)

Abstract:

In today's world of social media and politics, detection of false news is a crucial area for research. There may be difficulties because there aren't as many resources, such as published literature and available datasets. In this work, we advise using machine learning techniques to spot fake news. Three different machine learning categorization methods are compared. Decision tree classifier, random forest classifier, and logistic regression were the three techniques we employed. We have obtained varying degrees of accuracy for each method. The obtained results are helpful in detecting the veracity of the specified news.

Keywords — Machine learning, fake news, data pre- processing, deep learning

I. INTRODUCTION

Any article, post, story, or publication that may be classified as false or truthful has gained significant attention from researchers all around the globe. Many analyses and investigations have been undertaken to ascertain the impact of any inaccurate or deceptive data on human beings when they have been exposed to it. Fake news or information can be used to sway someone's basic assumptions in favour of something that may not be real.

The worldwide pandemic scenario serves as the best illustration of bogus news. Many news stories have been fabricated in the past and are now being done so in order to confuse people, cause disruption, and lead them to believe incorrect information.

Fake news spread over social media platform in India convinced people to consume dung or cow wee in an effort to prevent infection, while in the country, artiodactyl wee- wee with lime was lauded as a coronavirus-preventing remedy. The researchers looked at a number of suggestions like consuming garlic, donning heat-generating socks for curing the potentially fatal conditions.

A. Types of False News

1) Pictorial-based:

Graphics, including altered photographs, videos, and mixtures of both, are frequently used as material in false news articles.

2) User-based:

This type of news is produced by fictional accounts. Generally, this kind of made-up news is targeted at certain audiences that may represent particular age groups, genders, cultures, and political allegiances.

3) Knowledge-based:

These articles offer scientific explanations for certain unsolved problems, giving users the impression that they are true. For instance, natural remedies for the body's excessive blood sugar levels.

4) Stance-based:

It serves as a true example of utterances that alter their meaning and aim in some way.

II. LITERATURE REVIEW

Many authors have done preliminary research on "fake" news detection on twitter in the past. However, this topic became more well-known

during the US Presidential election, and everyone tried their hardest to come up with some better answers for this classification. On "The Atlantic" news, a brief history was presented [1]. Some writers also used a machine learning strategy that was implemented as a software system [2]. Mykhailo Granik et al.'s study [3] demonstrates how simple it is to spot fake news using a Naive Bayes classifier. This strategy was examined using a sample of Facebook news posts, and it was then applied to a software solution. In addition to three significant mainstream political news websites (Politico, CNN, and ABC News), they come from three significant left- and right-leaning Facebook pages. They were successful in achieving a classification accuracy of around 74%. False news is classified with somewhat less precision. This skewness may be caused by the fact that false news makes up just 4.9% of the sample.

The authors of [4] recommended deception detection utilising the benchmark data set "LIAR" with clearly enhanced effectiveness in detecting false postings and news in their study. The authors argued in favour of using corpora for political NLP research, position categorization, opinion mining, and rumour identification.

In order to address a variety of issues, including as the lack of accuracy, BotMaker's temporal lag, and the lengthy processing time required to handle hundreds of tweets in only one second, Himank Gupta et al. [5] built a framework based on various machine learning approaches. Using the HSpam14 dataset, they first collected 400,000 tweets. The 250,000 non-spam tweets and the 150,000 spam tweets are then further explained. In addition to the Top-30 phrases from Bag-of-phrases model 4 that are producing the most information gain, they also deduced a number of lightweight features. They performed around 18% better than the previous response and had an accuracy of 91.65%. Need for hoax detection has been introduced by the authors of [6]. They combined social content and news content techniques while using ML. The writers assert that the performance is superior to what has been written about it. It was implemented by the authors using a Facebook Messenger chatbot. Facebook Italian news postings were pulled from three separate databases. Both content-based strategies leveraging Boolean crowd sourcing algorithms and social and content signals were used. Three classifiers were used in Mandical et al.'s proposed fake news classification system in

order to accelerate the identification of false information. Eight datasets were implemented with this system [7].

Zhang *et al.* [8] observed the guidelines, procedures, and algorithms used for categorising fake and manufactured news articles, authors, and subjects from online social networks and assessing the reach and effectiveness of the accompanying efforts.

Bharadwaj and Shao [9] employed recurrent neural networks, a Naive Bayes classifier, and random forest classifiers with six feature extraction techniques on the kaggle.com false news dataset. Junaed *et al.* [10] provided a comprehensive performance study of several methodologies on three separate datasets.

This work ignored details like the source, author, or publication date that may have a significant influence on the outcome and instead concentrated on the language of the material and the mood it conveyed. Additionally, we included feelings while demonstrating our work in the detection process.

III. MEHODOLOGY

Three step methodology was employed in this investigation. In the first phase (pre-processing), the study used particular data filtering and cleaning techniques to extract semantic features from the raw data. The data was categorized using a stopword filter that eliminates prepositions. This study also employed non-English character removal technologies and HTML tags to filter the data and get rid of contaminants that weren't helpful for categorization. Semantic features were transformed into feature vectors during the second step (extracting numerical features), the final stage (classifiers) involved classifying the elements in the dataset using machine learning and deep learning classifiers. Separate applications of each of these techniques were made to the same dataset.

A. Used Dataset

The datasets we used for this analysis are available online as open source downloads. The information consists of both true and fraudulent news articles from various websites. The true news articles produced provide an accurate depiction of real occurrences, in contrast to fake news websites that make assertions that are not backed by the facts. Fact-checking websites like politifact.com and snopes.com may be used to manually verify many of the political assertions made in such

publications. Two distinct datasets were used in this research.

TRUE.CSV and FAKE.CSV: These datasets were utilized. They have three columns: text/keyword, statement, and label (false/true).

B. Pre-Processing

Data cleaning, formatting, and transformation are essential steps in the pre-processing phase of machine learning, which prepares data for use in a machine learning algorithm. The accuracy and performance of a machine learning model are strongly influenced by the quality of the training data. Data must thus be pre-processed to ensure that it is accurate, trustworthy, and relevant to the current problem. Data cleaning, which entails deleting any blank or incorrect numbers and resolving discrepancies in the data, is the initial stage of data pre-processing.

C. Algorithms used:

i. Logistic regression

The objective of classification problems, where machine learning is applied, is to predict a discrete outcome or label based on one or more input factors. In logistic regression, the weighted sum of the input variables is calculated, and the result is then run through a logistic function to get a probability value between 0 and 1. The final categorization choice is made using this probability. To narrow the discrepancy between expected probability and actual label, the model is trained by increasing the weights of the input variables. With the use of methods like one-vs-all or SoftMax regression, logistic regression—a common solution for binary classification issues—can also be expanded to tackle multi-class classification issues.

ii. Decision Tree

The decision tree classifier has become an established artificial intelligence technique for classification tasks. The data is recursively partitioned based on the input attributes, a tree-like model is built, and a set of decision rules are created, leading to a final classification determination. The leaves of the tree represent the final classification decisions, while each node represents a decision rule based on a specific attribute. Decision trees are a well-liked option for applications where comprehending the model is crucial because they are simple to grasp and display. They are capable of handling

categorical and continuous input data as well as binary and multiclass classification problems.

iii. Random Forest Classifier

A popular machine learning method for classifying tasks is Random Forest Classifier. It is an ensemble method that combines different decision trees to improve the model's overall accuracy and stability. During training, the random forest classifier creates a large number of decision trees and averages the forecasts of each tree to produce predictions. This method aids in improving the model's generalisation capabilities and decreasing overfitting. Additionally, random forests are capable of handling categorical and continuous input variables as well as binary and multiclass classification issues. Moreover, they can deal with missing data and outliers.

IV. WORKFLOW

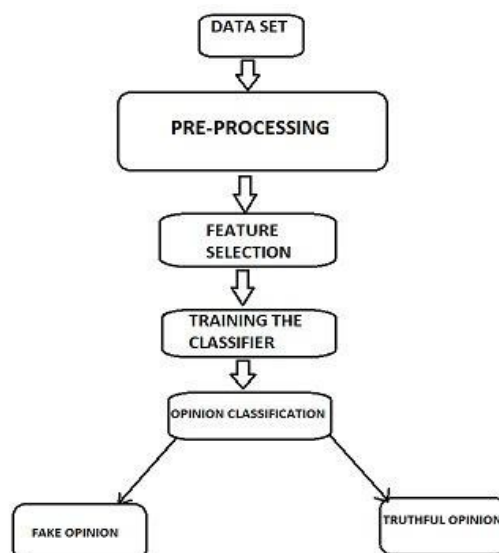


Fig. 1: Work Flow

The workflow in the Fig. 1 describes the standard process for identifying false news. The following steps make up the process:

Data gathering: The initial phase is gathering news items from numerous websites and social media platforms.

Text preprocessing: After the news items are gathered, any extraneous text, such as HTML elements, URLs, and stop words, is removed. This procedure is crucial since it aids in lowering the data's noise level and improving its suitability for analysis.

Feature extraction: The preprocessed news items are used to extract features in this stage. The characteristics of the news piece known as features

can be utilized to differentiate between authentic and false news. The frequency of specific terms, the article's length, and the tone of the piece are a few prominent criteria used to identify false news.

Model selection: A machine learning model is selected to classify the news articles as real or fake once the attributes are extracted. Several models, including decision trees, logistic regression, and random forest classifier, can be used for this.

Model education: A collection of labelled data is used to train the chosen model once it has been chosen. Data that has been labelled as being either true or false news is referred to as labelled data. This stage is crucial because it teaches the model how to appropriately identify fresh news pieces.

Evaluation of the model: Using a set of test data, the model is assessed following training. The test data is used to gauge how effectively the model generalizes to new situations because it differs from the training data.

V. EXPERIMENT

We used Python as a programming language throughout the experiment. Steps are listed as follows:

1: Importing different libraries required for building a model, Fig. 2.

Pandas – Used to Read and write for SSV files.

NumPy – It helps in dealing with numerical tasks.

Matplotlib – It helps in displaying in graph format.

Seaborn – It is used for visualization.

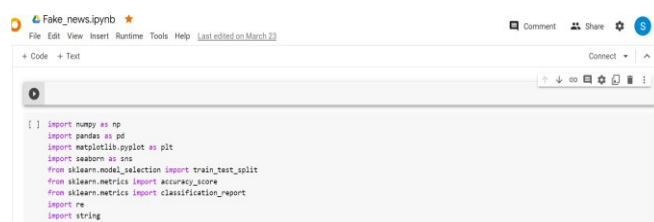


Fig. 2: Importing

Libraries 2: Data Pre-processing

In text mining, texts are regarded as unstructured data and may include a range of impurities, including HTML components, single letters, advertisements, non-English characters, numbers, and apostrophes. Data preparation is similar to text mining in this regard. Because of this, it is particularly challenging to express textual data in natural language. Unstructured data may be transformed using a variety of methods into structured data that computers can handle. The

classification dataset was cleaned in this study using the stopwords strategy. Stopword technology, which eliminates certain worthless words (such as the, in, a, an, with, etc.), is a typical approach applied for text classification, information retrieval, and data filtering. In other words, things that don't function as keywords yet impact classification are deleted. HTML tags were eliminated using the Python Standard Library's (remove_tags) function. The preprocess text feature then eliminated all non-English characters, Fig. 3.



Fig. 3: Data Preprocessing

3: MODEL BUILDING

Testing a machine learning model's performance on data that it has never seen before is essential. Due to this, the data was split into a training set and a test set.

The training set, which is frequently a larger portion of the data, is used to train the model. The model generates new predictions based on the values of the independent variables (attributes) and their corresponding dependent variables (labels) in the training set.

After model training, we evaluate the model's performance on an unfamiliar test set. To determine how effectively the model generalises to new data, we create predictions using the independent variables in the test set and then compare those predictions to the actual values of the dependent variables in the test set.

Scikit-learn's train_test_split function makes it simple to divide the data into training and test sets. The data is first divided into the desired proportions before being randomly shuffled using the programme. By default, training uses 80% of the data, while testing uses 20%.

The independent variables (attributes) and dependent variables (labels) are sent in as different arrays when we run `train_test_split`. The four arrays that the function returns are `X_train`, `X_test`, `y_train`, and `y_test`. The independent and dependent variables for the training set are included in `X_train` and `y_train`, whereas the independent and dependent variables for the test set are contained in `X_test` and `y_test`.

In conclusion, separating the data into a training set and a test set enables us to train and test our model on several sets of data, preventing overfitting and enabling a more precise evaluation of the model's performance.

```
[ ] from sklearn.tree import DecisionTreeClassifier
DT = DecisionTreeClassifier()
DT.fit(xv_train,y_train)
```

```
DecisionTreeClassifier
DecisionTreeClassifier()
```

```
[ ] pred_dt = DT.predict(xv_test)
```

```
[ ] DT.score(xv_test,y_test)
```

```
0.9958110516934047
```

```
[ ] print(classification_report(y_test,pred_lr))
```

	precision	recall	f1-score	support
0	0.99	0.98	0.99	5920
1	0.98	0.99	0.99	5300
accuracy			0.99	11220
macro avg	0.99	0.99	0.99	11220
weighted avg	0.99	0.99	0.99	11220

Fig. 4: Model Building using Decision tree classifier

```
[ ] from sklearn.linear_model import LogisticRegression
LR = LogisticRegression()
LR.fit(xv_train,y_train)
```

```
LogisticRegression
LogisticRegression()
```

```
[ ] pred_lr = LR.predict(xv_test)
```

```
[ ] LR.score(xv_test,y_test)
```

```
0.9854189483066
```

```
[ ] print(classification_report(y_test,pred_lr))
```

	precision	recall	f1-score	support
0	0.99	0.98	0.99	5920
1	0.98	0.99	0.99	5300
accuracy			0.99	11220
macro avg	0.99	0.99	0.99	11220
weighted avg	0.99	0.99	0.99	11220

Fig. 5: Model Building using Logistic Regression

```
[ ] from sklearn.ensemble import RandomForestClassifier
RF = RandomForestClassifier(random_state = 0)
RF.fit(xv_train,y_train)
```

```
RandomForestClassifier
RandomForestClassifier(random_state=0)
```

```
[ ] pred_rf = RF.predict(xv_test)
```

```
[ ] RF.score(xv_test,y_test)
```

```
0.987433250802109
```

```
[ ] print(classification_report(y_test,pred_rf))
```

	precision	recall	f1-score	support
0	0.99	0.99	0.99	5920
1	0.99	0.99	0.99	5300
accuracy	0.99	0.99	0.99	11220
macro avg	0.99	0.99	0.99	11220
weighted avg	0.99	0.99	0.99	11220

Fig. 6: Model Building using Random Forest

4: PREDICTING USING THE MODEL

Once we have separated our data into training and test sets, we may use the training set to train a variety of machine learning algorithms, including logistic regression, decision trees, and random forest classifier. The relationship between the independent (attributes) and dependent (label) variables, which these algorithms may learn from the training set, may then be utilised to predict results for newly discovered data.

When utilising a trained machine learning algorithm to forecast a new data point, we input the values of the independent variables, and the algorithm returns a predicted value for the dependent variable. The association that the algorithm discovered during the training phase provides the basis for this anticipated value.

The technique used for logistic regression describes the connection between the independent and dependent variables as a logistic function. This function produces a probability value, which is defined in order to classify the data point as belonging to one of two groups.

Depending on the independent variable values from the training set, the method for decision tree classifiers learns a decision tree. Following the decision tree's route according to the values of the independent variables, additional data points are then classified using this decision tree.

An ensemble of decision trees is created for random forest classifiers using the approach, with each tree being trained on an undetermined subset of the independent variables and an undetermined percentage of the training data. The algorithm combines the predictions from each tree when forecasting a new data point to produce a final forecast.

It's important to remember that the test set should first go through the same preprocessing steps as the training set before predictions can be made. This makes sure that the test set's input data are in the same format as the data used to train the algorithm.

The screenshot shows a Jupyter Notebook titled "Fake News, pytorch". The code is as follows:

```

File Edit View Insert Runtime Tools Help All changes saved

+ Code + Text

1 # def output_label(x):
2     if x == 0:
3         return "Fake News"
4
5     elif x == 1:
6         return "News is absolutely TRUE!"
7
8
9 # def neural_testing(model):
10     testing_data_loader = torch.utils.data.DataLoader(fake_news_loader, batch_size=batch_size, shuffle=False)
11     num_fake_news = 0
12     num_true_news = 0
13     for i, (inputs, labels) in enumerate(testing_data_loader):
14         outputs = model(inputs)
15         pred_outputs = torch.argmax(outputs, dim=-1)
16         for j in range(len(pred_outputs)):
17             if pred_outputs[j] == 0:
18                 num_fake_news += 1
19             else:
20                 num_true_news += 1
21     return print("Fake News Prediction: {} \t True News Prediction: {}".format(num_fake_news, num_true_news))
22
23
24 # # #
25
26 # # #
27
28 # # #
29
30 # # #
31
32 # # #
33
34 # # #
35
36 # # #
37
38 # # #
39
40 # # #
41
42 # # #
43
44 # # #
45
46 # # #
47
48 # # #
49
50 # # #
51
52 # # #
53
54 # # #
55
56 # # #
57
58 # # #
59
60 # # #
61
62 # # #
63
64 # # #
65
66 # # #
67
68 # # #
69
70 # # #
71
72 # # #
73
74 # # #
75
76 # # #
77
78 # # #
79
80 # # #
81
82 # # #
83
84 # # #
85
86 # # #
87
88 # # #
89
90 # # #
91
92 # # #
93
94 # # #
95
96 # # #
97
98 # # #
99
100 # # #
101
102 # # #
103
104 # # #
105
106 # # #
107
108 # # #
109
110 # # #
111
112 # # #
113
114 # # #
115
116 # # #
117
118 # # #
119
120 # # #
121
122 # # #
123
124 # # #
125
126 # # #
127
128 # # #
129
130 # # #
131
132 # # #
133
134 # # #
135
136 # # #
137
138 # # #
139
140 # # #
141
142 # # #
143
144 # # #
145
146 # # #
147
148 # # #
149
150 # # #
151
152 # # #
153
154 # # #
155
156 # # #
157
158 # # #
159
160 # # #
161
162 # # #
163
164 # # #
165
166 # # #
167
168 # # #
169
170 # # #
171
172 # # #
173
174 # # #
175
176 # # #
177
178 # # #
179
180 # # #
181
182 # # #
183
184 # # #
185
186 # # #
187
188 # # #
189
190 # # #
191
192 # # #
193
194 # # #
195
196 # # #
197
198 # # #
199
200 # # #
201
202 # # #
203
204 # # #
205
206 # # #
207
208 # # #
209
210 # # #
211
212 # # #
213
214 # # #
215
216 # # #
217
218 # # #
219
220 # # #
221
222 # # #
223
224 # # #
225
226 # # #
227
228 # # #
229
230 # # #
231
232 # # #
233
234 # # #
235
236 # # #
237
238 # # #
239
240 # # #
241
242 # # #
243
244 # # #
245
246 # # #
247
248 # # #
249
250 # # #
251
252 # # #
253
254 # # #
255
256 # # #
257
258 # # #
259
260 # # #
261
262 # # #
263
264 # # #
265
266 # # #
267
268 # # #
269
270 # # #
271
272 # # #
273
274 # # #
275
276 # # #
277
278 # # #
279
280 # # #
281
282 # # #
283
284 # # #
285
286 # # #
287
288 # # #
289
290 # # #
291
292 # # #
293
294 # # #
295
296 # # #
297
298 # # #
299
300 # # #
301
302 # # #
303
304 # # #
305
306 # # #
307
308 # # #
309
310 # # #
311
312 # # #
313
314 # # #
315
316 # # #
317
318 # # #
319
320 # # #
321
322 # # #
323
324 # # #
325
326 # # #
327
328 # # #
329
330 # # #
331
332 # # #
333
334 # # #
335
336 # # #
337
338 # # #
339
340 # # #
341
342 # # #
343
344 # # #
345
346 # # #
347
348 # # #
349
350 # # #
351
352 # # #
353
354 # # #
355
356 # # #
357
358 # # #
359
360 # # #
361
362 # # #
363
364 # # #
365
366 # # #
367
368 # # #
369
370 # # #
371
372 # # #
373
374 # # #
375
376 # # #
377
378 # # #
379
380 # # #
381
382 # # #
383
384 # # #
385
386 # # #
387
388 # # #
389
390 # # #
391
392 # # #
393
394 # # #
395
396 # # #
397
398 # # #
399
400 # # #
401
402 # # #
403
404 # # #
405
406 # # #
407
408 # # #
409
410 # # #
411
412 # # #
413
414 # # #
415
416 # # #
417
418 # # #
419
420 # # #
421
422 # # #
423
424 # # #
425
426 # # #
427
428 # # #
429
430 # # #
431
432 # # #
433
434 # # #
435
436 # # #
437
438 # # #
439
440 # # #
441
442 # # #
443
444 # # #
445
446 # # #
447
448 # # #
449
450 # # #
451
452 # # #
453
454 # # #
455
456 # # #
457
458 # # #
459
460 # # #
461
462 # # #
463
464 # # #
465
466 # # #
467
468 # # #
469
470 # # #
471
472 # # #
473
474 # # #
475
476 # # #
477
478 # # #
479
480 # # #
481
482 # # #
483
484 # # #
485
486 # # #
487
488 # # #
489
490 # # #
491
492 # # #
493
494 # # #
495
496 # # #
497
498 # # #
499
500 # # #
501
502 # # #
503
504 # # #
505
506 # # #
507
508 # # #
509
510 # # #
511
512 # # #
513
514 # # #
515
516 # # #
517
518 # # #
519
520 # # #
521
522 # # #
523
524 # # #
525
526 # # #
527
528 # # #
529
530 # # #
531
532 # # #
533
534 # # #
535
536 # # #
537
538 # # #
539
540 # # #
541
542 # # #
543
544 # # #
545
546 # # #
547
548 # # #
549
550 # # #
551
552 # # #
553
554 # # #
555
556 # # #
557
558 # # #
559
560 # # #
561
562 # # #
563
564 # # #
565
566 # # #
567
568 # # #
569
570 # # #
571
572 # # #
573
574 # # #
575
576 # # #
577
578 # # #
579
580 # # #
581
582 # # #
583
584 # # #
585
586 # # #
587
588 # # #
589
590 # # #
591
592 # # #
593
594 # # #
595
596 # # #
597
598 # # #
599
600 # # #
601
602 # # #
603
604 # # #
605
606 # # #
607
608 # # #
609
610 # # #
611
612 # # #
613
614 # # #
615
616 # # #
617
618 # # #
619
620 # # #
621
622 # # #
623
624 # # #
625
626 # # #
627
628 # # #
629
630 # # #
631
632 # # #
633
634 # # #
635
636 # # #
637
638 # # #
639
640 # # #
641
642 # # #
643
644 # # #
645
646 # # #
647
648 # # #
649
650 # # #
651
652 # # #
653
654 # # #
655
656 # # #
657
658 # # #
659
660 # # #
661
662 # # #
663
664 # # #
665
666 # # #
667
668 # # #
669
670 # # #
671
672 # # #
673
674 # # #
675
676 # # #
677
678 # # #
679
680 # # #
681
682 # # #
683
684 # # #
685
686 # # #
687
688 # # #
689
690 # # #
691
692 # # #
693
694 # # #
695
696 # # #
697
698 # # #
699
700 # # #
701
702 # # #
703
704 # # #
705
706 # # #
707
708 # # #
709
710 # # #
711
712 # # #
713
714 # # #
715
716 # # #
717
718 # # #
719
720 # # #
721
722 # # #
```

VI. RESULT

A. Random Forest

B. Decision Tree

C. Logistic Regression

A bar chart titled "Comparison of Accuracy Scores" comparing the accuracy of three models. The y-axis is labeled "Accuracy" and ranges from 0.0 to 1.0 in increments of 0.2. The x-axis is labeled "Models" and has three categories: "Model 1", "Model 2", and "Model 3". All three models have an accuracy of 1.0, represented by blue bars reaching the top of the chart.

Model	Accuracy
Model 1	1.0
Model 2	1.0
Model 3	1.0

Table I: Accuracy of models

Classifier	Accuracy
Random Forest	98%
Logistic Regression	98%
Decision Tree	99%

The effects of spreading erroneous information are always negative and harmful to society. The distinction between fake and true news continues to be widely misunderstood in society. Readers are frequently confused by false information, which causes them to get confused and behave improperly. This is why everyone's concern is unfounded. As a result, our paper can now utilise certainty to assess if the provided news is false or not. By considering the philosophy of the paper, people may at least confirm whether the news they are viewing is accurate. This paper discussed three approaches viz. random forest, decision tree and logistic regression. Their accuracies are shown in the Table I.

[1] R. Meyer. The Grim Conclusions of the Largest-Ever Study of Fake News ,<https://www.theatlantic.com/technology/archive/2018/03/largest-study-ever-fake-news-mit-twitter/555104/>

[2] Z. Sharf. and A. Us Saeed. Twitter News Credibility Meter, International Journal of Computer Applications, 83(6), 49-51, 2013.

[3] M. Granik and V. Mesyura, "Fake news detection using naive Bayes classifier," 2017 IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON), Kiev, 2017, pp. 900-903.

[4] Wang, W. Y. (2017). " liar, liar pants on fire": A new benchmark dataset for fake news detection. arXiv preprint arXiv:1705.00648.

[5] H. Gupta, M. S. Jamal, S. Madisetty and M. S. Desarkar, "A framework for real-time spam detection in Twitter," 2018 10th International Conference on Communication Systems & Networks (COMSNETS), Bengaluru, 2018, pp. 380-383

[6] Della Vedova, M. L., Tacchini, E., Moret, S., Ballarin, G., DiPierro, M., & de Alfaro, L. (2018). Automatic Online Fake News Detection Combining Content and Social Signals. In 2018 22nd Conference of Open Innovations Association (FRUCT), pp. 272-279. IEEE

[7] M. Rahul R., et al. "Identification of Fake News Using Machine Learning." 2020 IEEE International Conference on Electronics, Computing and Communication Technologies (CONECCT). IEEE, 2020.

[8] Zhang, J., Cui, L., Fu, Y., & Gouza, F. B. (2018). Fake news detection with deep diffusive network model. arXiv preprint arXiv:1805.08751.

[9] P. Bharadwaj, and Z. Shao. "Fake news detection with semantic features and text mining." International Journal on Natural Language Computing (IJNLC), Vol 8, 2019.

[10] Junaed Younus Khan, Md Khondaker, Tawkat Islam, Anindya Iqbal, and Sadia Afroz. A benchmark study on machine learning methods for fake news detection. arXiv preprint arXiv:1905.04749, 2019.