

Software Defect Prediction using Deep Learning: Systematic Literature Review

Swati Rai¹, Dr. Pinaki Ghosh², Dr. Gourav Shrivastava³, Dr. Kirti Jain⁴
¹Research Scholar, ²Head and Professor, ³Associate Professor, ⁴Associate Professor
SSAC, SAGE University, Bhopal, India

Abstract - Defect prediction is one of the key challenges in software development and programming language research for improving software quality and reliability. The problem in this area is to properly identify the defective source code with high accuracy. Developing a fault prediction model is a challenging problem, and many approaches have been proposed throughout history. The recent breakthrough in machine learning technologies, especially the development of deep learning techniques, has led to many problems being solved by these methods. Our survey focuses on the deep learning techniques for defect prediction. We analyze the recent works on the topic, study the methods for automatic learning of the semantic and structural features from the code, discuss the open problems and present the recent trends in the field.

Keywords: defect prediction; anomaly detection; program analysis; code understanding; neural networks; deep learning

How to cite this article: Swati Rai, Dr. Pinaki Ghosh, Dr. Gourav Shrivastava, Dr. Kirti Jain
"Software Defect Prediction using Deep Learning: Systemic Literature Review" Published in
International Journal of Scientific Modern Research and Technology (IJSMT), ISSN: 2582-8150,
Volume-12, Issue-1, Number 4, July 2023, pp.19-31, URL: www.ijsmrt.com/wp-content/uploads/2026/09/IJSMT-23070104.pdf

Copyright © 2023 by author (s) and International Journal of Scientific Modern Research and Technology Journal. This is an Open Access article distributed under the terms of the Creative Commons Attribution License (CC BY 4.0) (<http://creativecommons.org/licenses/by/4.0/>)



IJSMT-23070104

I. INTRODUCTION

According to the IEEE Standard Classification for Software Anomalies [1], a software defect is "an imperfection or deficiency in a work product where that work product does not meet its requirements or specifications and needs to be either repaired or replaced". Software defects can cause different problems. Common ways to find software defects are manual testing and code review. The main drawback of these methods is that they are quite expensive in terms of time and effort.

The automatic approaches to the Software Defect Prediction (SDP) would allow one to reduce the costs and improve quality of the software projects. Thus, Software Defect Prediction is an important problem in the fields of the software engineering and

programming language research. The task is to identify the defective code with high accuracy (in terms of the precision and recall). The development and breakthrough of machine learning led to the fact that many tasks can be solved by these methods. Recent advances in the fields of artificial neural networks and machine learning, as well as the increasing power of the modern computers (such as supercomputers based on GPUs with AI accelerating modules), allowed new concepts, such as deep learning, to emerge.

The main idea is that an artificial neural network with multiple layers is capable of progressively extracting the higher-level features from the original data to solve complex problems. For the problem of software defect prediction, the researchers have proposed the representation-learning algorithms to learn semantic

representations of programs automatically and use this representation to identify the defect-prone code. Using these implicit features shows better results than the previous approaches based on the explicit features, such as the code metrics [2].

The software defect prediction is a rapidly developing field, and the state-of-the-art surveys on the topic [3–5] do not sufficiently cover the recent works describing the cuttingedge techniques. For example, recent advances in the related fields of Natural Language Processing (NLP) provided new powerful tools such as Transformer language models. These techniques were later successfully applied to the software engineering tasks. The goal of our survey is to describe these latest achievements taking into account the newest primary studies published in 2019–2021. We hope that this survey can be useful for researchers and practitioners in the software defect prediction, code understanding and other related fields.

Some semantic defects are hard to find using only source code. For example, in [6], the bytecode of Kotlin programs is processed to detect the so called compiler-induced anomalies, which arise only in the compiled bytecode. Another example is presented in [7] where to expose the program behavior, the assembly code (generated from the C source code by the compiler) is used to learn the defect features.

Nevertheless, the source code remains the main source of data for the defect prediction. In this survey, our main interest lies in techniques devoted to analyze the source code. Usually, the process of developing the model for the defect prediction consists of the following steps (see Figure 1):

1. Prepare the dataset by collecting the source code samples from repositories of the software projects (or choose the suitable existing dataset).
2. Extract features from the source code.
3. Train the model using the train dataset.
4. Test the model using the test dataset and assess the performance using the quality metrics.

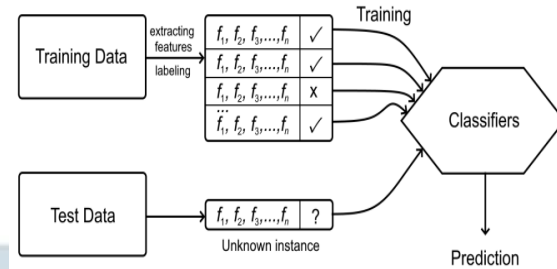


Figure 1: Scheme of the process of constructing the defect prediction model.

The survey is structured as follows: Section 2 briefly describes the methodology of our survey. Section 3 presents the overview on the various deep learning techniques applied to the defect prediction. In Section 4, we outline the main difficulties of the problem. Section 5 presents the study of the latest trends in the techniques and methods for defect prediction. Section 6 concludes the study and offers our vision on the future developments on the field.

II. RESEARCH QUESTIONS

To summarize the work of our survey, let us formulate the following research questions:

- RQ1. What deep learning techniques have been applied to software defect prediction?
- RQ2. What are the key factors contributing to the difficulty of the problem?
- RQ3. What are the trends in the primary studies on the use of deep learning for the software defect prediction?

A. RQ1. What Techniques Have Been Applied to This Problem?

In order to work with the source code, we need to have its representation. On the one hand, this representation should be simple as a vector, since most machine learning algorithms work with vectors. On the other hand, the representation should contain all the necessary information. The numerical vector representing the source code is called an “embedding”.

There are different ways to represent the source code. Moreover, we need different granularities for different tasks, for example, for code completion we need token-level embedding and for function clone detection we need function embedding.

For the software defect prediction problem, various levels of granularity are used, such as sub-system, component, file/class, method and change (see [8,9] for more info on various code embeddings). One way is to create the vector from the hand-crafted features. This approach assumes that an expert invents a set of features and selects best of them (e.g., [10,11]). Usually, these features include the statistical characteristics of code, such as its size, code complexity, code churn or process metrics. Another way is to create the numerical vector by processing the source code. One way to represent the code is a sequence of elements.

Usually, they are code tokens or characters [12]. The neural networks based on the sequences are usually trained to predict the subsequent element. Another approach to build the representation of the source code is the abstract syntax trees (AST) [13]. The nodes of the tree correspond to the statement and operators, and the leaves represent the operands and values. The tree-based models are trained to predict the code by generating new nodes taking into account the existing tree structure. The most common approach to defect prediction is to use some classification algorithm to divide the source code into two categories: defect code and correct one (e.g., [14]).

However, the approaches based on the hand-crafted features usually do not sufficiently capture the syntax and semantics of the source code. Most traditional code metrics cannot distinguish code fragments if these fragments have the same structure and complexity but implement a different functionality. For example, if we switch several lines in the code fragments, traditional features, such as the number of lines of code, number of function calls and number of tokens, would remain the same (see [2]). Thus, the semantic information is more important for defect prediction than these metrics.

Modern approaches are usually based on extracting the implicit structural, syntax and semantic feature from

the source code rather than using the explicit hand-crafted ones. The most popular deep learning techniques for software defect prediction are: Deep Belief Networks (DBN), Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM), and Transformer architecture.

(i) Deep Belief Networks

Deep Belief Network [15] generative models are based on a multilevel neural network. This network contains one input layer, one output layer and multiple hidden layers. The output layer generates a feature vector representing the data fed to the input layer. Each layer consists of the stochastic nodes. The important feature of the DBN is that the nodes are only connected to the nodes in the adjacent layers but not to the nodes within the same layer as shown in Figure 2.

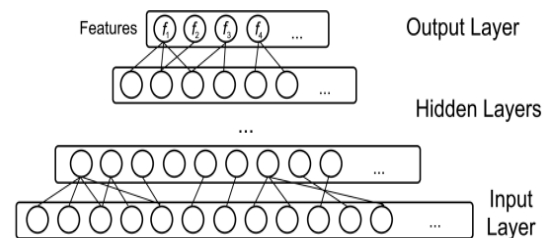


Figure 2: Architecture of the Deep Belief Network.

Perhaps one of the first works combining AST with the deep learning is [16]. The authors propose the approach for software defect prediction on a changes level. The DBN (which is fed by the traditional code metrics) generates the new expressive features and use them in classical machine learning classifiers. They extract the relations from the traditional code metrics, such as number of modified modules, directories and files, added and deleted lines, and several features related to the developer's experience. Later, the authors proposed the "TLEL" approach [17] based on the decision tree and ensemble learning for classification.

The works of Wang et al. [2,18] also use the DBN, but in a different manner. For predicting the defects on the basis of the code semantics, the authors have developed a DBN to automatically learn a semantic features from the source code. As the input for the network, the programs' AST and source code changes

are used for the cases of file-level and change-level prediction, respectively. Then, the authors use the classical machine learning classifiers and extracted features to classify source code files whether they are buggy or clean. The main drawback of the DBN is that it does not sufficiently capture the context of the code elements, such as the order of statement execution and function calls.

(ii) Long Short-Term Memory

The Long Short-Term Memory [19] is a subtype of the recurrent neural network specialized for processing the data sequences. The LSTM network consists of LSTM units (see Figure 3). The key element of the unit is a memory cell, which allows the unit to store the values for a short, as well as, for a long time intervals. This provides the LSTM-based models the ability to capture the long-range context information from the source code.

The LSTM-based model was used in work [11] for learning both the semantic and syntactic features of code. The proposed approach represents the code as a sequence of code tokens, which is fed into a LSTM system to transform code into a feature vector and a token state representing the semantic information of the token. Later the Tree-LSTM model was developed using the AST representation as input [20].

A neural bug finding technique is proposed in [21]. The authors train a neural network on examples of the defective and correct code, and then use the resulting binary classifier for bug detection. To prepare a labeled dataset, the authors use the existing static bug detection software to identify the specific kind of bugs. The code is represented as a tokens sequence and converted to a real-value vector by using the one-hot encoding for each token. Then, a bi-directional network with LSTM is used as model.

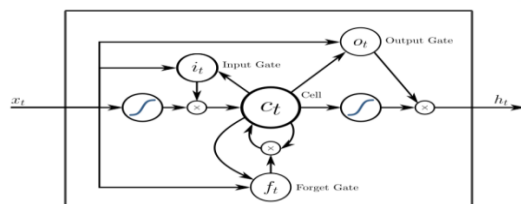


Figure 3: Scheme of the LSTM unit.

In [22], the authors propose a model for defect prediction on the base of AST path pair representation. To process the code, the path in the AST is extracted as combination of symbol sequence and control sequence. These sequences are fed to a Bi-LSTM network to generate a path vector. Then, all the vectors are combined using the global attention technique to generate the vector for the entire code fragment. These final embedding representations are used for classification.

(iii) Convolutional Neural Networks

The Convolutional Neural Networks [23] are a type of neural network specialized for processing the data with a mesh-like structure. This network is characterized by two important features. Firstly, the local connection pattern between the units is repeated over the entire network. It allows the network to capture the short-term structural context of the source code. Secondly, each unit have the same parameters. It allows the network to learn the information on the code element irrespective of its position in the code. The scheme of general CNN is shown in Figure 4.

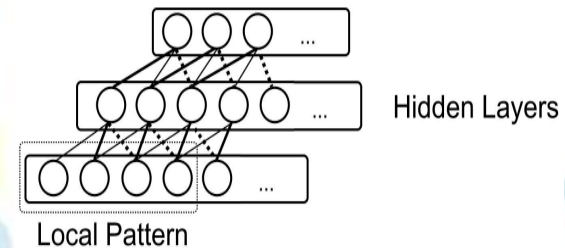


Figure 4: Architecture of the Convolutional Neural Network.

Reference [24] presents the model based on the CNN architecture. Based on the program's AST, the token vectors are extracted and converted to numerical vectors. Then, these vectors are fed into a CNN. After that, the combination of the extracted semantic and structural features and code metrics is used for software defect prediction applying the logistic regression.

A deep learning model to predict defects on the basis of the commit messages and code changes is developed in [25]. This model is based on the CNN. It uses the convolutional network layers for processing

the code changes and commit text and the feature combination layer to fuse these two embedding vectors into a single one. Another deep learning-based model for defect prediction is proposed in [26]. The training of the neural network utilizes the triplet loss technique and the weighted cross-entropy loss technique. The random forest is used as a classifier.

In [27], the features learning technique based on CNN is proposed. This model extract features from token vectors in the AST of the code and learns the transferable joint features. Combining these deep-learning-generated features with the hand-crafted ones allows the model to perform the cross-project defect prediction. Later, the authors propose a new treebased convolutional network to perform this task [28]. It uses the tree-based continuous bag-of-word for encoding the AST nodes to be fed into CNN.

(iv) Transformer Models

Recently, the big success of pre-trained contextual representations in the NLP, for example, [29], led to a rise of attempts to apply these techniques to source code. Usually, these models are based on the multi-layer Transformer architecture [30] shown in Figure 5. They are pre-trained using the massive unlabeled corpora of programs with the self-supervised objectives, such as masking language modelling and next sentence prediction [31,32]. After the pre-training phase, the model can be fine-tuned for specific tasks using the supervised techniques.

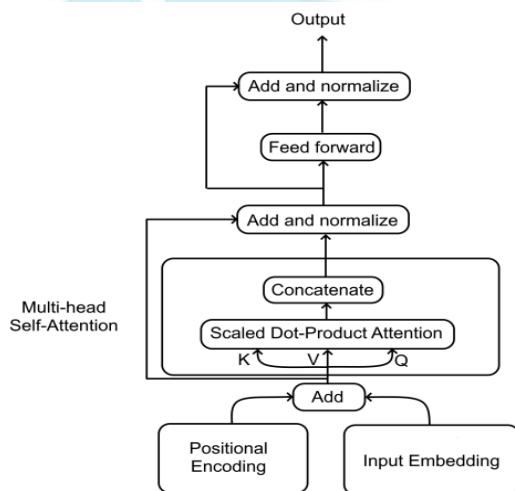


Figure 5: Architecture of the multi-layer transformer.

The authors of [33] state that the approaches based on the traditional complexity metrics are useless since there is no need for a tool to tell the engineer that longer and more complex code is more defect-prone. The methods of learning features from the source code do not guarantee capturing semantic and syntactical similarity, and very similar source codes can have very different features. These features can correlate with defects rather than directly cause them. In contrast, the authors propose an approach based on the self-attention transformer encoder to the semantic defect prediction. The matrix representing the defectiveness of each token in the fragment is generated. Attention and layer normalization are used as a regularization technique. The resulting model provides the defect prediction with the semantic highlight of defective code regions.

The CuBERT model is presented in [31]. The authors use a corpus of Python files from the GitHub to create a benchmark for evaluating code embeddings on five classification tasks and a program repair task. They train their model and compare it with various other models including the BiLSTM and Transformer. It is shown that the CuBERT outperforms the baseline models consistently.

A bimodal language model called CodeBERT is presented in [32]. It is based on the multilayer bidirectional Transformer neural architecture. To prepare the data, the natural language text is represented as a sequence of words, and the source code is presented as a sequence of tokens. The output of the CodeBERT model is a contextual vector learned from the natural language and source code, as well as the aggregated sequence. The resulting model efficiently solves the problems of both code to the documentation and natural language code search.

Work [34] presents a multi-layer bidirectional transformer architecture GraphCode-BERT, which utilizes three components as input: the source code, paired comments and data flow graph. Data flow graph represents relations between variables, for example, where the value of a variable comes from. This allows the model to consider the code structure for code representation. For pre-training tasks, the traditional masked language modeling, as well as the edge prediction and node alignment of data flow graph were

used. It supports several downstream code-related tasks including the code clone detection, code translation and code refinement.

(v) Other Networks

In [35], a software defect prediction technique based on stacked denoising autoencoders model is presented. The stacked denoising autoencoder is used to extract higher level features from the traditional metrics. The two-stage ensemble learning is used for classification. To address the class imbalance, the authors use the ensemble learning strategy. Later, the feature selection algorithm was applied to this method to address the feature redundancy problem [36].

A model for the software defect prediction was constructed in work [37] on the base of the Siamese parallel fully-connected networks. This model utilizes the paired parallel Siamese networks architecture and the deep learning approach. The network produces the high-level features that are used for classification. To address the imbalance between the minority and majority classes, the network takes into account the cost-sensitivity features.

The neural forest networks are used to learn feature representations in [38]. To perform a classification, a decision forest is used. It also guides the learning of the neural network. In [39], a new deep forest model is proposed for the software defect prediction. To detect the essential defect features, it uses the cascade

learning strategy, which consists in reforming a set of the random forest classifiers into a layered network.

The graph neural network to predict the software defects is constructed in work [40]. It extracts the semantics and context features from the AST of the code fragments. To capture the defect-related information from the source code, the ASTs for the buggy and fixed version of a fragment are constructed and pruned using the community detection algorithm, which extracts the defect-related subtree. Then, the Graph Neural Network is used to capture the latent defect information.

B. RQ2. What Are the Key Factors Contributing to Difficulty of the Problem?

The problem of software defect prediction is considered very complex and very challenging for the machine learning models based on the neural networks.

(i) Lack of Data

One of the difficulties is lack of available large labelled datasets devoted to the defect prediction. To alleviate this problem, one can utilize the pre-trained contextual embeddings. This technique consists in pre-training the language model on a massive corpora of unlabeled source code using the self-supervised objectives, such as masked language modelling, next sentence prediction and replaced token detection.

Table 1: presents the popular unlabeled code datasets suitable for this task.

Dataset	Content	Size	Used in Tasks
Bigquery github repos [3]	Python source code	4 M files	Pre-training CuBERT model
Py150 [4]	Python source code, AST	8423 repos, 149,993 files	Fine-tuning CuBERT model

Js150 [5]	Javascript source code, AST	150,000 source files	Code Summarization; Defect Prediction
Datasets for [6]	Java source code	9500 projects, 16 M samples in the largest one	Code summarization
GitHub Java Corpus [8]	Java source code	11,000 projects	Language Modelling
CodeNN Dataset [9]	C# source code and summaries	66,015 fragments	Code Captioning
Dataset for [6]	Kotlin source code, AST, bytecode	47,751 repos, 932,548 files, 4,044,790 functions	Anomaly detection, defect prediction
Dataset for [10]	C# source code	29 projects, 2.9 M lines of code	Variable Misuse detection

The pre-trained model may then be fine-tuned for the defect prediction using much smaller labeled datasets. Table 2 presents a list of publicly available datasets

devoted to the defect prediction. Usually, such datasets include pairs of correct and defective code fragments

Table 2: List of labeled datasets.

Dataset	Content	Size	Used in Tasks
SEIP Lab Software Defect Prediction Data [11]	Complexity metrics	5 subsequent releases of 3 projects from the Java Eclipse community	Data collection and linking
PROMISE Software Engineering Repository [12]	Numeric metrics; reported defects (false/true)	15,000 modules	Defect prediction

NASA Defect Dataset [13]	Numeric metrics; reported defects (false/true)	51,000 modules	Defect prediction
REPD datasets [15]	Numeric metrics, semantic features, reported defects	10,885 fragments in the largest one	Defect prediction
GPHR [16]	Java code and metrics	3526 pairs of fragments, buggy and fixed, code metrics	Defect prediction
BugHunter [17]	Java source code; metrics; fix-inducing commit; number of reported bugs	159 k pairs for 3 granularity levels (file/class/method), 15 projects	Analyzing the importance of complexity metrics
GitHub Bug DataSet [18]	Java source code; code metrics; number of reported bugs and vulnerabilities	15 projects; 183 k classes	Bug prediction
Unified Bug Dataset [19]	Java source code; code metrics; number of reported bugs	47,618 classes; 43,744 files	Bug prediction
Neural Code Translator Dataset [20]	Pairs of buggy and fixed abstracted method-level fragments 46 k pairs of small fragments	(under 50 tokens), 50 k pairs of medium fragments (under 100 tokens)	Code refinement

As with the other factors affecting the difficulty of constructing datasets, we can highlight that the distribution of the classes in the real code projects is often imbalanced. Usually, there are fewer buggy files or methods in a project than the correct ones. This may lead to the situation where the common classifiers would correctly detect the major class (correct code)

and ignore the much smaller class of the defect-prone code. This will lead to bad performance of the model.

To address this imbalance, several oversampling methods are proposed. The authors constructed hybrid approaches. It is based on the Synthetic Minority Over-Sampling Technique (SMOTE and SMOTUNED) for preparing the datasets and ensemble

approaches for classifying the defective and correct code. In [22], the authors take into account the proportion of the correct and defective code in each project in the dataset. To balance the classes, they duplicate the elements of the smaller class.

(ii) Lack of Context

Another problem is the complexity of the context for the code. Unlike the natural texts, the code element may depend on another element located far away, maybe, even in another code fragment. Moreover, it is often hard to say if the code element is defective without considering its context. If dataset consists of the pairs of bugged and fixed code fragments, it is often hard to extract the essence of defect.

Approaches based on the Transformer networks were aimed to NLP problems where data display a great deal of locality of reference. Most information about a token can be derived from its neighboring tokens. Thus, most such models represent the source code as a sequence of tokens.

The traditional Transformer architectures based on self-attention matrices do not scale well because of quadratic complexity. Usually, they are designed to handle the input sequences with limited length (usually, 512 or 1024 tokens). Therefore, their applicability to understanding the context of the source code is limited.

There are several modifications to the Transformer architecture that improve its ability to comprehend long sequences. These approaches alleviate the problem of limited length of the input, giving the Transformers the potential to work with a complex context of the source code. Another approach is to capture the structural and global relations on the code, combining the sequence-based and graph-based models for code representation [34]. Thus, representing the code context is essential in the software defect prediction.

C. RQ3. What Are the Trends in the Primary Studies on the Use of Deep Learning for the Software Defect Prediction?

The earliest works, such as [16], utilize the deep learning techniques trying to extract the implicit features from the traditional explicit features (such as code metrics). The main drawback of this approach is that these traditional features usually cannot capture the semantic difference between the correct and defective code. Therefore, the combination of these features would also fail to do this [24].

Later approaches [20,25] use the generic or tailored deep learning techniques to extract the semantic and syntactic features directly from the source code, usually, from the abstract syntax trees. These deep learned features are used in combination with the traditional ones in the machine classifiers to produce the accurate defect prediction.

Modern software development often prioritizes writing the human-readable source code. This includes using the meaningful names for the functions and variables and writing the code documentation in natural language. This leads to a situation where we can extract the semantic information from the source code using the techniques originally intended for the NLP, such as the pre-trained language representations such as BERT [7].

Learning useful models with supervised setting is often difficult because labelled data are usually limited. Thus, many unsupervised approaches have been proposed recently to utilize the large unlabelled datasets that are more readily available. Usually, this means that pre-training is performed with automatic supervisions without manual annotation of the samples. Then, the model may be fine-tuned for the specific task using much smaller supervised data [31].

The most recent techniques in software engineering are based on using the general purposed pre-trained models for programming languages [34]. These models learn to “understand” the source code from unlabelled datasets using the self-supervised objectives. A large corpus of source code is used for pre-training. Usually, the objective is the Masked Language Modelling where at some positions the tokens are masked out and the model must predict the original token [32]. Utilizing these techniques alleviates the need for the task-specific architectures

and training on large labelled datasets for each task separately.

III. CONCLUSIONS

One of the major challenges in modern software engineering is predicting defective code. Recent developments in the field of machine learning, especially the multi-layered neural networks and deep learning algorithms, provide powerful techniques, which utilize learning algorithms for representations of the source code that captures semantic and structural information.

This survey presents the latest research progress in software defect prediction using the deep learning techniques, such as the Transformer architectures. We formulate the main difficulties of the defect prediction problem as lack of data and complexity of context and discuss the ways to alleviate these problems.

Taking into account the latest trends in the machine learning techniques for the software defect prediction problem, we believe that progress in this field will be achieved largely due to the implementation of the following ideas.

- To reduce the requirements for the size of the labelled datasets, one should use the self-supervised training on large corpora of the unlabelled data. In addition, it is necessary to use the unlabelled data for the pre-training of related tasks and to contribute to the fact that the trained models will have a deeper and more comprehensive understanding of the source code. This, in the turn, will allow one to find the deeper defects.
- To leverage the latest advances in the machine learning techniques in the natural language processing in the programming languages, we are already seeing the successful migration of these methods to solve various code understanding problems. For example, optimization of the self-attention mechanism for the transformers will allow one to use them for long sequences, which, in the turn, will lead to a more complete consideration of the code context for finding the defects.

- Often a defect is not limited to a single line of code or one function, and there are various ways to fix it. For example, a bug can be fixed either inside the function or at calling this function. Thus, the defect ceases to have specific coordinates inside the source file. In addition, not being an explicit defect, a line of code can become defective at a certain point in time. A changed context may lead to the fact that the purpose of the code changes, and, therefore, the old implementation no longer corresponds to the new requirements or specifications.

All this leads to a blurring of the concept of a defect. Thus, we come to the concepts of “potentially defective” code or “strange” code. In this regard, as promising problems, we want to note the task of finding an atypical (or anomalous) code and the task of the code refinement. These task require good representations of the code and code changes, taking into account the specifics of the source code, such as structure and context.

It is difficult to state which of the state-of-the-art models performs in the best way. There are no universally accepted standard benchmarks for the problem and different researchers utilize different performance metrics and use different data. Thus, the experimental results from the primary works cannot be directly compared. The existing comparative studies such as show that while the state-of-the-art deep learning techniques usually perform better than standard deep learning and traditional metrics-based ones (achieving the increase of F1 from 60% up to 80% in some cases). None of the approaches achieves a consistently high performance in terms of recall, precision and accuracy sufficient for the practical application. Thus, the defect prediction problem remains an open one.

REFERENCES

- [1] Jitimon Angskun, Suda Tipprasert and Thara Angskun, “Big data analytics on social networks for realtime depression detection”, Journal of Big Data, 2022.
- [2] Md. Rafidul Hasan Khan; Umme Sunzida Afroz; Abu Kaisar Mohammad Masum; Sheikh Abujar; Syed Akhter Hossain, “Sentiment Analysis

from Bengali Depression Dataset using Machine Learning”, 11th International Conference on Computing, Communication and Networking Technologies (ICCCNT), 2020.

[3] Tanna D, Dudhane M, Sardar A. Deshpande K, Deshmukh N., “Sentiment analysis on social media for emotion classification”, In: International Conference on intelligent computing and control systems, 2020.

[4] Md Kamrul Hasan, Hasan Mahmud, Ahmed Al Marouf, “Comparative Analysis of Feature Selection Algorithms for Computational Personality Prediction From Social Media”, IEEE Transactions on Computational Social Systems, 2020.

[5] Yang X, McEwen R, Ong LR, Zihayat M., “A big data analytics framework for detecting user-level depression from social networks”, Int J Inf Manag. 2020.

[6] Lyua YW, Chow JC-C, Hwang J-J, “Exploring public attitudes of child abuse in mainland China: a sentiment analysis of China’s social media”, Weibo. Child Youth Serv Rev. 2020.

[7] Chen B, Cheng L, Chen R, Huang Q, Phoebe Chen Y-P. Deep neural networks for multiclass sentiment classification. In: IEEE 20th International Conference on high performance computing and communications, IEEE 16th International Conference on Smart City, IEEE 4th International Conference on Data Science and Systems 2018; pp. 854–59.

[8] Sethi M, Pande S, Trar P, Soni P. Sentiment identification in COVID-19 specific tweets. In: International Conference on electronics and sustainable communication systems (ICESC 2020), pp. 509–16, <https://doi.org/10.1109/ICESC48915.2020.9155674>.

[9] Kundale JU, Kulkarni NJ. Language independent multi-class sentiment analysis. In: 5th International Conference on computing communication control and automation (ICCUBEA), 2019; pp. 1–7, <https://doi.org/10.1109/ICCUBEA47591.2019.9128383>.

[10] Ruz GA, Henriquez PA, Mascareno A. Sentiment analysis of Twitter data during critical events through Bayesian networks classifiers. Future Gener Comput Syst. 2020;106:92–104.

[11] Yang X, McEwen R, Ong LR, Zihayat M. A big data analytics framework for detecting user-level depression from social networks. Int J Inf Manag. 2020;54:102141.

[12] Tao X, Dharmalingam R, Zhang J, Zhou X, Li L, Gururajan R. Twitter analysis for depression on social networks based on sentiment and stress. In: 6th International Conference on behavioral, economic and socio-cultural computing, 2019; pp. 1-4, <https://doi.org/10.1109/BESC48373.2019.8963550>.

[13] Tanna D, Dudhane M, Sardar A. Deshpande K, Deshmukh N. Sentiment analysis on social media for emotion classification. In: International Conference on intelligent computing and control systems (ICICCS 2020), pp. 911–15, <https://doi.org/10.1109/ICICCS48265.2020.9121057>.

[14] Arora P, Arora P. Mining Twitter data for depression detection. In: IEEE International Conference on signal processing and communication (ICSC), 2019; pp. 186–89, <https://doi.org/10.1109/ICSC45622.2019.8938353>.

[15] Chen Y, Zhou B, Zhang W, Gong W, Sun G. Sentiment analysis based on deep learning and its application in screening for perinatal depression. In: IEEE Third International Conference on data science in cyberspace. 2018; pp. 451–6. <https://doi.org/10.1109/DSC.2018.00073>.

[16] Uddin AH, Bapery D, Arif ASM. Depression analysis from social media data in Bangla language using long short term memory (LSTM) recurrent neural network technique. In: International Conference on computer, communication, chemical, materials and electronic engineering (IC4ME2), 11–12 July, 2019; pp. 1-4, <https://doi.org/10.1109/IC4ME247184.2019.9036528>.

[17] Cheng L-C, Tsai S-L. Deep learning for automated sentiment analysis of social media. In:

IEEE/ACM International Conference on advances in social networks analysis and mining. 2019; pp. 1001–4. <https://doi.org/10.1145/3341161.3344821>.

[18] Al Asad N, Pranto MAM, Afreen S, Islam MM. Depression detection by analyzing social media posts of user. In: IEEE International Conference on signal processing, information, communication & systems (SPICSCON). Dhaka, Bangladesh, 2019; pp. 13–17, <https://doi.org/10.1109/SPICSCON48833.2019.9065101>.

[19] Lyua YW, Chow JC-C, Hwang J-J. Exploring public attitudes of child abuse in mainland China: a sentiment analysis of China's social media Weibo. Child Youth Serv Rev. 2020;116:102520.

[20] Abid F, Li C, Alam M. Multi-source social media data sentiment analysis using bidirectional recurrent convolutional neural networks. Comput Commun. 2020;157:102–15.

[21] Hammou BA, Lahcen AA, Mouline S. Towards a real-time processing framework based on improved distributed recurrent neural network variants with fastText for social big data analytics. Inf Process Manag. 2020;57:102122.

[22] Tadessi MM, Lin H, Xu B, Yang L. Detection of depression related posts in reddit social media forum. IEEE Access. 2019;7:44883–93. <https://doi.org/10.1109/ACCESS.2019.2909180>.

[23] Trotzek M, Koitka S, Friedrich CM. Utilizing neural networks and linguistic metadata for early detection of depression indications in text sequences. IEEE Trans Knowl Data Eng. 2018;32:588–601.

[24] Tariq S, Akhtar N, Afzal H, Khalid S, Mufti MR, Hussain S, Habib A, Ahmad G. A novel co-training based approach for the classification of mental illnesses using Social media posts. IEEE Access. 2019;7:166165–72. <https://doi.org/10.1109/ACCESS.2019.2953087>.

[25] Rao G, Zhang Y, Zhang L, Cong Q, Feng Z. MGL-CNN: a hierarchical posts representations model for identifying depressed individuals in online

forums. IEEE Access. 2020;8:32395–403. <https://doi.org/10.1109/ACCESS.2020.297373>.

[26] Syarif I, Ningtias N, Badriyah T. Study on mental disorder detection via social media mining. In: IEEE. 2019; pp. 1–6. <https://doi.org/10.1109/CCCS.2019.8888096>.

[27] Hussain J, Satti FA, Afzal M, Khan WA, Bilal SM, Muhammad ZA, Hafz FA, Hur T, Bang J, Kim J-I, Park GH, Seung H, Lee S. Exploring the dominant features of social media for depression detection. J Inf Sci. 2019;46:1–21.

[28] Katchapakirin K, Wongpatikaseree K, Yomaboot P, Kaewpitakkun Y. Facebook social media for depression detection in the Thai community. In: 15th International Joint Conference on computer science and software engineering (JCSSE), 2018; pp. 1–6, <https://doi.org/10.1109/JCSSE.2018.8457362>.

[29] Yazdavar AH, Mahdavinjad MS, Bajaj G, Romine W, Sheth A, Monadjemi AH, Thirunarayan K, Meddar JM, Myers A, Pathak J, Hitzler P. Multimodal mental health analysis in social media. 2020. <https://doi.org/10.1371/journal.pone.0226248>.

[30] Islam MR, Kabir MA, Ahmed A, Kamal ARM, Wang H, Ulhaq A. Depression detection from social network data using machine learning techniques. Health Inf Sci Syst. 2018;6:1–12.

[31] Shen G, Jia J, Nie L, Feng F, Zhang C, Hu T, Chua T-S, Zhu W. Depression detection via harvesting social media: a multimodal dictionary learning solution. In: Twenty-Sixth International Joint Conference on artificial intelligence (IJCAI-17) 2017; pp. 3838–844.

[32] Kumar A, Sharma A, Arora A. Anxious depression prediction in real-time social data. In: International Conference on advanced engineering, science, management and technology—2019 (ICAESMT19).

[33] Nalinde PB, Shinde A. Machine learning framework for detection of psychological disorders at OSN. Int J Innov Technol Explor Eng (IJITEE). 2019;8(11), (ISSN: 2278-3075).

- [34] Tajuddin M, Kabeer M, Misbahuddin M. Analysis of social media for psychological stress detection using ontologies. In: Fourth International Conference on inventive systems and control (ICISC 2020) IEEE Xplore Part Number: CFP20J06-ART; ISBN: 978-1-7281-2813-9.
- [35] Baheti RR, Kinariwala S. Detection and analysis of stress using machine learning techniques. *Int J Eng Adv Technol (IJEAT)*. 2019; 9(1), (ISSN: 2249–8958).
- [36] Ahmad S, Asghar MZ, Alotaibi FM, Awan I. Detection and classification of social media-based extremist affiliations using sentiment analysis techniques. *Human Centric Comput Inf Sci*. 2019;24:1–23.
- [37] Cornn K. Identifying depression on social media. 2019. <https://web.stanford.edu/>.
- [38] Jabreel M, Moreno A. A deep learning-based approach for multi-label emotion classification in tweets. *MDPI Appl Sci*. 2019;9(6):1123. 33. Bouzazi M, Ohtsuki T. A pattern-based approach for multi-class sentiment analysis in Twitter. *IEEE Access*. 2017;5:20617–39. <https://doi.org/10.1109/ACCESS.2017.2740982>.
- [39] Rosa RL, Schwartz GM, Ruggiero WV, Rodriguez DZ. A knowledge-based recommendation system that includes sentiment analysis and deep learning. *IEEE Trans Ind Inf*. 2019;15(4):2124–35. <https://doi.org/10.1109/TII.2018.2867174>.
- [40] Yang L, Li Y, Wang J, Sherrarat RS. Sentiment analysis for E-commerce product reviews in chinese based on sentiment lexicon and deep learning. *IEEE Access*. 2020;8:23522–30. <https://doi.org/10.1109/ACCESS.2020.2969854>.
- [41] Sadr H, Pedram MM, Teshnehlal M. Multi-view deep network: a deep model based on learning features from heterogeneous neural networks for sentiment analysis. *IEEE Access*. 2020;8:86984–97. <https://doi.org/10.1109/ACCESS.2020.2992063>.
- [42] Chen F, Ji R, Jinsong S, Cao D, Gao Y. Predicting microblog sentiments via weakly supervised multi-modal deep learning. *IEEE Trans Multimed*. 2018;20(4):997–1007. <https://doi.org/10.1109/TMM.2017.2757769>.