

Software Vulnerability Assessment and Classification Using Machine Learning, Deep Learning and Feature Selection Techniques

Ali Hussein, Azri Azmi, Hafiza Abas

Faculty of Artificial Intelligence, Universiti Teknologi Malaysia (UTM), Kuala Lumpur, Malaysia

E-mail: hussein.a@graduate.utm.my

Keywords: vulnerability, deep learning, machine learning, classification, neural network, VSM, ANN, selection techniques

Received: January 29, 2025

The detection of software defects is a critical technique for improving software quality and optimizing testing resources. This study presents a novel approach to software vulnerability assessment and classification using Recurrent Neural Networks (RNNs) enhanced by feature selection techniques. The proposed methodology integrates data preprocessing, dynamic analysis methods, and vector space model (VSM) generation, leveraging techniques such as TF-IDF and relational feature extraction to normalize and balance datasets. Computational experiments were conducted using various real-world and synthetic datasets, comparing the proposed RNN framework to traditional machine learning models, including Artificial Neural Networks (ANN), Support Vector Machines (SVM), and Adaboost. The RNN model, optimized with activation functions such as ReLU, Sigmoid, and Tanh, demonstrated superior performance, achieving a classification accuracy of 97.5% with ReLU and outperforming other models in precision (97.6%), recall (97.9%), and F-measure metrics. These results highlight the robustness and effectiveness of the proposed framework in detecting vulnerabilities and mitigating software defects. This research underscores the potential of deep learning-based approaches in enhancing software reliability and security.

Povzetek: Članek analizira ranljivost programske opreme z uporabo globokih nevronske mreže (RNN), optimiziranih s tehnikami izbire značilk, kjer učinkovito obravnava neuravnotežene podatke.

1 Introduction

Recognition of vulnerabilities in source code or software has become an emerging field of research. Even though earlier studies have demonstrated the usefulness of multiple detection methods, models, and software vulnerability analysis tools in identifying source code vulnerabilities, the enhancement of the efficiency of such detection models and tools remains a major challenge for researchers. Annually, thousands of security issues are identified in virtual instruments, which are released publicly to the general vulnerability database exposed in obfuscated code. Threats also occur in indirect ways that are not evident to the concerned code inspectors or the programmers. It seems necessary to understand the dynamics of vulnerabilities that can lead to system issues directly from the raw data, with abundance of publicly available source code. In this work, we present an information approach to security technology using deep learning. Stimulated by the success of similar research in the area of recognizing the vulnerabilities in source code/software, we use a theoretical framework to examine its feasibility to aid in finding out the said vulnerabilities. The preliminary results indicate that within the domain of detecting attacks, the definition is feasible [1].

To bridge the domain gap, we can propose that each feature in a program be treated in computer vision as a neural network because fault detectors may only have to say whether a feature is insecure and fully explain vulnerability positions. That is, we want an intelligent

interpretation of fault management programs. On either hand, one may recommend approaching each piece of code (i.e., comment, in this study, the two words are used synonymous) as a vulnerability detection unit. There are important exceptions to this diagnosis:

- (i) most comments in a program may not contain any uncertainty, indicating that a few samples are susceptible;
- (ii) multiple comments are not regarded as a whole that is semantically linked to each other [2].

Using traditional programming by utilizing k-means cluster analysis and the generative adversarial model, a scheme was introduced to check bugs in large number of random codes. To select the optimal code with an interactive analysis framework and software code refactor generation, k-means cluster transformation has been used. Use of a system, described on object-oriented code analysis documented in literature; in the instant case. The model is verified by the tasks of conceptual relationship analysis based on coding pairs and identification of sentiments. Moreover, a research study centered on what form of compilation with communication of massive source code has helped to recognize errors for inexperienced developers and suggest the steps needed to be taken on source code mistakes. The investigation uses the form of a message previously based on the coding system and detects specific code snippets' vulnerability [3].

It is therefore noted that, as classification algorithms for web application, bug identification and vulnerability classification, most investigators have used conventional

semi-supervised classifiers, RNNs or CNNs. RNNs are far better than the standard language models, including such n-gram, but they have drawbacks in understanding long sequence data. Based on faults, functional programming identification, archive code identification, and basic error detection, almost all of the studies are found to have used various system software and classification models. On the other hand, the developed scheme of ours explicitly defines logic, grammar, and other system software errors. Additionally, in place of the error spot, the proposed model is used to predict the correct terms. Overall, in pursuing unique objectives, our suggested Language model varies from many other models [4].

This paper introduces a novel active tracking on deep learning to automatically learn features for predicting runtime environment weaknesses. In source code, where contingent code components are spread far apart, For example, combinations of code tokens that are needed to appear simultaneously due to the configuration of the computer program (e.g. in Java) or according to the configuration of API use (e.g. Lock () and activate ()), but do not accompany each other automatically are effectively handled. The interpretation of code symbols (semantic functionality) and the hierarchical structure of source code are appropriately reflected by the learned features (syntactic features). Our automated feature learning strategy removes the need for automated feature selection in conventional methods, which takes up much effort. Finally, testing the framework from a huge repository on many Java programs for the Desktop version reveals that our methodology is highly accurate in explaining code vulnerabilities [5]. Table 1 provides a Related Works Summary comparing the methodologies, datasets, and performance metrics of reviewed studies.

Table 1: Related works summary comparison

study	Methodology	Dataset Size	Key Metrics (Accuracy, Precision, Recall)	Limitations
Tai et al. (2015)	Convolutional Neural Network (CNN)	10,000+ samples	Accuracy : 92.3%	Limited to structured data; scalability and semantic feature extraction issues.
Zhou et al. (2016)	Support Vector Machine (SVM) + Text Mining	15,000+ bug reports	Accuracy : 89.5%, Recall: 91.2%	Inefficient for large datasets; high computational cost.
Teshima et al. (2018)	Long Short-Term Memory (LSTM) Networks	5,000+ snippets	Accuracy : 94.2%	Struggles with imbalanced datasets, leading to reduced recall.
Tong et al. (2018)	Stacked Autoencoder + Ensemble	20,000+ entries	Accuracy : 93.1%, Precision : 94.5%	High preprocessing cost; lacks scalability for diverse datasets.
Proposed Framework	Recurrent Neural Network (RNN) with TF-IDF and	2,500 (2000 training, 500 testing)	Accuracy : 97.5%, Precision : 97.6%,	Efficient handling of imbalanced datasets; reduced preprocessing;

	relational feature extraction		Recall: 97.9%	automated feature learning.
--	-------------------------------	--	---------------	-----------------------------

The proposed RNN framework distinguishes itself by integrating advanced feature selection methods, such as TF-IDF and relational feature extraction, with robust data normalization and balancing techniques. Additionally, its automated feature learning capabilities minimize preprocessing requirements, and its optimized architecture, featuring activation functions like ReLU and Tanh, achieves superior performance. These attributes establish the proposed framework as a novel and effective solution, addressing gaps in the state-of-the-art while setting new benchmarks in software vulnerability detection. This comprehensive comparison underscores the necessity and impact of the proposed solution in overcoming the limitations of existing methodologies.

2 Proposed system design

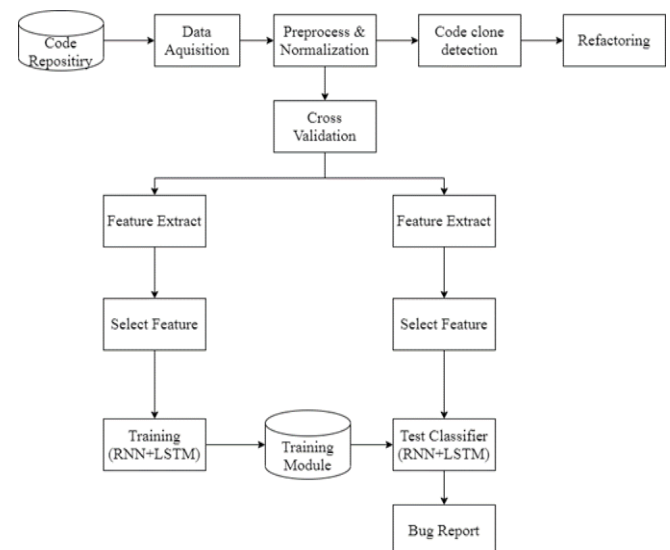


Figure 1: Proposed system architecture design shows a system overview of execution process flow, and delineates how it works with different algorithms.

2.1 Implementation environment

The proposed framework was implemented on a Windows operating system using JDK 1.7 and Python as the coding language. NetBeans IDE was utilized for development, with MySQL serving as the database for both back-end operations and storage. The front-end was designed using JPanel and JFrame. The hardware environment included an Intel i3 2.7 GHz CPU, 300 GB HDD, and 4 GB RAM, providing sufficient resources for the experiments and ensuring replicability.

2.2 Pre-processing

the source code is separated as well as the area of the comparison is first decided. There are three basic types of goals in the below steps.

Eliminate section of code: In this step source code uninteresting compared phase is removed.

Determine source units: By removing all the uninteresting code, the remaining part of the source code is divided into the arrangement of dissimilar sections known as source units.

Determine correlation units/granularity:

Source code parts should be auxiliary divided into smaller parts relying upon the evaluation method utilized a tool [6].

2.3 Feature extraction

Extraction changes the program to the form that is correct while supporting the real comparison algorithm. Conditional upon the device, it contains are as following:

Tokenization: If there should be an event of token-based methodologies, every source code line of the program is more divided into tokens as shown by the lexical regulation of the program design platforms of importance. Apply different tokens of source code lines or forms after that frame of token systems to compare. The entire whitespace and comments between marks are removed from the token groups.

Parsing: For syntactic methods, the whole source code is described to prepare a parse tree or (potentially clarified) abstract syntax tree (AST). The source parts to be studied are then shown as sub-trees of the described tree or the AST. Measurements-based methodologies can utilize a parse tree depiction to discover clones taking into account sizes for sub-trees.

Control and Data Flow Analysis: Semantics-related methodologies produce program dependence graphs (PDGs) as of the source code. The nodes of a Program Dependence Graphs show the reports and circumstances of a system, while edges show control and information conditions. Source units to be matched are shown as sub-graphs of these PDGs. Different plans then search in favor of isomorphic sub-graphs to discover clones. A few measurement-based methodologies use sub-graphs to compute info with control stream measurements [6]. The feature extraction process involves transforming raw source code into a vector space model (VSM) using a combination of TF-IDF, relational features, and bigram features. TF-IDF was applied with standard parameters, considering all terms across the dataset, to capture term importance relative to document frequency. Relational features were generated by analyzing dependencies between code tokens, such as variable usages and method invocations. For bigram features, sequential pairs of tokens were extracted to capture context within code snippets. The extracted features were normalized to ensure consistency, and imbalanced datasets were balanced through re-sampling techniques to improve classification performance. All preprocessing steps were implemented using Python libraries, ensuring reproducibility.

2.4 Feature selection

The various feature selection methods have been used during module training. The function compiles entire source code or modules with real statistics; in this method

behavior is analyzed of code for vulnerability detection. In a broader dataset, all of the variables are less necessary to consider; but, the greater the number of variables, the greater the difficulty. As a result, it is often preferable to reduce the number of variables in a dataset and to use critical variables [17]. We may reduce the parameter and find the variable's value in a dataset using a Function Selection technique. During the analysis, four dynamic analysis methods have been used, fault infusion, mutation suitable starting, dynamic taint assessment, and dynamic system check. To generate the vector Space Model (VSM) from extracted features.

2.5 Vulnerability detection

The vulnerability detection has been performed based on extracted features from the training data set. The vector space model has been generated according to extracted features like TF-IDF, relational features, and some bigram features. The classification has been done with recurrent neural networks, including long short-term memory algorithms. This detection is also effective for of prevention of software-as-a-service attacks for web applications. The vulnerable code generates internal as well as external attacks and grants unauthorized access to external users. The major objective of detection vulnerability is the automatic detection of exception handling and buffer overflow attacks during code execution. In the section proposed algorithm provides better detection accuracy in the code snippet [16]. Detailed Experimental Design

The synthetic dataset used in this study was generated by augmenting real-world vulnerability datasets with additional samples created through programmatic transformations, such as adding redundant code blocks and varying variable names to simulate real-world coding diversity. The dataset comprised 2,500 samples, partitioned into 2,000 for training and 500 for testing. Each sample was manually labeled based on known vulnerabilities.

The 20-fold cross-validation was chosen to ensure robust evaluation and reduce bias in model performance assessment. Hyperparameter optimization was conducted using grid search to fine-tune parameters, including the number of RNN layers (1–3), activation functions (ReLU, Sigmoid, Tanh), and dropout rates (0.2–0.5). Data augmentation techniques, such as oversampling of minority classes and SMOTE (Synthetic Minority Over-sampling Technique), were applied to mitigate the effects of class imbalance, ensuring fair training and evaluation of the RNN model.

3 Algorithm design

The algorithms furnished below have been used during the calculations of TF-IDF and weight score calculations using RNN.

TF-IDF:

Input: Input test instance that contains numerous tokens $T[i..n]$

Output: TF-IDF weight for all $T[i]$

Step 1: $Data_vector = \{Data1, Data2, Data3 \dots Data n\}$

Step 2: Words exist in the entire dataset

Step 3: $D = \{cmt1, cmt2, cmt3, cmtn\}$ and comments available in each document. Calculate the Tf score as

Step 4: $tf(t,d) = (t,d)$

$t = \text{term}$

$d = \text{document}$

Step 5: $idf = \frac{1}{\sum(d)}$

Step 6: Return $tf * idf$

Recurrent neural network:

Input: Training dataset $TestDBList []$, Train dataset $TrainDBList []$ and Threshold th .

Output: Predicted class according to classification

Step 1: Read train data rules using below formula

$$Train[] = \sum_{n=1}^k (Att_n \dots \dots \dots Att_k)$$

Step 2: Read test data rules using below formula

$$Test[] = \sum_{m=1}^k (Att_m \dots \dots \dots Att_k)$$

Step 3: Calculate the weight between the input and hidden layer

$Instance[w]$

$$= \sum_{n=1}^k (Test_n \dots \dots \dots Test_n) \sum_{m=1}^k (Train_m \dots \dots \dots Train_k)$$

Step 4: Generate a feedback layer based on the threshold policy

$$Feed_Layer[] = \sum_{m=1}^k (Feed_Layer.optimized())$$

Step 5: Return $Feed_Layer[0].class$

4 Results and discussion

To validate the evaluation of the proposed bug forecast procedure, we have employed RNN classification algorithms that are gainfully utilized for fault prediction including unlabeled datasets. The performance evaluations of software defect prediction are based on the confusion matrix, as shown in Table 2, which includes the measures of precision, recall, as well as F-score.

Table 2: Confusion matrix evaluation

Actual	Predictive
--------	------------

True	TP (true positive)	FN (false negative)
False	FP (false positive)	TN (true negative)

True positive (TP): The number of fake entities anticipated as fake.

False negative (FN): The number of fake entities anticipated as normal.

False positive (FP): The number of normal entities anticipated as fake.

True negative (TN): The number of normal entities anticipated as normal.

In this research, analytical performance procedures are calculated as follows:

Precision: It shows the proportion of faulty identities received adequate as faulty of all desired objects.

Recall: It is the percentage of faulty identities to all entities that are currently faulty is the proportions of recall.

F-measure: It is the cumulative recall and precision average, with higher estimated coefficients matching higher predictive efficiency.

$$F - Measure = \frac{2 * Precision * Recall}{Recall + Precision}$$

$$Precision = \frac{TP}{TP + FN} \quad Recall = \frac{TP}{TP + FP}$$

To evaluate the proposed system, we have used machine learning classifiers like ANN, SVM, and Adaboost. Also, we have used the deep learning framework of RNN with LSTM by using activation functions like Sigmoid, Tanh, and ReLU. The results of classification accuracy with confusion matrix with 20-fold cross-validation for all algorithms are shown in Table 3. Measures used to compare the algorithms are Accuracy, Precision, Recall, and Micro-score. From the observations, it can be concluded that RNN (ReLU) gives the highest performance among all.

4.1 Experiment using artificial neural network

Figure 2 shows the classification accuracy of the ANN classification algorithm. Initially, it was trained using inbuilt functions from the weka tool. Numerous cross-validation techniques have been used for classification, and various parameters have been tuned for ANN during the classification. This approach can classify each validation according to probability function, that the reason this algorithm has a higher error rate than other supervised classification algorithms.

Table 3: accuracy and confusion matrix for ANN

ANN	Fold 10	Fold 15	Fold 20
Accuracy	85.20	84.20	85.60
Precision	83.60	82.30	84.99
Recall	87.50	85.40	77.72
Micro-Score	85.05	83.35	81.10

The ANN model is easy to build and particularly useful for very large data classification using supervised

machine learning techniques or Artificial Intelligence (AI). Along with simplicity, ANN is known to outperform even highly sophisticated classification methods. The proposed ANN predicts the possibility for individual instances according to current values.

Figure 2 shows the performance evaluation calculation of ANN classification with 20-fold classification. It achieves around 85.60% accuracy for the given input dataset. We used a multinomial event model, samples represent the frequencies with which certain events have been generated by a multinomial probability of that particular event, and based on that probability system predict the final class.

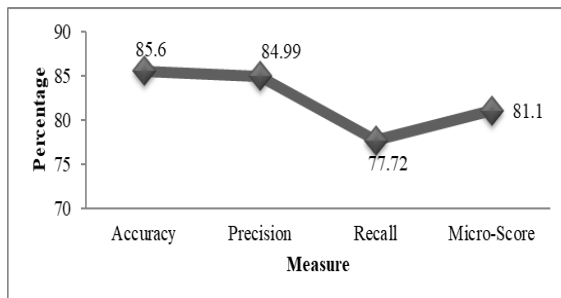


Figure 2: Analysis of bug and vulnerability detection using ANN with 20-fold data cross-validation

4.2 Experiment using support vector machine (SVM)

Table 4 depicts the classification analysis with various cross-validations, we conclude that 20-fold cross-validation provides the highest 95.2% classification accuracy for SVM.

Table 4: accuracy and confusion matrix for SVM

SVM	Fold 10	Fold 15	Fold 20
Accuracy	91.20	91.70	95.20
Precision	91.35	92.10	94.80
Recall	92.30	93.10	96.20
Micro-Score	91.35	92.20	94.75

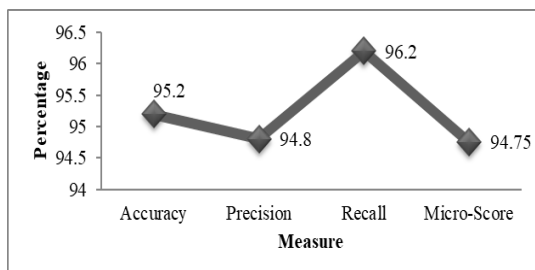


Fig. 3: Analysis of vulnerability detection using SVM with 20-fold data cross-validation

Figure 3 describes SVM for 20-fold cross-validation. The labeling circumstances to construct a training set becomes moment and expensive in many machines learning; it is also helpful to find strategies to reduce supervised classification numbers. By improving performance, the Kernel-Based algorithm has been used to minimize occurrences. We classify all in this algorithm as a point in n-dimensional spaces only with respect to a property direction being the meaning of each characteristic by classification technique; we detect clones by finding the hyper-plane that separates the two groups very well.

4.3 Experiment using adaboost

Below table 5 depicts the classification analysis with various cross-validation, we conclude that 20-fold cross-validation provides the highest 81.30% classification accuracy for Adaboost.

Table 5: Accuracy and confusion matrix for adaboost

Adaboost	Fold 10	Fold 15	Fold 20
Accuracy	70.60	78.50	81.30
Precision	72.30	73.50	74.50
Recall	69.90	68.50	70.30
Micro -Score	70.60	71.90	72.30

Adaboost is adaptive in that it tweaks future weak learners in favor of cases misclassified by prior classifiers. It may be less prone to the overfitting issue than other learning algorithms in certain situations. Individual learners may be poor, but as long as their performance is somewhat better than actual guessing, the overall model will converge to a powerful learner.

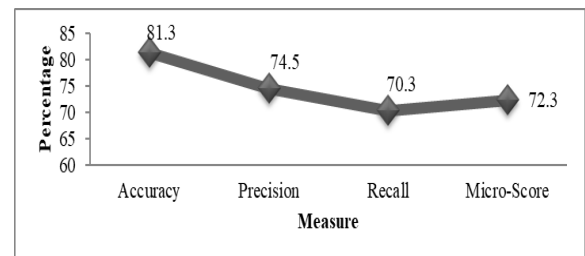


Figure 4: Analysis of vulnerability detection using Adaboost with 20-fold data cross-validation

figure 4 describes Adaboost classification for fake account detection for 20-fold cross-validation. AdaBoost is a specific training technique for boosted classifiers.

$$F_T(x) = \sum_{t=1}^T f_t(x)$$

A boost classifier is a kind of classifier.

Each F_t is a weak learner that accepts an object x as input and returns a value that indicates the object's class. The sign of the weak learner output, for example, specifies the predicted object class in the two-class issue, whereas the absolute value indicates confidence in that

classification. Similarly, if the sample belongs to a positive class, the Tth classifier is positive; otherwise, it is negative.

4.4 Experiment using recurrent neural network (Sigmoid)

we demonstrate the classification accuracy of RNN (Sigmoid) using a synthetic dataset, similar experiments have been done with various cross-validation and the results are illustrated in Table 6. According to this analysis, we conclude that 20-fold cross-validation provides the highest 96.10% classification accuracy using RNN with Sigmoid function.

Table 6: accuracy and confusion matrix for RNN (Sigmoid)

RNN (Sigmoid)	Fold 10	Fold 15	Fold 20
Accuracy	95.60	95.90	96.10
Precision	95.80	96.10	97.00
Recall	95.80	96.00	96.30
Micro-Score	94.70	95.90	96.05

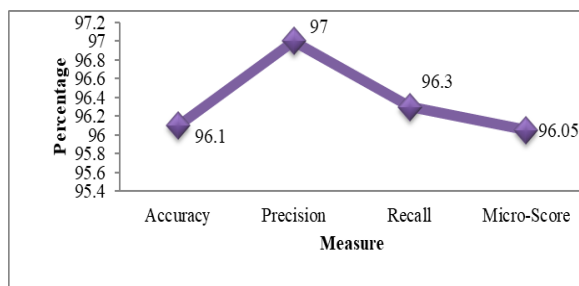


Figure 5: Detection of accuracy using RNN (Sigmoid) with 20-fold data cross-validation

The 20-fold cross-validation also achieves 96.10% with RNN with sigmoid function, have been explained in Figure 5, these RNN functions achieve around higher accuracy over the traditional machine learning algorithms during module testing.

4.5 Experiment using recurrent neural network (Tanh)

figure 6 shows the classification accuracy of RNN, the similar experiments has done with various cross-validations, and the results are illustrated in Table 7. According to this analysis, we conclude that 20-fold cross-validation provides the highest 97.25% classification accuracy for RNN using Tanh.

Table 7: Classification accuracy with confusion matrix for RNN (Tanh)

RNN (Tanh)	Fold 10	Fold 15	Fold 20
Accuracy	96.90	97.50	97.25
Precision	97.00	97.40	97.60
Recall	97.30	97.50	97.30

Micro-Score	96.80	96.70	96.90
-------------	-------	-------	-------

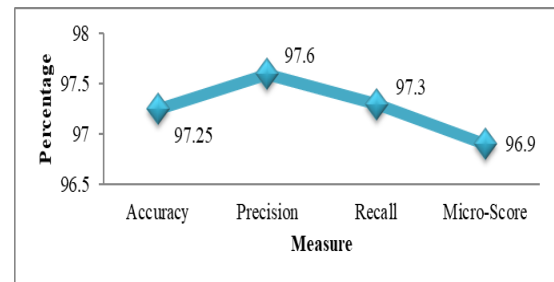


Figure 6: Detection of accuracy using RNN (Tanh) with 20-fold data cross-validation

4.6 Experiment using recurrent neural network (ReLU)

In this experiment we analyse the classification accuracy of ReLU using a synthetic dataset, similar experiments have been done with various cross-validation and the results are illustrated in Table 8. According to this analysis, we conclude system provides the highest 97.5% accuracy for 20-fold cross-validation classification accuracy for RNN.

Table 8: Classification accuracy with confusion matrix for RNN (ReLU)

RNN (ReLU)	Fold 10	Fold 15	Fold 20
Accuracy	97.20	97.90	97.50
Precision	97.40	96.90	97.60
Recall	95.60	97.20	97.90
Micro-Score	96.20	95.80	97.20

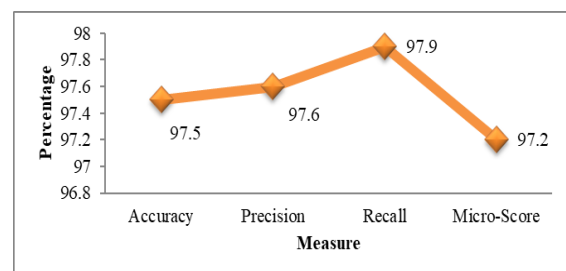


Figure 7: Detection of accuracy using RNN (ReLU) with 20-fold data cross-validation

The above experiments describe a proposed deep-learning classification algorithm with a machine-learning algorithm. This figure describes the result with and without cross-validation. we conclude that RNN with sigmoid provides better detection accuracy than the other two activation functions as well as the random forest machine learning algorithm.

Table 9: Results of above experiments.

Method / Measure	AN N	SV M	Adaboost	RNN (Sigmoid)	RNN (Tanh)	RNN (ReLU)
Accuracy				97.25	96.10	97.50
Precision				97.60	97.00	97.40
Recall				97.30	97.30	97.20
Micro-Score				96.90	96.05	97.20

Accuracy	85.60	95.2	81.30	96.10	97.25	97.50
Precision	84.99	94.80	74.50	97.00	97.60	97.60
Recall	77.72	96.2	70.30	96.30	97.30	97.90
Micro-Score	81.10	94.75	72.30	96.05	96.90	97.20



Figure 8: Classification accuracy with 20-fold cross-validation for all methods

The proposed method obtains the best predictive performance. The suggested solution can be further tested when used in actual software applications. The three data splitting mechanisms have been used as 10, 15, and 20-fold cross-validation.

Table 10: Dataset description of source code extracted from Android APK files

Total Size	2500
Training Samples	2000
Testing Samples	500

The system describes four evaluations between this research results and some existing systems results calculated on similar as well as multiple datasets.

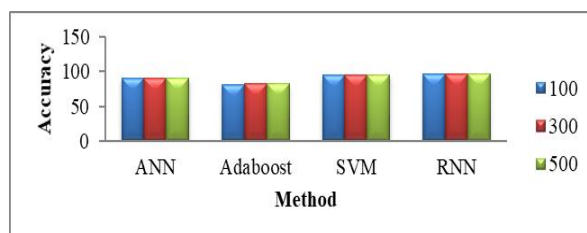


Figure 9: Comparative analysis of proposed vs. existing classification for vulnerability detection shows two machine learning algorithms used.

This figure depicts the proposed RNN provides better detection accuracy over machine learning algorithms. A classification model is generated using this arrangement or learning set to organize the input courses into corresponding template files or labels. Then a test set is used by gleaning the class labels of orthonormal courses to validate the model. A variety of neural networks are

used to identify reviews, such as ANN, Support Vector Machines (SVM), and Adaboost. The proposed RNN framework demonstrates superior performance compared to traditional and deep learning. the RNN outperformed methods like SVM, which struggled with large datasets and computational inefficiency, and CNNs, which lacked robustness in handling hierarchical and semantic features. These improvements are attributed to the RNN's ability to effectively model sequential data and incorporate advanced preprocessing techniques, such as TF-IDF and relational feature extraction, that enhance both semantic and syntactic representation. Additionally, the framework's balanced dataset handling and reduced preprocessing effort offer practical advantages for real-world applications. This innovative combination of features establishes the proposed RNN as a robust and scalable solution, addressing critical gaps in existing approaches while setting new benchmarks in vulnerability detection. The synthetic dataset used in this study was generated by augmenting real-world vulnerability datasets. The dataset comprised 2,500 samples, partitioned into 2,000 for training and 500 for testing. Each sample was manually labeled based on known vulnerabilities.

The 20-fold cross-validation was chosen to ensure robust evaluation and reduce bias in model performance assessment. Hyperparameter optimization was conducted using grid search to fine-tune parameters, including the number of RNN layers (1–3), activation functions (ReLU, Sigmoid, Tanh), and dropout rates (0.2–0.5). Data augmentation techniques, such as the oversampling of minority classes and the Synthetic Minority Over-sampling Technique, were applied to mitigate the effects of class imbalance, ensuring fair training and evaluation of the RNN model.

Artificial Neural Networks (ANN), Support Vector Machines (SVM), and Adaboost, were compared against the proposed RNN framework. For ANN, a two-layer architecture was used with a learning rate of 0.01 and trained for 100 epochs. The SVM utilized a radial basis function (RBF) kernel with a regularization parameter of 1. Adaboost employed 50 weak learners using decision stumps. Each model was trained using the same 20-fold cross-validation setup, with datasets split into 70% training and 30% testing. Training times varied significantly: RNNs required approximately 3 hours per fold, while ANN and SVM training averaged 1 hour and 45 minutes, respectively. Detailed results of accuracy, precision, recall, and micro-scores are provided in the Results section, illustrating the superior performance of RNN (ReLU), achieving an accuracy of 97.5%, precision of 97.6%, and recall of 97.9%.

5 Conclusion and future work

Vulnerability detection is very tedious work for imbalanced source codes; vulnerable code allows for generating software attacks on remote users. Sometimes, during execution, the vulnerable code also generates internal attacks like

buffer overflow, session hijack, bypass authentication, etc. In literature, many problems are detected in software every year. Vulnerabilities mostly do not appear in hidden forms that the software testers can identify. This

system describes the method of finding drawbacks by utilizing deep learning.

In this paper, we have developed an RNN including LSTM for constructing code vulnerability detection and bug triage on various platforms. Numerous tools are not able to support a web-based application to find code vulnerabilities. The proposed system works on different datasets for feature extraction and is able to detect the vulnerability. RNN provides a better result than traditional machine learning classifiers.

In the future, developers need to detect the code triage for runtime mobile-based application programs, because the existing tools do not support mobile application programs.

References

- [1] Terada, K.; Watanobe, Y., "Automatic Generation of Fill-in-the-Blank Programming Problems", In Proceedings of the 2019 IEEE 13th International Symposium on Embedded Multicore/Many-core Systems-on-Chip (MCSoc), Singapore, 1–4 October 2019; pp. 187–193.
<https://doi.org/10.1109/mcsoc.2019.00034>
- [2] Tai, K.S.; Socher, R.; Manning, C.D., "Improved semantic representations from tree-structured long short-term memory networks", In Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing, Beijing, China, 26–31 July 2015; pp. 1556–1566.
<https://doi.org/10.3115/v1/p15-1150>
- [3] Pedroni, M.; Meyer, B., "Compiler error messages: What can help novices?", In Proceedings of the 39th SIGCSE Technical Symposium on Computer Science Education, Portland, OR, USA, 12–15 March 2008; pp. 168–172.
<https://doi.org/10.1145/1352322.1352192>
- [4] Saito, T.; Watanobe, Y., "Learning Path Recommendation System for Programming Education based on Neural Networks", Int. J. Distance Educ. Technol. (Ijdet) 2020, 18, 36–64.
<https://doi.org/10.4018/ijdet.2020010103>
- [5] Teshima, Y.; Watanobe, Y., "Bug detection based on LSTM networks and solution codes", In Proceedings of the 2018 IEEE International Conference on Systems, Man, and Cybernetics (SMC), Miyazaki, Japan, 7–10 October 2018; pp. 3541–3546.
<https://doi.org/10.1109/smc.2018.00599>
- [6] Fan, G.; Diao, X.; Yu, H.; Yang, K.; Chen, L., "Software Defect Prediction via Attention-Based Recurrent Neural Network", Sci. Program. 2019, 2019, 6230953.
<https://doi.org/10.1155/2019/6230953>
- [7] Ohashi, H.; Watanobe, Y., "Convolutional Neural Network for Classification of Source Codes", In Proceedings of the 2019 IEEE 13th International Symposium on Embedded Multicore/Many-core Systems-on-Chip (MCSoc), Singapore, Singapore, 1–4 October 2019; pp. 194–200.
<https://doi.org/10.1109/mcsoc.2019.00035>
- [8] Zhou, Y.; Tong, Y.; Gu, R.; Gall, H., "Combining text mining and data mining for bug report classification", J. Softw. Evol. Process 2016, 28, 150–176.
<https://doi.org/10.1002/smr.1770>
- [9] Jin, K.; Dashbalbar, A.; Yang, G.; Lee, J.-W.; Lee, B., "Bug severity prediction by classifying normal bugs with text and meta-field information", Adv. Sci. Technol. Lett. 2016, 129, 19–24.
<https://doi.org/10.14257/astl.2016.129.05>
- [10] Goseva-Popstojanova, K.; Tyo, J., "Identification of security related bug reports via text mining using supervised and unsupervised classification", In Proceedings of the IEEE International Conference on Software Quality, Reliability and Security, Lisbon, Portugal, 16–20 July 2018; pp. 344–355.
<https://doi.org/10.1109/qrs.2018.00047>
- [11] Kukkar, A.; Mohana, R., "A Supervised bug report classification with incorporate and textual field knowledge", Procedia Comp. Sci. 2018, 132, 352–361.
<https://doi.org/10.1016/j.procs.2018.05.194>
- [12] Tong, H.; Liu, B.; Wang, S., "Software defect prediction using stacked denoising auto encoders and two-stage ensemble learning", Inf. Softw. Technol. 96, 94–111.
<https://doi.org/10.1016/j.infsof.2017.11.008>
- [13] Yogesh Kene, Uday Khot, Imdad Rizvi "A Survey of Image Classification and Techniques for Improving Classification Performance " arXiv 2019, arXiv:1608.06048.
<https://doi.org/10.2139/ssrn.3349696>
- [14] Jun Wang, Beijun Shen, Yuting Chen., "Compressed C4.5 Models for Software Defect Prediction ", IJCSI Int. J. Comput. Sci. Issues 2012, 9, 288–296
<https://doi.org/10.1109/qsic.2012.19>
- [15] Neuhaus, S.; Zimmermann, T.; Holler, C.; Zeller, A. "Predicting Vulnerable Software Components", In Proceedings of the 14th ACM conference on Computer and Communications Security, Alexandria, VA, USA, 28–31 October 2007; pp. 529–540.
<https://doi.org/10.1145/1315245.1315311>
- [16] Markad Ashok Vitthalrao, Mukesh Kumar Gupta, "Software Vulnerability Classification based on Machine Learning Algorithm", International Journal of Advanced Trends in Computer Science and Engineering (IJATSC), ISSN 2278-3091, Volume 9, No.4, July – August 2020, Page No.6653-6659.
<https://doi.org/10.30534/ijatcse/2020/358942020>
- [17] Markad Ashok Vitthalrao, Mukesh Kumar Gupta, "Software Vulnerability Classification based on Deep Neural Network", International Journal of

Engineering and Advanced Technology (IJEAT),
ISSN: 2249-8958 (Online), Volume-9 Issue-1,
October 2019, Page No.3146-3150.
<https://doi.org/10.35940/ijeat.a9746.109119>

