



FAST: Boosting Uncertainty-based Test Prioritization Methods for Neural Networks via Feature Selection

DOI:

[10.1145/3691620.3695472](https://doi.org/10.1145/3691620.3695472)

Document Version

Accepted author manuscript

[Link to publication record in Manchester Research Explorer](#)

Citation for published version (APA):

Chen, J., Wang, J., Zhang, X., Sun, Y., Kwiatkowska, M., Chen, J., & Cheng, P. (2024). FAST: Boosting Uncertainty-based Test Prioritization Methods for Neural Networks via Feature Selection. In *39th IEEE/ACM International Conference on Automated Software Engineering* <https://doi.org/10.1145/3691620.3695472>

Published in:

39th IEEE/ACM International Conference on Automated Software Engineering

Citing this paper

Please note that where the full-text provided on Manchester Research Explorer is the Author Accepted Manuscript or Proof version this may differ from the final Published version. If citing, it is advised that you check and use the publisher's definitive version.

General rights

Copyright and moral rights for the publications made accessible in the Research Explorer are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

Takedown policy

If you believe that this document breaches copyright please refer to the University of Manchester's Takedown Procedures [<http://man.ac.uk/04Y6Bo>] or contact openresearch@manchester.ac.uk providing relevant details, so we can investigate your claim.



FAST: Boosting Uncertainty-based Test Prioritization Methods for Neural Networks via Feature Selection

Jialuo Chen^{*}
Zhejiang University
Hangzhou, China
chenjialuo@zju.edu.cn

Jingyi Wang[†]
Zhejiang University
Hangzhou, China
wangjyee@zju.edu.cn

Xiyue Zhang
University of Oxford
Oxford, United Kingdom
xiyue.zhang@cs.ox.ac.uk

Youcheng Sun
University of Manchester
Manchester, United Kingdom
youcheng.sun@manchester.ac.uk

Marta Kwiatkowska
University of Oxford
Oxford, United Kingdom
marta.kwiatkowska@cs.ox.ac.uk

Jiming Chen
Zhejiang University
Hangzhou, China
cjm@zju.edu.cn

Peng Cheng[†]
Zhejiang University
Hangzhou, China
lunarheart@zju.edu.cn

Abstract

Due to the vast testing space, the increasing demand for effective and efficient testing of deep neural networks (DNNs) has led to the development of various DNN test case prioritization techniques. However, the fact that DNNs can deliver high-confidence predictions for incorrectly predicted examples, known as the over-confidence problem, causes these methods to fail to reveal high-confidence errors. To address this limitation, in this work, we propose FAST, a method that boosts existing prioritization methods through guided **FeAture SelecTion**. FAST is based on the insight that certain features may introduce noise that affects the model's output confidence, thereby contributing to high-confidence errors. It quantifies the importance of each feature for the model's correct predictions, and then dynamically prunes the information from the noisy features during inference to derive a new probability vector for the uncertainty estimation. With the help of FAST, the high-confidence errors and correctly classified examples become more distinguishable, resulting in higher APFD (Average Percentage of Fault Detection) values for test prioritization, and higher generalization ability for model enhancement. We conduct extensive experiments to evaluate FAST across a diverse set of model structures on multiple benchmark datasets to validate the effectiveness, efficiency, and scalability of FAST compared to the state-of-the-art prioritization techniques.

^{*}Work done while at the University of Oxford.

[†]Co-corresponding authors.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

Conference'17, July 2017, Washington, DC, USA

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-x-xxxx-xxxx-x/YY/MM

<https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>

CCS Concepts

• **Software and its engineering** → **Software testing and debugging**; • **Computing methodologies** → **Neural networks**.

Keywords

Deep Neural Networks, Test Input Prioritization.

ACM Reference Format:

Jialuo Chen, Jingyi Wang, Xiyue Zhang, Youcheng Sun, Marta Kwiatkowska, Jiming Chen, and Peng Cheng. 2024. FAST: Boosting Uncertainty-based Test Prioritization Methods for Neural Networks via Feature Selection. In . ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>

1 Introduction

Deep neural networks (DNNs) have demonstrated remarkable performance in diverse safety-critical domains such as self-driving cars [33] and malware detection [47]. Such domains have a low tolerance for mistakes. Compared to traditional code-driven software, DL-based systems adopt a data-driven programming paradigm, and thus introduce reliability challenges related to data faults (e.g., mispredicted data samples). Therefore, it is essential to conduct a high-quality test for DNNs to expose the hidden vulnerabilities as much as possible. However, testing for DNNs generally requires massive labeled data, while building test oracles by manually labeling for a large test set can be costly [2], especially for tasks that require domain-specific knowledge (e.g., medical images).

To increase the efficiency of DNN testing and reduce the labeling cost, it is essential to prioritize the test cases and identify a set of high-quality examples to reveal the prediction faults. A variety of test case prioritization methods [7, 28, 37, 43, 48] have been proposed based on carefully designed testing metrics, such as neuron coverage (NC) [18, 25, 31], surprise adequacy (SA) [15, 39], and prediction uncertainty [3, 7, 19]. Recent studies [2, 19, 41] have validated that the prioritization methods based on uncertainty achieved dominant results and less overhead. These methods (e.g., DeepGini [7]) calculate the uncertainty score directly from the model's probability output, with the assumption that test cases

with higher uncertainty scores (closer to the decision boundary) are more likely to be prediction errors. However, DNNs can be very confident on the wrongly predicted examples [10, 24] (the over-confidence problem), which causes these errors to be naturally overlooked by existing uncertainty-based methods. In practical applications, especially those heavily dependent on high-confidence decisions, over-confident errors can be more troublesome. Recently, nearest neighbors smoothing (NNS) [2] has been proposed to help mitigate this problem by averaging the outputs with a set of closely matched samples, showing notable improvements. However, its performance largely depends on the quality of the neighbors and can be computationally costly due to the search algorithm.

Since DNNs make the final prediction of an input based on the extracted features (the outputs of neurons), there should be connections between certain characteristics of hidden features and the prediction results, whether correct or erroneous. Despite considerable efforts within the software engineering community to characterize the behaviors of neurons, such as neuron activation coverage metrics [18, 25, 31], several studies [6, 7, 36, 44] have found that these measurements are not necessarily even negatively indicative of bug-revealing capabilities. We suppose that existing metrics focus solely on analyzing output values of intermediate layers; however, due to the high non-linearity of DNNs, their actual contribution to final predictions has not yet been clearly examined. Intuitively, different features are combined together to describe specific groups of examples (e.g., ‘cat’ or ‘dog’), each contributing differently to the class prediction. By examining the intermediate features of the model and empirically measuring each feature’s contribution, we show that there is a certain portion of ‘redundant’ features having minimal impact on the correctly predicted examples. On the contrary, such features can undesirably introduce noisy information aggregated to the model’s final output, thereby leading to high-confidence errors.

To effectively reveal such high-confidence errors, we propose a simple yet effective method for test input prioritization, namely FAST. It prioritizes test inputs by integrating guided feature selection during inference, so as to remove the potential noise incorporated in the output mispredictions for uncertainty measure. Specifically, FAST selectively prunes a subset of ‘noisy’ features from a given feature layer based on their contribution measurements, then uses the remaining features to derive the probability vector for uncertainty estimation. As a result, the original high-confidence errors will tend to be given a lower confidence score, making them easier to prioritize based on the calibrated probability vector, as illustrated in Fig. 1. FAST positions itself as a plug-and-play technique that can be easily applied with existing uncertainty-based test prioritization methods. We extensively evaluated FAST on 7 benchmark datasets (including image, text, and audio data) with 9 different DNN structures (including both CNN and RNN structures) against 10 prioritization baselines from different types. The results confirm the effectiveness, efficiency, and scalability of FAST compared to the state-of-the-art methods. In terms of the overall prioritization performance, FAST achieves 3.19% higher APFD (Average Percentage of Fault Detection) values than uncertainty-based methods and 1.78% higher than recent work (NNS [2]) on average. Specifically, for small selection budgets, FAST exposes over 13.63% more errors

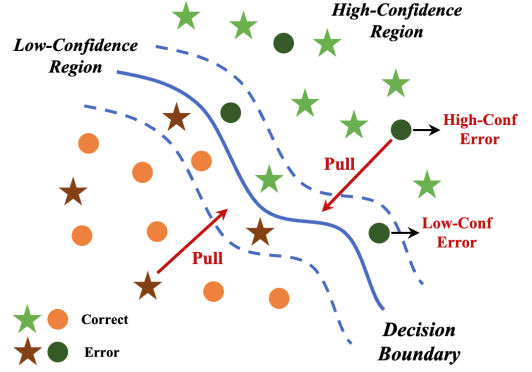


Figure 1: Illustration of the key limitation of existing uncertainty-based methods, which focus primarily on the test cases close to the decision boundary while neglecting the high-confidence errors. FAST helps to expose such high-confidence errors by dynamically suppressing their confidence, pulling them towards the low-confidence region.

than uncertainty-based methods and 8.16% more than NNS, respectively. In terms of model enhancement guided by test prioritization, FAST is effective in enhancing the model performance by rectifying more prediction errors through retraining (with 5% selection budget), achieving an accuracy improvement of 3.47% on average, which is 13.36% higher than NNS. Notably, FAST only incurs a slight additional time cost compared to the uncertainty-based methods, making it more efficient than most existing methods.

In summary, we make the following contributions.

- We propose and implement a novel prioritization method FAST through feature selection (noisy feature pruning). With FAST, more high-confidence errors can be exposed compared to traditional uncertainty-based methods.
- We extensively evaluate FAST on multiple common benchmark datasets with various model structures, validating the superior performance in prioritizing test cases compared to previous work. As a lightweight and generic method, FAST also shows high efficiency and scalability.
- We release FAST at [1] to facilitate future studies in this area.

2 Background

2.1 Deep Neural Networks

We focus on deep neural networks (DNNs) for classification. Typically, a DNN classifier [9] is a decision function $f : X \rightarrow Y$ mapping an input $x \in X$ to a label $y \in Y = \{1, 2, \dots, C\}$, where C is the number of classes. It comprises of L layers: $\{f^1, f^2, \dots, f^{L-1}, f^L\}$, where f^1 is the input layer, f^L is the probability output layer, and $\{f^2, \dots, f^{L-1}\}$ are the hidden layers. Each layer f^l can be denoted by a collection of neurons: $\{n_{l,1}, n_{l,2}, \dots, n_{l,N_l}\}$, where N_l is the total number of neurons at that layer. Each neuron is a computing unit that computes its output by applying a linear transformation followed by a nonlinear operation to its input (i.e., output from the precedent layer). We use $n_{l,i}(x)$ to denote the function that returns the output of neuron $n_{l,i}$ for a given input

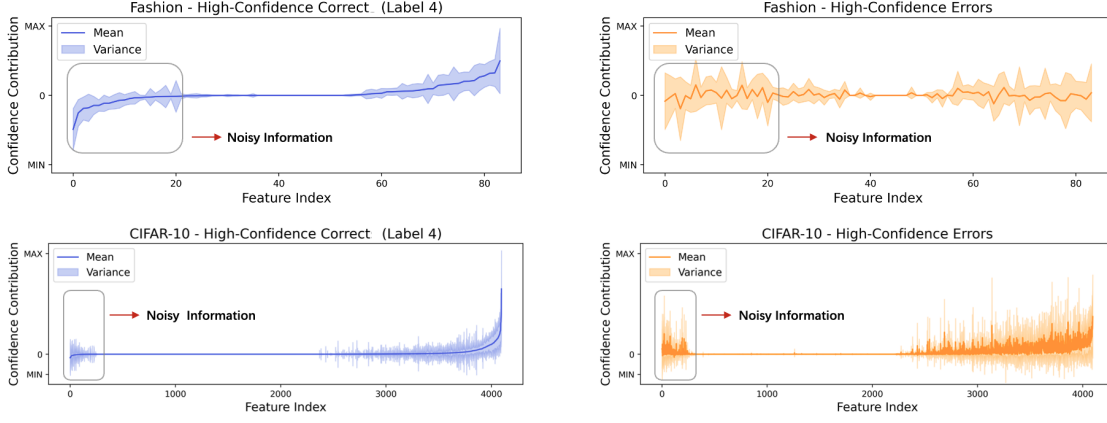


Figure 2: The feature contribution to the output confidence for the FASHION model (LeNet-5) and the CIFAR-10 model (VGG-16). The features are sorted in an increasing order based on their average contribution over the correctly classified examples with high confidence. The key facts are: 1) the most important features at the higher end significantly impact both correctly and incorrectly classified samples with high confidence, and 2) while the features at the lower end contribute the least to high-confidence correct classifications, they can still significantly affect high-confidence misclassifications.

$x \in X$. Then, we have the output vector of layer f^l ($2 \leq l \leq L$): $f^l(x) = \langle n_{l,1}(x), n_{l,2}(x), \dots, n_{l,N_l}(x) \rangle$. Specifically, we use $p(x)$ to denote the output probability vector of the final layer, then, the predicted label $f(x)$ is computed as $f(x) = \arg \max p(x)$.

2.2 Test Case Prioritization

To improve the DNN testing efficiency and reduce the labeling cost, a variety of test case prioritization methods have been proposed. For a given model f to test and a test suite T , the purpose of test case prioritization is to prioritize the test cases in T in a guided way, and when the test stops at a certain point, it is expected that the executed test cases (often selected according to the rank) will expose as many faults as possible. For a classification model, a test case $x \in T$ is identified as a fault when the model’s predicted label $f(x)$ is inconsistent with its ground-truth label. The core of test case prioritization is to quantify the possibility of a test case being a fault-trigger. Based on the required information, existing prioritization methods can be classified into two classes: 1) internal-based and 2) prediction-based (or uncertainty-based).

Internal-based Prioritization. As the model prediction is based on the propagation of layer outputs, it is natural to analyze the model internal status (i.e., hidden feature outputs) to provide evidence for judging a potential error. Multiple coverage metrics such as Neuron Activation Coverage (NAC) [25] and Neuron Boundary Coverage (NBC) [18] have been proposed to depict the covered neuron status for guiding the testing procedure. A test case that achieves a higher coverage is regarded as a more valuable input and will be prioritized. Another group of metrics utilizing the model internal is Surprise Adequacy (SA) [15], which measures how surprising an input is, i.e., how far is a given test case from the training set. Specifically, Likelihood-based Surprise Adequacy (LSA) utilizes the kernel density estimation (KDE) to obtain a density function from the training instances, which allows the direct estimation of the likelihood of a test case being a fault trigger. Distance-based

Surprise Adequacy (DSA) measures the distance between the test case and the closest example in the training dataset. However, several studies have found that all the above coverage-based criteria are not necessarily or even negatively indicative of fault-revealing capabilities [4, 6, 7, 44].

Uncertainty-based Prioritization. Instead of utilizing the internal outputs, uncertainty-based methods solely rely on the model’s final outputs, i.e., the probability vector $p(x)$. The intuition behind is that, if a DNN model predicts a test instance with nearly equal probability values on all candidate classes, indicating the model is uncertain (or confused) on predicting this test case, then it is more likely to be an incorrect prediction. We briefly introduce representative uncertainty metrics that are used in uncertainty-based test prioritization methods. DeepGini [7] quantifies the likelihood of a test case x being incorrectly predicted as: $Gini(x) = 1 - \sum_{i=1}^C p_i^2(x)$, where $p_i(x)$ is the prediction probability of class i . A high Gini index impurity represents high uncertainty. MaxP [19] directly uses the highest probability output as: $MaxP(x) = \max p_i(x)$ as the metric. Margin [27] is another way for quantifying prediction uncertainty, which is the difference between the highest and the second highest prediction probabilities, and a low margin represents a high uncertainty. After quantifying the uncertainty, the test cases with high uncertainty (close to the decision boundary) will be prioritized. Compared to internal-based methods, uncertainty-based methods are more lightweight and achieve better results [14, 41], as they do not require recording a heavy profile of neuron activations.

Over-confidence Problem of DNNs. Recent studies [2, 7, 41] have demonstrated the dominance of simple uncertainty-based methods in prioritizing the faults of DNNs. However, they can be further improved. Modern DNN models could suffer from the over-confidence problem, in that they tend to give high probability scores to the wrong predictions, due to model capacity, insufficient training data, or other factors. In such cases, the uncertainty-based methods will fall short of exposing and prioritizing these high-confidence faults.

As illustrated in Fig. 1, uncertainty-based methods will only prioritize the test cases close to the decision boundary (in the dotted low-confidence region), while neglecting the remote errors that are far away from the decision boundary. In this work, we propose FAST, a method to improve the uncertainty-based prioritization methods to uncover more high-confidence errors.

3 Methodology

In this section, we first present the motivation behind and then give details of each part of our prioritization method FAST.

3.1 Motivation

Considering the fact that the final prediction results are made based on a group of extracted features, we investigate the over-confidence problem faced by existing prioritization methods from the perspective of features. For a given class, different features (representing different learned patterns) have different levels of ‘contribution’ to the predictions. Instead of utilizing the numerical output values of the neurons to quantify the contribution (e.g., activation value [25] and activation frequency [34]), we measure it in a more intuitive way, that is, the observable difference on the model’s final output with and without that feature (we give details in Section 3.4). In other words, the feature whose removal causes a significant drop in confidence will be regarded as contributing more, and otherwise as having minimal influence on the final prediction. Then, for a selected feature layer of the model, we calculate the feature-wise contribution for a set of high-confidence correctly and wrongly predicted examples separately. We use correct predictions and errors in the following for simplicity.

As in Fig. 2, there are clear gaps between the feature contributions for supporting the model’s correct predictions (indicated by the blue color), with the lowest-ranked features making the least contributions. However, these same low-contributing features could potentially yield a higher contribution towards misleading the predictions (indicated by the orange color), as they undesirably bring noisy information that is aggregated to the model’s final output, further causing high-confidence errors. These features also show higher variance. Note that the magnitude of contribution can be influenced by the DNN structure, e.g., the VGG network shows higher sparsity than the smaller one, LeNet, but the patterns are similar across the models. Motivated by this observation, we propose a simple yet effective feature selection mechanism, FAST, to selectively drop the identified noisy features, to avoid them passing any information to influence the following decision process. That is, FAST tends to give lower confidence scores to the original high-confidence errors, while the correct predictions are influenced more subtly. As a consequence, the high-confidence errors become more distinguishable from the correct ones, further boosting the fault-revealing capability of simple uncertainty-based methods.

3.2 Overview of FAST

Figure 3 presents an overview of the proposed FAST approach, which emphasizes feature selection. Given a model and an unlabeled test suite for prioritization, existing uncertainty-based methods like DeepGini [7] directly calculate the uncertainty score from the output probability vector, as depicted in the dashed blue box.

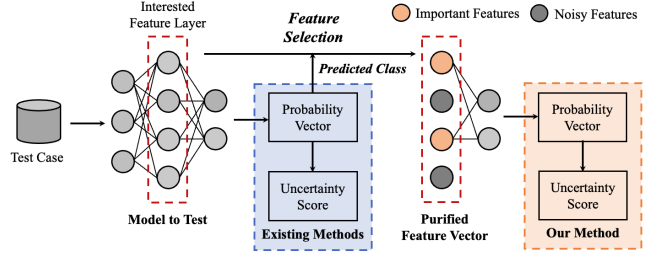


Figure 3: An overview of FAST framework. It selectively drops a set of noisy features during inference to derive the probability vector for uncertainty estimation.

In contrast, FAST focuses on performing feature selection at an intermediate layer before calculating the uncertainty score. Only a part of important features will be passed to the following layers, and other redundant features will be discarded. Then the purified feature vector will be passed to the next layer and derive a new probability vector for the uncertainty calculation, as illustrated in the orange box. Based on the scores, the unlabeled test suite can be ranked for prioritization. Specifically, FAST utilizes an output mask to achieve such a feature selection procedure without the need to modify the model structure or parameters and can be easily applied to any intermediate layer as a plug-in. Note that the additional cost of FAST lies in the propagation of partial layers (usually very few layers), thus only yielding a slight additional run-time cost compared to existing prioritization methods such as DSA [15] and NNS [2] that require massive distance computations.

3.3 Feature Selection

The feature selection will be performed on a specific hidden layer l of the given DNN. According to the definitions in the background (see Section 2), layer l encodes the inputs to a feature space with dimension N_l . We denote the feature vector of this layer for a given test input x as $f^l(x)$, which is then passed to the following layers to perform the final prediction. For convolutional layers, each kernel will be considered a basic feature element, while for fully-connected layers, each neuron is the basic element. For C classes, we construct a feature mask $M = [M_1, M_2, \dots, M_C]$ with $N_l \times C$ dimensions, where C is the number of classes and N_l is the number of features. M is a binary matrix (only 0 and 1), and M_c identifies the redundant features for each class. Specifically, the i^{th} element in M_c equal to 0 indicates the i^{th} feature in f^l is redundant for class c and should be omitted, whereas a value of 1 indicates that the feature is important and should be preserved.

As illustrated in Fig. 3, based on the originally predicted class $c = \arg \max p(x)$, the corresponding mask M_c is then applied to the feature vector $f^l(x)$ to obtain the purified feature vector as:

$$\hat{f}^l(x) = f^l(x) \odot M_c \quad (1)$$

where $\odot$ denotes the element-wise multiplication. This operation effectively prevents information from noisy features, which will have consistent zero output values, from being passed to subsequent layers. The modified vector $\hat{f}^l(x)$ is then used to derive the new probability vector $\hat{p}(x)$ through forward propagation by the layers

behind l . This feature selection process can be easily applied to any hidden layer with the use of an output mask, without the need to modify the model structure. It is noteworthy that different layers extract different semantic features, where shallow layers tend to extract more common features while deep layers focus on more complex features [45]. Based on our empirical studies, we find that the deep layer is a good choice for the feature selection, which effectively enlarges the differences between high-confidence erroneous and correct predictions. We further elaborate on the influence of layer locations in our experiments.

3.4 Contribution Measurement

To construct the feature mask M , we need to quantify the importance (contribution) of each feature for a specific class. The software engineering community has put a significant effort into devising a variety of neuron-level metrics to guide the testing procedure of DNNs in a more targeted way. For instance, DeepInspect [34] utilizes the neuron activation frequency for measuring how the neurons interact with the data from different classes. The neuron that is more frequently activated (larger than the threshold) is regarded as more important. Moreover, other strategies such as variance-based [15] and gradient-based [30] are also proposed to help measure the importance of neurons. These quantifiers can be naturally used for our purpose. However, due to the highly non-linear interactions of DNNs, it is still unclear how important a neuron (feature) is for the model's final prediction, as a larger output value or frequency does not mean a greater influence on the model's final output.

Instead of relying on existing static internal analyses, we adopt a direct approach inspired by the Shapley Value [8] from the machine learning community to examine the feature importance based on the prediction outcomes. It calculates the overall contribution of neurons by adding them to every possible subnetwork. However, applying such a neuron-wise approach is challenging in our setting, due to its significant computational overhead, which leads to exponential time complexity, making it difficult to meet high-efficiency requirements. Thus, we conduct the testing procedure in a one-by-one manner, similarly to elimination in program analysis, by directly assessing the individual contribution of features to the original model, thereby achieving linear time complexity. Specifically, for the i^{th} feature f_i^l of layer l , we empirically calculate its contribution for a given class c by comparing the model's output difference with and without feature f_i^l as:

$$\text{Contribution}(f_i^l, D_c) = \text{Conf}(f, D_c) - \text{Conf}(f^*, D_c) \quad (2)$$

where f is the original model and f^* is the model variant where the output of the feature $f_i^l(x)$ is consistently set to zero while keeping the other features unchanged. The function Conf measures the model's average confidence scores over a dataset D_c as:

$$\text{Conf}(f, D_c) = \frac{1}{|D_c|} \sum_{x \in D_c} p_c(x) \quad (3)$$

where p_c is the probability score for class c , and D_c is a clean dataset consisting of correctly predicted examples with the ground-truth label being class c (we sample from the training dataset). We carefully exclude uncertain examples (with low confidence),

Algorithm 1: FAST(f, f^l, r, T, λ)

Input: Model to test f , feature layer f^l , selection rate r , unlabeled test suite T , uncertainty metric λ

Output: Prioritization order of T .

```

1 Initialize score list  $Q$ 
  // Preparation Step
2 for each class  $c \in Y$  do
3   Obtain contribution vector  $V_c$  for  $f^l$ 
4   Obtain binary mask  $M_c$  from  $V_c$  ( $r\%$  of  $M_c$  are set 0)
5 end
  // Prioritization Step
6 for  $x \in T$  do
7   Get original predict label  $c$  for  $x$ 
8   Get original feature vector  $f^l(x)$ 
9   Get purified feature vector  $\hat{f}^l(x) = f^l(x) \odot M_c$ 
10  Get probability vector  $\hat{p}(x)$  using  $\hat{f}^l(x)$ 
11  Get uncertainty score  $s = \lambda(\hat{p}(x))$ 
12  Append  $s$  to list  $Q$ 
13 end
14 return  $\text{ArgSort}(Q)$ 

```

thereby focusing on a clear decision-making process and enabling a targeted analysis of features important or redundant for accurate predictions for a class. Specifically, the threshold of 0.9 is set based on empirical observations with the aim to identify the features important or irrelevant for correctly predicting each class, as high-confidence examples are regarded as representative examples of the class. The intuitive measurement here is that a significant drop in final confidence (a large contribution), when a feature is removed, underscores its importance to the model's ability to correctly predict class c . Note that FAST is flexible and allows integrating other candidate importance measurements to obtain the feature mask, and we elaborate on the impact of different quantification schemes in our experiments.

After quantifying the contribution of features using Equation 2, we get a contribution vector V_c of the features for each class c . To construct the binary mask M_c , we select the top- k features with the lowest contribution values as redundant, setting their corresponding mask values to 0 and marking the remaining, more crucial features with 1. We define a hyper-parameter $r = \frac{k}{N_l}$ to denote the percentage of features selected to drop. A low r indicates a smaller fraction of redundant features will be pruned and a larger fraction of features are preserved. When $r = 0$, the newly derived output $\hat{p}(x)$ becomes the same as the original one $p(x)$. We discuss the influence of the pruning rate later in experiments.

3.5 FAST for Test Prioritization

The complete procedure of our FAST method for test case prioritization is detailed in Algorithm 1. FAST is designed to be a generic tool compatible with all existing uncertainty-based methods such as DeepGini [7]. Given a selected feature layer f^l , we first measure the importance of features for each class using a specific quantifier (Line 3), and we then prepare the binary feature mask M_c from V_c ,

setting $r\%$ of it to 0 (Line 4). During the prioritization process (Line 6-13), for the unlabeled test suite T , the objective is to prioritize T such that faults appear at the forefront. For each test case $x \in T$, the predicted label c is first determined (Line 7). The original feature vector $f^l(x)$ is then taken (Line 8). According to the mask M_c , only a part of important features are preserved as $\hat{f}^l(x)$ (Line 9). Finally, we get a new probability vector $\hat{p}(x)$ based on $\hat{f}^l(x)$, and we calculate the uncertainty of the test case using the given uncertainty measure λ as input (Line 11). Finally, the test samples are ranked by their uncertainty scores (line 14). Compared to traditional uncertainty-based methods, the additional cost with FAST involves constructing the critical feature mask M and executing partial layer propagation (behind the selected feature layer f^l). Since the mask construction is a one-time effort, FAST is considered efficient.

3.6 Comparison with Similar Work

Here, we provide a qualitative comparison between FAST and two similar approaches, all aiming to obtain a justified probability vector for distinguishing more prediction faults from the correct ones. FAST is similar to Monte-Carlo Dropout (MC-Dropout) [40] in its approach of feature dropping in the hidden space. MC-Dropout utilizes the dropout functionality (which is often used in model training) during inference, running the DNN multiple times (e.g., t times) while randomly dropping a proportion of neuron activations. The output probability is then averaged over the t runs. In contrast, FAST is guided by a critical feature mask, which helps to drop noisy features in a targeted manner without needing multiple runs for a single test case, thereby achieving much better effectiveness and efficiency across experiments. More recently, Nearest Neighbor Smoothing (NNS) [2] has been proposed to improve uncertainty-based prioritization methods. NNS adjusts the probability vector by interpolating its original vector with those of its neighbors (typically using K-NN to find the neighbors in the test suite) as:

$$\hat{p}(x) = \alpha \cdot p(x) + \frac{(1-\alpha)}{k} \sum_{x' \in x_{KNN}} p(x') \quad (4)$$

where x' denotes the found neighbors, α is the balance factor, k is the number of neighbors, and $\hat{p}(x)$ is the derived new probability output over the original test case and its k neighbors. However, NNS requires iteratively calculating a heavy distance matrix with the test suite for the K-NN algorithm, which can be particularly costly for large models. In contrast, FAST is more efficient. Once the feature mask is obtained as a one-time effort, it incurs only a slight computational cost during inference. We provide a quantitative comparison of FAST with the two methods in experiments.

4 Evaluation

4.1 Experimental Setup

4.1.1 Datasets and Models. Following recent work [2, 41], we apply FAST to four benchmark datasets commonly used in the literature, i.e., MNIST [17], FASHION [42], CIFAR-10 [16], and SVHN [23]. For each dataset, we adopt two different DNN model structures for evaluation, where LeNet [17], ResNet [11], and VGG [29] are standard and popular structures. Conv-6 and Conv-8 [15] are the models composed of multiple convolutional layers. Additionally, we

extend our evaluation to include three datasets representing high-dimensional image, audio, and text data: Imagenette [13], GSCmd [38], and Imdb [20], respectively.

4.1.2 Test Suite for Prioritization. We prepare both clean and noisy test suites to evaluate the test prioritization methods. In the clean setting, we split the original dataset following [2] and then use the unlabeled clean test samples for evaluation. This is to test how well a prioritization method works on a test suite with the same distribution as the training data and whether it can expose faults earlier. Considering that the data processed by the DNN in a practical environment can be diverse, we use corrupted datasets that are manipulated with a range of modifiers inspired by real-world input corruptions such as noise and blur. As discussed in [2, 41], corrupted data provides model-independent and realistic data compared to adversarial examples. Specifically, we use MNIST-C [22], FASHION-C [41], and CIFAR-C [12] as provided in the literature. For SVHN, since no corrupted datasets existed, we apply multiple corruption actions following [12] to create an SVHN-C version for a comprehensive evaluation.

4.1.3 Test Prioritization Baselines. We assess the effectiveness and efficiency of FAST by comparing it with 10 baselines from four different groups, including three coverage-based methods: NAC [25], NBC [18], and OBSAN [5], an improved version of NBC by incorporating class-wise neuron boundary information; two surprise adequacy (SA)-based methods: DSA [15] and LSA [15], which also use class-wise information for computation; three uncertainty-based methods: DeepGini [7], MaxP [19], and Margin [27]; and two post-hoc methods as introduced in Section 3.6: MC-Dropout [40] and NNS [2] (state-of-the-art). The uncertainty-based and post-hoc methods are designed to be class-agnostic. We use the same uncertainty calculation method, DeepGini, for the base of evaluated post-hoc methods. We configure each baseline according to the default settings reported in its respective paper. Specifically, for MC-Dropout, we use $t = 50$, and for NNS, we set $\alpha = 0.5$ and $k = 10$. We also consider the pure random prioritization baseline.

4.1.4 Evaluation Metrics. We use Average Percentage of Fault Detection (APFD) [26] to evaluate the overall performance of a test case prioritization method as:

$$APFD = 1 - \sum \frac{TF_i}{n * m} + \frac{1}{2 * n} \quad (5)$$

where n denotes the total number of test cases in the test suite, m denotes the total number of faults exposed in the DNN under the test suite, and TF_i denotes the position of the first test in the test suite that exposes fault i . The value of APFD ranges from 0 to 1, with higher values representing higher fault detection efficiency.

While APFD is a good measure of the overall performance of test case prioritization methods, it is not suitable for comparing different methods within a given budget, i.e., test case selection based on the prioritization results. Therefore, we also use Test Relative Coverage (TRC) [2] to evaluate the performance of a test case prioritization method for a given budget as:

$$TRC = \frac{\#SelectedFaults}{\min(\#SelectBudget, \#TotalFaults)} \quad (6)$$

TRC measures how far a prioritization method is from the ideal case in a given budget. The value of TRC ranges from 0 to 1, where a higher value represents a higher proportion of faults exposed within the selection budget.

4.1.5 Experimental Environment. Our experiments were conducted on a computational server with an Intel(R) Xeon(R) CPU E5-2680 v4 @ 2.40GHz, an NVIDIA GTX 2080Ti GPU, and 256GB of RAM. The operating system used was Ubuntu 16.04.5 LTS. We conducted all experiments using Python 3.7.4 and Tensorflow 2.6.0.

4.2 Research Questions

In the following, we evaluate FAST through extensive experiments and answer five research questions.

4.2.1 RQ1: Is FAST more effective in test case prioritization?

To answer this question, we evaluate the effectiveness of FAST with representative baselines in prioritizing test cases in the clean and noisy setting, respectively.

1) Effectiveness on Clean Data. Table 1 displays the APFD values obtained by each test prioritization method on clean data, with the highest scores in bold. Across all models tested, FAST achieves the highest APFD values, averaging 85.42%, which is 3.0% higher than the uncertainty-based methods (e.g., DeepGini). This improvement is more than 2.7 times the improvement observed with NNS (about 1.1%). This suggests the superior capability of FAST in prioritizing misclassifications compared to state-of-the-art methods.

Next, we examine different types of prioritization methods. Generally, uncertainty-based methods outperform coverage-based and surprise-adequacy-based methods, a trend consistent with the findings in the replication study [41]. With the APFD values for random ordering around 50%, two of the three coverage-based methods, i.e., NAC and NBC, perform even worse than the random baseline. This performance can be attributed to the high non-linearity of DNNs, which limits the connections between neuron output magnitudes and prediction faults. For the SA-based methods, DSA consistently achieves higher values than LSA, still performing slightly worse than DeepGini in 7 out of 8 cases, and about 1.4% worse on average. LSA shows large fluctuations in performance compared to random, which we attribute to model complexity and the influence of kernel density estimation (KDE) used in LSA, which is heavily affected by the number of features. For instance, in a hidden layer of the VGG network, which contains thousands of features, the performance of LSA becomes close to random ordering.

As expected, the three uncertainty-based methods, DeepGini, MaxP, and Margin, exhibit similar performance, with their average APFD values all around 83%. Interestingly, the MC-Dropout method, a post-hoc baseline that averages multiple predictions through random dropout, causes a slight decline (about 0.75%) compared to the original uncertainty-based methods. This drop is likely because MC-Dropout occasionally drops some key features, which can confuse the decision process. In contrast, FAST, by selectively dropping features guided by contribution measurements, achieves stable improvements over uncertainty-based methods and consistently outperforms them. Although the improvement of NNS over DeepGini is noticeable, it is sometimes subtle when compared to FAST. We hypothesize that NNS tends to select neighboring test

cases with similar feature patterns, resulting in outputs that are too similar, which minimally influences the averaged probability outputs, and the difference between errors and correct may not be significantly enlarged in certain tasks.

While the APFD value serves as an effective overall performance metric, it does not precisely reveal how many errors can be exposed within a given selection budget. Fig. 4 shows the TRC scores obtained by each method under various budgets, ranging from 1% to 100% of the test suite. As TRC reaches 100% when the entire test suite is selected, it becomes less meaningful for detailed comparisons. Our focus is therefore on the TRC values at smaller budgets, where a higher TRC score indicates that more faults can be uncovered at a lower cost. The TRC curves for each method exhibit a distinctive turning point; before this turning point, differences between methods are more pronounced but diminish as the selection budget increases. FAST consistently achieves higher TRC values than all other baselines across the selection budget levels, indicating that it can expose more errors with the same budget compared to other methods. This advantage is particularly notable at small selection budgets. To quantify this, we calculated the average TRC values before the turning point for each method. FAST shows substantial improvements of 12.92%, 14.46%, 9.54%, and 17.58% over the uncertainty-based methods, and 7.76%, 9.08%, 6.84%, and 8.97% higher than NNS across the four datasets, respectively. These results highlight a significant enhancement in the practical application of FAST for test case selection compared to existing methods.

2) Effectiveness on Noisy Data. Table 2 shows a comparison of the APFD values obtained by each method on noisy data. The overall trend is similar to what we observe with clean data, though all evaluated methods show a certain drop in APFD when compared to their performance on clean data (Table 1), indicating that noisy data is generally harder for test case prioritization. Nevertheless, FAST achieves the highest APFD values in 7 out of 8 cases, with the only exception being the MNIST dataset, where NNS surpasses FAST. The improvement of FAST over the uncertainty-based methods is more than 3% on average. Additionally, we note that DSA surpasses the uncertainty-based methods in 3 cases, and the improvement of NNS on noisy data is more pronounced than on clean data (about 1.7%), indicating a noticeable advantage when dealing with data from different distributions.

Answer to RQ1: FAST consistently improves the performance of uncertainty-based test case selection methods and outperforms the state-of-the-art.

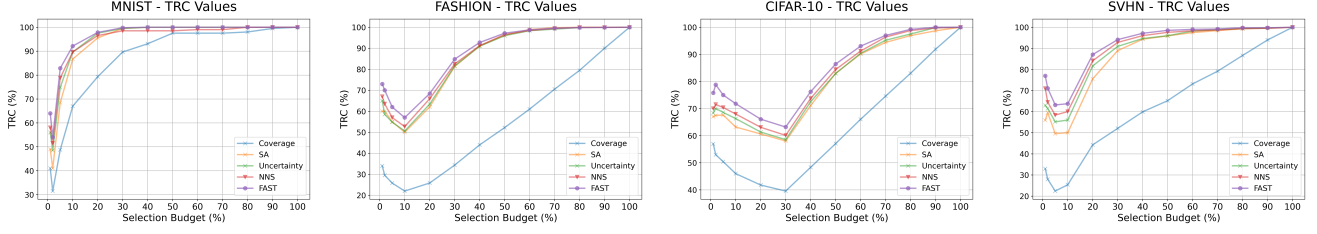
4.2.2 RQ2: Is FAST more efficient?

To answer this question, we report the total time cost of each method for prioritizing the same test suite (10K samples) in Table 3, inclusive of any necessary preparation time. The time cost for FAST consists of two parts, the feature mask preparation (S1) and running FAST during inference for prioritization (S2). We include the number of features N_f and split the overall time cost of FAST into two parts as S1/S2.

FAST demonstrates marked efficiency in prioritizing test inputs, requiring less than 2 minutes on average across various models. This efficiency is over 3 times and 9 times greater than that of other post-hoc methods, NNS and MC-Dropout, respectively. MC-Dropout

Table 1: The APFD values (capability to expose mis-classifications) on clean data. The higher APFD value the better.

Dataset	Model	Random	Coverage-based			SA-based		Uncertainty-based			Post-Uncertainty		
			NAC	NBC	OBSAN	DSA	LSA	DeepGini	MaxP	Margin	MC-Dropout	NNS	FAST
MNIST	LeNet-5	50.64	51.86	41.62	70.41	92.06	75.97	92.88	92.91	92.78	91.83	93.54	94.23
	Conv-6	52.45	40.86	46.84	88.03	93.15	78.55	94.84	94.86	94.89	94.60	95.52	96.37
FASHION	LeNet-5	49.73	36.57	50.03	52.79	79.44	71.33	81.24	81.46	81.49	79.31	82.03	83.36
	Conv-8	49.20	39.95	48.89	54.53	82.01	56.26	82.45	82.48	82.42	81.65	82.79	83.85
CIFAR-10	ResNet-18	49.84	47.22	49.76	55.48	71.10	50.26	72.04	72.08	71.93	71.49	72.58	75.11
	VGG-16	50.65	44.48	48.69	69.89	71.74	50.01	71.07	71.00	70.92	71.28	72.26	75.34
SVHN	ResNet-34	49.93	34.21	51.06	63.12	82.60	50.75	83.84	83.96	84.04	83.48	85.60	86.84
	WRN-28	51.47	41.83	48.32	66.14	82.50	52.16	85.21	85.50	85.66	84.92	86.63	88.30
Average		50.49	42.12	48.15	65.05	81.83	60.66	82.95	83.03	83.02	82.32	83.87	85.43

**Figure 4: The TRC values on clean data. The higher TRC value the better.****Table 2: The APFD values (capability to expose mis-classifications) on noisy data. The higher APFD value the better.**

Dataset	Model	Random	Coverage-based			SA-based		Uncertainty-based			Post-Uncertainty		
			NAC	NBC	OBSAN	DSA	LSA	DeepGini	MaxP	Margin	MC-Dropout	NNS	FAST
MNIST	LeNet-5	49.77	54.78	41.65	61.45	83.34	67.87	83.34	83.36	83.13	82.58	86.15	85.88
	Conv-6	48.43	39.30	49.64	76.47	89.94	66.02	90.67	90.74	90.70	90.11	92.40	93.24
FASHION	LeNet-5	50.20	52.23	53.74	53.11	57.17	57.24	56.42	56.37	55.99	56.88	58.01	58.49
	Conv-8	50.11	48.07	52.52	52.73	57.62	53.67	58.39	58.29	57.84	58.65	59.14	60.56
CIFAR-10	ResNet-18	50.07	47.90	50.96	51.70	60.71	51.59	61.38	61.34	61.17	61.52	61.83	63.85
	VGG-16	49.68	45.75	49.02	59.08	63.33	52.79	62.55	62.49	62.35	62.84	63.00	64.83
SVHN	ResNet-34	50.26	40.25	50.39	59.16	77.76	44.12	77.66	77.61	77.47	78.00	78.28	80.57
	WRN-28	50.34	45.10	47.36	64.55	77.87	42.09	79.06	79.05	78.97	79.19	80.33	81.44
Average		49.86	46.67	49.41	59.78	70.97	54.42	71.18	71.16	70.95	71.22	72.39	73.61

Table 3: Time cost (seconds) for test cases prioritization (per 10K samples).

Dataset	Model	Coverage-based			SA-based		Uncertainty-based			Post-Uncertainty			FAST	
		NAC	NBC	OBSAN	DSA	LSA	DeepGini	MaxP	Margin	MC-Dropout	NNS	FAST	S1/S2	N_l
MNIST	LeNet-5	20.0s	33.5s	34.5s	916.9s	180.6s	<5s	<5s	<5s	872.5s	228.9s	21.6s	16.1s/5.5s	84
	Conv-6	21.6s	40.5s	43.8s	1108.0s	166.2s	<5s	<5s	<5s	995.1s	264.3s	82.5s	64.5s/18.0s	256
FASHION	LeNet-5	16.5s	30.6s	31.3s	1015.4s	322.9s	<5s	<5s	<5s	880.4s	244.5s	29.3s	23.1s/6.2s	84
	Conv-8	24.9s	43.2s	47.3s	1008.4s	201.2s	<5s	<5s	<5s	916.4s	288.1s	35.0s	26.2s/8.8s	128
CIFAR-10	ResNet-18	33.8s	59.1s	68.4s	1584.2s	365.8s	<5s	<5s	<5s	1133.8s	395.9s	126.2s	99.0s/27.2s	512
	VGG-16	65.2s	119.2s	130.7s	1985.1s	620.8s	<5s	<5s	<5s	1833.7s	642.8s	306.1s	273.0s/33.1s	4096
SVHN	ResNet-34	37.5s	55.6s	70.2s	1730.5s	450.7s	<5s	<5s	<5s	1212.6s	422.6s	174.3s	134.6s/39.7s	512
	WRN-28	41.4s	67.1s	80.5s	1331.2s	329.7s	<5s	<5s	<5s	1306.3s	455.3s	203.2s	155.7s/47.5s	640
Average		32.6s	56.1s	63.3s	1335.0s	329.7s	<5s	<5s	<5s	1143.9s	367.8s	122.3s	99.0s/23.3s	—

involves randomly sampling multiple stochastic predictions for the same test input, while NNS requires computing the heavy distance matrix between the test sample and the entire dataset to find the neighbors, both resulting in significantly higher computational costs compared to FAST. The primary cost for FAST is the feature contribution estimation step (for obtaining mask M), which varies depending on the target model structure, taking from 20 seconds for smaller models to about 5 minutes for larger models such as VGG-16. Generally, the feature contribution assessment (S1) comprises the majority of the overall time cost. It is noteworthy that the estimation step of FAST is a one-time effort, after which the additional inference cost for FAST involves merely the information propagation of few layers (influenced by the layer location). The time overhead incurred for the contribution assessment mainly depends on the number of features N_l of the selected layer. As the assessment procedure for each feature is independent, it can be further accelerated using parallel processing.

In comparison across different method types, undoubtedly, pure uncertainty-based methods emerge as the most efficient, requiring less than 5 seconds (with GPU acceleration), as they do not require any preparatory work. The time cost here reflects the actual time needed for the DNN to make the final predictions. Coverage-based methods incur considerably higher costs than uncertainty-based methods, as they need to extract fine-grained information from the internals and maintain a heavy neuron map, yet their performance does not justify these costs. SA-based methods take even longer than coverage-based methods. Specifically, DSA is the most time-consuming among the evaluated methods as it involves traversing all training data to identify the closest neighbor sample.

Answer to RQ2: FAST is more efficient than the SA-based prioritization methods and other post-hoc methods.

4.2.3 RQ3: What factors influence the performance of FAST?

To understand how FAST performs under different settings, we conduct an ablation study for FAST on three main influencing factors: the feature selection strategy, the feature selection (pruning) rate, and the layer location where FAST is applied.

1) Feature Selection Strategy. The core of our feature selection procedure is the importance quantification strategy of each feature regarding the predictions (as detailed in Section 3.4). Here, we replace the measurement method used in FAST with five representative strategies popular in the SE community. The Output strategy [25] supposes that features with higher output values have stronger influence on the model decision, and thus will prune the features with the lowest average output values across the validation dataset. The Activation Frequency strategy [34] quantifies how frequently a neuron (feature) is activated (larger than a threshold) during inference, considering rarely activated features as unimportant. The Variance strategy [15] assumes that features with lower variance contain less information, thus contributing less to the predictions. The Gradient strategy [30] calculates the gradients of the model predictions with respect to the feature outputs to measure the contribution. We also consider the Random strategy, where features will be selected and pruned randomly. For consistency, we set the same feature pruning rate ($r = 5\%$) for all methods.

Fig. 5 shows the results of FAST using different feature selection strategies. Notably, our method (represented by orange bars) demonstrates significantly higher APFD values compared to all baseline strategies, indicating the pruned features have stronger connections with the prediction errors. This distinction allows errors to become more distinguishable from correct predictions after pruning the identified noisy features. The performance of baseline methods varies considerably. For instance, the Gradient-based strategy achieves comparable performance to ours on the MNIST dataset but performs the worst on the FASHION dataset. In some cases, other strategies even underperform the random baseline. This suggests that the features selected by these baseline strategies have limited influence on the model's final predictions. We speculate that these methods quantify the features based on static outputs (except for the gradient-based, which incorporates the model's outputs in its calculation), often assuming that higher activation values or frequencies indicate greater contributions to final decisions. However, this may not hold true due to the high non-linearity of DNNs, which involves complex layer-by-layer propagation. In contrast, our method adopts a more intuitive approach by dynamically pruning and directly testing the contribution of each feature with respect to the observable prediction results, which leads to better performance in the feature selection step across the tested models. Moreover, experiments in RQ1 have already validated that metrics based on output values (e.g., NAC and NBC) have limited or even negative correlations with DNN faults. The results underline the superiority of our contribution-based method over static statistical methods in feature selection for complex DNNs. Although our method requires a higher computational cost than the baselines, it plays an indispensable role in the effectiveness of FAST.

2) Feature Selection Rate. Next, we vary the feature selection (pruning) rate and evaluate the corresponding effects of FAST on the APFD values. We also present the results of baseline uncertainty-based methods, i.e., DeepGini and NNS, for a better comparison.

Fig. 6 summarizes the trends of FAST with different pruning rates. We observe that when r is relatively small (below 5%), the APFD score increases steadily as r increases across datasets (MNIST is an exception that achieves optimal performance at an early stage). This finding aligns with the insights from Fig. 2 that there is a portion of noisy features contributing to the high-confidence errors. As more of the noisy information is discarded, the separability between correct and erroneous predictions increases. However, as r continues to increase, the APFD values drop significantly, suggesting a trade-off between the selection rate and the effectiveness of FAST. We suppose that pruning a larger portion of features will inevitably block some essential information influencing the decisions on correct examples, leading to decreased confidence in these classifications and making it harder to distinguish the errors from correct predictions. Note that the optimal r appears to vary depending on the specific characteristics of the model. Based on the visualized results, a rate between 3% to 10% is generally suitable for FAST, consistently yielding higher APFD values than the baselines.

3) Feature Selection Layer. We further explore the impact of applying FAST to different layer locations within the model, specifically examining the shallow, middle, and deep layers. We divide the model into three chunks, each consisting of 1/3 of the total

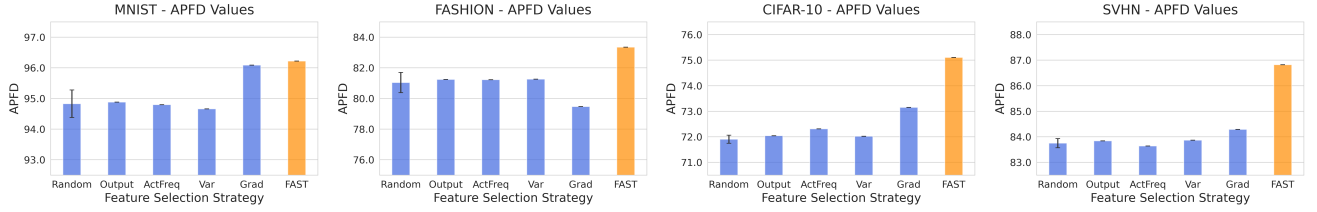


Figure 5: The performance of FAST with different feature selection strategies.

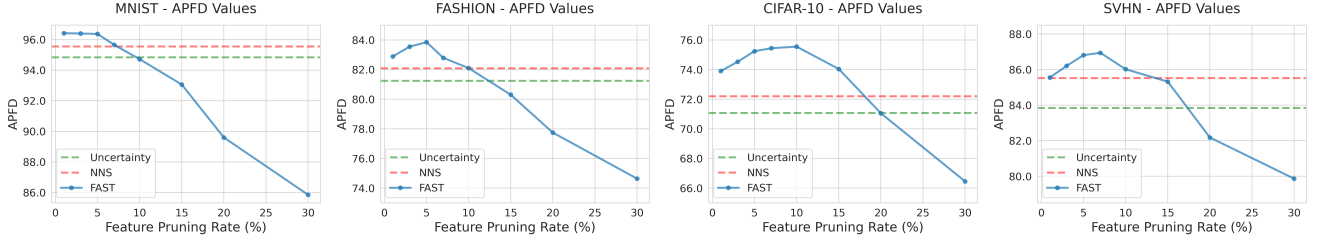


Figure 6: The performance of FAST with different feature selection rates.

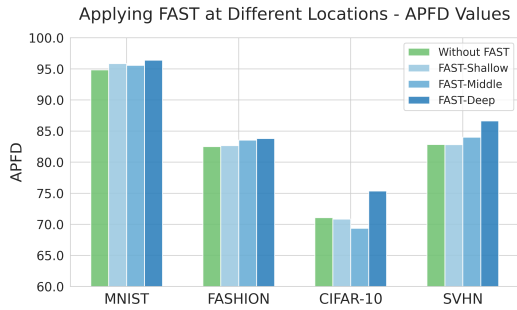


Figure 7: The performance of FAST with different layers.

Table 4: The APFD values on clean data.

Dataset	Model	Random	DeepGini	NNS	FAST
Imdb	LSTM-4	49.74	69.65	69.69	70.47
GSCmd	LSTM-6	50.53	88.49	88.50	89.40
Imagenette	VGG-16	49.82	65.85	66.41	67.53

number of layers, and we select one layer from each block to apply FAST. Specifically, the selected layer indices are 1/3/5, 1/4/7, 2/8/14, 4/16/32 for the Conv-6, Conv-8, VGG-16, ResNet-34 models, respectively. We also present the results without FAST (i.e., pure uncertainty calculation) for a better comparison.

Fig. 7 shows the results. It is evident that applying FAST at the deep layer consistently achieves the best performance across all datasets. Conversely, when FAST is applied to the shallower layers, it sometimes underperforms compared to the baseline without using FAST. This indicates that the feature selection in these layers can sometimes hurt the model’s normal decision-making process for correct predictions. This observation aligns with the understanding that shallow layers tend to extract more abstract features, while

deeper layers encode higher-level, more complex features [45], which contain richer information for FAST to distinguish the errors more effectively. Therefore, the deep layer is the preferred location for applying FAST to optimize performance.

Answer to RQ3: FAST’s performance can be influenced by several factors. Configuring a suitable setting is important for FAST to achieve satisfying performance.

4.2.4 RQ4: Is FAST scalable to high-dimensional input and other data domains?

To investigate the scalability of FAST, we conduct validation experiments on three additional datasets, the Imagenette dataset, which comprises high-dimensional images of $224 \times 224 \times 3$ pixels; the GSCmd dataset [38], which includes audio commands such as ‘left’ and ‘right’; and the Imdb dataset [20] containing textual reviews for sentiment classification. We perform necessary pre-processing for the audio and text data, and use the RNN model structures (LSTM [46] and fully-connected layers) for model training. The results are presented in Table 4.

Compared to the results on image data, the improvements of NNS and FAST over uncertainty-based methods in the text and audio domains are relatively subtle. We hypothesize that this is due to the more abstract and entangled nature of features extracted in these domains, which affects both NNS and FAST to a certain extent. Nevertheless, FAST still demonstrates a non-trivial improvement over the uncertainty-based methods. This further validates the scalability of FAST to handle complex data and model structures.

Answer to RQ4: FAST is scalable to high-dimensional input and promising for other data domains. Compared to other baselines, it still improves APFD values more effectively.

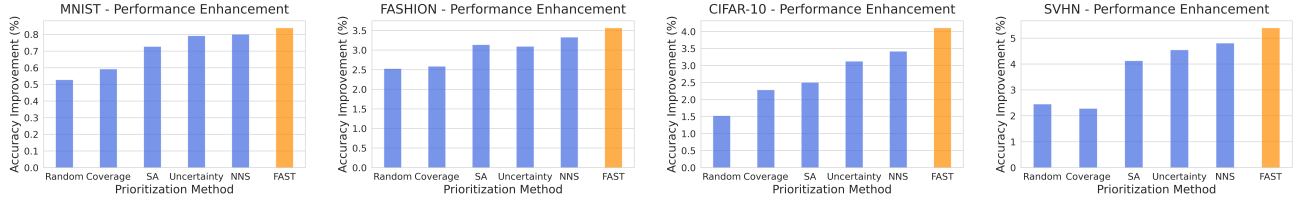


Figure 8: The model enhancements guided by test prioritization methods through retraining (with 5% selection budget).

4.2.5 RQ5: Is FAST able to guide the retraining for model enhancement?

Prioritizing the test suite and selecting the valuable part to enhance the DNN model performance (e.g., accuracy) by retraining is an important application scenario in DNN testing. To evaluate the effectiveness of FAST as well as other prioritization methods, we select a set of test cases with the same selection budget to retrain the models. For each model, we use the same training settings.

Fig. 8 shows the performance enhancement results of different prioritization methods. We can observe that FAST consistently obtains the highest accuracy improvements across all scenarios. Compared to the random baseline, most other prioritization methods also demonstrate certain level of effectiveness, except for the coverage-based methods. For all the retrained models, FAST achieves 3.47% accuracy improvements on average, which is 13.36% higher than NNS. This indicates that FAST can expose more hidden errors and rectify them than the baseline methods.

Answer to RQ5: FAST can guide the retraining procedure to help improve the model performance.

4.3 Discussion

In this work, we proposed FAST, a lightweight test prioritization designed to enhance the efficiency of DNN testing. The key idea behind FAST is to selectively drop certain noisy features to obtain an alternative probability vector for uncertainty measurement. As a generic method, FAST can be easily applied to various types of DNN structures and support different uncertainty-based metrics, not limited to DeepGini. We also examined the sensitivity of FAST to different hyperparameters through necessary experiments. Some hyperparameters, such as the threshold used for class-wise analysis, are currently set based on empirical observations to identify features that are important or irrelevant for correctly predicting each class. We plan to investigate automatic parameter tuning to optimize performance. Additionally, the overconfidence issue explored in this work can be mitigated to some extent by existing model calibration methods [21, 32, 35] that involve modifying the model structure or training process. We believe that FAST, like other test prioritization methods, is complementary to model calibration techniques. FAST functions as a plug-and-play module, requiring no modifications to the models themselves.

4.4 Threats to Validity

The external threat to validity primarily arises from the datasets, model structures, and baselines used in our study. Following recent

replication studies [2, 41], we considered commonly used public datasets in the literature to ensure a broad evaluation scope. Additionally, we validated the scalability of FAST to high-dimensional image data and other data domains including audio and text. To mitigate threats from DNN models, we employed various well-known architectures with different capacities to limit the impact of model dependency to some extent. Since our method is generic and independent of datasets and model structures, it can be adapted with minor adjustments for further applications. Regarding the prioritization baselines, while there are other works with different approaches [37, 48], we selected multiple representative methods from different categories and demonstrated the effectiveness and efficiency of FAST. The internal threat could arise from the implementation of FAST and the baselines used for comparison. To reduce this threat, we utilized available implementations from the authors' released code or re-implemented them according to the original papers. We carefully checked the correctness of our code. The construct threat lies in the hyper-parameter settings of the experiments. For running the baselines, we followed the standard settings in existing works. Additionally, we conducted necessary ablation studies to provide guidance on setting hyper-parameters for optimal performance of FAST.

5 Conclusion

In this paper, we propose a novel and effective DNN test prioritization method FAST to improve testing efficiency and reduce labeling costs. The key idea behind FAST is to prune noisy information from the irrelevant features which could contribute to the high-confidence errors, which sharpens the uncertainty estimation and improves differentiation between high-confidence correct and erroneous predictions. Extensive experiments on multiple benchmark datasets with various model structures validate the effectiveness, efficiency, and scalability of FAST. Our results demonstrate that FAST consistently enhances the performance of existing uncertainty-based test case selection methods and outperforms state-of-the-art work. As a plug-and-play technique, FAST can be easily integrated with different uncertainty measures.

Acknowledgements

This work was supported by the Key R&D Program of Zhejiang (Grant No. 2022C01018), the National Science Foundation of China Program (Grant No. 62102359, 62293511, 62088101) and the Program of China Scholarship Council (Grant No. 202306320425). MK and XZ received partial support from ELISA: European Lighthouse on Secure and Safe AI project (Grant No. 101070617 under UK guarantee) and the ERC under the European Union's Horizon 2020 research and innovation program (FUN2MODEL, Grant No. 834115).

References

- [1] [n. d.]. <https://github.com/Testing4AI/FAST>.
- [2] Shenglin Bao, Chaofeng Sha, Bihuan Chen, Xin Peng, and Wenyun Zhao. 2023. In defense of simple techniques for neural network test case selection. In *Proceedings of the 32nd ACM SIGSOFT International Symposium on Software Testing and Analysis*. 501–513.
- [3] Taejoon Byun, Vaibhav Sharma, Abhishek Vijayakumar, Sanjai Rayadurgam, and Darren Cofer. 2019. Input prioritization for testing neural networks. In *2019 IEEE International Conference On Artificial Intelligence Testing (AITest)*. IEEE, 63–70.
- [4] Jialuo Chen, Jingyi Wang, Xingjun Ma, Youcheng Sun, Jun Sun, Peixin Zhang, and Peng Cheng. 2022. QuoTe: Quality-oriented Testing for Deep Learning Systems. *ACM Transactions on Software Engineering and Methodology* (2022).
- [5] Yanzuo Chen, Yuanquan Yuan, and Shuai Wang. 2023. OBSan: An Out-Of-Bound Sanitizer to Harden DNN Executables.. In *NDSS*.
- [6] Yizhen Dong, Peixin Zhang, Jingyi Wang, Shuang Liu, Jun Sun, Jianye Hao, Xinyu Wang, Li Wang, Jinsong Dong, and Ting Dai. 2020. An Empirical Study on Correlation between Coverage and Robustness for Deep Neural Networks. In *2020 25th International Conference on Engineering of Complex Computer Systems (ICECCS)*. IEEE, 73–82.
- [7] Yang Feng, Qingkai Shi, Xinyu Gao, Jun Wan, Chunrong Fang, and Zhenyu Chen. 2020. DeepGini: Prioritizing Massive Tests to Enhance the Robustness of Deep Neural Networks. In *Proceedings of the 29th ACM SIGSOFT International Symposium on Software Testing and Analysis (ISSTA 2020)*. Association for Computing Machinery, New York, NY, USA, 177–188.
- [8] Amirata Ghorbani and James Y Zou. 2020. Neuron shapley: Discovering the responsible neurons. *Advances in Neural Information Processing Systems* 33 (2020), 5922–5932.
- [9] Ian Goodfellow, Yoshua Bengio, and Aaron Courville. 2016. *Deep learning*. MIT press.
- [10] Chuan Guo, Geoff Pleiss, Yu Sun, and Kilian Q Weinberger. 2017. On calibration of modern neural networks. In *International conference on machine learning*. PMLR, 1321–1330.
- [11] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 770–778.
- [12] Dan Hendrycks and Thomas Dietterich. 2019. Benchmarking neural network robustness to common corruptions and perturbations. *arXiv preprint arXiv:1903.12261* (2019).
- [13] Jeremy Howard. 2019. *The Imagenette dataset*.
- [14] Qiang Hu, Yuejun Guo, Maxime Cordy, Xiaofei Xie, Wei Ma, Mike Papadakis, and Yves Le Traon. 2021. Towards exploring the limitations of active learning: An empirical study. In *2021 36th IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE, 917–929.
- [15] Jinhan Kim, Robert Feldt, and Shin Yoo. 2019. Guiding deep learning system testing using surprise adequacy. In *2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE)*. IEEE, 1039–1049.
- [16] Alex Krizhevsky, Geoffrey Hinton, et al. 2009. Learning multiple layers of features from tiny images. (2009).
- [17] Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. 1998. Gradient-based learning applied to document recognition. *Proc. IEEE* 86, 11 (1998), 2278–2324.
- [18] Lei Ma, Felix Juefei-Xu, Fuyuan Zhang, Jiyuan Sun, Minhui Xue, Bo Li, Chunyang Chen, Ting Su, Li Li, Yang Liu, et al. 2018. Deepgauge: Multi-granularity testing criteria for deep learning systems. In *Proceedings of the 33rd ACM/IEEE International Conference on Automated Software Engineering*. ACM, 120–131.
- [19] Wei Ma, Mike Papadakis, Anestis Tsakmalis, Maxime Cordy, and Yves Le Traon. 2021. Test selection for deep learning systems. *ACM Transactions on Software Engineering and Methodology (TOSEM)* 30, 2 (2021), 1–22.
- [20] Andrew Maas, Raymond E Daly, Peter T Pham, Dan Huang, Andrew Y Ng, and Christopher Potts. 2011. Learning word vectors for sentiment analysis. In *Proceedings of the 49th annual meeting of the association for computational linguistics: Human language technologies*. 142–150.
- [21] Lassi Meronen, Martin Trapp, Andrea Pilzer, Le Yang, and Arno Solin. 2024. Fixing overconfidence in dynamic neural networks. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*. 2680–2690.
- [22] Norman Mu and Justin Gilmer. 2019. Mnist-c: A robustness benchmark for computer vision. *arXiv preprint arXiv:1906.02337* (2019).
- [23] Yuval Netzer, Tao Wang, Adam Coates, Alessandro Bissacco, Baolin Wu, Andrew Y Ng, et al. 2011. Reading digits in natural images with unsupervised feature learning. In *NIPS workshop on deep learning and unsupervised feature learning*, Vol. 2011. Granada, Spain, 7.
- [24] Anh Nguyen, Jason Yosinski, and Jeff Clune. 2015. Deep neural networks are easily fooled: High confidence predictions for unrecognizable images. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 427–436.
- [25] Kexin Pei, Yinzhi Cao, Junfeng Yang, and Suman Jana. 2017. Deepxplore: Automated whitebox testing of deep learning systems. In *Proceedings of the 26th Symposium on Operating Systems Principles*. ACM, 1–18.
- [26] Gregg Rothermel, Roland H. Untch, Chengyun Chu, and Mary Jean Harrold. 2001. Prioritizing test cases for regression testing. *IEEE Transactions on software engineering* 27, 10 (2001), 929–948.
- [27] Tobias Scheffer, Christian Decomain, and Stefan Wrobel. 2001. Active hidden markov models for information extraction. In *International symposium on intelligent data analysis*. Springer, 309–318.
- [28] Weijun Shen, Yanhui Li, Lin Chen, Yuanlei Han, Yuming Zhou, and Baowen Xu. 2020. Multiple-boundary clustering and prioritization to promote neural network retraining. In *Proceedings of the 35th IEEE/ACM International Conference on Automated Software Engineering*. 410–422.
- [29] Karen Simonyan and Andrew Zisserman. 2014. Very deep convolutional networks for large-scale image recognition. *arXiv preprint 1409.1556* (2014).
- [30] Jeongju Sohn, Sungmin Kang, and Shin Yoo. 2023. Arachne: Search-Based Repair of Deep Neural Networks. *ACM Transactions on Software Engineering and Methodology* 32, 4 (2023), 1–26.
- [31] Youcheng Sun, Min Wu, Wenjie Ruan, Xiaowei Huang, Marta Kwiatkowska, and Daniel Kroening. 2018. Concolic Testing for Deep Neural Networks. In *Proceedings of the 33rd ACM/IEEE International Conference on Automated Software Engineering (ASE 2018)*. Association for Computing Machinery, New York, NY, USA, 109–119.
- [32] Sunil Thulasidasan, Gopinath Chennupati, Jeff A Birmes, Tanmoy Bhattacharya, and Sarah Michalak. 2019. On mixup training: Improved calibration and predictive uncertainty for deep neural networks. *Advances in neural information processing systems* 32 (2019).
- [33] Yuchi Tian, Kexin Pei, Suman Jana, and Baishakhi Ray. 2018. Deeptest: Automated testing of deep-neural-network-driven autonomous cars. In *Proceedings of the 40th international conference on software engineering*. 303–314.
- [34] Yuchi Tian, Ziyuan Zhong, Vicente Ordonez, Gail Kaiser, and Baishakhi Ray. 2020. Testing DNN image classifiers for confusion & bias errors. In *Proceedings of the acm/ieee 42nd international conference on software engineering*. 1122–1134.
- [35] Deng-Bao Wang, Lei Feng, and Min-Ling Zhang. 2021. Rethinking calibration of deep neural networks: Do not be afraid of overconfidence. *Advances in Neural Information Processing Systems* 34 (2021), 11809–11820.
- [36] Jingyi Wang, Jialuo Chen, Youcheng Sun, Xingjun Ma, Dongxia Wang, Jun Sun, and Peng Cheng. 2021. RobOT: Robustness-oriented testing for deep learning systems. In *2021 IEEE/ACM 43rd International Conference on Software Engineering (ICSE)*. IEEE, 300–311.
- [37] Zan Wang, Hanmo You, Junjie Chen, Yingyi Zhang, Xuyuan Dong, and Wenbin Zhang. 2021. Prioritizing test inputs for deep neural networks via mutation analysis. In *2021 IEEE/ACM 43rd International Conference on Software Engineering (ICSE)*. IEEE, 397–409.
- [38] Pete Warden. 2018. Speech commands: A dataset for limited-vocabulary speech recognition. *arXiv preprint arXiv:1804.03209* (2018).
- [39] Michael Weiss, Riddhi Chakraborty, and Paolo Tonella. 2021. A review and refinement of surprise adequacy. In *2021 IEEE/ACM Third International Workshop on Deep Learning for Testing and Testing for Deep Learning (DeepTest)*. IEEE, 17–24.
- [40] Michael Weiss and Paolo Tonella. 2021. Fail-safe execution of deep learning based systems through uncertainty monitoring. In *2021 14th IEEE conference on software testing, verification and validation (ICST)*. IEEE, 24–35.
- [41] Michael Weiss and Paolo Tonella. 2022. Simple techniques work surprisingly well for neural network test prioritization and active learning (replicability study). In *Proceedings of the 31st ACM SIGSOFT International Symposium on Software Testing and Analysis*. 139–150.
- [42] Han Xiao, Kashif Rasul, and Roland Vollgraf. 2017. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747* (2017).
- [43] Xiaoyuan Xie, Pengbo Yin, and Songqiang Chen. 2022. Boosting the revealing of detected violations in deep learning testing: A diversity-guided method. In *Proceedings of the 37th IEEE/ACM International Conference on Automated Software Engineering*. 1–13.
- [44] Shenao Yan, Guan hong Tao, Xuwei Liu, Juan Zhai, Shiqing Ma, Lei Xu, and Xiangyu Zhang. 2020. Correlations between deep neural network model coverage criteria and model quality. In *Proceedings of the 28th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering*. 775–787.
- [45] Jason Yosinski, Jeff Clune, Anh Nguyen, Thomas Fuchs, and Hod Lipson. 2015. Understanding neural networks through deep visualization. *arXiv preprint arXiv:1506.06579* (2015).
- [46] Yong Yu, Xiaosheng Si, Changhua Hu, and Jianxun Zhang. 2019. A review of recurrent neural networks: LSTM cells and network architectures. *Neural computation* 31, 7 (2019), 1235–1270.
- [47] Zhenlong Yuan, Yongqiang Lu, Zhaoguo Wang, and Yibo Xue. 2014. Droid-sec: deep learning in android malware detection. In *Proceedings of the 2014 ACM conference on SIGCOMM*. 371–372.
- [48] Haibin Zheng, Jinyin Chen, and Haibo Jin. 2023. CertPri: certifiable prioritization for deep neural networks via movement cost in feature space. In *2023 38th IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE, 1–13.