

# Themis: Automatic and Efficient Deep Learning System Testing with Strong Fault Detection Capability

Dong Huang\*

University of Hong Kong  
dhuang@cs.hku.hk

Tsz On Li\*

University of Hong Kong  
toli2@cs.hku.hk

Xiaofei Xie

Singapore Management University  
xiaofei.xfxie@gmail.com

Heming Cui

University of Hong Kong  
heming@cs.hku.hk

**Abstract**—Deep Learning Systems (DLSs) have been widely applied in safety-critical tasks such as autopilot. However, when a perturbed input is fed into a DLS for inference, the DLS often has incorrect outputs (i.e., faults). DLS testing techniques (e.g., DeepXplore) detect such faults by generating perturbed inputs to explore data flows that induce faults. Since a DLS often has infinitely many data flows, existing techniques require developers to manually specify a set of activation values in a DLS’s neurons for exploring fault-inducing data flows. Unfortunately, recent studies show that such manual effort is tedious and can detect only a tiny proportion of fault-inducing data flows.

In this paper, we present Themis, the first automatic DLS testing system, which attains strong fault detection capability by ensuring a full coverage of fault-inducing data flows at a high probability. Themis carries a new workflow for automatically and systematically revealing data flows whose internal neurons’ outputs vary substantially when the inputs are slightly perturbed, as these data flows are likely fault-inducing. We evaluated Themis on ten different DLSs and found that on average the number of faults detected by Themis was 3.78X more than four notable DLS testing techniques. By retraining all evaluated DLSs with the detected faults, Themis also increased (regained) these DLSs’ accuracies on average 14.7X higher than all baselines.

\*Tsz On Li and Dong Huang contributed equally to this work.

## I. INTRODUCTION

Deep Learning Systems (DLSs) have been widely applied in safety-critical tasks such as autopilot and smart cities [7], [11], [27], [36], [43]. However, when a DLS (e.g., an autopilot system) is deployed in a real-world environment (e.g., a crowded city), the DLS’s input (e.g., a road condition image of the crowded city) is often perturbed by environmental noise such as raindrops and fog effects, causing incorrect values and disaster [11], [21]–[23], [27], [36], [43]. The incorrect output value of a DLS caused by perturbation on the DLS’s input is defined as the DLS’s fault [16].

Since a DLS’s faults greatly undermine the DLS’s reliability, a DLS must be systematically tested in order to detect as many faults as possible [11], [27], [36], [43]. These faults are essential for further developing a DLS (i.e., improving the DLS’s accuracy by retraining the DLS with perturbed inputs which lead to faults). The rationale of existing DLS testing is inspired by conventional software testing. In conventional software testing [39], [51], a testing technique generates a set of inputs (known as the *test set*) to explore diverse data

flows of a software (e.g., statements or branches), especially the data flows which likely lead to the software’s faults such as exceptions or crashes (these data flows are called “fault-inducing data flows”). When all data flows of the software have been explored (i.e., the test coverage metric adopted by the testing technique, such as code coverage, reaches 100%), the testing process is considered completed, and the testing technique stops.

Similarly, existing DLS testing techniques (e.g., DeepXplore [43], DeepGauge [36], DeepImportance [11], and Surprise Adequacy for Deep Learning System [27]) generate a test set (a set of perturbed inputs) to explore a DLS’s data flows, where the DLS’s data flow is defined as a set of numerical values: each of the numerical value is one of the DLS’s neurons’ output value (a DLS’s neuron is a “function”, and the output value of a neuron is called “activation value” [36], [43]) corresponding to a DLS’s input. A DLS testing instance generates noise test set by perturbing test set with the same type of environmental noise (e.g., raindrop effect) of various levels of noise strength (e.g., various raindrop densities).

However, since the activation values of a DLS’s neurons are discrete numbers from negative infinity to positive infinity, generating a test set to make a DLS’s neurons output all possible sets of activation values is prohibitively inefficient. To mitigate such inefficiency in DLS testing, for each type of environmental effects (noise), existing techniques (DeepXplore [43], DeepGauge [36], DeepImportance [11], and Surprise Adequacy for Deep Learning System [27]) require a DLS’s developer to manually specify likely sets of fault-inducing activation values in all neurons.

For example, DeepXplore [43] requires developers to manually partition a neuron’s outputs into two segments (i.e., outputs larger than a threshold and outputs smaller than a threshold, where the threshold has to be manually specified by the developers). Then, DeepXplore [43] generates test inputs to cover each of the two segments for all neurons. Nevertheless, recent work [16] showed that manually specifying the threshold is error-prone, and such manual effort is often overly coarse-grained, identifying a DLS’s all activation values that may induce faults for each type of environmental effects is fundamentally difficult, and not all values will induce faults on all neurons. Worse, existing DLS testing techniques [11], [43]

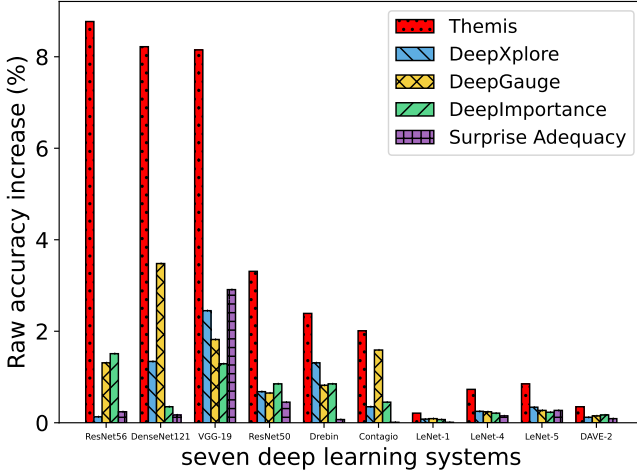


Fig. 1: Increase in a DLS’s accuracy by retraining the DLS with faults detected by Themis and the baselines. The increases in DLS’s accuracy brought by Themis were on average 14.7X more than the increases in DLS’s accuracy brought by baselines.

often can’t attain strong fault detection capability (strong fault detection capability means the correlation between the error rate of a DLS and the number of the DLS’s faults detected is larger than 0.7).

We believe that the root cause behind the limitations of the existing techniques is that the triggering condition and the total number of a DLS’s fault-inducing data flows are unknown [11], [16], [27], [31], [36], [43]. Hence, existing techniques inevitably require manual expert effort in inferring a set of fault-inducing data flows, in order to guide the exploration of a DLS’s data flows and the computation of test coverage (i.e., the proportion of the fault-inducing data flows being explored from the DLS during testing). Unfortunately, despite these advancements [11], [16], [27], [31], [36], [43], in principle, manual effort is far away from guaranteeing that these techniques can explore a full coverage of fault-inducing data flows from a DLS, while the full coverage is an essential and sufficient condition for a testing technique to attain strong fault-detection capability [16], [31]. Overall, an automated DLS testing technique that can theoretically explore a full coverage of fault-inducing data flows is highly desirable but vacant.

The main observation of the paper is that, a fault-inducing data flow in a DLS is often *sensitive* to perturbation on a DLS’s input. We denote the data flow with respect to a DLS  $M$  and an input  $I$  as  $DF(M, I)$ . For a DLS ( $M$ ), a clean input ( $I$ ) and its perturbed input ( $I'$ ), we consider a data flow  $DF(M, I')$  *sensitive* to perturbation if the difference of  $DF(M, I')$  and  $DF(M, I)$  (denoted as  $DF(M, I') - DF(M, I)$ ) is large, given that  $I'$  and  $I$  just have slight difference (e.g.,  $I'$  and  $I$  differ by light raindrops). Intuitively, sensitive  $DF(M, I')$  inevitably causes  $M(I')$  and  $M(I)$  to be different, confirmed in our evaluation (Figure 4b).

With this observation, we present Themis, a DLS test-

ing technique that systematically explores fault-inducing data flows guided by data flows’ sensitivity: the difference between  $DF(M, I')$  and  $DF(M, I)$  with respect to the difference between  $I'$  and  $I$ . Themis leverages math optimization techniques (e.g., Gradient Descent) to adjust the intensity of the perturbation (e.g., raindrops’ densities) on  $I$ , in order to generate a new test set of  $I'$ , which maximizes  $DF(M, I') - DF(M, I)$ . Hence, Themis maximizes the likelihood that  $DF(M, I')$  will lead to a fault. With this new workflow, Themis explores faults from a DLS without the necessity of either manual effort or exploring all data flows of a DLS. One major challenge for Themis to achieve strong fault detection capability is how Themis infers the coverage of fault-inducing data flows being explored by Themis’s generated test set during testing (i.e., how Themis computes the test coverage metric). This is because the actual number of a DLS’s fault-inducing data flows on all inputs is unknown.

To tackle this challenge, we summarize existing AI theories [7], [32], [46], [57] to show that, a DLS’s  $DF(M, I') - DF(M, I)$  on all inputs often follows normal distribution: for all clean inputs, most noise added to these inputs will lead to similar influence on  $M$ ’s final outputs, and only a tiny portion of noise will lead to outlying influence. Hence, once Themis’s workflow infers that  $DF(M, I') - DF(M, I)$  converges to a normal distribution driven by Themis’s generated test set, according to these theories, Themis has explored a statistically full coverage of fault-inducing data flows (i.e., only a tiny portion of fault-inducing noise and their data flows were missed by Themis), and Themis’s testing instance can complete. Our paper derives a proof (§IV-B) to show that Themis’s workflow can achieve statistically full coverage (i.e., achieving the “fault detection capability with high probability” in this paper).

We implemented Themis on Mindspore and compared Themis against four recent and notable DLS testing techniques (Deepxplore [43], DeepGauge [36], DeepImportance [11] and Surprise Adequacy for Deep Learning System [27]), which have been deployed in the Mindspore security framework and have been evaluated by several AI developers. We evaluated Themis and these baselines on ten popular Deep Learning models (e.g., LeNet [30], VGG [49], ResNet [18]) trained on six public datasets (Cifar10 [2], ImageNet [9], Driving [53], Contagio/VirusTotal [8], Drebin [4] and MNIST [29]), which cover a complete set of datasets evaluated by the baselines [11], [27], [36], [43]. Evaluation shows that:

- Themis consistently achieved strong fault detection capability. Themis achieved a high correlation [16] (0.70 to 0.95) for all the DL models, while the baselines’ correlation varied from -0.89 to 0.59 (Table II).
- Themis detected 3.78X more faults than the baselines, when Themis’s and the baselines’ test coverage reach 100% .
- By retraining the DL models with faults detected in testing, Themis increased (regained) the DL models’ accuracy by 0.21% to 8.77%, while baselines increased the DL models’ accuracy by 0.01% to 3.48%. Overall, Themis increased the DL models’ accuracy on average 14.7X higher than the baselines (Figure 1).

- Themis has been integrated into the Mindspore security framework, enhancing the testing of hundreds to thousands of DNNs in real-world scenarios<sup>1</sup>.

Our main contribution is Themis, the first automatic DLS testing technique which can empirically attain strong fault detection capability for perturbed inputs. The key novelty is Themis’s new workflow for systematically exploring a DLS’s fault-inducing data flows and computing its test coverage metric without manual effort, by leveraging our observation that most fault-inducing data flows are sensitive. Our theoretical analysis (§IV-B) shows that Themis can explore a full coverage of fault-inducing data flows at high probability (95%), so Themis empirically attained strong fault detection capability for all the evaluated DLSs (Table II).

## II. BACKGROUND AND RELATED WORK

### A. Deep Learning Testing

DeepXplore [43] proposes the first white-box coverage criteria, i.e., Neuron Coverage, which calculates the percentage of activated neurons. A differential testing approach is proposed to detect the errors by increasing NC. DeepGauge [34] then extends NC and proposes a set of more fine-grained coverage criteria by considering the distribution of neuron outputs from training data. Inspired by the coverage criteria in traditional software testing, some coverage metrics [11], [20], [33], [35], [50], [52], [56] are proposed. DeepCover [50] proposes the MC/DC coverage of DNNs based on the dependence between neurons in adjacent layers. DeepCT [33] adopts the combinatorial testing idea and proposes a coverage metric that considers the combination of different neurons at each layer. DeepMutation [35] adopts the mutation testing into DL testing and proposes a set of operators to generate mutants of the DNN. Furthermore, Sekhon et al. [48] analyzed the limitation of existing coverage criteria and proposed a more fine-grained coverage metric that considers both the relationships between two adjacent layers and the combinations of values of neurons at each layer. Based on the neuron coverage, DeepPath [54] initially proposes the path-driven coverage criteria, which considers the sequentially linked connections of the DNN. IDC [11] adopts the interpretation technique to select the important neurons in one layer. Based on the training data, it then groups the activation values of important neurons into a set of clusters and uses the clusters to measure the coverage. Kim et al. [26] proposed the coverage criteria that measure the surprise of the inputs. The assumption is that surprising inputs introduce more diverse data such that more behaviors could be tested. Surprise metric measures the surprise score by considering all neuron outputs of one or several layers. It is still unclear how the surprise coverage (calculated from some layers) is related to the decision logic. Xie et al. [55] propose Neuron Path Coverage, a novel, and interpretable coverage criterion aimed at characterizing the decision logic of models.

### B. Existing DLS testing are unautomated

DLS testing techniques (e.g., Deepexplore [43], DeepGauge [36], DeepImportance [11], and Surprise Adequacy for Deep Learning System [27]) are proposed to detect a DLS’s faults. A testing instance includes a pretrained DLS  $M$  (e.g., an autopilot system), an arbitrary input  $I$  fed to  $M$  for inference (e.g., a road condition image), and the same type of environmental noise with arbitrary noise strength  $E(\theta)$  (e.g., raindrop effect  $E$ , with arbitrary raindrop size  $\theta$ , where  $\theta$  is within a given range such as the real-world raindrop’s sizes) which may be present on  $I$ . For any  $M(I + E(\theta))$  which is different from  $M(I)$ , the  $M(I + E(\theta))$  is considered a DLS’s fault.

Note that an incorrect  $M(I + E(\theta))$  is considered as a fault instead of a failure because existing DLS testing techniques aim to test deep learning models of a DLS rather than the entire DLS (i.e., both deep learning models and the software code of a DLS). Since in existing DLSs (e.g., industrial autonomous systems such as Autoware [25] and Apollo [1]), deep learning models are intermediate components, the incorrect output of these deep learning models are regarded as the “fault” of a DLS [16].

These techniques generate perturbed inputs to make a DLS’s neurons output diverse activation values, in order to explore the DLS’s data flows (§I). Specifically, denote  $N_i, i = 1, \dots, n$  as the  $i$ th neuron of a DLS, and  $N_i(I + E(\theta)) \in \mathbb{R}$  as the activation values of the  $i$ th neuron corresponding to input  $I + E(\theta)$ . A DLS’s data flow is defined as a set of numerical values, where each numerical value is each neuron’s activation of the DLS (i.e., a data flow is defined as  $\{N_i(I + E(\theta))\}$ ). Existing work generates diverse  $I + E(\theta)$ , in order to trigger diverse data flows.

Since activation values are discrete numbers from negative infinity to positive infinity, existing techniques proposed various heuristic rules for a DLS’s developers to manually specify fault-inducing data flows. For instance, DeepXplore [43] requires a DLS’s developer to divide the activation values of neurons into two segments (e.g., values larger than zero and values smaller than zero), such that DeepXplore only needs to generate at minimal two perturbed inputs (we denote these inputs as  $I_1 + E_1(\theta_1)$  and  $I_2 + E_2(\theta_2)$ ), to cover the segments (e.g.,  $N_i(I_1 + E_1(\theta_1)) \geq 0 \forall i$  and  $N_i(I_2 + E_2(\theta_2)) > 0 \forall i$ ). Other DLS testing techniques also proposed similar heuristic rules for developers to specify fault-inducing data flows. DeepGauge [36] requires the user to divide activation values into several segments, where the segments that contain activation values towards zero have smaller intervals (because most inputs of a DLS make the neurons’ output activation values close to zero). DeepImportance [11] requires a DLS developer to specify fault-inducing activation values for “important” neurons only (the “important” neurons are identified via a popular technique in DL called “Layer-wise Relevance Propagation” [41]). Surprise Adequacy for Deep Learning System [27] requires a DLS’s developer to specify fault-inducing activation values for neurons in softmax layers only.

<sup>1</sup>Themis: [https://github.com/mindspore-ai/mindarmour/blob/master/mindarmour/fuzz\\_testing/sensitivity\\_convergence\\_coverage.py](https://github.com/mindspore-ai/mindarmour/blob/master/mindarmour/fuzz_testing/sensitivity_convergence_coverage.py)

### C. Fault-inducing data flows have to be automatically identified

Overall, these techniques allow a DLS to be tested within a reasonable time (e.g., several minutes). However, substantial studies [16], [31] showed that manual effort in specifying fault-inducing data flows often causes DLS testing to have unsatisfactory fault detection capability because fault-inducing data flows are often unknown (see §I). To automatically identify fault-inducing data flows, our observation is that for any  $I$  and  $E(\theta)$ , the  $N_i(I + E(\theta))$ ,  $1 \leq i \leq n$  which induces a fault usually has a large difference with  $N_i(I)$ ,  $1 \leq i \leq n$  (i.e., the value of  $\sum_{i=1}^n |N_i(I + E(\theta)) - N_i(I)|$  is large, we call this value as the *sensitivity* of a DLS's data flow on perturbed inputs). It is because large sensitivity often results in a large difference between  $M(I + E(\theta))$  and  $M(I)$  (i.e.,  $M(I + E(\theta))$  is likely to be a fault). This observation is also confirmed in our evaluation (§VI-B). Based on this observation, we propose an input generation technique guided by *sensitivity* (Sensitivity Maximizing Fuzzer in §III).

### D. Computing test coverage automatically

Test coverage metric measures the proportion of fault-inducing data flows explored. Ideally, when test coverage metric reaches 100%, all fault-inducing data flows of a DLS are explored [16], [31]. However, since the actual number of fault-inducing data flows in a DLS is unknown, proposing an accurate test coverage metric is an open challenge [16], [31]. Existing testing techniques [7], [11], [27], [36], [43] mitigate this challenge by requiring developers to specify a set of fault-inducing data flows. However, recent studies show that such manual effort is tedious and error-prone (§II-B).

To solve this challenge, Themis leverages the sensitivity of a DLS's data flow to infer the test coverage. We summarize the existing theories to show that for any DLS,  $N_i(I + E(\theta)) - N_i(I)$  follows a normal distribution. First, existing theories show that DNN is an ordinary differential equation [7], [32], and DNN's neurons are differential operators of an ordinary differential equation. Second, when random noise is present on an ordinary differential equation's input, the variation on the output values of the ordinary differential equation's differential operators often follows a normal distribution [46], [57]. By combining the first and the second theories, we can derive that when random noise is present on a DLS's input, the variation of the DLS's activation values (i.e.,  $N_i(I + E(\theta)) - N_i(I)$  defined in §II-B) often follows a normal distribution (we denote  $ND_i$  as the normal distribution, and we call  $\hat{ND}_i$  as "sensitivity distribution" in Figure 2). We also confirmed these observations in our evaluation (Figure 3b).

Since  $\hat{ND}_i$  consists of the frequency of all  $N_i(I + E(\theta)) - N_i(I)$  values including the fault-inducing data flows  $\{N_i(I + E(\theta))\}$ , the condition that  $\hat{ND}_i$  being identified as  $ND_i$  implies all fault-inducing data flows are also identified (we carry a more detailed analysis in §IV-B). Based on this observation, we proposed a test coverage metric based on the distribution of  $ND_i$  (Sensitivity Convergence Coverage in §III).

## III. OVERVIEW

This section presents the architecture and workflow of Themis. Figure 2 shows Themis's major components (namely Sensitivity Calculator, Sensitivity Convergence Coverage and Sensitivity Maximizing Fuzzer) and these components' workflow. To ease discussion, this section uses the notations defined in §II-B: given a DLS ( $M$ ), a set of inputs ( $I$ ) for the DLS to perform inference and environmental noise to be applied on  $I$  ( $E(\theta)$ ), Themis aims to detect the  $M(I + E(\theta))$  which is not equal to  $M(I)$  (such an  $M(I + E(\theta))$  is a fault).

To do so, Themis first perturbs each input of  $I$  with environmental noise  $E(\theta)$  (for each input in  $I$ ,  $\theta$  is a randomly chosen value within a given range). Then, Themis feeds both  $I$  and  $I + E(\theta)$  to **Sensitivity Calculator (Phase 1)**, to compute the difference between the outputs of each  $M$ 's neurons with respect to these inputs (i.e.,  $N_i(I + E(\theta)) - N_i(I)$  defined in §II-B, which are the  $i$ th neuron's outputs with respect to  $I$  and  $I + E(\theta)$  respectively). Sensitivity Calculator also computes  $M(I)$  and  $M(I + E(\theta))$ , and reports faults to DLS developers (i.e.,  $M(I + E(\theta))$  which is not equal to  $M(I)$ ).

Note that Themis needs to compute both  $N_i(I)$  and  $N_i(I + E(\theta))$ , while existing DLS testing techniques only compute  $N_i(I + E(\theta))$ , even though all these techniques aim to trigger fault-inducing  $N_i(I)$ . It is because Themis leverages the theoretical implication behind  $N_i(I + E(\theta)) - N_i(I)$  to guide the detection of fault-inducing  $N_i(I + E(\theta))$ , in order to automatically identify fault-inducing data flows (see §II-C). On the other hand, existing techniques require fault-inducing  $N_i(I + E(\theta))$  to be specified by DLS developers, which has been proven tedious and error-prone (§II-D).

Sensitivity Calculator then passes  $N_i(I + E(\theta)) - N_i(I)$  values to **Sensitivity Convergence Coverage (Phase 2)**, which infers a distribution from the values (i.e.,  $\hat{ND}_i$  defined in §II-D) for each neuron (i.e., each  $i \in n$ ). Sensitivity Convergence Coverage then determines whether  $\hat{ND}_i$  converges to a normal distribution and returns a test coverage value accordingly. Specifically, for each neuron in a DLS, Sensitivity Convergence Coverage feeds the  $N_i(I + E(\theta)) - N_i(I)$  values to Monte Carlo Markov Chain [12] (MCMC, a popular statistical technique for inferring probability distributions from a set of numerical values).

MCMC then outputs both  $\hat{ND}_i$  (the distribution inferred by MCMC) and the Monte Carlo Standard Error (MCSE, a metric to measure the potential error between the inferred distribution and the ground-truth normal distribution) value corresponding to  $\hat{ND}_i$ . Sensitivity Convergence Coverage thus determines whether  $\hat{ND}_i$  converges to a normal distribution based on the MCSE value: If MCSE equals zero within a confidence interval  $\alpha$  (Themis's default value of  $\alpha$  is 95%, a standard value of setting a confidence interval [3], [19]), then Sensitivity Convergence Coverage considers  $\hat{ND}_i$  converges to a normal distribution at a high probability (i.e., 95%). Otherwise,  $\hat{ND}_i$  is considered not converged.

Sensitivity Convergence Coverage then computes the test coverage metric as the proportion of neurons whose  $\hat{ND}_i$  con-

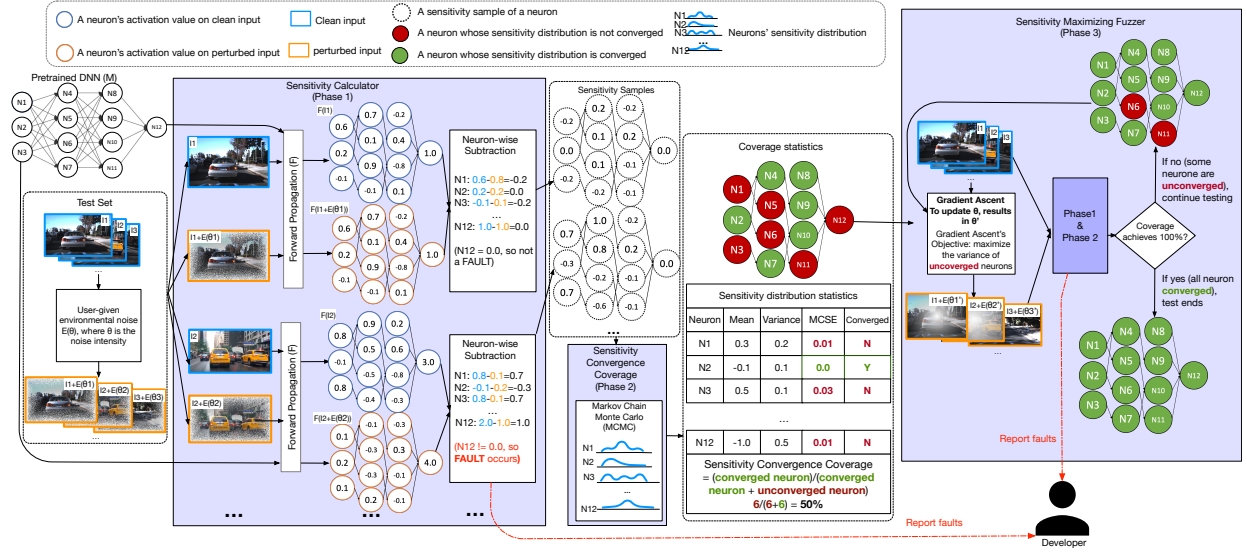


Fig. 2: Themis’s architecture. Themis’s components are shaded in purple. The DNN’s neurons (e.g., N2) are colored in green if their MCSE value (as shown in the “Sensitivity distribution statistics”) equals 0.0.

verges. For instance, in Figure 2, after Sensitivity Convergence Coverage computes  $\hat{N}D_i$  for the twelve neurons, six neurons among the twelve neurons (i.e., N2, N4, N7, N8, N9 and N10 which are colored in green) have MCSE as 0.0 (i.e., their  $\hat{N}D_i$  are converged), while the rest of the six neurons of the DLS have MCSE large than 0.0 (i.e., their  $\hat{N}D_i$  are not converged). Hence, Sensitivity Convergence Coverage computes the test coverage as 50%.

Sensitivity Convergence Coverage addresses an open challenge of DLS testing: accurately inferring the proportion of faults undetected from a DLS (i.e., computing a test coverage metric), even the total number of the DLS’s faults is unknown. Specifically, Sensitivity Convergence Coverage does so by inspecting  $\hat{N}D_i$ ’s convergence condition, which is correlated with the proportion of the DLS’s faults being detected (§II-D). Our theoretical analysis showed that when Sensitivity Convergence Coverage reaches 100%, all fault-inducing data flows are identified (i.e., explored) at high probability (95%). In contrast, existing works mitigate this challenge by requiring DLS developers to specify a set of fault inducing data flows, such that the test coverage metric is computed as the specified fault inducing data flows being explored. d

After that, Sensitivity Convergence Coverage passes the neurons whose  $\hat{N}D_i$  is converged to Themis’s **Sensitivity Maximizing Fuzzer (Phase 3)**. Sensitivity Maximizing Fuzzer then leverages math optimization techniques (e.g., Gradient Ascent) to generate a new test set, to explore unexplored fault-inducing  $N_i(I + E(\theta))$ . Specifically, Sensitivity Maximizing Fuzzer iteratively adjusts  $\theta$  of  $E$  (we denote the adjusted  $\theta$  as  $\theta'$ ), such that  $I + E(\theta')$  maximizes the sum of  $N_i(I + E(\theta)) - N_i(I)$  for all neurons whose  $\hat{N}D_i$  are not converged (i.e., N2, N4, N7, N8, N9 and N10 in Figure 2). Sensitivity Maximizing Fuzzer does so to ensure that the generated test set can explore new fault-inducing data flows (i.e.,  $N_i(I + E(\theta))$  which has a large difference

| Dataset (I) | Description                            | DNN (M)             | Number of neurons |
|-------------|----------------------------------------|---------------------|-------------------|
| MNIST       | Hand-written digit classification      | LeNet-1 [10], [30]  | 7206              |
|             |                                        | LeNet-4 [10], [30]  | 69362             |
|             |                                        | LeNet-5 [10], [30]  | 107786            |
| Contagio    | Malware classification in PDF Files    | <200,200> [43]      | 35410             |
| Drebin      | Malware classification in Android apps | <200,10> [14], [43] | 15230             |
| ImageNet    | General Image classification           | VGG-19 [49]         | 14888             |
|             |                                        | ResNet-50 [18]      | 16168             |
| Udacity     | Road condition classification          | DAVE-2 [58]         | 1560              |
| Cifar10     | General Image classification           | ResNet56 [18]       | 532490            |
|             |                                        | DenseNet121 [24]    | 563210            |

TABLE I: Datasets and DNNs for evaluating Themis, which covers the complete set of datasets evaluated by baselines.

with  $N_i(I)$ ). Sensitivity Maximizing Fuzzer then feeds the generated  $I + E(\theta')$  to Sensitivity Calculator for computing  $N_i(I + E(\theta')) - N_i(I)$ , and thus Themis enters another iteration of the testing.

#### IV. THEMIS’S RUNTIME

##### A. The challenge in designing Themis’s workflow

One major novelty of Themis is a new workflow to compute test coverage metric, by computing  $\hat{N}D_i$  for all of a DLS’s neurons (§II-B and §III). However, realising such a workflow is challenging, because  $\hat{N}D_i$  for all of a DLS’s neurons may cause Themis unscalable to a DLS’s size.

Computing  $\hat{N}D_i$  for just one neuron of a DLS is already computationally expensive, because it requires Themis’s Sensitivity Convergence Coverage to run MCMC on thousands of sensitivity samples (§III). Hence, for DLSs of large sizes (e.g., ResNet56 which consists of more than five hundred thousand neurons, see Table I), the strawman approach of brutal force computing would require more than hours to conduct the testing, inefficient compared to the existing techniques [11], [27], [36], [43], which only takes several minutes to complete testing. On the other hand, computing  $\hat{N}D_i$  for only a subset of a DLS’s neurons may cause Themis to miss substantial faults. Because  $\hat{N}D_i$  of different neurons have

different convergence rate ( $\hat{N}D_i$ 's convergence rate of a DLS neuron is essential for Themis to compute test coverage, see §III), computing  $\hat{N}D_i$  for a subset of a DLS's neurons may cause Themis to compute test coverage incorrectly (e.g., stop testing too soon or too late).

To tackle this challenge, we present Sensitivity Sampler, a sampling technique used by Themis's Sensitivity Convergence Coverage to compute  $\hat{N}D_i$  for a subset of a DLS's neurons, in order to precisely approximate the  $\hat{N}D_i$ 's convergence rate of all the DLS's neurons. Our observation behind Sensitivity Sampler is that  $\hat{N}D_i$ 's convergence rate is inversely proportional to  $\hat{N}D_i$ 's variance (the larger the variance, the slower the convergence rate, as pointed out by a classic statistic theory called Chebyshev's inequality [40]). Hence, for a DLS of any size, Themis samples a constant number (by default one thousand, see §IV-B) of neurons for computing  $\hat{N}D_i$  based on the variance of  $N_i(I + E(\theta)) - N_i(I)$  values of each neuron.

### B. Themis's algorithm

Algorithm 1 shows the algorithm of Themis, including Themis's three components (§III): Sensitivity Calculator (Phase 1), Sensitivity Convergence Coverage (Phase 2) and Sensitivity Maximizing Fuzzer (Phase 3), as well as Sensitivity Sampler (§IV-A). When  $M$ ,  $\{I\}$ , and  $E(\theta)$  are fed into Themis, Themis starts testing process. Specifically, for  $I$  in  $\{I\}$ , Themis's Sensitivity Calculator (line 2-12) first computes  $N_i(I)$  and  $N_i(I + E(\theta))$  (line 5-6), which are defined in §II-C. Sensitivity Calculator then computes the difference (denoted as  $D_i$ ) between  $N_i(I)$  and  $N_i(I + E(\theta))$  (line 8- 10).  $D_i$  are *sensitivity samples* (§III). Meanwhile, Sensitivity Calculator reports faults to the developer for any  $M(I)$  and  $M(I + E(\theta))$  with different DLS's outputs.

Then, Themis's Sensitivity Convergence Coverage infers  $\hat{N}D_i$  of the DLS's neurons based on  $D_i$  (line 2- 15). Specifically, Sensitivity Sampler computes the variance of each item in  $D_i$ , and sort these items according to their variance. Then Sensitivity Sampler selects  $k$  (by default one thousand) items from all the items in an evenly-spaced manner (i.e., select the item for every  $\lfloor \frac{|D_i|}{k} \rfloor$  items). The default value of  $k$  (i.e., one thousand) is selected based on Mann-Witney statistic [5], which points out that one thousand samples from a sorted list is sufficient to accurately approximate all the remaining values in the sorted list.

After Sensitivity Sampler computes  $\hat{N}D_i$  for each sampled neurons (line 14), Sensitivity Sampler computes *sensitivity convergence* (the test coverage of Themis) as the number of converged  $\hat{N}D_i$  over  $k$  (line 15). Specifically, the convergence criterion is whether Monte Carlo Standard Error (MCSE) equals 0.0, a standard criterion which indicates the probability of convergence identified by MCMC [12], [45]. By having MCSE equal 0.0, each batch of samples drawn by MCMC have almost identical probability distributions. We also evaluated the choice of MCSE value and our evaluation result (Figure 5b) shows that MCSE equal 0.0 allowed Sensitivity Sampler to precisely infer  $\hat{N}D_i$ 's convergence rate of all neurons. After Sensitivity Sampler computed *sensitivity convergence*,

### Algorithm 1: Themis's entire workflow

---

**Input:**  $\{I\}$ : raw inputs,  $M$ : DNN to be tested,  $n$ : the number of neurons of  $M$ ,  $E(\theta)$ : User-specified perturbation,  $t$ : threshold value for MCMC to determine convergence,  $k$ : sample size for Sensitivity Sampler,  $c$ : desired sensitivity coverage

---

```

1 Function Themis ( $\{I\}, E(\theta), M, n$ ):
2   Initialise List sensitivitySample; Initialise List
   Faults
3   while True do
4     /* Sensitivity Calculator (Phase 1)
       begins */
5     for each  $I \in \{I\}$  do
6        $\{N_i(I)\}_{i=1}^n \leftarrow M(\{I\})$ 's activation values
7        $\{N_i(I + E(\theta))\}_{i=1}^n \leftarrow M(I + E(\theta))$ 's activation
       values
8       Initialise Array  $\{Diff_i\}_{i=1}^n$ 
9       for each  $j \in [1, \dots, n]$  do
10         $Diff_j = |N_j(I + E(\theta)) - N_j(I)|$ 
11        sensitivitySample.insert(Diff)
12        if  $M(I) \neq M(I + E(\theta))$  then
13          Faults.insert(I, E(θ))
14      /* Sensitivity Convergence Coverage
        (Phase 2) begins */
15       $\{Diff_i\}_{i=1}^k \leftarrow$ 
        sensitivitySampler(sensitivitySample, k)
16       $\{\hat{N}D_i\}_{i=1}^k \leftarrow \text{MCMC}(\{Diff_i\}_{i=1}^k)$ 
17      sensitivityCoverage  $\leftarrow$  proportion of  $\{\hat{N}D_i\}$  that
        has MCSE == 0.0
18      /* Sensitivity Maximizing Fuzzer
        (Phase 3) begins */
19      if sensitivityCoverage <  $c$  then
20        UncovNeuron  $\leftarrow$  unconverged neuron
21        obj = maximizer(UncovNeuron)
22         $I + E(\theta') \leftarrow \text{gradientAscent}(\text{obj}, I + E(\theta))$ 
23      else
24        break
25  return
26 Function sensitivitySampler(sensitivitySample, k):
27   Initialise Array  $\{variance_i\}_{i=1}^n$ 
28   for each  $j \in [1, \dots, n]$  do
29     Variancej =
30       compute_variance(sensitivitySample[j])
31   Dsort = sort  $\{Variance_i\}_{i=1}^n$  based on variance
32   for each  $j \in [1, \dots, k]$  do
33     selectedj = variancej *  $\lfloor \frac{|D_i|}{k} \rfloor$ 
34   return  $\{selected_i\}_{i=1}^k$ 

```

---

Themis's Sensitivity Maximizing Fuzzer (line 18- line 19) performs fuzzing in order to explore the unconverged  $\hat{N}D_i$  (i.e., coverage-guided fuzzing).

**Analysis of Themis's fault detection capability.** Themis has strong fault detection capability because Themis theoretically can explore all fault-inducing data flows at high probability (95%). We first explain how we derive this theoretical property of Themis. As discussed in §II-D, we summarize two existing theories [7], [32], [46], [57] and conclude that when random noise is present on a DLS's input, the variation of a DLS's activation values (i.e.,  $N_i(I + E) - N_i(I)$  defined in §II-B) often follows a normal distribution (i.e.,  $\hat{N}D_i$  defined in §II-B). Hence, Themis examines whether the generated

test set explores a full coverage of the fault-inducing data flows (i.e.,  $N_i(I + E)$  which has large difference with  $N_i(I)$ ) by examining whether a normal distribution can be inferred from the  $N_i(I + E) - N_i(I)$  values associated to the test set (i.e., whether the test set explores a full coverage of  $N_i(I + E) - N_i(I)$  values).

Note that even when distribution inferred from the  $N_i(I + E) - N_i(I)$  values (i.e.,  $\hat{N}D_i(I + E)$ ) converges to a normal distribution, Themis is not 100% guaranteed to identify all fault-inducing data flows. It is because the  $N_i(I + E) - N_i(I)$  values can coincidentally form a normal distribution different from  $ND_i$ . Hence, Themis probabilistically (rather than deterministically 100%) identifies a full coverage of fault-inducing data flows, and the corresponding probability depends on the confidence interval (by default 95%) adopted by Themis when inferring  $ND_i$  (§III).

After we justify the reason why Themis theoretically can explore all fault-inducing data flows at high probability (95%), we show that this theoretical property of Themis allows Themis to have strong fault detection capability. As discussed in §IV-B, for any DLS having more number of faults (i.e., has lower accuracy), the DLS's  $\hat{N}D_i$  often has greater variance. This implies that these DLSs require Themis to generate more  $I + E(\theta)$  to identify the DLS's  $\hat{N}D_i$  (any statistical distribution with greater variance requires more samples to identify the distribution, according to Chebyshev's inequality [40]). Since with more  $I + E(\theta)$ , Themis can detect more faults (the number of  $I + E(\theta)$  generated by Themis is the upper limit of the number of faults can be detected by Themis). Hence, Themis detects more faults from a DLS which has lower accuracy (i.e., Themis has strong fault detection capability).

## V. IMPLEMENTATION

We implemented Themis using PyTorch 1.8.1 [42]. Themis's implementation consists of around 9,791 lines of Python code. We adopted PYMC3 [47], a popular Python package which realises MCMC, to perform MCMC in Themis. Specifically, Themis's Sensitivity Calculator (§III) calls PYMC3.DATA to load the samples of sensitivity, and constructs the prior distributions by calling PYMC3.NORMAL. Then, Themis's Sensitivity Convergence Coverage calls PYMC3.SAMPLE to perform MCMC. Sensitivity Convergence Coverage examines the convergence rate (i.e., MCSE, §IV) by ARVIZ.SUMMARY [28], a popular Python package for analysis of Bayesian models. Finally, Themis's Sensitivity Maximizing Fuzzer is realised with textscAutograd [37] library, which efficiently performs gradient ascent on  $\hat{N}D_i$ (§IV).

## VI. EVALUATION

### A. Experiment Setup

Our evaluation was done on a computer equipped with twenty CPU cores and four Nvidia RTX2080TI graphic cards. Inspired by [17] that analyze the correlation of increases in neuron coverage and ASR, we compared Themis with the state-of-the-art DLS testing techniques: DeepeXplore [43], DeepGauge [36], DeepImportance [11], and Surprise Ade-

quacy [27]. During the testing process, we will mutate test cases with perturbations in III and require the NC (DeepXplore), KMNC (DeepGauge), IDC (DeepImportance), and LSA (Surprise Adequacy) to select test cases to increase the coverage rate. For NC, we set the threshold as 0.5, For KMNC, we follow the configuration [36] and set  $k$  as 1,000. For LSA, we set the upper bound as 2,000. For IDC, we set the  $m$  as 12.

Table I shows our evaluated datasets and models, which covers a complete set of six datasets evaluated by four baselines. To be fair and comprehensive, our evaluated models covered **all** models evaluated by four notable DL testing works [11], [27], [36], [43] that have been deployed in Mindspore security framework. The models we evaluated include VGG [49], ResNet [18], DenseNet [24], etc. We believe the architecture of these models covers all basic building blocks (such as convolution layers and residual connections) of most modern real-world deployed AI applications. These DNN models all achieve high precision. Hence, finding faults (i.e., a DNN's incorrect outputs) from these well-studied and well-tested DNN will be valuable.

We choose the correlation metric [16], [31] as the main evaluation metric to compare the fault detection capability between Themis and the baselines. The value of the correlation metric is between -1.0 and 1.0 and the latest study [16] points out that a high correlation (e.g., 0.7) implies strong fault detection capability. We also evaluated the increased accuracy of DLS after retraining the DLS with the detected faults, to evaluate the quality of faults detected by a DLS testing technique. Finally, we evaluated the test time cost for Themis and the baselines to achieve 100% test coverage<sup>2</sup>.

To study Themis and the baselines' fault detection capability, we applied four well-studied perturbations (CW [6], FGSM [13], PGD [38], and Gaussian noise) to each dataset of the evaluated DNN models in Table I. These perturbations are evaluated by Themis's baselines [11], [36] because these perturbations are common in real world DLS applications. Hence, faults detected from a DLS under these perturbations greatly promote the reliability of a DLS.

The values of the magnitude of these perturbations in our evaluation are listed in Table III. We follow Themis's baselines to set these values [11], [27], [36], [43]. The evaluation questions are as follows:

- §VI-B: How is Themis's fault detection capability compared to the baselines?
- §VI-C: How is the improvement on a DLS's accuracy attained by Themis?
- §VI-D: How is Themis's efficiency compared to baselines?
- §VI-E: What are the factors that affect Themis's fault detection capability?
- §VI-F: What are Themis's limitations?

<sup>2</sup>In this paper, we follow the Mindspore fuzzing testing setup for our experiments, where the testing will be stopped when the coverage reaches 100%.

| Dataset(DNN)          | CW     |             |            |              |                   | FGSM     |             |            |              |                   |
|-----------------------|--------|-------------|------------|--------------|-------------------|----------|-------------|------------|--------------|-------------------|
|                       | Themis | Deep-Xplore | Deep-Gauge | Deep-Import. | Surprise Adequacy | Themis   | Deep-Xplore | Deep-Gauge | Deep-Import. | Surprise Adequacy |
| MNIST (LeNet-1)       | 0.79   | -0.58       | 0.25       | 0.08         | -0.21             | 0.77     | 0.21        | 0.35       | -0.35        | 0.33              |
| MNIST (LeNet-4)       | 0.72   | -0.88       | 0.01       | -0.62        | -0.5              | 0.74     | 0.26        | 0.34       | 0.24         | -0.18             |
| MNIST (LeNet-5)       | 0.71   | 0.35        | -0.41      | -0.17        | 0.74              | 0.83     | -0.02       | -0.02      | 0.18         | -0.05             |
| Contagio (<200,200>)  | 0.71   | 0.33        | -0.07      | 0.25         | -0.87             | 0.77     | 0.71        | -0.04      | 0.1          | 0.38              |
| Drebin (<200, 10>)    | 0.70   | 0.11        | 0.18       | 0.78         | -0.37             | 0.81     | -0.06       | 0.08       | -0.19        | 0.2               |
| ImageNet (VGG-19)     | 0.79   | -0.77       | -0.07      | 0.12         | -0.29             | 0.86     | 0.3         | -0.62      | -0.01        | 0.74              |
| ImageNet (ResNet-50)  | 0.71   | 0.4         | 0.33       | -0.06        | -0.09             | 0.74     | 0.3         | 0.1        | -0.05        | 0.1               |
| Udacity (DAVE-2)      | 0.72   | -0.16       | 0.31       | 0.76         | 0.21              | 0.84     | 0.77        | 0.4        | 0            | -0.08             |
| Cifar10 (ResNet56)    | 0.86   | 0.35        | -0.36      | 0.4          | 0.21              | 0.85     | 0.37        | 0.01       | -0.63        | 0.08              |
| Cifar10 (DenseNet121) | 0.88   | 0.25        | -0.7       | 0.05         | -0.01             | 0.78     | -0.51       | 0.72       | 0.33         | 0.39              |
| Dataset(DNN)          | PGD    |             |            |              |                   | GAUSSIAN |             |            |              |                   |
|                       | Themis | Deep-Xplore | Deep-Gauge | Deep-Import. | Surprise Adequacy | Themis   | Deep-Xplore | Deep-Gauge | Deep-Import. | Surprise Adequacy |
| MNIST (LeNet-1)       | 0.72   | -0.89       | 0.01       | 0.39         | 0.23              | 0.7      | 0.2         | -0.21      | -0.06        | 0.71              |
| MNIST (LeNet-4)       | 0.77   | -0.03       | -0.54      | 0.37         | 0.19              | 0.72     | 0.18        | -0.07      | -0.6         | -0.68             |
| MNIST (LeNet-5)       | 0.71   | 0.01        | -0.04      | -0.23        | 0.76              | 0.89     | -0.17       | -0.09      | 0.75         | 0                 |
| Contagio (<200,200>)  | 0.78   | -0.13       | 0.3        | 0.23         | -0.09             | 0.77     | -0.63       | 0.18       | 0.37         | 0.28              |
| Drebin (<200, 10>)    | 0.71   | 0.4         | -0.02      | -0.1         | -0.57             | 0.82     | 0.26        | -0.32      | 0.36         | 0.11              |
| ImageNet (VGG-19)     | 0.75   | 0.09        | 0.1        | -0.13        | -0.19             | 0.84     | 0.79        | 0.28       | 0.1          | -0.03             |
| ImageNet (ResNet-50)  | 0.78   | -0.09       | -0.56      | 0.02         | -0.05             | 0.86     | -0.16       | -0.17      | -0.09        | -0.17             |
| Udacity (DAVE-2)      | 0.71   | 0.31        | 0.07       | -0.71        | -0.74             | 0.8      | 0.4         | 0.34       | -0.54        | 0.08              |
| Cifar10 (ResNet56)    | 0.8    | -0.1        | 0.09       | -0.08        | 0.36              | 0.95     | -0.57       | 0.11       | 0.06         | 0.35              |
| Cifar10 (DenseNet121) | 0.86   | -0.04       | 0.25       | -0.56        | -0.73             | 0.91     | 0.22        | -0.14      | 0.31         | 0.4               |

TABLE II: Correlation between the number of faults identified by DLS testing and the error rate of a DLS. Correlation larger than 0.7 (colored in green) is considered strong [16]. Correlation less than 0.0 is colored in red.

| Variables (Perturbations) | Magnitude Value |      |      |      |      |
|---------------------------|-----------------|------|------|------|------|
| $c$ Confidence (CW [6])   | 10              | 20   | 30   | 40   | 50   |
| $\epsilon$ (FGSM [13])    | 0.1             | 0.2  | 0.3  | 0.4  | 0.5  |
| $\epsilon$ (PGD [38])     | 0.1             | 0.2  | 0.3  | 0.4  | 0.5  |
| $\sigma$ (Gaussian noise) | 0.01            | 0.02 | 0.03 | 0.04 | 0.05 |

TABLE III: Perturbations adopted in our evaluation.

### B. Themis's Fault Detection Capability Result

We first investigate the fault detection capability of Themis and baselines. Specifically, we perturbed raw inputs (e.g., road condition images) with various noise intensities as shown in Table III. Then, for each set of inputs perturbed by the same noise intensity (e.g., PGD with epsilon as 0.1), we fed these perturbed inputs into a DLS and computed the DLS's error rate (i.e., the proportion of the perturbed inputs being wrongly classified by the DLS) and the number of faults detected by Themis until Themis's coverage metric reached 100%. Finally, the fault detection capability is computed as the correlation between all pairs of a DLS's error rate and the number of faults detected (e.g., the pairs of DLS's error rate and the number of faults detected corresponding to PGD's epsilon as 0.1, 0.2, ..., 0.5 as shown in Table III). Then, the inference accuracy of these DNNs and the number of faults detected by Themis or the baselines are used to compute the fault detection capability of

Themis and the baselines.

Table II shows the fault detection capability of Themis and the baselines on all datasets. Themis's correlation was larger than 0.7 for all the evaluated datasets, while baselines rarely achieved 0.7 correlation or even had negative correlation (i.e., detected fewer faults for DNN which has lower accuracy). It is because the baselines tested a DLS deterministically, while Themis tested a DLS probabilistically (§II-C). We further inspect why Themis achieved strong fault detection capability in Figure 3b. As discussed in §IV-B, the condition for Themis to achieve strong fault detection capability is that Themis precisely computes  $\hat{N}D_i$ , which guides Themis to attain full fault coverage. Hence, we inspect whether Themis precisely computes  $\hat{N}D_i$ , by inspecting the consistency of  $\hat{N}D_i$ 's means computed by Themis: Themis is supposed to compute the same  $\hat{N}D_i$ 's mean for the same DLS on different sets of test inputs because theoretically, there is only one  $\hat{N}D_i$  for each DLS's neuron for the same set of test inputs (§II-C).

Figure 3b shows that Themis computed  $\hat{N}D_i$ 's mean with less than 10% deviation, so according to conventional standards in statistics [15], Themis successfully computed  $\hat{N}D_i$ . Themis can do this because theoretically,  $\hat{N}D_i$  could be identified by Bayesian analysis, and Themis's MCMC component could identify it. Figure 4b also confirms our observation that greater sensitivity ( $N_i(I + E(\theta)) - N_i(I)$ ) values resulted in

higher number of faults.

### C. Increase in accuracy obtained by Themis

One main purpose of DLS testing is to retrain a DLS with detected faults to increase the DLS's accuracy. Figure 1 shows the increase in DLS's accuracy brought by Themis and the baselines. Overall, Themis increased the DLS's accuracy higher than baselines for all datasets. The main reason is that Themis identified more faults than the baselines (Figure 3a). Specifically, in our evaluation, we let both Themis and the baselines keep generating inputs until their coverage metric reached 100%. By doing so, Themis on average generated 21399 test inputs, DeepXplore [43] generated 3774 test inputs, DeepGuage [36] generated 18501 test inputs, DeepImportance [11] generated 20052 test inputs, and Surprise Adequacy [27] generated 26757 test inputs. Overall, the test inputs generated by Themis consisted of much more faults than the baselines.

We found that Themis increased the accuracy of Cifar10 (densenet121) under PGD perturbation the most (i.e., 11.56%). It is because densenet121 is vulnerable to PGD perturbation [44], so it suffered from the greatest accuracy loss under this perturbation. Since Themis detected more faults from DLS which had lower accuracy (Table II), Themis detected more faults for this DNN than other DNNs, which allowed Themis to increase the DNN's accuracy more than Themis did on the other DNNs. This also implies Themis is valuable to real-world DLSs, which are often trained with limited datasets and have moderate accuracy (especially on safety-critical tasks, see §I).

### D. Themis's testing time cost

We then study the efficiency of Themis. Figure 4a shows the test time cost of Themis and the baselines. Themis on average had 27.1% more testing time cost than baselines. It is because Themis was probabilistic and hence required more computation than deterministic approaches (see IV-A). Nevertheless, Themis identified 3.78X more faults than baselines, so the testing time cost is worthwhile.

To identify the source of Themis's time cost, we break down Themis's time cost, which is comprised of three main components: Sensitivity Calculator, Sensitivity Convergence Coverage, and Sensitivity Maximizing Fuzzer (III). Table IV shows Themis's time cost in Sensitivity Calculator, Sensitivity Convergence Coverage and Sensitivity Maximizing Fuzzer with or without Themis's Sensitivity Sampler. From the table, we can see that Themis's performance overhead was mainly from Sensitivity Maximizing Fuzzer, which iteratively performs MCMC (known to be time-consuming). Without Themis's sampler component, Themis's time on Sensitivity Maximizing Fuzzer was enormous (more than one hour), because Themis had to perform MCMC on all DNN's neurons. With Themis's sampler, Themis could accurately compute the coverage metric based on the results of one thousand DNN neurons ( §IV-A).

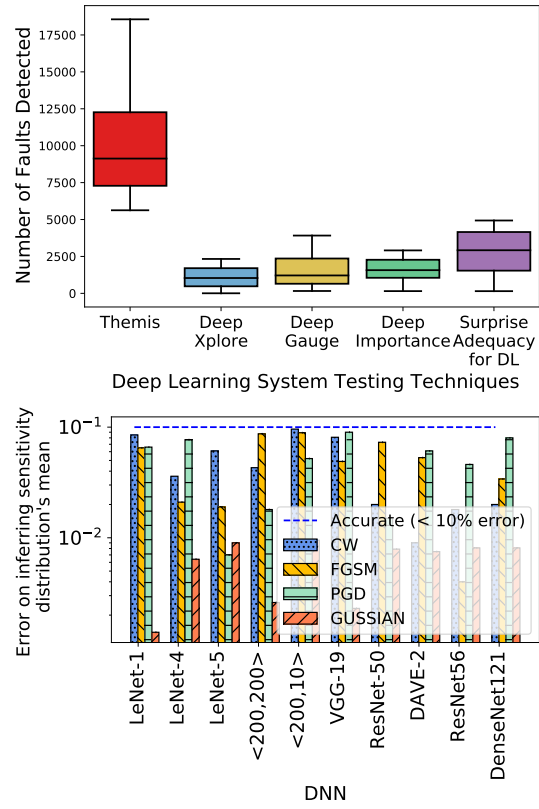


Fig. 3: Above(a): Number of faults detected by Themis and the baselines when their test coverage metric reaches 100%. Bottom(b) The mean of  $\hat{N}\hat{D}_i$  inferred by Themis for each evaluated DNN. Themis is considered accurate in inferring the mean if the variation of the mean is below 10% [15] .

### E. Sensitivity studies on Themis's parameters

We study the sensitivity of Themis's parameters (i.e., threshold  $t$  and Sensitivity Sampler's sampler size). Figure 5 shows Themis's coverage variation on these values. Specifically, Figure 5a shows that Themis's coverage was the same as the ground truth when the sample size equaled one thousand, which was coherent with Mann-Witney Test theory [5] (§IV-B). We could also observe that when the sample size increased, Themis's coverage rate was closer to the ground-truth. It is because statistically, with more samples, we could approximate the ground-truth distribution better. Nevertheless, the increase in precision is diminishing when the sample size grows larger. Hence, setting the sample size as one thousand is desirable, as justified theoretically (§IV-B) and empirically (Figure 5a). Figure 5b shows the variation of Themis's coverage against a threshold value ( $t$ ). The figure shows that setting  $t$  as zero (i.e., the default value) allowed Themis to achieve the ground-truth coverage value.

### F. Limitations and Threat to Validity

#### a) limitations

Themis has two main limitations. First, Themis requires extra testing time (20%) than the baselines, because Themis probabilistically tests a DNN (§VI-D). Nevertheless, even

| Dataset(DNN)          | with Sensitivity Sampler |                           |                         |                                    | without Sensitivity Sampler |                           |                         |                                    |
|-----------------------|--------------------------|---------------------------|-------------------------|------------------------------------|-----------------------------|---------------------------|-------------------------|------------------------------------|
|                       | Total<br>Test Time       | Sensitivity<br>Calculator | Sensitivity<br>Coverage | Sensitivity Fuzzer<br>(iterations) | Total<br>Test Time          | Sensitivity<br>Calculator | Calculator<br>Estimator | Sensitivity Fuzzer<br>(iterations) |
| MNIST (LeNet-1)       | 160                      | 27                        | 24                      | 109 (5)                            | 11171                       | 24                        | 2006                    | 9141 (5)                           |
| MNIST (LeNet-4)       | 143                      | 22                        | 11                      | 110 (10)                           | 2434                        | 23                        | 217                     | 2194 (11)                          |
| MNIST (LeNet-5)       | 208                      | 59                        | 34                      | 115 (4)                            | 1704                        | 56                        | 379                     | 1269 (4)                           |
| Contagio (<200,200>)  | 201                      | 49                        | 9                       | 143 (15)                           | 14320                       | 50                        | 856                     | 13414 (16)                         |
| Drebin (<200,10>)     | 211                      | 27                        | 17                      | 167 (10)                           | 12690                       | 27                        | 1139                    | 11524 (12)                         |
| ImageNet (VGG-19)     | 202                      | 49                        | 23                      | 130 (6)                            | 10273                       | 50                        | 1533                    | 8690 (6)                           |
| ImageNet (ResNet-50)  | 131                      | 57                        | 2                       | 72 (32)                            | 5470                        | 51                        | 162                     | 5257 (32)                          |
| Udacity (DAVE-2)      | 202                      | 33                        | 7                       | 162 (24)                           | 2055                        | 33                        | 80                      | 1942 (25)                          |
| Cifar10 (ResNet56)    | 241                      | 39                        | 24                      | 178 (7)                            | 18036                       | 39                        | 2159                    | 15838 (7)                          |
| Cifar10 (DenseNet121) | 155                      | 22                        | 23                      | 110 (5)                            | 1613                        | 21                        | 270                     | 1322 (6)                           |

TABLE IV: Themis’s testing time breakdown (all the numbers are in the unit of seconds). Those total testing time costs comparable to the baselines are colored in green.

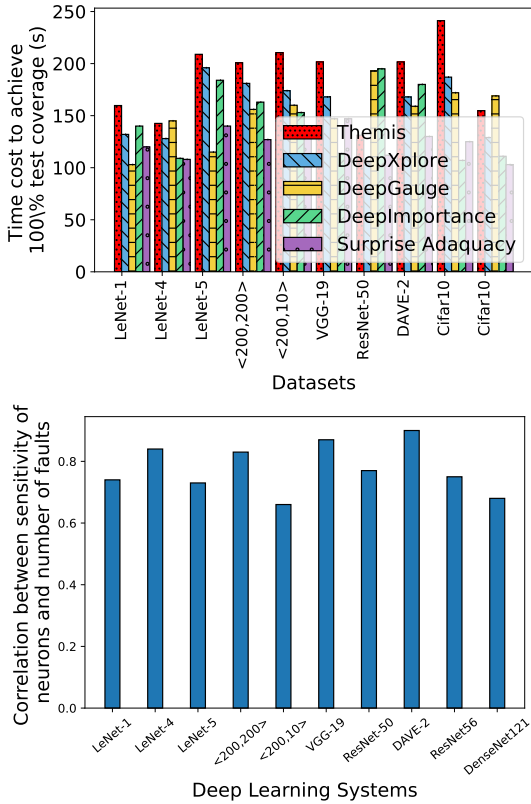


Fig. 4: Above (a): Average time taken by DLS testing to complete testing (achieve 100% test coverage). Bottom (b): correlation between sensitivity of neurons and the error rate.

though Themis has the additional testing time, Themis still completes the testing within several minutes, which is comparable with related work [11], [27], [36], [43]. Besides, Themis detected 3.78X more faults and enhanced the accuracy of DNN on average 14.7X times more than baselines. Hence, Themis’s additional testing time cost is worthwhile. The second limitation is that although Themis increased a DLS’s accuracy by retraining the DLS with faults detected by Themis, Themis does not guarantee these faults would be eliminated

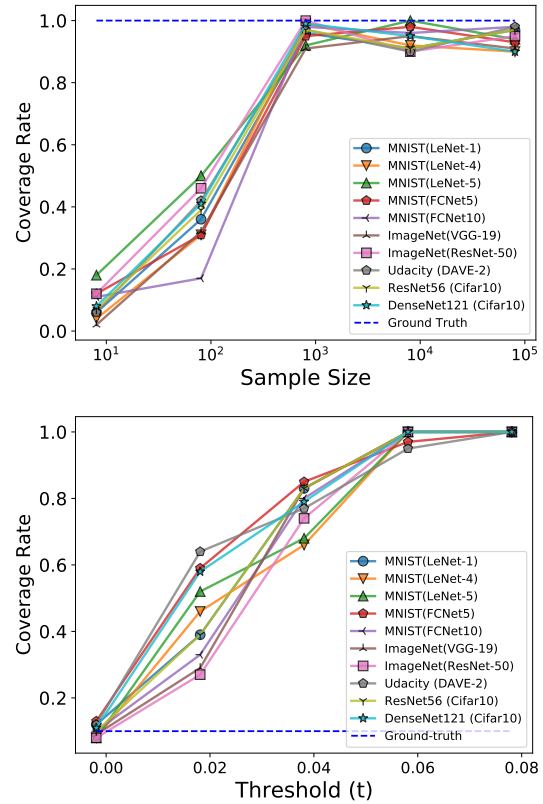


Fig. 5: Above (a): Themis’s coverage values for different sample sizes. Bottom (b): Themis’s coverage values for different threshold values.

from the DLS after retraining. Due to the randomness nature of sampling techniques, there is a significant probability that sensitivity samples converge to a normal distribution while a tiny portion of fault-inducing data flows is not covered (§IV-B): the sensitivity samples can coincidentally converge to a normal distribution different from the ground-truth; hence, Themis probabilistically (95%, rather than deterministically 100%) covers the fault-inducing data flows in a DLS. Indeed, how a DLS’s faults can be eliminated with a guarantee is still

an open challenge for all DLS testing techniques [16].

#### b) Threats to Validity

The accuracy of our experimental measures and the relevance of the theoretical concepts tested are crucial for construct validity. A significant concern is whether the chosen datasets and deep learning models adequately represent the complexity of real-world scenarios. To mitigate this, we utilized a variety of well-recognized datasets (MNIST, Contagio, Drebin, ImageNet, Udacity, and CIFAR10) and deep learning systems (LeNet, VGG, ResNet, DAVE, DenseNet), each with different architectures. This diverse range ensures our findings are not limited to specific models or scenarios, thus enhancing the robustness of our construct validity.

### VII. CONCLUSION

This paper presents Themis, an automated testing technique that addresses the critical challenge of fault detection in Deep Learning Systems (DLSs) used in safety-critical applications. Traditional DLS testing methods, though inspired by software testing principles, fall short in dealing with the complexities of deep learning models, particularly their sensitivity to input perturbations. Themis overcomes these limitations by automatically exploring fault-inducing data flows, significantly reducing the reliance on manual testing efforts. This novel approach is based on the key observation that most fault-inducing data flows in DLSs are sensitive to slight changes in input. The effectiveness of Themis is demonstrated through rigorous evaluation, showcasing its superior performance in fault detection when compared to existing DLS testing techniques. It not only achieves a higher correlation in detecting faults but also enhances the overall accuracy of DLS models. Consequently, Themis contributes significantly to improving the reliability and efficiency of DLS applications in real-world scenarios. This advancement highlights the importance of continued innovation in AI and machine learning, ensuring these technologies meet the evolving demands of safety-critical systems in our increasingly digital world.

### VIII. ACKNOWLEDGEMENTS

The work is supported in part by National Key R&D Program of China (2022ZD0160201), HK RIF (R7030-22), HK ITF (GHP/169/20SZ), the Huawei Flagship Research Grant in 2023, HK RGC GRF (Ref.: 17208223 and 17204424), the HKU-CAS Joint Laboratory for Intelligent System Software, and the Shanghai Artificial Intelligence Laboratory, the National Research Foundation, Singapore, and the Cyber Security Agency under its National Cybersecurity R&D Programme (NCRP25-P04-TAICeN). Any opinions, findings and conclusions or recommendations expressed in this material are those of the author(s) and do not reflect the views of National Research Foundation, Singapore and Cyber Security Agency of Singapore.

### REFERENCES

- [1] Baidu Apollo team (2017), Apollo: Open Source Autonomous Driving, howpublished = <https://github.com/apolloauto/apollo>, note = Accessed: 2019-02-11.
- [2] Cifar10 model in keras.
- [3] Douglas G Altman and J Martin Bland. How to obtain the confidence interval from a p value. *Bmj*, 343, 2011.
- [4] Daniel Arp, Michael Spreitzenbarth, Malte Hubner, Hugo Gascon, Konrad Rieck, and CERT Siemens. Drebin: Effective and explainable detection of android malware in your pocket. In *Ndss*, volume 14, pages 23–26, 2014.
- [5] ZW Birnbaum. On a use of the mann-whitney statistic. In *Volume 1 Contribution to the Theory of Statistics*, pages 13–18. University of California Press, 2020.
- [6] Nicholas Carlini and David Wagner. Towards evaluating the robustness of neural networks. In *2017 IEEE Symposium on Security and Privacy (SP)*, pages 39–57. IEEE, 2017.
- [7] Ricky T. Q. Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. In S. Bengio, H. Wallach, H. Larochelle, K. Grauman, N. Cesa-Bianchi, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018.
- [8] Contagio. Contagio, pdf malware dump. 2010.
- [9] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pages 248–255. Ieee, 2009.
- [10] Li Deng. The mnist database of handwritten digit images for machine learning research [best of the web]. *IEEE Signal Processing Magazine*, 29(6):141–142, 2012.
- [11] Simos Gerasimou, Hasan Ferit Eniser, Alper Sen, and Alper Cakan. Importance-driven deep learning system testing. In *2020 IEEE/ACM 42nd International Conference on Software Engineering (ICSE)*, pages 702–713. IEEE, 2020.
- [12] Charles J Geyer. Practical markov chain monte carlo. *Statistical science*, pages 473–483, 1992.
- [13] Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. *arXiv preprint arXiv:1412.6572*, 2014.
- [14] Kathrin Grosse, Nicolas Papernot, Praveen Manoharan, Michael Backes, and Patrick McDaniel. Adversarial perturbations against deep neural networks for malware classification. *arXiv preprint arXiv:1606.04435*, 2016.
- [15] Bradley Harding, Christophe Tremblay, and Denis Cousineau. Standard errors: A review and evaluation of standard error estimators using monte carlo simulations. *The Quantitative Methods for Psychology*, 10(2):107–123, 2014.
- [16] Fabrice Harel-Canada, Lingxiao Wang, Muhammad Ali Gulzar, Quanquan Gu, and Miryung Kim. Is neuron coverage a meaningful measure for testing deep neural networks? In *Proceedings of the 28th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, pages 851–862, 2020.
- [17] Fabrice Harel-Canada, Lingxiao Wang, Muhammad Ali Gulzar, Quanquan Gu, and Miryung Kim. Is neuron coverage a meaningful measure for testing deep neural networks? *Proceedings of the 28th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, 2020.
- [18] Kaiping He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [19] David W Hosmer and Stanley Lemeshow. Confidence interval estimation of interaction. *Epidemiology*, pages 452–456, 1992.
- [20] Qiang Hu, L. Ma, Xiaofei Xie, Bing Yu, Yang Liu, and Jianjun Zhao. Deepmutation++: A mutation testing framework for deep learning systems. *2019 34th IEEE/ACM International Conference on Automated Software Engineering (ASE)*, pages 1158–1161, 2019.
- [21] Dong HUANG and Qingwen Bu. Adversarial feature map pruning for backdoor. In *The Twelfth International Conference on Learning Representations*, 2024.
- [22] Dong Huang, Qingwen Bu, Yichao Fu, Yuhao Qing, Bochong Xiao, and Heming Cui. Neuron sensitivity guided test case selection for deep learning testing, 2023.
- [23] Dong Huang, Qingwen Bu, Yahao Qing, Yichao Fu, and Heming Cui. Feature map testing for deep neural networks, 2023.
- [24] Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely connected convolutional networks. In *Proceedings*

- of the IEEE conference on computer vision and pattern recognition, pages 4700–4708, 2017.
- [25] Shinpei Kato, Shota Tokunaga, Yuya Maruyama, Seiya Maeda, Manato Hirabayashi, Yuki Kitsukawa, Abraham Monrroy, Tomohito Ando, Yusuke Fujii, and Takuya Azumi. Autoware on board: Enabling autonomous vehicles with embedded systems. In *2018 ACM/IEEE 9th International Conference on Cyber-Physical Systems (ICCPs)*, pages 287–296. IEEE, 2018.
  - [26] Jinhan Kim, Robert Feldt, and Shin Yoo. Guiding deep learning system testing using surprise adequacy. *2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE)*, pages 1039–1049, 2018.
  - [27] Jinhan Kim, Robert Feldt, and Shin Yoo. Guiding deep learning system testing using surprise adequacy. In *2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE)*, pages 1039–1049. IEEE, 2019.
  - [28] Ravin Kumar, Colin Carroll, Ari Hartikainen, and Osvaldo Martin. Arviz: a unified library for exploratory analysis of bayesian models in python. *Journal of Open Source Software*, 4(33):1143, 2019.
  - [29] Yann LeCun, Bernhard E Boser, John S Denker, Donnie Henderson, Richard E Howard, Wayne E Hubbard, and Lawrence D Jackel. Hand-written digit recognition with a back-propagation network. In *Advances in neural information processing systems*, pages 396–404, 1990.
  - [30] Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
  - [31] Zenan Li, Xiaoxing Ma, Chang Xu, and Chun Cao. Structural coverage criteria for neural networks could be misleading. In *2019 IEEE/ACM 41st International Conference on Software Engineering: New Ideas and Emerging Results (ICSE-NIER)*, pages 89–92. IEEE, 2019.
  - [32] Zichao Long, Yiping Lu, Xianzhong Ma, and Bin Dong. Pde-net: Learning pdes from data. In *International Conference on Machine Learning*, pages 3208–3216. PMLR, 2018.
  - [33] L. Ma, Felix Juefei-Xu, Minhui Xue, Bo Li, Li Li, Yang Liu, and Jianjun Zhao. Deepct: Tomographic combinatorial testing for deep learning systems. *2019 IEEE 26th International Conference on Software Analysis, Evolution and Reengineering (SANER)*, pages 614–618, 2019.
  - [34] L. Ma, Felix Juefei-Xu, Fuyuan Zhang, Jiyuan Sun, Minhui Xue, Bo Li, Chunyang Chen, Ting Su, Li Li, Yang Liu, Jianjun Zhao, and Yadong Wang. Deepgauge: Multi-granularity testing criteria for deep learning systems. *2018 33rd IEEE/ACM International Conference on Automated Software Engineering (ASE)*, pages 120–131, 2018.
  - [35] L. Ma, Fuyuan Zhang, Jiyuan Sun, Minhui Xue, Bo Li, Felix Juefei-Xu, Chao Xie, Li Li, Yang Liu, Jianjun Zhao, and Yadong Wang. Deepmutation: Mutation testing of deep learning systems. *2018 IEEE 29th International Symposium on Software Reliability Engineering (ISSRE)*, pages 100–111, 2018.
  - [36] Lei Ma, Felix Juefei-Xu, Fuyuan Zhang, Jiyuan Sun, Minhui Xue, Bo Li, Chunyang Chen, Ting Su, Li Li, Yang Liu, et al. Deepgauge: Multi-granularity testing criteria for deep learning systems. In *Proceedings of the 33rd ACM/IEEE International Conference on Automated Software Engineering*, pages 120–131, 2018.
  - [37] Dougal Maclaurin. *Modeling, inference and optimization with composable differentiable procedures*. PhD thesis, 2016.
  - [38] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. *arXiv preprint arXiv:1706.06083*, 2017.
  - [39] Valentin Jean Marie Manès, HyungSeok Han, Choongwoo Han, Sang Kil Cha, Manuel Egele, Edward J Schwartz, and Maverick Woo. The art, science, and engineering of fuzzing: A survey. *IEEE Transactions on Software Engineering*, 2019.
  - [40] Albert W Marshall and Ingram Olkin. Multivariate chebyshev inequalities. *The Annals of Mathematical Statistics*, pages 1001–1014, 1960.
  - [41] Grégoire Montavon, Alexander Binder, Sebastian Lapuschkin, Wojciech Samek, and Klaus-Robert Müller. Layer-wise relevance propagation: an overview. *Explainable AI: interpreting, explaining and visualizing deep learning*, pages 193–209, 2019.
  - [42] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. *arXiv preprint arXiv:1912.01703*, 2019.
  - [43] Kexin Pei, Yinzhao Cao, Junfeng Yang, and Suman Jana. Deepxplore: Automated whitebox testing of deep learning systems. In *proceedings of the 26th Symposium on Operating Systems Principles*, pages 1–18, 2017.
  - [44] Arash Rahnema, Andre T Nguyen, and Edward Raff. Robust design of deep neural networks against adversarial attacks based on lyapunov theory. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8178–8187, 2020.
  - [45] Samik Raychaudhuri. Introduction to monte carlo simulation. In *2008 Winter simulation conference*, pages 91–100. IEEE, 2008.
  - [46] Christian P Robert, Peter J Bickel, P Diggle, S Fienberg, K Krickeberg, I Olkin, N Wermuth, and S Zeger. *Discretization and MCMC convergence assessment*, volume 135. Springer Science & Business Media, 1998.
  - [47] John Salvatier, Thomas V Wiecki, and Christopher Fonnesbeck. Probabilistic programming in python using PyMC3. *PeerJ Computer Science*, 2:e55, apr 2016.
  - [48] Jasmine Sekhon and Cody Harrison Fleming. Towards improved testing for deep learning. *2019 IEEE/ACM 41st International Conference on Software Engineering: New Ideas and Emerging Results (ICSE-NIER)*, pages 85–88, 2019.
  - [49] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
  - [50] Yousheng Sun, Min Wu, Wenjie Ruan, Xiaowei Huang, M. Kwiatkowska, and Daniel Kroening. Concolic testing for deep neural networks. *2018 33rd IEEE/ACM International Conference on Automated Software Engineering (ASE)*, pages 109–119, 2018.
  - [51] Michael Sutton, Adam Greene, and Pedram Amini. *Fuzzing: brute force vulnerability discovery*. Pearson Education, 2007.
  - [52] Yuchi Tian, Kexin Pei, Suman Sekhar Jana, and Baishakhi Ray. Deeptest: Automated testing of deep-neural-network-driven autonomous cars. *2018 IEEE/ACM 40th International Conference on Software Engineering (ICSE)*, pages 303–314, 2017.
  - [53] Udacity-challenge. Using deep learning to predict steering angles. 2016.
  - [54] Dong Wang, Ziyuan Wang, Chunrong Fang, Yanshan Chen, and Zhenyu Chen. Deeppath: Path-driven testing criteria for deep neural networks. *2019 IEEE International Conference On Artificial Intelligence Testing (AITest)*, pages 119–120, 2019.
  - [55] Xiaofei Xie, Tianlin Li, Jian Wang, L. Ma, Qing Guo, Felix Juefei-Xu, and Yang Liu. Npc: Neuron path coverage via characterizing decision logic of deep neural networks. *ACM Transactions on Software Engineering and Methodology (TOSEM)*, 31:1 – 27, 2022.
  - [56] Xiaofei Xie, L. Ma, Felix Juefei-Xu, Minhui Xue, Hongxu Chen, Yang Liu, Jianjun Zhao, Bo Li, Jianxiong Yin, and S. See. Deephunter: a coverage-guided fuzz testing framework for deep neural networks. *Proceedings of the 28th ACM SIGSOFT International Symposium on Software Testing and Analysis*, 2019.
  - [57] Stephan Zheng, Yang Song, Thomas Leung, and Ian Goodfellow. Improving the robustness of deep neural networks via stability training. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4480–4488, 2016.
  - [58] Yi Zhou, Li Liu, Ling Shao, and Matt Mellor. Dave: A unified framework for fast vehicle detection and annotation. In *European Conference on Computer Vision*, pages 278–293. Springer, 2016.