



Original Software Publication

GIFTbench: Generative image fuzz testing benchmark

Maryam Maryam ^a, Matteo Biagiola ^{b,*}, Andrea Stocco ^c, Vincenzo Riccio ^a^a University of Udine, Italy^b University of St. Gallen and Università della Svizzera italiana, Switzerland^c Technical University of Munich & fortiss, Germany

ARTICLE INFO

Keywords:

Software testing

Generative AI

Search-based software engineering

ABSTRACT

GIFTbench is a modular framework for testing Deep Learning image classifiers that combines Generative AI with genetic algorithms. Its architecture integrates pretrained generative models with a user-friendly Gradio interface, enabling automated, reproducible, and interpretable robustness testing. Supporting VAE, GAN, and Diffusion models, GIFTbench generates test inputs by perturbing latent representations to expose misbehaviors of the classifier under test. By automating test input generation and reducing the need for manual coding, GIFTbench accelerates experimentation and facilitates comparative evaluation of both classifiers and generative models. Designed for researchers and practitioners, it enables reproducible assessment of image classifiers, while supporting studies on classifier vulnerabilities, mutation strategies, and the role of generative models in robustness testing.

Metadata

Table 1

Code metadata.

Code metadata description	Please fill in this column
Current code version	v1.2.1
Permanent link to code/repository used for this code version	https://doi.org/10.5281/zenodo.16966421
Permanent link to Reproducible Capsule	https://hub.docker.com/r/maryam483/giftbench
Legal Code License	MIT
Code versioning system used	git
Software code languages, tools, and services used	Python 3.10, PyTorch, diffusers, Gradio
Compilation requirements, operating environments and dependencies	CUDA-enabled, Docker, other required packages listed in <code>requirements.txt</code>
If available, link to developer documentation/manual	https://github.com/deeptestai/genai_tigs/blob/tool/README.md
Support Email for questions	maryam@spes.uniud.it

1. Motivation and significance

Deep Neural Network (DNN) based image classifiers have become integral components of software systems, also in safety-critical domains such as healthcare and autonomous driving. On standard benchmark datasets, these classifiers often outperform traditional

* Corresponding author.

E-mail addresses: maryam@spes.uniud.it (M. Maryam), matteo.biagiola@usi.ch (M. Biagiola), andrea.stocco@tum.de (A. Stocco), vincenzo.riccio@uniud.it (V. Riccio).

<https://doi.org/10.1016/j.scico.2026.103502>

Received 29 August 2025; Received in revised form 13 April 2026; Accepted 27 April 2026

Available online 6 May 2026

0167-6423/© 2026 The Authors. Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

approaches and even human experts [1–3]. However, these benchmarks do not fully capture the diversity and unpredictability of real-world conditions encountered during operation. As a result, DNNs struggle to generalize when exposed to new or slightly perturbed inputs, raising concerns about their robustness and reliability in practice [4,5].

The gap between training data and real-world inputs highlights the need for systematic testing approaches. A major challenge for software testers is to generate test inputs that accurately reflect real-world operating conditions and trigger misclassifications, i.e., unexpected behaviors in which the predicted labels deviate from the expected ones.

To address this challenge, researchers have proposed several *Test Input Generators* (TIGs), i.e., tools that automatically produce synthetic images for assessing the quality of DNN classifiers [6–9]. Recent advances in TIGs exploit the power and creativity of distribution-aware Generative AI (GenAI) models [9–17], which learn the input data distribution in the form of a latent space, a lower-dimensional representation of the input space that captures the key features of the problem domain [18]. By manipulating latent representations, GenAI models can generate novel inputs that are both diverse and semantically meaningful, providing more realistic test cases than traditional approaches [9].

GenAI-based TIGs adopted a variety of architectures, ranging from simpler models such as Variational AutoEncoders (VAEs) to more sophisticated Generative Adversarial Networks (GANs). More recently, diffusion models have emerged as state-of-the-art generative approaches, achieving impressive results but at the cost of increased complexity and computational demands for training. However, it is challenging to assess the specific contribution of different GenAI models, since existing TIGs are often influenced by confounding factors such as variations in testing algorithms, the absence of standardized training and hyperparameter tuning, and limited support for recent innovations like diffusion models. Unlike earlier test generators based on diverse input generation paradigms, GIFTbench supports comparative testing with multiple GenAI backends, including diffusion models, within a single configurable framework. In the literature, similar challenges related to the fair and systematic comparison of alternative DNN testing techniques have been addressed through the proposal of unified benchmarking frameworks, e.g., for test prioritization [19], generation [9], and validation [20], including relevant safety-critical domains such as autonomous driving [21,22]. However, such benchmarks do not readily support the systematic comparison or practical use of recent GenAI models. Moreover, they often lack accessible interfaces, making them difficult to adopt for non-expert users.

To this end, we propose **GIFTbench** (Generative Image Fuzz Testing Benchmark), a framework that combines search-based test generation with different state-of-the-art GenAI models to enable automated testing of DNN classifiers through latent space manipulation. GIFTbench allows researchers and practitioners to analyze classifier behavior with inputs crafted to induce misclassifications. To make experimentation more accessible, GIFTbench integrates a lightweight web-based interface built with Gradio [23], a Python framework that allows users to interact with the tool, configure experiments, and visualize results without requiring extensive coding effort.

GIFTbench has already enabled a fair comparison of test generation capabilities across different GenAI models, performed in our prior work [17]. In a large-scale empirical study, we evaluated three representative architectures (VAEs, GANs, and diffusion models) across four datasets of increasing complexity. Our study revealed several key insights into GenAI-based test generation. In particular, we observed that diffusion models achieve superior performance on complex tasks, but at the cost of substantially higher inference time (up to 10× slower than alternative models). We also found that applying larger perturbations can accelerate test generation without compromising input validity. GIFTbench includes predefined datasets, pretrained generative models, default classifiers, and parameter configurations that allow users to reproduce the experimental setting of our study [17] with minimal setup.

By providing an intuitive and modular integration of search-based testing techniques with GenAI models, GIFTbench supports the advancement of research on the quality assurance of DNNs, while also helping practitioners identify the generative approach best suited to their classification tasks. A permanent archived release of GIFTbench is available on Zenodo [24].

2. The GIFTbench framework

GIFTbench is a modular and user-friendly framework for testing image classifiers using GenAI models. Its architecture integrates (1) input generation through GenAI models, (2) search-based optimization for latent space manipulation, and (3) evaluation and visualization.

At its core, GIFTbench automatically generates diverse and perturbed inputs to evaluate the robustness of classifiers. It currently supports 3 state-of-the-art GenAI architectures: VAEs [25], GANs [18] and Diffusion Models [26]. These are combined with an evolutionary search mechanism based on a Genetic Algorithm, which systematically explores input variations to produce misbehavior-inducing test cases. The full specification of this search procedure, including initialization, fitness evaluation, selection, crossover, mutation, and stopping criteria, is presented in our prior research paper [17]. Our framework is highly configurable, allowing users to control key testing parameters such as the testing budget (i.e., the number of evolutionary iterations) and the magnitude of perturbations applied to latent vectors. To improve accessibility, GIFTbench provides an interactive interface built with Gradio [23], enabling experimentation and visualization of results without requiring programming expertise.

2.1. Software architecture

GIFTbench follows a layered architecture that separates the user interface from the core engine, as depicted in Fig. 1.

The **User Interface Layer** serves as the system's front end, implemented with Gradio [23]. It allows users to interact with the platform via a simple web interface [27]. Through this layer, users can: (1) select datasets (e.g., MNIST [28], ImageNet [29]); (2) configure pretrained GenAI models (e.g., VAEs, diffusion models) and classifiers (default or user-supplied); (3) adjust search

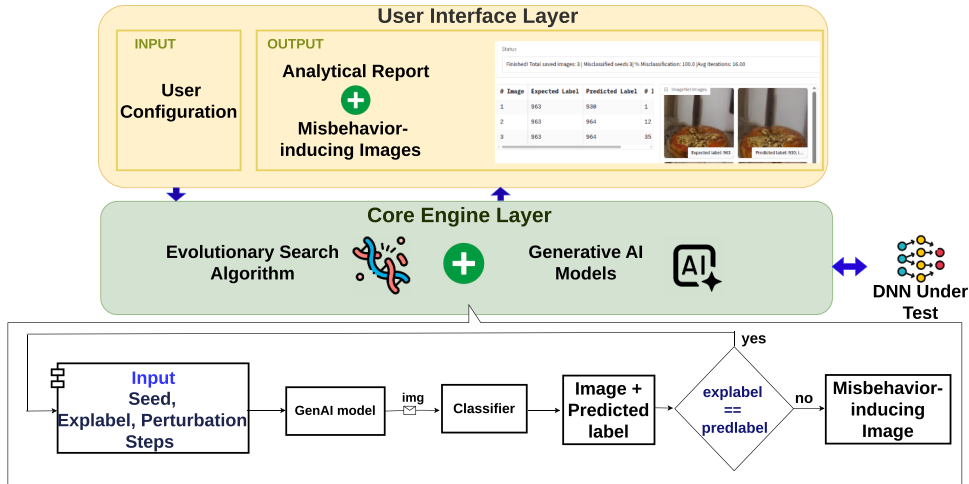


Fig. 1. Layered architecture of GIFTbench, showing the user interface, core engine, and the iterative search process in the latent space.

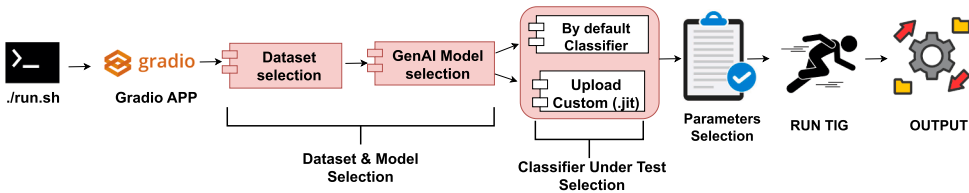


Fig. 2. GIFTbench workflow.

parameters using sliders and toggles; and (4) launch and monitor test campaigns. Key search parameters are exposed by the UI to support controlled experimentation without requiring users to modify the framework internals. The interface automatically verifies dataset-model compatibility (ensuring, for instance, that the chosen generative model was trained on the selected dataset), and enforces classifier-specific constraints such as input dimensions or file formats (e.g., .jit TorchScript models [30]). We integrated such compatibility checks in the interface to prevent invalid experiment configurations early. Once a test session is executed, results are visualized directly in the interface, including generated images, performance metrics, and summary reports.

The **Core Engine Layer** is the heart of GIFTbench’s test generation process. It integrates a genetic algorithm with the selected GenAI model to test the target classifier. In particular, this layer generates and evolves candidate inputs in the latent space with the goal of inducing misclassifications. Although it interacts closely with the User Interface, the Core Engine is designed to remain independent and easily extensible to new datasets, generative models, and classifiers, thus supporting reuse and adaptation. The lower part of Fig. 1 illustrates its iterative test generation process: starting from a latent seed, an expected label, and a perturbation budget, the test generator progressively modifies the latent representation until it produces an input that causes the classifier’s prediction to differ from the expected label, thereby exposing a robustness weakness.

2.2. Software functionality

Fig. 2 illustrates the end-to-end workflow of GIFTbench. Executing `./run.sh` sets up the Docker-based Gradio app and launches the user interface. From there, the user proceeds with the **dataset selection**, i.e., either *MNIST*, *SVHN*, *CIFAR-10*, or two classes of *ImageNet* (i.e., pizza and teddy bear). Once a dataset is selected, the user must perform **GenAI model selection**, by choosing one of the three generative models (VAE, GAN, or *Diffusion Model*). The user then performs **classifier under test selection** by choosing either the dataset-specific default classifier or uploading a custom classifier. A custom classifier is a user-supplied model whose architecture and training are fully defined by the user. The model is uploaded as a single TorchScript (.jit) file, i.e., a PyTorch-serialized, executable representation of the model produced via `torch.jit`. TorchScript improves usability by encapsulating user-defined classifiers into a single, self-contained executable artifact, which eliminates the need for dependency resolution and custom code integration within the GIFTbench codebase. We adopted TorchScript as the interface for custom classifiers to avoid ad hoc integration of model-specific dependencies. To ensure compatibility, the uploaded classifier must follow the same dataset-specific input configuration (e.g., image resolution, channel format, and preprocessing) as the selected dataset, while the internal model architecture and training procedure remain entirely user-defined. Through **parameters selection**, the user can initialize the testing process with one of the default configurations from our empirical study [17], or customize the parameters’ sliders, which is particularly useful when testing custom classifiers. The number of test images can be specified via the `Images_to_Sample` parameter. The underlying Genetic Algorithm is also customizable: population size (default 25), number of iterations (default 250), and perturbation size (low/high). Then, the

user can **run the TIG**. Our TIG algorithm manipulates the latent vectors in two steps: first, an initial perturbation to generate the base population, and then iterative adjustments to trigger misclassifications guided by the classifier's responses. Users can toggle perturbation size (low/high) to observe its effects on model robustness and TIG efficiency. Each perturbed image is evaluated by the classifier, and its misclassification likelihood is used as fitness score. The loop continues until a misclassification is detected or the iteration budget is exhausted.

The Gradio **output** interface provides real-time monitoring and results visualization. Users can inspect side-by-side galleries of original and perturbed images, compare expected and predicted labels, and track metrics through a status bar summarizing misclassification rate and average iteration count. A detailed tabular report logs image IDs, expected/predicted labels, and iteration counts, while all generated images can be exported as a .zip archive with metadata-enriched filenames for downstream analysis. This interactive execution and reporting flow makes GIFTbench a practical and efficient tool for evaluating classifier robustness in real time.

3. Using GIFTbench

3.1. Installation and setup

GIFTbench is packaged as a dockerized Python application for easy deployment across different environments. Detailed hardware and network requirements are available in the online user guide [31]. There are two ways to run GIFTbench: (1) using the pre-built Docker image from Docker Hub or (2) building and running from source code (recommended for developers).

3.1.1. Running from docker hub

The pre-built container image provides the fastest way to reproduce our results without installing any dependencies locally.

```
docker pull maryam483/giftbench:v1.2.0
docker run --name giftbench-running --gpus all -p 7860:7860 \
  maryam483/giftbench:v1.2.0
```

Listing 1. Running GIFTbench from Docker Hub

These commands pull the versioned image and launch GIFTbench. On the first run, the container automatically downloads the required pretrained GenAI models. Internet access is therefore required only once, while subsequent runs reuse the cached models.

3.1.2. Building from source

The repository can be cloned from GitHub, and the application can be launched with a single command. To run the tool, users must have Docker already installed.

```
git clone https://github.com/deeptestai/genai_tigs.git
cd genai_tigs
git checkout tool
./run.sh
```

Listing 2. Cloning and running the GIFTbench tool from the tool branch

3.1.3. Accessing the interface

When launched, Gradio provides two possible endpoints:

- **Local URL:** <http://localhost:7860>, accessible directly from the user's browser on the same machine.
- **Public URL:** a temporary link of the form <https://abcdef12345.gradio.live>, automatically generated by Gradio to share sessions. This link expires after 72 h or when the container is stopped.

3.2. Supported datasets and default classifiers

Table 2 establishes the experimental context by summarizing the default classifiers provided in GIFTbench, each trained on one of the four supported datasets (i.e., MNIST, SVHN, CIFAR-10, ImageNet). Users who wish to reproduce our experiments [17] can simply select the corresponding predefined classifier from the dataset-specific options in the interface. Alternatively, GIFTbench allows users to evaluate their own classifiers by selecting the custom classifier option from the drop-down menu, enabling robustness testing of third-party or proprietary models without modifying the framework.

3.3. Test generation

The test generation process can be configured and executed through the user interface shown in **Fig. 3**. After selecting the dataset, GenAI model, classifier, and configuring the genetic algorithm parameters, the user can start the test generation by clicking the Run TIG button.

Table 2
Datasets' input size and classifiers' training accuracy.

Dataset	Image Size	Classifier	Accuracy (%)
MNIST	28x28x1	deepconv [11]	99.46
SVHN	32x32x3	VGGNET [11]	95.20
CIFAR-10	32x32x3	VGGNET [11]	86.38
ImageNet	224x224x3	VGG19bn [32]	75.00

TEST INPUT GENERATORS – (VAE / GAN / Diffusion)

MNISTSVHNCIFAR-10ImageNet (class-pizza)ImageNet (class-teddy)

Model Type

☐ VAE (pizza)

☐ GAN (pizza)

☒ DM (pizza)

Generations

250

10050

Population Size

25

1050

Best Left (Selection)

10

520

Perturbation Size

Low

Initial Perturbation Size

Low

Images to Sample

3

1100

Classifier Under Test

VGG19bn

Prompt (only for DM)

Run ImageNet (class-Pizza) TIGStop

Fig. 3. Configuration of GIFTbench parameters.

Table 3
Misclassified images generated by different models for the ImageNet Pizza class.

Model

Gradio Results

Status

Finished! Total saved images: 3 | Misclassified seeds 3 | % Misclassification: 100.0 | Avg Iterations: 16.00

#	Image	Expected Label	Predicted Label	#
1		963	930	1
2		963	964	12
3		963	964	35

ImageNet Images

Expected label: 963

Predicted label: 930; L...

VAE

Status

Finished! Total saved images: 10 | Misclassified seeds 9 | % Misclassification: 90.0 | Avg Iterations: 147.90

#	Image	Expected Label	Predicted Label	#
1		963	957	192
2		963	963	256
3		963	892	90
4		963	567	176

ImageNet Images

Expected label: 963

Predicted label: 957; L...

GAN

Status

Finished! Total saved images: 4 | Misclassified seeds 4 | % Misclassification: 100.0 | Avg Iterations: 132.50

#	Image	Expected Label	Predicted Label	#
6		963	964	223
7		963	962	1
8		963	567	145
9		963	962	161

ImageNet Images

Expected label: 963

Predicted label: 964; L...

DM



Fig. 4. Representative images generated by GIFTbench for the ImageNet *pizza* class. In all cases, the expected label is *pizza*, while the classifier predicts *trifle*.

3.4. Illustrative example: Testing an ImageNet classifier

A complete evaluation of the GenAI models integrated in GIFTbench is presented in our earlier work [17], where VAE, GAN, and Diffusion Models were benchmarked for their effectiveness in generating valid, misclassification-inducing test inputs across four datasets.

In this paper, we provide a demonstration revisiting a small subset of that setup to illustrate how GIFTbench facilitates test generation and model evaluation. We configured GIFTbench to test a default VGG19bn ImageNet classifier, focusing on the Pizza class. As shown in Fig. 3, the Gradio interface is set up for generating 10 seeds, perform 250 iterations of the genetic algorithm with a population of size 25, and adopting a small perturbation extent in the latent space. By applying the same configuration to all three GenAI model types, we ensure a fair comparison of their effectiveness.

Table 3 demonstrates how GIFTbench operationalizes hypotheses about model robustness, misclassification sensitivity, and the impact of input perturbation extent. After each run, the status bar reports the number and percentage of misclassified seeds, together with the average number of perturbation iterations required. In our illustrative runs (not statistically significant due to the small number of seeds), the GAN model achieved the highest misclassification rate, with 9 out of 10 seeds triggering a misclassification within the test budget.

After executing the search-based testing process, GIFTbench enables users to inspect misclassification-inducing images and analyze failure patterns and erroneous decision boundaries of the classifier under test. However, misclassification-inducing inputs should be interpreted with caution, since some may be invalid or may not preserve the seed's class [9]; only validated inputs provide reliable support for robustness analysis or retraining [33]. This inspection can be aided by state-of-the-art input validators [9,10,34,35] or explanation tools [36,37]. Fig. 4 shows four representative examples in which inputs expected to be classified as “pizza” are instead predicted as “trifle”. These failures suggest that the classifier is confused by top-view images of pizza whose visual characteristics (e.g., color and texture) resemble those of cream- and red berry-based desserts. Such visual evidence supports the diagnosis of systematic weaknesses in the model. The validated misclassified inputs identified by GIFTbench can then be used to augment the training set, enabling targeted retraining aimed at improving the classifier's robustness.

4. Expected impact and significance

GIFTbench provides a novel experimental environment for studying the interplay between generative AI models, latent space perturbations, and classifier vulnerabilities under controlled conditions. It supports systematic comparisons of classifier robustness and generative model effectiveness across diverse architectures and mutation settings.

Compared to prior robustness testing frameworks, GIFTbench offers the most comprehensive integration of evolutionary test generation with GenAI-based input synthesis within a unified framework. Its automated workflows minimize manual coding effort while ensuring reproducibility and extensibility. We anticipate that the tool will serve as an accelerator for researchers, students, and practitioners, enabling them to conduct reproducible robustness evaluations through an intuitive interface that lowers entry barriers and facilitates both experimentation and analysis. At the same time, its modular and extensible design makes it a promising foundation for reuse in early-stage practitioner workflows.

Usage Scenarios. In the following, we summarize the intended usage scenarios and the required user actions at increasing levels of expertise and coding effort.

1. **Benchmarking alternative GenAI-based test generators on a target classifier.** Users can compare VAE-, GAN-, and Diffusion-based test generation under the same search budget and perturbation settings, evaluating their effectiveness and efficiency in generating misclassification-inducing inputs.
2. **Exploring search-based testing by modifying configuration parameters.** Users can explore different search settings (e.g., perturbation budgets, population size) by simply interacting with the Gradio interface controls. This enables experimentation without modifying the framework's source code.
3. **Testing custom classifiers on supported datasets.** GIFTbench allows users to plug in their classifiers with minimal effort. Thanks to the use of TorchScript jit-compatible interfaces [30], classifiers can be uploaded via the Gradio interface without changes to the core framework.

Table 4

Characteristics of the generative AI models: latent vector (LV) size, training time until convergence, and average inference time.

Dataset	Model	LV size	t_{train} (min)	t_{infer} (ms)
MNIST	VAE [38]	400	6	0.27
	GAN [39,40]	100	9	0.70
	DM [41]	16,384	405	960.68
SVHN	VAE [42]	800	93	4.07
	GAN [39,40]	100	86	1.75
	DM [41]	16,384	572	1213.49
CIFAR-10	VAE [42]	1024	423	2.51
	GAN [39,40]	100	450	1.73
	DM [41]	16,384	362	1903.29
ImageNet	VAE [43]	512	2521	11.92
	GAN [44]	128	21,600	20.68
	DM [41]	16,384	30	1945.77

- 4. Extending GIFTbench to new datasets.** Applying the framework to a new dataset is mostly handled at source code level, requiring training dataset-specific generative models, which is standard practice in generative robustness testing. To lower the entry barrier, we: provide training scripts and links to reference model checkpoints for all included datasets; provide a user guide with explicit instructions on how to train and register new dataset-specific generative models [31]; rely on a standard Gradio interface, which users can easily extend once a new dataset and corresponding models are registered.
- 5. Advanced extensions by expert users.** Advanced researchers and practitioners may transparently integrate alternative search or testing algorithms by modifying the code of the engine layer and substituting the test runners in the main `GIFTbench.py` file. These changes do not affect the Gradio interface, allowing advanced experimentation while preserving usability for non-expert users.

Known limitations and trade-offs. GIFTbench currently focuses on image classification tasks and does not yet support other computer vision problems such as object detection or segmentation, nor data modalities beyond images (e.g., text or video). As shown in Table 4, which reports results of our empirical comparison [17], higher-resolution inputs (i.e., ImageNet) and larger generative AI models (i.e., diffusion models) introduce higher inference times. Regarding training costs, smaller generative models are more suitable for simpler datasets, while the considered diffusion models can be efficiently fine-tuned on complex datasets.

5. Conclusions and future work

We presented GIFTbench, a framework that combines generative AI models with evolutionary algorithms to systematically assess the quality of DNN-based image classifiers. By integrating pretrained GenAI models with an intuitive interface, the tool enables automated, reproducible, and visually interpretable robustness testing across multiple datasets.

Future work will focus on enhancing the flexibility and extensibility of the framework, enabling seamless integration of custom classifiers, datasets, and generative models, thus broadening its applicability for both research and practice. Moreover, extending GIFTbench with additional fitness functions and exploration strategies will open avenues for mutation analysis [45,46], targeted test generation [5,15], deeper input space and behavioral boundary exploration [47,48], and increased input diversity [49].

CRedit authorship contribution statement

Maryam Maryam: Conceptualization, Data curation, Investigation, Methodology, Project administration, Software, Validation, Visualization, Writing – original draft, Writing - review & editing; **Matteo Biagiola:** Supervision; **Andrea Stocco:** Conceptualization, Methodology, Project administration, Supervision, Writing – original draft, Writing – review & editing; **Vincenzo Riccio:** Conceptualization, Data curation, Investigation, Methodology, Project administration, Resources, Software, Supervision, Validation, Visualization, Writing – original draft, Writing – review & editing.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgment

This work was supported by the project PNRR M4 C2 I1.3 “SEcurity and RiGhts in the CyBerSpace (SERICS)” (PE00000014 - CUP H73C2200089001, D33C22001300002) PE7, funded by Next-Generation EU. Matteo Biagiola is partially supported by Fondo Istituzionale per la Ricerca granted by Università della Svizzera italiana (USI).

References

- [1] D. Sarwinda, R.H. Paradisa, A. Bustamam, P. Anggia, Deep learning in image classification using residual network (resnet) variants for detection of colorectal cancer, *Procedia Comput. Sci.* 179 (2021) 423–431. <https://doi.org/10.1016/j.procs.2021.01.025>
- [2] G. Lou, Y. Deng, X. Zheng, M. Zhang, T. Zhang, Testing of autonomous driving systems: where are we and where should we go?, in: *Proceedings of the 30th ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, Association for Computing Machinery, New York, NY, USA, 2022, pp. 31–43. <https://doi.org/10.1145/3540250.3549111>
- [3] S. Tang, Z. Zhang, Y. Zhang, J. Zhou, Y. Guo, S. Liu, S. Guo, Y.-F. Li, L. Ma, Y. Xue, Y. Liu, A survey on automated driving system testing: landscapes and trends, 2023. 2206.05961
- [4] A. Guerriero, R. Pietrantuono, S. Russo, Operation is the hardest teacher: estimating DNN accuracy looking for mispredictions, in: *Proceedings of the 43rd International Conference on Software Engineering*, ICSE '21, IEEE Press, 2021, pp. 348–358. <https://doi.org/10.1109/ICSE43902.2021.00042>
- [5] T. Zohdinasab, V. Riccio, P. Tonella, Deepatash: focused test generation for deep learning systems, in: *Proceedings of the 32nd ACM SIGSOFT International Symposium on Software Testing and Analysis*, Association for Computing Machinery, New York, NY, USA, 2023, pp. 954–966. <https://doi.org/10.1145/3597926.3598109>
- [6] J.M. Zhang, M. Harman, L. Ma, Y. Liu, Machine learning testing: survey, landscapes and horizons, *IEEE Trans. Software Eng.* 48 (1) (2020) 1–36.
- [7] V. Riccio, G. Jahangirova, A. Stocco, N. Humbatova, M. Weiss, P. Tonella, Testing machine learning based systems: a systematic mapping, *Empir. Software Eng.* 25 (6) (2020) 5193–5254.
- [8] H.B. Braiek, F. Khomh, On testing machine learning programs, *J. Syst. Software* 164 (2020) 110542.
- [9] V. Riccio, P. Tonella, When and why test generators for deep learning produce invalid inputs: an empirical study, in: *2023 IEEE/ACM 45th International Conference on Software Engineering (ICSE)*, IEEE, 2023, pp. 1161–1173.
- [10] S. Dola, M.B. Dwyer, M.L. Soffa, Distribution-aware testing of neural networks using generative models, in: *2021 IEEE/ACM 43rd International Conference on Software Engineering (ICSE)*, IEEE, 2021, pp. 226–237.
- [11] S. Kang, R. Feldt, S. Yoo, Search-based image space navigation for dnn image classifier test input generation, in: *Proceedings of the IEEE/ACM 42nd International Conference on Software Engineering Workshops*, Sinval, 2020, pp. 521–528.
- [12] I. Dunn, H. Pouget, D. Kroening, T. Melham, Exposing previously undetectable faults in deep neural networks, in: *Proceedings of the 30th ACM SIGSOFT International Symposium on Software Testing and Analysis*, 2021, pp. 56–66.
- [13] A. Aleti, Software testing of generative ai systems: challenges and opportunities, in: *2023 IEEE/ACM International Conference on Software Engineering: Future of Software Engineering (ICSE-FoSE)*, 2023.
- [14] S. Dola, R. McDaniel, M.B. Dwyer, M.L. Soffa, Cit4dnn: generating diverse and rare inputs for neural networks using latent space combinatorial testing, in: *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering*, 2024, pp. 1–13.
- [15] O.W. I. A. Abdellatif, X. Chen, G. Merabishvili, V. Riccio, S. Kacianka, A. Stocco, Targeted deep learning system boundary testing, *ACM Trans. Software Eng. Methodol.*, Association for Computing Machinery, New York, NY, USA, Just Accepted (2025). <https://doi.org/10.1145/3771557>
- [16] L. Baresi, D.Y.X. Hu, A. Stocco, P. Tonella, Efficient domain augmentation for autonomous driving testing using diffusion models, in: *Proceedings of 47th International Conference on Software Engineering*, ICSE '25, IEEE, 2025.
- [17] M. Maryam, M. Biagiola, A. Stocco, V. Riccio, Benchmarking generative AI models for deep learning test input generation, in: *2025 IEEE Conference on Software Testing, Verification and Validation (ICST)*, IEEE, 2025, pp. 174–185.
- [18] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, Y. Bengio, Generative adversarial networks, *Commun. ACM* 63 (11) (2020) 139–144.
- [19] C. Birchler, S. Klikovits, M. Fazzini, S. Panichella, ICST tool competition 2025-self-driving car testing track, in: *2025 IEEE Conference on Software Testing, Verification and Validation (ICST)*, IEEE, 2025, pp. 801–804.
- [20] J. Zhang, J. Keung, X. Ma, X. Li, Y. Xiao, Y. Li, W.K. Chan, Enhancing valid test input generation with distribution awareness for deep neural networks, in: *2024 IEEE 48th Annual Computers, Software, and Applications Conference (COMPSAC)*, IEEE, 2024, pp. 1095–1100.
- [21] M. Biagiola, S. Klikovits, J. Peltomäki, V. Riccio, Sbt tool competition 2023-cyber-physical systems track, in: *IEEE/ACM International Workshop on Search-Based and Fuzz Testing*, IEEE, 2023, pp. 45–48.
- [22] S.C. Lambertenghi, H. Leonhard, A. Stocco, Benchmarking image perturbations for testing automated driving assistance systems, in: *2025 IEEE Conference on Software Testing, Verification and Validation (ICST)*, IEEE, 2025, pp. 150–161.
- [23] A. Abid, A. Abdalla, A. Abid, D. Khan, A. Alfozan, J. Zou, Gradio: hassle-free sharing and testing of ML models in the wild, 2019. 1906.02569
- [24] M. Maryam, M. Biagiola, A. Stocco, V. Riccio, Giftbench, GIFTbench web-based application: Release v1.2.0, Zenodo, 2025. <https://doi.org/10.5281/zenodo.19521006>
- [25] D.P. Kingma, M. Welling, Auto-encoding variational bayes (2022). 1312.6114
- [26] J. Sohl-Dickstein, E.A. Weiss, N. Maheswaranathan, S. Ganguli, Deep unsupervised learning using nonequilibrium thermodynamics (2015). 1503.03585
- [27] Gradio, Gradio: build machine learning demos in python, 2025, <https://www.gradio.app>. Accessed: 2025-05-30.
- [28] L. Deng, The mnist database of handwritten digit images for machine learning research, *IEEE Signal Process. Mag.* 29 (6) (2012) 141–142. [best of the web].
- [29] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, L. Fei-Fei, ImageNet: a large-scale hierarchical image database, in: *2009 IEEE Conference on Computer Vision and Pattern Recognition*, 2009, pp. 248–255.
- [30] PyTorch Team, TorchScript: a way to create serializable and optimized models from pytorch code, 2018, <https://pytorch.org/docs/stable/jit.html>. Accessed: 2025-05-30.
- [31] M. Maryam, M. Biagiola, A. Stocco, V. Riccio, GIFTbench user guide, 2025, https://deeptestai.github.io/genai_tigs/user-guide/. accessed: 2026-04-04.
- [32] PyTorch Contributors, Source code for torchvision.models.VGG, 2024, https://pytorch.org/vision/0.12/_modules/torchvision/models/vgg.html#vgg19. Accessed: 2024-01-01.
- [33] X. Chen, M. Biagiola, V. Riccio, M. d'Amorim, A. Stocco, XMutant: XAI-based fuzzing for deep learning systems, *Empir. Software Eng.* 31 (4) (2026) 102.
- [34] D. Ghobari, M.H. Amini, S. Park, S. Nejati, M. Sabetzadeh, et al., Test input validation for vision-based dl systems: an active learning approach, in: *2025 IEEE/ACM 47th International Conference on Software Engineering: Software Engineering in Practice (ICSE-SEIP)*, IEEE, 2025, pp. 630–640.
- [35] M. Maryam, M. Biagiola, P. Tonella, V. Riccio, DEEPNAQQAL: human-aligned automated validation of test inputs for deep learning, in: *2026 IEEE Conference on Software Testing, Verification and Validation (ICST)*, IEEE, 2026.
- [36] H. Fahmy, F. Pastore, M. Bagherzadeh, L. Briand, Supporting deep neural network safety analysis and retraining through heatmap-based unsupervised learning, *IEEE Trans. Reliab.* 70 (4) (2021) 1641–1657.
- [37] T. Zohdinasab, V. Riccio, P. Tonella, An empirical study on low-and high-level explanations of deep learning misbehaviours, in: *2023 ACM/IEEE International Symposium on Empirical Software Engineering and Measurement (ESEM)*, IEEE, 2023, pp. 1–11.
- [38] S. Kang, R. Feldt, S. Yoo, VAE training code for MNIST, 2020, <https://github.com/coinse/SINVAD/blob/master/vae/train.py>. Accessed: August 1, 2023.
- [39] PyTorch, DCGAN example in pytorch, 2025. <https://github.com/pytorch/examples/blob/main/dcgan/main.py>.
- [40] Zokovie, cDCGAN-MNIST, 2025, <https://github.com/zokovi/cDCGAN-MNIST>. Accessed: September 3, 2023.
- [41] B. Maltais, Kohya SS - LoRA GUI code, 2025, https://github.com/bmaltais/kohya_ss/blob/master/kohya_gui/lora_gui.py. Accessed: November 30, 2023.
- [42] S. Kang, R. Feldt, S. Yoo, VAE convolutional training code, 2023, https://github.com/coinse/SINVAD/blob/master/vae/train_conv.py. Accessed: September 29, 2023.
- [43] M. Bhandari, Training VAE on ImageNet with PyTorch, 2024, <https://www.kaggle.com/code/maunish/training-vae-on-imagenet-pytorch>. Accessed: April 2024.
- [44] A. Brock, J. Donahue, K. Simonyan, Large scale GAN training for high fidelity natural image synthesis, in: *International Conference on Learning Representations*, 2019.

- [45] V. Riccio, N. Humbatova, G. Jahangirova, P. Tonella, DEEPMETIS: augmenting a deep learning test set to increase its mutation score, in: 2021 36th IEEE/ACM International Conference on Automated Software Engineering (ASE), IEEE, 2021, pp. 355–367.
- [46] A. Abbasishahkoo, M. Dadkhah, L. Briand, D. Lin, TEASMA: a practical methodology for test adequacy assessment of deep neural networks, IEEE Trans. Software Eng. 50 (12), IEEE Computer Society, Los Alamitos, CA, USA (2024) 3307–3329 <https://doi.org/10.1109/TSE.2024.3482984>.
- [47] V. Riccio, P. Tonella, Model-based exploration of the frontier of behaviours for deep learning system testing, in: Proceedings of the 28th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering, Association for Computing Machinery, New York, NY, USA, 2020, pp. 876–888. <https://doi.org/10.1145/3368089.3409730>
- [48] T. Zohdinasab, V. Riccio, A. Gambi, P. Tonella, DeepHyperion: exploring the feature space of deep learning-based systems through illumination search, in: Proceedings of the 30th ACM SIGSOFT International Symposium on Software Testing and Analysis, 2021, pp. 79–90.
- [49] Z. Aghababaeian, M. Abdellatif, L. Briand, M. Bagherzadeh, et al., Black-box testing of deep neural networks through test case diversity, IEEE Trans. Software Eng. 49 (5) (2023) 3182–3204.