

Evading Deep Learning-Based Malware Detectors via Obfuscation: A Deep Reinforcement Learning Approach

Brian Etter*

*Department of Management Information Systems
University of Arizona
Tucson, USA
etterb@arizona.edu*

James Lee Hu*

*Department of Management Information Systems
University of Arizona
Tucson, USA
jameshu@arizona.edu*

Mohammadreza Ebrahimi

*School of Information Systems and Management
University of South Florida
Tampa, USA
ebrahimim@usf.edu*

Weifeng Li

*Department of Management Information Systems
University of Georgia
Athens, USA
weifeng.li@uga.edu*

Xin Li

*Department of Computer Science
University of Arizona
Tucson, USA
xinli2@arizona.edu*

Hsinchun Chen

*Department of Management Information Systems
University of Arizona
Tucson, USA
hsinchun@arizona.edu*

Abstract—Adversarial Malware Generation (AMG), the generation of adversarial malware variants to strengthen Deep Learning (DL)-based malware detectors has emerged as a crucial tool in the development of proactive cyberdefense. However, the majority of extant works offer subtle perturbations or additions to executable files and do not explore full-file obfuscation. In this study, we show that an open-source encryption tool coupled with a Reinforcement Learning (RL) framework can successfully obfuscate malware to evade state-of-the-art malware detection engines and outperform techniques that use advanced modification methods. Our results show that the proposed method improves the evasion rate from 27%-49% compared to widely-used state-of-the-art reinforcement learning-based methods.

Index Terms—Adversarial Robustness, Reinforcement Learning, Adversarial Malware Variants, Adversarial Malware Generation, Obfuscation.

I. INTRODUCTION

Detection and identification of malware is one of the top priorities in cybersecurity. However, as malware authors write more sophisticated malware, existing malware detectors become less efficient [11]. At a high level, malware detectors typically fall into one of two categories: traditional signature-based detectors and, more recently, detectors incorporating Deep Learning (DL) architectures [11], [7]. While showing

promise in the early detection of new malware variants, DL-based malware detectors are vulnerable to adversarial attacks, which involve the deliberate insertion of benign code alongside the malicious code in order to fool the detector into misclassifying the malware as a benign file, thus evading detection [9]. It has been shown that adversarially generated malware samples that successfully evade detection can be leveraged to retrain and enhance, or robustify, the detector, improving subsequent performance against future attacks [12]. To this end, Adversarial Malware Generation (AMG) aims to robustify malware detectors through automated generation of, and training on, crafted adversarial malware variants [3].

The majority of current AMG studies focus on small additive or editing actions with very few exploring full-file obfuscation [13]. We expect that these small perturbations ultimately decrease the covert nature of adversarial variants, with the full potential of AMG unrealized. Thus, changing the makeup of the entire malware file through full-file obfuscation has the potential to contribute to the AMG field. Additionally, this would better reflect the practices of real-world hackers, who often obfuscate an entire file through the use of powerful editing actions [14], [15].

Compression and encryption are techniques used by real-world hackers and can be used in generating evasive variants [16]. A multitude of tools and methods can apply such actions in the obfuscation of malware [2], generating a large action space. However, considering all possible combinations of these

*: Corresponding author

Acknowledgments: This material is based upon work supported by the National Science Foundation (NSF) under the Secure and Trustworthy Cyberspace (1936370) program.

TABLE I
SELECTED RECENT SIGNIFICANT AMG STUDIES

Year	Author(s)	Detector Type	Focus	Attack Method	Modification
2022	Song et al. [1]	DL-based	Attack	Reinforcement Learning	Additive
2021	Javaheri et al. [2]	Signature-based	Defense	Genetic Algorithm	Editing
2021	Ebrahimi et al. [3]	Signature-based	Attack	Deep Language Modeling	Editing
2020	Demetrio et al. [4]	DL-based	Attack	Genetic Algorithm	Additive
2019	Park et al. [5]	DL-based	Attack	Dynamic Programming	Additive
2019	Castro et al. [6]	Signature-based	Attack	Random Perturbations	Additive
2019	Suciu et al. [7]	DL-based	Attack	Benign Sequence Append	Additive
2019	Rosenberg et al. [8]	DL-based	Defense	GAN	Additive
2018	Anderson et al. [9]	Signature-based	Attack	Deep Reinforcement Learning	Editing
2018	Dey et al. [10]	Signature-based	Attack	Genetic Programming	Additive

Note: GAN: Generative Adversarial Network

tools, methods, and their parameters leads to a larger combinatorial action space. It is worth noting that not all modifications necessarily lead to an increase in evasive capabilities [13]. To address this, Reinforcement Learning (RL) provides a promising framework to search for effective action sequences, or strings of modifying actions, that result in evasive malware variants [17]. In this study, we aim to present a deep RL-based framework coupled with prevailing open-source obfuscation tools for conducting effective adversarial malware generation.

The rest of this paper is organized as follows. First, we conduct a literature review of adversarial malware generation, obfuscation, and reinforcement learning. We then present the details of the components of the proposed framework. Next, we compare the proposed method against extant state-of-the-art AMG methods. Lastly, we highlight several promising future directions.

II. LITERATURE REVIEW

In this work, we review four areas of the literature. First, the AMG landscape is surveyed and used as a foundation for the study. Second, we explore obfuscation methodologies and tools with an emphasis on open-source availability and capacity to automate. Third, RL is reviewed as a means to automate the generation of adversarial samples and identify evasive variants. Fourth, we explore Deep Q-Networks (DQN) and their role in model effectiveness and efficiency.

A. Adversarial Malware Generation

Malware detectors, by and large, fall into one of two categories, signature-based or DL-based. Signature-based detectors utilize libraries containing code known to be malicious and look for these sequences of code in the target file [9], [3], [2], [6], [10]; whereas DL-based detectors detect malicious files based on a learned representation of bytes [11], [7], [1], [4], [5], [8]. Despite their ability in identifying new variants, DL-based detectors tend to be more vulnerable to adversarial variants.

Adversarial Malware Generation (AMG) is the practice of adversarially modifying malware samples with the goal of being misclassified as benign by DL-based malware detectors. The ultimate goal of AMG studies is to strengthen detectors by using the adversarially crafted samples that evade detection

to retrain the model, thus robustifying the detector against future adversarial attacks [9]. We summarize several significant recent works based on the detector type, focus (attack or defense), attack method, and modification strategy in Table I.

The AMG methods generate adversarial variants through modifying the original malware files. Such adversarial modifications fall into two categories of actions: additive or editing. Additive actions are predominantly the appending of benign bytes, or sequences, to a malicious file [7], [1], [10], [8]. Random perturbations [6], insertion of benign bytes [4], [10], and insertion of dummy code [5] are also common; the latter tend to utilize more advanced DL-based models such as Generative Adversarial Networks (GAN), or Genetic or Dynamic Programming to facilitate the injection of benign bytes [2], [4], [5], [10]. The editing actions utilized rarely encompass more than a few sections of the malware file, and those that do are limited to file compression, or packing, through UPX [9], [3]. Due to the prevalence of UPX's use in malware obfuscation however, most detectors see the act of compression/packing by UPX as malicious without scrutinizing the actual file, resulting in higher occurrences of false positives [4], [18].

The key observations from Table I are that there is a prevalence of additive actions [7], [1]- [8] and a sparsity of editing actions [9], [3]. In particular, the additive actions are often applied using methods such as Genetic Programming (GP), Dynamic Programming (DP), Generative Adversarial Networks (GANs), etc., and may only be possible to implement by seasoned hackers or software engineers [2], [4], [5], [10], [8]. In these studies, the majority of actions are designed to only modify small parts of the file or append a small number of bytes to avoid breaking its functionality, leaving the majority of the file's makeup unchanged. However, hackers tend to use obfuscation techniques such as encryption, packing, or encoding to conceal the entirety of a file [19]. The use of open-source software to obfuscate the entire malware file may help to advance AMG techniques by leveraging readily available tools used in real-world scenarios. Thus, the inclusion of this approach may afford enhanced robustification of DL-based detectors [20].

TABLE II
SELECTED MAJOR OBFUSCATION TOOLS FROM LITERATURE

Reference	OS / Commercial	API / CLI	Modification	Tool
Aghakhani et al. [18]	Commercial	No	Compression	Obsidium, PECompact, PELock
			Encryption	Themida
	Open-Source	Yes	Compression	Petite, UPX
Anderson et al. [9]	Open-Source	Yes	Compression	UPX
Bergenholtz et al. [21]	Commercial	No	Compression	Obsidium, PECompact, PELock
			Encryption	Themida
	Open-Source	Yes	Compression	Petite, UPX
Halls [22]	Open-Source	Yes	Encryption	Darkarmour
Nichol [23]	Open-Source	Yes	Encryption	Gobfuscate
Metasploit Project [24]	Open-Source	Yes	Multiple	MSFVenom

B. Obfuscation

The functionality of a portable executable (PE) file is extremely sensitive to changes in the code, with the change of one byte being enough to render the malware file corrupt [1]. Because of this, the majority of editing actions, with the exception of compression by UPX, are designed to preserve the overall makeup of the malware file [9]. However, given the prevalence of malware file encryption in the real-world [14], [15], we expect that adding this approach will extend the potential of AMG research to better emulate real-world scenarios. Thus, through the inclusion of readily available open-source tools that are used in real-world scenarios, DL-based detectors may experience heightened levels of robustification and be less susceptible to adversarial attacks [20]. Additionally, this approach would better reflect the actions of real-world hackers [14], [15]. Tools and methods found in the current research encompass a small portion of the spectrum for obfuscation tools [9], [18], [21]. Resources and subject matter experts (SMEs) outside of academic research also provide valuable insights about such commonly used tools and methods [14], [15], [20].

The literature largely documents tools that fall into two categories, open-source and commercial, with open-source tools typically found on GitHub, and commercial tools obtained directly from the developer. The ability to automate the obfuscation process through scripting and the command line interface (CLI) requires access to an application programming interface (API). Most tools found in extant literature utilize compression, encryption, or a combination of these two methods to obfuscate malware. Table II summarizes major open-source and commercial grade tools used for obfuscation in literature or by the hacker community.

Three key points can be drawn from Table II. First, within the commercial tools their API is not readily available with the main functionality locked behind a paywall or graphical user interface (GUI). The API call is of particular importance for automation purposes (interaction with the RL agent). Second, extant AMG research primarily documents tools utilized by the academic and/or research communities; with very few references to tools used by the hacker community outside of UPX (packing/compression) [11], [9]. There are several non-AMG studies done on packing tools and their ability to increase evasion of malware detectors ranging from commercial grade

(paid or freemium) to open-source [18], [21]. However, due to the prevalence of packing in malware obfuscation, as noted previously, most detectors see the act of compression/packing, such as the use of the UPX packing tool, as malicious without scrutinizing the actual file. This is because UPX leaves specific artifacts in the code of a compressed PE file, which typically cause it to be flagged as potentially malicious [4], [18], resulting in higher occurrences of false positives.

```
usage: darkarmour.py [-h] [-f FILE] [-e ENCRYPT] [-S SHELLCODE] [-b]
                  [-d] [-u] [-j] [-r] [-s] [-k KEY] [-l LOOP]
                  [-o OUTFILE]

optional arguments:
  -h, --help            show this help message and exit
  -f FILE, --file FILE  file to crypt, assumed as binary if not
                        told otherwise
  -e ENCRYPT, --encrypt ENCRYPT
                        encryption algorithm to use (xor)
  -S SHELLCODE, --shellcode SHELLCODE
                        file containing the shellcode, needs to be
                        in the 'msfvenom -f raw' style format
  -b, --binary          provide if file is a binary exe
  -d, --dll            use reflective dll injection to execute
                        the binary inside another process
  -u, --upx            pack the executable with upx
  -j, --jmp            use jmp based pe loader
  -r, --runpe          use runpe to load pe
  -s, --source          provide if the file is c source code
  -k KEY, --key KEY    key to encrypt with, randomly generated if
                        not supplied
  -l LOOP, --loop LOOP number of levels of encryption
  -o OUTFILE, --outfile OUTFILE
                        name of outfile, if not provided then
                        random filename is assigned

kali@kali:~/github/darkarmour$ python3 darkarmour.py -f /home/kali/
adversarialai/Ransomware/fe003e8b216d8ea480351ae040ae99d142945bed62
d4ccdbaf94e1cfa2b6cb13 -e xor -j -b -o /home/kali/adversarialai/Tes
tOut/notmalware -l 2
```

Fig. 1. Interface of Darkarmour

Lastly, in our research, we considered recommendations from a Red Team SME who endorsed three tools: Darkarmour [22], Gobfuscate [23], and MSFVenom [24]. After a thorough evaluation, we opted for Darkarmour as the most suitable choice. This decision is underpinned by Darkarmour's simplicity of implementation, open-source accessibility, and availability through the CLI with an associated API. While Gobfuscate and MSFVenom received commendation from the SME, practical challenges, such as configuration and environment issues with Gobfuscate and misalignment with our research requirements in the case of MSFVenom as well as, led us to select Darkarmour as our preferred tool. Additionally,

the Petite tool was excluded from consideration due to its absence of recent updates or active development, which did not align with our criteria. Accessing a tool's API via the CLI is crucial for automating and seamlessly integrating it into various workflows. This capability is demonstrated in Figure 1 through Darkarmour's CLI, featuring a rich set of parameters suitable for automation. With a large number of available tools and parameters, RL has emerged as an ideal medium for navigating this action space and identifying optimal combinations of actions in the generation of evasive variants.

C. Deep Reinforcement Learning

Reinforcement Learning (RL) provides a useful framework to conduct goal-directed learning or optimization [25]. The goal of RL is to find the best action sequence, from a given set of actions, for a given input and for the agent to learn a (near) optimal policy; in this case the best set of modifications (actions) for a given malware file [10].

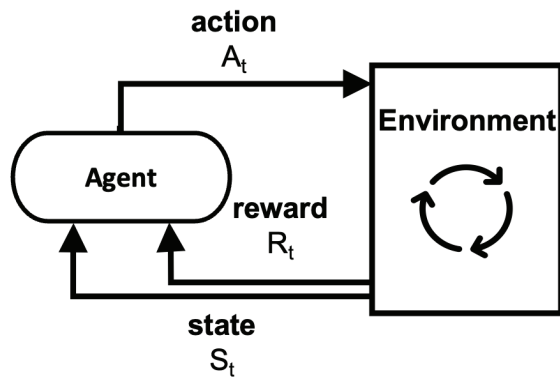


Fig. 2. RL agent-environment interaction

Recently, RL has been shown to be capable of systematically choosing action sequences to conduct AMG [9], [3], [17]. Figure 2 provides a conceptual overview of the RL at time t , where the agent applies a set of actions A_t to the environment generating the state S_t and reward R_t pair which are then used by the agent to determine the next action or set of actions to apply, or whether the optimal configuration has been reached. The agent-environment interaction is composed of the following components:

- **Agent:** Selects actions that maximize future rewards.
- **Environment:** Generates the states and reward based on the actions taken by the agent at each step.
- **State:** Current configuration of the agent in the environment.
- **Actions:** A set of options, or moves, available to the agent.
- **Reward:** A positive, negative, or zero value returned after each action is applied.

In the AMG scenarios, the primary objective is to develop techniques that can robustify malware detectors against adversarial attacks [9], [3], [13], [4]. This requires the generation of evasive malware samples that can avoid detection

by the detectors. However, attackers need to minimize their interactions with the environment (detector) to avoid detection and possible countermeasures by the defender [8]. Accurately emulating the defender, for instance, by implementing a query limit, is crucial for the generation of evasive adversarial variants. RL provides an effective approach to operate under such limitations. To this end, many RL approaches adopt Deep-Q Networks (DQN), which have been shown to be highly efficient and thus require fewer interactions with the environment [9], [3], [17].

Deep Reinforcement Learning (DRL) is the combination of RL and deep neural networks (DNN), to learn policies for decision-making in complex scenarios [26], [27]. Deep Q-Network (DQN), a DRL algorithm, was developed by Mnih et al., 2013, and has shown exceptional success in a wide range of applications including video games [26], [28], malware evasion [9], [3], [17], and penetration testing [27]. DQN uses a multi-layered deep neural network to estimate the optimal action value, or Q-function $Q(s, a)$ for a given state [9], [28].

Traditional RL methods utilize a tabular Q-learning method, consisting of a lookup table of historical actions, but such methods often experience performance issues with large state/search spaces [27], [29]. DQN attempts to solve this issue by replacing the lookup table with a DNN, in particular a deep convolutional neural network [30]. The implementation of DQN by Mnih et al. outperformed top benchmark RL methods in the Atari domain and performed at a level comparable to professional human gamers [30]. DQN has demonstrated exceptional success in other applications including video games, malware detection, and penetration testing [9], [27], [30]. In the AMG applications, thanks to its sample efficiency, we expect that DQN is a viable choice to decrease the frequency of interactions with the malware detector, minimizing the risk of countermeasures implemented by the defender [9], [17].

III. RESEARCH GAPS AND QUESTIONS

Based on our literature review, the following research gaps have been identified. First, most current AMG attack methods focus on small additive or editing actions, not complete obfuscation, potentially limiting the ability to robustify DL-based malware detectors. Second, it is unclear from a methodological perspective if RL can be used in conjunction with obfuscation tools to modify malware resulting in functional adversarial variants. To address the identified gaps, the following research question is proposed:

- How can RL, in combination with obfuscation, be used to obfuscate malware and generate evasive functional variants?

Motivated by this question, we propose OBFU-mal, a deep RL-based framework to automate the obfuscation of malware for adversarial variant generation.

IV. PROPOSED METHOD

We first introduce the threat model under which OBFU-mal operates in line with previous AMG studies [3], [1], [31], [32]. Then, we present our proposed OBFU-mal architecture

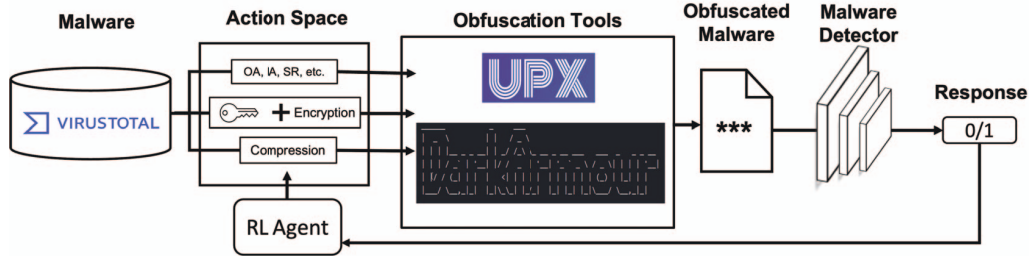


Fig. 3. The OBFU-mal architecture

(Figure 3). Third, we detail the malware testbed and extended action space containing obfuscating actions. Finally, we present the malware detectors used in OBFU-mal’s evaluation.

A. Threat Model

Following [1], [31], [32], our threat model is composed of three main components: adversarial goal, adversarial capability, and adversarial knowledge:

- **Adversarial Goal:** The adversary aims to evade DL-based malware detectors by modifying a known malicious file such that it is recognized as benign.
- **Adversarial Capability:** The adversary is capable of using open-source tools to obfuscate the whole malware file (through encryption). These obfuscation tools can also be utilized by novice hackers with little experience.
- **Adversarial Knowledge:** The adversary operates in a black-box setting [9], in which no knowledge of the malware detection model is assumed. The only available information to the adversary is whether the modified malware was detected or not.

B. The OBFU-mal Architecture

OBFU-mal’s architecture closely emulates the techniques employed by real-world hackers to develop advanced malware. The OpenAI-based gym-malware environment, developed by Anderson et al. [9], serves as the foundation for OBFU-mal, providing a testing ground for generating adversarial samples. Real-world hackers adapt their strategies to evade detection, making this an ideal environment for testing the robustness of malware detectors. OBFU-mal employs an RL agent with a DQN algorithm, mirroring real-world hackers’ trial-and-error approach to refining malware. This agent utilizes an extended action space replicating the tactics of malicious actors and incorporates open-source tools for modifying malware samples. This process is illustrated in Figure 3, and proceeds as follows. First, the Malware sample is introduced to the RL agent. Then, the RL agent applies actions from the extended action space summarized in Table III. Next, the modified malware sample is assessed by the detector. The feedback from the detector indicates detection or evasion. The process iterates as the next sample is introduced to the agent, which applies actions based on DQN feedback, and the detector evaluates the adversarially generated malware sample.

This approach closely mimics the diverse set of techniques employed by real-world hackers to create obfuscated and adversarial variants. The iterative process, where the detector

evaluates the adversarially generated malware samples, reflects the persistent nature of real-world hackers, who continually adjust their techniques based on feedback. As such, OBFU-mal provides a more accurate representation of how real-world hackers try and attack malware detectors.

C. Action Space

The extended action space, shown in Table III, details the actions used in select recent AMG studies and includes those added in this study [9], [3], [17], [1].

TABLE III
EXPANDED ACTION SPACE IN OBFU-MAL

Action	Modification	Description
Overlay Append	Additive	Append bytes to the end of malware exe
Imports Append	Additive	Add an entry to the import table
Section Rename	Edit	Changes section’s name in malware exe
Remove Signature	Edit	Unlink digital signature from certification table
Remove Debug	Additive	Unlink debug section from header
Section Append	Additive	Add a new section to the malware exe
Break Checksum	Edit	Set file’s checksum
Change Timestamp	Edit	Change / set timestamp
UPX Pack (Compress)	Obfuscate	Compress malware exe
Darkarmour XOR EL1	Obfuscate	Apply one XOR encryption loop
Darkarmour XOR EL2	Obfuscate	Apply two XOR encryption loops
Darkarmour XOR EL3	Obfuscate	Apply three XOR encryption loops

The actions used in extant research only impact small portions of the file with the exception of UPX which operates on the whole file. The inclusion of UPX compression, which affects the entire file, mirrors its common use in both legitimate and malicious software. However, malware compressed by UPX tends to be categorized as malicious by most malware detectors [18], leading to a decrease in evasion rate. OBFU-mal introduces the actions ”Darkarmour XOR EL1/2/3” that utilize one or more loops of eXclusive OR (XOR) encryption, a boolean operator used in cryptography to obfuscate the malware file. The XOR loops are applied by the open-source tool Darkarmour which modifies the malware in a way such that the execution does not require bytes touching the disk

TABLE IV
STATE-OF-THE-ART BENCHMARK AMG METHODS

Method Category	Method Selected	Description	Reference
RL	MAB	Append to file and edit sections	Song et al. [1]
RL	AMG-VAC	Append to file and edit sections	Ebrahimi et al. [3]
CLM	MalGPT	Appends benign, file-specific perturbations	Hu et al. [32]
GA	GAMMA	Append to file, file-specific perturbations	Demetrio et al. [4]
Feature Append	Benign Append	Appends benign bytes to end of malware file	Castro et al. [6]
Feature Append	Enhanced BFA	Strategically append bytes to file; intent to lower conf. score.	Chen et al. [33]
Feature Append	Random Append	Bytes randomly appended to malware file	Suciu et al. [7]
RL	ACER	Append to file and edit sections	Anderson et al. [9]
RL	DDQN	Append to file and edit sections	Hasselt et al. [28]

[22]. XOR encryption adds intricacy to the code, making it more challenging for detectors to recognize patterns in the malware.

V. EVALUATION

A. Experimental Design

To test the utility and effectiveness of OBFU-mal, we conducted three experiments. First, the evasiveness of OFBU-mal-generated variants was tested against two well-known malware detectors: MalConv and LightGBM (LGBM)/EMBER. Results are compared against a total of nine state-of-the-art benchmark AMG tactics detailed in Table IV. Attack method category includes RL, Genetic Algorithm (GA), Causal Language Model (CLM), and feature append; attack methods include Multi-Armed Bandit (MAB) and AMG-Variational Actor Critic (VAC). For RL-based methods, to ensure a fair comparison, a five-query limit was implemented during attacks on the detector. Moreover, in our second experiment, a qualitative analysis of the action sequences that generated evasive variants was performed, and results were contrasted against related work [3]. In our third experiment, an ablation analysis was conducted for OBFU-mal to assess the contribution of each component of the framework.

A repository of 3,456 malware samples from VirusTotal was used during the experiment, detailed in Table VI. These malware samples were modified using OBFU-mal, and evaluated against MalConv and LGBM/EMBER, two well-known malware detectors. MalConv, developed by the Laboratory for

Physical Sciences, is a DL-based detector utilizing a deep convolutional neural network (CNN) which takes the first 2M bytes of a sample as its input [11]. LGBM/EMBER, developed by Endgame, is a signature-based detector which utilizes a gradient-boosted decision tree (GBDT) and is trained on 2,381 features extracted from the malware binary [9].

TABLE VI
COMPOSITION OF MALWARE TESTBED

Malware Category	Description	# of Files
Botnet	A network of bots connected through the internet	526
Ransomware	Encrypts data and files, restricting access and usage until decrypted	900
Rootkit	Grants admin rights to malware authors	731
Spyware	Allows malware authors to covertly steal personal data	640
Virus	Corrupts files on the host system	659
Total	-	3,456

Performance is measured using evasion rate, which is consistent with prior literature [1], [3], [9]. Specifically, evasion rate, E , is defined in the following equation:

$$E = M_e / M_t$$

Where M_e denotes the number of samples that evaded detection, and M_t denotes the total number of generated adversarial samples.

TABLE V
EVASION RATE RESULT OF OBFU-MAL SAMPLES AGAINST MALCONV BENCHMARK AMG TACTICS

Detector	Method	Botnet	Ransomware	Rootkit	Spyware	Virus	Average
MalConv	ACER	37.07%	25.33%	29.41%	56.09%	44.76%	35.04%
	Random Append	2.47%	3.78%	4.73%	2.50%	1.88%	2.79%
	BFA	1.14%	0.11%	3.77%	1.88%	2.43%	1.27%
	Enhanced-BFA	21.86%	14.44%	3.77%	11.50%	12.29%	15.86%
	GAMMA	3.51%	1.79%	0%	1.86%	1.35%	1.70%
	MalGPT	25.86%	20.33%	24.53%	22.97%	28.38%	24.51%
	MAB-malware	52.93%	48.38%	45.23%	50.99%	55.96%	50.70%
	OBFU-mal (Ours)	69.58%	56.89%	57.80%	66.04%	75.42%	65.15%
LGBM/EMBER	DDQN	23.00%	44.33%	39.12%	19.80%	27.77%	28.44%
	ACER	30.99%	60.11%	27.51%	26.87%	62.82%	37.18%
	Enhanced-BFA	3.02%	4.44%	4.73%	5.34%	6.16%	3.90%
	AMG-VAC	48.29%	65.22%	61.15%	29.53%	82.40%	51.67%
	MAB-malware	4.21%	11.25%	4.89%	10.98%	24.13%	11.09%
	OBFU-mal (Ours)	67.32%	92.44%	81.48%	88.07%	63.46%	79.20%

TABLE VII
COMMONLY OCCURRING ACTION SEQUENCES APPLIED BY OBFU-MAL

Method	Evasive Action Sequences	# of Occurrences
OBFU-mal VS MalConv	Change TDS → Overlay Append	94
	Overlay Append → Overlay Append	53
	Overlay Append → Darkarmour XOR EL2	38
	Change TDS → Darkarmour XOR EL2	38
	Change TDS → Overlay Append → Darkarmour XOR EL2	30
OBFU-mal VS LGBM/EMBER	Overlay Append → Darkarmour XOR EL2	106
	Change TDS → Darkarmour XOR EL2	69
	Change TDS → Overlay Append → Darkarmour XOR EL2	58
	Imports Append → Darkarmour XOR EL1	57
	Overlay Append → Darkarmour XOR EL3	52

B. Experiment Results

In the first experiment, the evasiveness of OBFU-mal-generated samples was compared against samples generated by the identified benchmark AMG methods. Table V summarizes the results of OBFU-mal against benchmark methods attacking the MalConv and LGBM/EMBER malware detectors.

In table V we detail the evasion rates obtained by OBFU-mal generated samples. Against the MalConv detector OBFU-mal achieved an overall evasion rate of 65.15% which is a percentage-point increase of 14.45%, 30.11%, and 40.64% over MAB-malware, ACER, and MalGPT respectively. Samples interrogated by the LGBM/EMBER detector achieved an overall evasion rate of 79.20%, a percentage-point increase of 27.53%, 42.02%, and 68.11% against AMG-VAC, Enhanced-BFA, and MAB-malware respectively.

Within the table, we also observed variations in OBFU-mal's performance across different malware types. Notable examples include an increased evasion rate of virus files compared to a decreased evasion rate of ransomware files when attacking MalConv. While outside the scope of this paper, we can offer an educated guess regarding the cause of this performance variance: file complexity. Virus files are typically simple scripts that exploit known vulnerabilities. As such, the binary is easily and massively changed with OBFU-mal's obfuscating actions, rendering the file evasive. Meanwhile, malware such as ransomware contains complex sequences that cannot easily be obfuscated, thus being caught by MalConv's byte-level scrutiny. Interestingly, we find the exact opposite phenomenon when attacking LGBM/EMBER. This could be due to the feature-based nature of LGBM/EMBER, making it more attentive to minute alterations in distinctive attributes of the malware, such as the encryption key of ransomware, while being less sensitive to generic signs of obfuscation used by most software products, signs which may be more represented when obfuscating a simpler virus file. However, these are all conjectures, and further research into why these variations exist is exciting but outside the scope of this paper. Regardless

of these variations, OBFU-mal outperformed all benchmark AMG methods across all categories, with the exception of viruses in the case of LGBM/EMBER against the AMG-VAC method, in our research testbed.

In the second experiment, a qualitative analysis of action sequences was conducted to identify sequences that generated the most evasive malware variants. Table VII presents the most commonly occurring sequences used in the evasion of MalConv and LGBM/EMBER as applied by OBFU-mal. Obfuscation actions Darkarmour XOR EL1, EL2, and EL3 were observed in the majority of evasive action sequences. This appears to indicate that the newly introduced obfuscation actions are effective in the generation of evasive variants; the aforementioned actions largely appear as the final action in the sequences they were included in.

Lastly, in the third experiment, the impact of including obfuscation in the OBFU-mal framework is assessed through an ablation study. The individual contributions of obfuscation alone [22], RL by itself [9], and the combined impact of obfuscation applied alongside the RL actions, developed by Anderson et al. [9], against MalConv were evaluated. Table VIII illustrates the contribution of each component assessed in the experiment. RL outperformed obfuscation in only one malware category, spyware; however, the combination of obfuscation and RL outperforms both individual methods. Through the integration of obfuscation into the RL framework enabled by OBFU-mal, we see an average of 17.82% increase in evasion rates.

Overall, the experiments show that the inclusion of obfuscation actions, as applied by OBFU-mal, significantly improves the evasion of malware against malware detectors, not limited to DL-based, and outperforms benchmark methods. Additionally, the utilization of relatively simple open-source software is sufficient to evade most state-of-the-art detectors, and access to advanced ML- or DL-based methods may not be necessary to create evasive variants.

TABLE VIII
ABLATION STUDY OF OBFU-MAL COMPONENTS

Method	Botnet	Ransomware	Rootkit	Spyware	Virus	Average
RL without obfuscation [9]	37.07%	25.33%	29.41%	56.09%	44.76%	35.04%
Obfuscation alone [22]	53.61%	42.44%	45.96%	40.63%	56.45%	46.5%
RL with obfuscation (Ours)	69.58%	56.89%	57.80%	66.04%	75.42%	64.32%

VI. CASE STUDY

The majority of extant AMG research is focused on evading open-source malware detectors such as MalConv and LGBM/EMBER, with commercial and pseudo-commercial detectors understudied. One such detector is ClamAV (<https://www.clamav.net>), a pseudo-commercial open-source malware detector that also provides endpoint security. Through the previously documented experiments, the results suggest that attackers can create adversarial variants against open-source DL-based malware detectors, and in turn, those samples are likely to be evasive against ClamAV (i.e., high-quality adversarial variants are transferable). From this, we can test the efficacy of OBFU-mal in a real-world scenario by testing adversarially crafted samples against ClamAV.

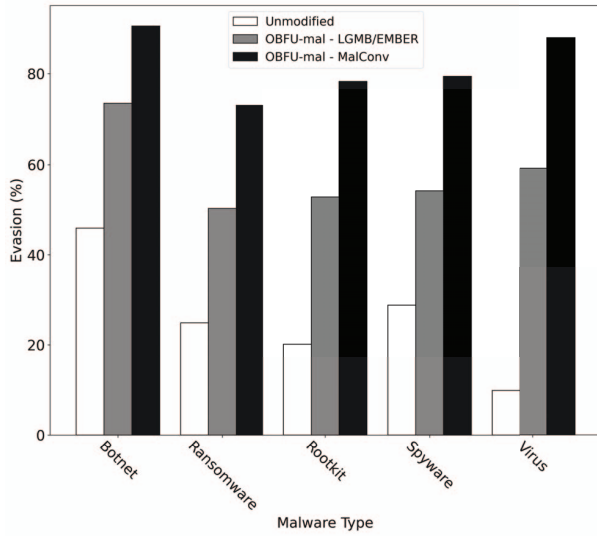


Fig. 4. Evasion rates against ClamAV by OBFU-mal

Figure 4 highlights OBFU-mal’s performance against ClamAV showing evasion rates for unmodified malware samples and malware modified against LGBM/EMBER and MalConv. Samples crafted against LGBM/EMBER and MalConv evade ClamAV 86.44% and 69.67% of the time respectively. These results show the transferability of the malware samples generated by OBFU-mal, as samples generated against both LGBM/EMBER and MalConv were still able to evade ClamAV at a higher rate than their unmodified counterparts despite the lack of knowledge of ClamAV during the sample generation process.

VII. IMPLICATIONS

Given OBFU-mal’s performance against established malware detectors, concerns about its ethics will naturally arise. However, previous adversarial literature has shown that a better understanding of an adversary will yield better defense against such adversaries [34]. Given that OBFU-mal seeks to model real-world malware adversaries and their obfuscation methods, we surmise that defenses built upon OBFU-mal’s generated adversarial malware samples will be better equipped against real-world hackers.

We can offer a potential way such a defense may be constructed: adversarial retraining. This refers to finetuning the initial detector (MalConv or LGBM for this study) on evasive adversarial malware samples (i.e., those assigned a benign label by the original detector) and their correct, malicious labels [35]. With this in mind, it becomes clear that adversarially retraining a detector with OBFU-mal samples forces a detector to learn the obfuscating tactics of real-world hackers, thus making them more robust in real-world attacks. While such a defense method could be promising, OBFU-mal is a method focused on offensive AMG and generating evasive adversarial malware samples. Thus, adversarial training and its effect on malware detectors are also outside the scope of this paper.

VIII. CONCLUSION AND FUTURE WORK

AMG studies provide valuable insight into improving anti-malware engines against the application of targeted modifications to malware. However, the vast majority of extant research mostly outlines subtle additive and editing techniques, real world tactics and methods such as obfuscation, encryption and utilization of open source tools are rarely reviewed.

Through the inclusion of obfuscation applied by open-source tools in AMG studies, there may be potential to open many additional avenues for the robustification of DL-based detectors. This work shows that an RL agent can be coupled with open-source obfuscation tools and generate evasive variants capable of evading DL- and signature-based detectors. Automation through RL allows for more efficient exploration of the large action space created when considering open-source tools.

Our proposed OBFU-mal framework, successfully outperformed the vast majority of benchmark methods through obfuscation. This shows that open source tools, requiring little to no advanced training, are sufficient in bypassing state-of-the-art detectors. Moreover, we were able to evade a pseudo-commercial antivirus tool, ClamAV, successfully demonstrating the generalizability of our method.

Promising future directions include using sophisticated source-code obfuscation methods enabled by reverse engineering malware binaries to enhance the AMG capabilities and the adversarial robustness of deep learning-based malware detectors. Another direction is the aforementioned exploration of OBFU-mal’s performance variation over certain malware categories, which may yield insight into the strengths and weaknesses of certain malware detectors. Lastly, OBFU-mal’s contribution to defensive operations is another exciting yet unexplored direction.

ACKNOWLEDGMENT

We would like to sincerely thank VirusTotal for providing the malware dataset and granting access to the corresponding APIs for functionality assessment.

This material is based upon work supported by the National Science Foundation (NSF) Secure and Trustworthy Cyberspace program (grant No. 1936370).

REFERENCES

- [1] W. Song, X. Li, S. Afroz, D. Garg, D. Kuznetsov, and H. Yin, "Mab-malware: A reinforcement learning framework for blackbox generation of adversarial malware," in *Proceedings of the 2022 ACM on Asia Conference on Computer and Communications Security*, p. 990–1003.
- [2] D. Javaheri, P. Lalbakhsh, and M. Hosseinzadeh, "A novel method for detecting future generations of targeted and metamorphic malware based on genetic algorithm," *IEEE Access*, vol. 9, p. 69951–69970.
- [3] M. Ebrahimi, J. Pacheco, W. Li, J. Hu, and H. Chen, "Binary black-box attacks against static malware detectors with reinforcement learning in discrete action spaces," in *2021 IEEE Security and Privacy Workshops (SPW)*, p. 85–91.
- [4] L. Demetrio, B. Biggio, G. Lagorio, F. Roli, and A. Armando, "Functionality-preserving black-box optimization of adversarial windows malware," *IEEE Transactions on Information Forensics and Security*, vol. 16, p. 3469–3478.
- [5] D. Park, H. Khan, and B. Yener, "Generation evaluation of adversarial examples for malware obfuscation," in *Proceedings - 18th IEEE International Conference on Machine Learning and Applications, ICMLA 2019*, p. 1283–1290. [Online]. Available: <https://doi.org/10.1109/ICMLA.2019.00210>.
- [6] R. Castro, C. Schmitt, and G. Rodosek, in *ARMED: How Automatic Malware Modifications Can Evade Static Detection? 2019 5th International Conference on Information Management (ICIM)*, p. 20–27.
- [7] O. Suci, S. Coull, and J. Johns, "Exploring adversarial examples in malware detection," in *2019 IEEE Security and Privacy Workshops (SPW)*. [Online]. Available: <https://doi.org/10.1109/spw.2019.00015>.
- [8] I. Rosenberg, A. Shabtai, Y. Elovici, and L. Rokach, "Defense methods against adversarial examples for recurrent neural networks," *ArXiv*.
- [9] H. Anderson, A. Kharkar, B. Filar, D. Evans, and P. Roth, "Learning to evade static pe machine learning malware models via reinforcement learning," *ArXiv*.
- [10] S. Dey, A. Kumar, M. Sawarkar, P. Singh, and S. Nandi, *EvadePDF: Towards Evading Machine Learning Based PDF Malware Classifiers*. ISEA-ISAP.
- [11] E. Raff, J. Barker, J. Sylvester, R. Brandon, B. Catanzaro, and C. Nicholas, "Malware detection by eating a whole exe," in *Workshops at the Thirty-Second AAAI Conference on Artificial Intelligence*.
- [12] M. Ebrahimi, N. Zhang, J. Hu, M. Raza, and H. Chen, "Binary black-box evasion attacks against deep learning-based static malware detectors with adversarial byte-level language model," *ArXiv*.
- [13] X. Ling, L. Wu, J. Zhang, Z. Qu, W. Deng, X. Chen, C. Wu, S. Ji, T. Luo, J. Wu, and Y. Wu, "Adversarial attacks against windows pe malware detection: A survey of the state-of-the-art," in *arXiv [cs.CR]*. *arXiv*. [Online]. Available: <http://arxiv.org/abs/2112.12310>.
- [14] C. Nachreiner, "How hackers hide their malware: The basics," *Dark Reading*, retrieved October 7, 2022, from. [Online]. Available: <https://www.darkreading.com/attacks-breaches/how-hackers-hide-their-malware-the-basics>.
- [15] P. Baltazar, "How and where do hackers hide malware? malwarefox," retrieved October 7, 2022, from. [Online]. Available: <https://www.malwarefox.com/how-and-where-do-hackers-hide-malware/>.
- [16] R. Badhwar, "Polymorphic and metamorphic malware. the ciso's next frontier."
- [17] Z. Fang, J. Wang, B. Li, S. Wu, Y. Zhou, and H. Huang, "Evading anti-malware engines with deep reinforcement learning," *IEEE Access*, vol. 7, p. 48867–48879.
- [18] H. Aghakhani, F. Gritti, F. Mecca, M. Lindorfer, S. Ortolani, D. Balzarotti, G. Vigna, and C. Kruegel, *When Malware is Packing' Heat: Limits of Machine Learning Classifiers Based on Static Analysis Features*. NDSS.
- [19] N. Malviya, "Simple malware obfuscation techniques. infosec resources." [Online]. Available: <https://resources.infosecinstitute.com/topic/simple-malware-obfuscation-techniques/>.
- [20] MITRE, "Obfuscated files or information, technique t1027 - enterprise — mitre attck®," retrieved November 2, 2021, from. [Online]. Available: <https://attack.mitre.org/techniques/T1027/>.
- [21] E. Bergenholtz, E. Casalicchio, D. Ilie, and A. Moss, "Detection of metamorphic malware packers using multilayered lstm networks," *Information and Communications Security*, p. 36–53.
- [22] D. Halls, "Darkarmour Computer software," <https://github.com/bats3c/darkarmour>, 2021.
- [23] A. Nichol, "gobfuscate: A go code obfuscator," <https://github.com/unixpickle/gobfuscate>, 2016, accessed: March 8, 2022.
- [24] M. Project. (2015) How to use msfvenom. <https://docs.metasploit.com/docs/using-metasploit/basics/how-to-use-msfvenom.html>. Accessed: Date.
- [25] R. Sutton and A. Barto, *Reinforcement Learning: An Introduction (Adaptive Computation and Machine Learning series*, 2nd ed. A Bradford Book.
- [26] V. Mnih, K. Kavukcuoglu, D. Silver, A. Graves, I. Antonoglou, D. Wierstra, and M. Riedmiller, "Playing atari with deep reinforcement learning," *arXiv* 2013, arXiv:1312.5602.
- [27] S. Zhou, J. Liu, D. Hou, X. Zhong, and Y. Zhang, "Autonomous penetration testing based on improved deep q-network," *Applied Sciences*, vol. 11, no. 19, p. 8823. [Online]. Available: <https://doi.org/10.3390/app11198823>.
- [28] H. Hasselt, A. Guez, and D. Silver, "Deep reinforcement learning with double q-learning," *ArXiv*.
- [29] Z. Alom, T. Taha, C. Yakopcic, S. Westberg, P. Sidike, M. Nasrin, M. Hasan, B. Essen, A. S. Awwal, and V. Asari, "A state-of-the-art survey on deep learning theory and architectures," *Electronics*, vol. 8, no. 3, p. 292.
- [30] V. Mnih, K. Kavukcuoglu, D. Silver, A. Graves, I. Antonoglou, D. Wierstra, and M. Riedmiller, "Human-level control through deep reinforcement learning," *Nature*, vol. 518, no. 7540, p. 529–533.
- [31] N. Carlini, A. Athalye, N. Papernot, W. Brendel, J. Rauber, D. Tsipras, I. Goodfellow, A. Madry, and A. Kurakin, on evaluating adversarial robustness. *arXiv preprint arXiv:1902.06705*.
- [32] J. Hu, M. Ebrahimi, and H. Chen, "Single-shot black-box adversarial attacks against malware detectors: A causal language model approach," in *2021 IEEE International Conference on Intelligence and Security Informatics (ISI)*, p. 1–6.
- [33] B. Chen, Z. Ren, C. Yu, I. Hussain, and J. Liu, "Adversarial examples for cnn-based malware detectors," *IEEE Access*, vol. 7, pp. 54 360–54 371, 2019.
- [34] W. Li and Y. Chai, "Assessing and enhancing adversarial robustness of predictive analytics: An empirically tested design framework," *Journal of Management Information Systems*, vol. 39, no. 2, pp. 542–572, 2022.
- [35] H. Rathore, A. Samavedhi, S. K. Sahay, and M. Sewak, "Robust malware detection models: learning from adversarial attacks and defenses," *Forensic Science International: Digital Investigation*, vol. 37, p. 301183, 2021.