

基于挤压激励网络的恶意代码家族检测方法*

申高宁^{1,2}, 陈志翔³, 王 辉³, 陈 姮^{1,2}

(1. 闽南师范大学 计算机学院, 福建 漳州 363000;

2. 数据科学与智能应用福建省高校重点实验室, 福建 漳州 363000;

3. 闽南师范大学 物理与信息工程学院, 福建 漳州 363000)

摘 要: 恶意代码已经成为威胁网络安全的重要因素。基于机器学习的恶意代码检测方法已经取得较好的效果, 但面对相似的恶意代码家族, 往往效果不佳。对此, 提出了一种基于挤压激励网络的检测算法, 由卷积神经网络(Convolutional Neural Network, CNN)与挤压和激励(Squeeze-and-Excitation, SE)模块构成。CNN先快速提取恶意代码的图像特征, SE模块对多通道特征图进行全局平均池化, 将全局信息压缩, 然后通过全连接层自适应学习, 并将每个通道特征图赋予不同的权重来表示不同的重要程度, 指导激励或抑制特征信息。实验结果表明, 该方法相对于传统机器学习方法有更好的检测效果, 与深度学习算法相比检测效果也有一定的提升且参数数量大大减少。

关键词: 恶意代码; 机器学习; 卷积神经网络; 挤压和激励网络

中图分类号: TP393

文献标识码: A

DOI: 10.19358/j.issn.2096-5133.2022.06.001

引用格式: 申高宁, 陈志翔, 王辉, 等. 基于挤压激励网络的恶意代码家族检测方法[J]. 信息技术与网络安全, 2022, 41(6): 1-9.

A family detection method for malicious code based on squeezed-and-excitation networks

Shen Gaoning^{1,2}, Chen Zhixiang³, Wang Hui³, Chen Heng^{1,2}

(1. School of Computer Science, Minnan Normal University, Zhangzhou 363000, China;

2. Key Laboratory of Data Science and Intelligent Applications, Zhangzhou 363000, China;

3. School of Physics and Information Engineering, Minnan Normal University, Zhangzhou 363000, China)

Abstract: Malicious code has become an important factor threatening cyber security. Machine learning-based malicious code detection methods have achieved good results, but often poorly in the face of similar malicious code families. In this paper, a detection algorithm based on extrusion excitation network was proposed, which consists of Convolutional Neural Network(CNN) and squeeze-and-excitation(SE) module. Firstly, the CNN quickly extracts the image features of the malicious code, and the SE module carries out global average pooling of multi-channel feature map to compress the global information, then learns adaptively through the full connection layer, and weights each channel feature graph to represent different degrees of importance, guiding motivating or suppressing the feature information. The experimental results show that the proposed method has a better detection effect compared with the traditional machine learning methods, and the detection effect is improved and the number of parameters is greatly reduced compared with the deep learning algorithm.

Key words: malicious code; machine learning; convolutional neural network; squeeze and excitation network

0 引言

在过去几年里随着互联网的飞速发展, 恶意代

码数量也呈爆发式增长。2020年瑞星“云安全”系统共截获病毒样本总量1.48亿个^[1], 病毒感染次数为3.52亿次, 病毒总体数量比2019年同期上涨43.71%, 恶意代码已经成为网络安全的重要威

* 基金项目: 国家自然科学基金(62001199); 漳州市自然科学基金(zz2020533)

胁之一^[2]。恶意软件作者经常会重用代码用来生成具有相似特征的其他恶意变体,而这些恶意变体通常可以归类为同一个恶意软件家族。因此,识别恶意软件家族的能力变得十分重要,通过对恶意代码的分类,可以更好防范恶意代码攻击。

近年来,恶意软件检测分类出现了静态分析和动态分析。静态分析侧重于统计特征,例如 API 调用、操作码序列等。Wang^[3]等人通过提取权限、硬件功能和接收者动作等 122 个特征,使用多种机器学习分类器进行训练和测试,并使用随机森林(Random Forest)分类器获得较高的分类准确率。动态分析则是使用虚拟的环境来分析恶意应用程序的行为^[4]。但是这些技术大多数需要提取大量特征,检测效率不高,对特征的选择需要一些专家知识,并且有一定的主观性。

为了降低特征工程成本和领域专家知识,一些研究人员使用可视化方法来解决恶意软件家族分类问题。例如,Nataraj 等人^[5]提出把恶意代码二进制文件转化为灰度图,然后利用 k 近邻算法对恶意代码进行分类,这种方法相比于之前未转换灰度图,直接分类的方法准确率有一定提高,但是该方法用 GIST 提取图片特征需要耗费大量时间,导致效率不高。

随着深度学习在图像分类领域的快速发展,有学者将深度学习引入到恶意代码检测领域。Choi 等人^[6]把恶意代码二进制文件转化为灰度图像,运用深度学习的技术,在 12 000 个样本中达到了 95.66% 的准确率。Su 等人^[7]用 light-weight DL 技术进行恶意代码家族分类,取得 94.00% 的成绩,但是他们提出的网络只对两类家族进行分类,有一定的局限性。Cui 等人利用卷积神经网络在图像分类的出色表现,并分别利用蝙蝠算法^[8]和 NSGA-II 算法^[9]处理恶意代码样本数量不均的问题,该方法准确率明显

高于传统机器学习方法,且算法复杂度较低。随着更深网络的提出,Rezende 等人提出将 VGG16 网络^[10]以 ResNet 网络^[11]运用在恶意代码检测分类上,该方法准确率有所提升,但是参数量变得巨大,分类效率有待提升。

基于上述方法产生的问题,本文提出了一种基于卷积神经网络^[12]的分类方法 SE-CNN,实现恶意代码家族分类。首先将恶意代码的二进制文件转化成灰度图得到灰度图像数据集,然后构建 SE-CNN 网络模型对灰度图像数据集进行训练,最后实现对恶意代码的检测分类。该方法采用 CNN 对灰度图像自动提取特征,解决了特征提取慢且耗时的问题;通过结合 SE 模块自适应学习通道重要程度信息,并赋予特征通道权重,从而激励有用特征信息,同时抑制无用信息,提升了模型分类准确率。实验结果表明,本文方法准确率高于传统机器学习方法,且参数量相较于先进的深度学习方法更低。

1 基于 SE-CNN 的分类检测算法

本文恶意代码检测模型框架如图 1 所示,首先将恶意代码二进制文件转化为灰度图像,然后将灰度图作为输入层,通过卷积神经网络层对图片特征进行提取,再经过挤压和激励网络模块对特征通道权重分配,接着通过丢弃(Dropout)层减少一定参数,从而防止模型过拟合,提升模型泛化能力,最后通过全连接层进行检测分类。具体模型参数如表 1 所示,其中 Conv 表示卷积层,Pool 表示池化层,Fc 表示全连接层。

1.1 恶意代码可视化

恶意代码可执行文件中的恶意行为不仅仅只存在于恶意代码中,还可能存在于其他数据中。如果只提取恶意代码进行检测,可能检测不全,如果全部提取又会增加时间消耗,还会降低准确度。本文将恶意代码可执行文件直接转换为灰度图像^[4],

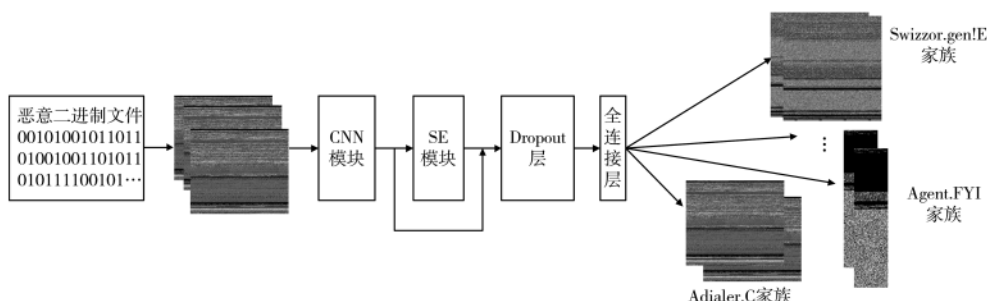


图 1 本文模型架构

表 1 模型参数

层	卷积核	输入图片	激活函数
Input		128×128×1	ReLU
Conv1	5×5	128×128×32	ReLU
Pool1		64×64×32	
Conv2	3×3	64×64×64	ReLU
Pool2		32×32×64	
Conv3	3×3	32×32×128	ReLU
Pool3		16×16×128	
Conv4	3×3	16×16×256	ReLU
Pool4		8×8×256	
Conv5	2×2	8×8×512	ReLU
SE		8×8×512	ReLU
Dropout		32 768	
Fc		25	Softmax

因此可以处理恶意代码可执行文件中的所有部分,同时相比于直接对文件进行特征提取,时间可有效缩短。具体方法如图 2 所示,本文把二进制可执行文件转化为二维的矩阵。

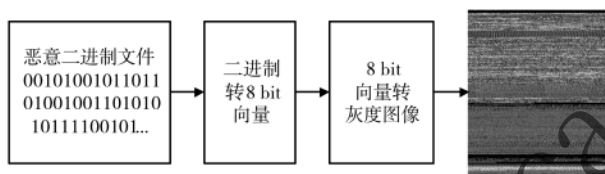


图 2 恶意代码可视化

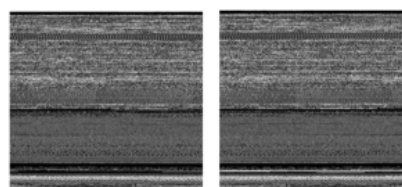
在二进制文件转换后的灰度图像中,相同的家族有着类似的灰度图纹理,不同家族的灰度图纹理具有一定的差别,如图 3 所示,Swizzor.gen!E 和 Swizzor.gen!I 家族灰度图有着相似的纹理。在一些算法模型上,对有着相似纹理的家族分类效果也不太理想。

1.2 卷积神经网络模块

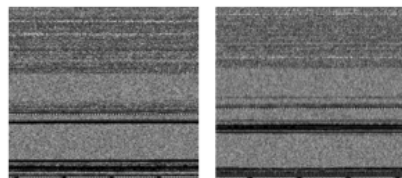
卷积神经网络是一种前馈神经网络^[12],广泛用于图像分类^[13]、语音识别等领域,它能以较少的提取特征信息。卷积神经网络对输入的图片提取特征,然后通过全连接层进行分类。

1.2.1 卷积层

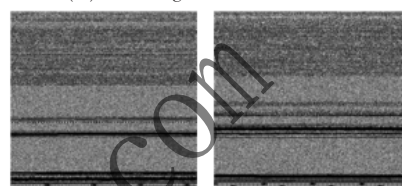
卷积层是对输入的特征图进行特征提取,通过卷积核和特征图做卷积运算,然后通过激活函数得到输出特征图。本文模型使用了线性整流函数(Rectified Linear Unit, ReLU)作为卷积层的激活函数,式(1)为 ReLU 函数的形式。ReLU 函数在 x 大于 0



(a) Adialer.C 家族灰度图



(b) Swizzor.gen!E 家族灰度图



(c) Swizzor.gen!I 家族灰度图

图 3 恶意代码家族灰度图

时,可以有效防止梯度弥散,加快计算。

$$f(x) = \begin{cases} 0, & x < 0 \\ x, & x \geq 0 \end{cases} \quad (1)$$

1.2.2 池化层

池化层又叫下采样层,经过卷积层后输出的特征图会被传递到池化层,然后进行特征选择和信息过滤。本文采取了极大池化(max pooling)方法,极大池化即输出区域内最大值。通过池化层能减少特征维度,降低过拟合。

1.2.3 全连接层

全连接层可以整合具有类别区分性的局部信息,在卷积神经网络中,通常会有一个或者一个以上的全连接层。本文模型最后是一个全连接层,用来对提取特征信息进行整合,然后将结果传递给 Softmax 分类器进行分类。式(2)为 Softmax 函数。通过 Softmax 函数可以使分类的概率范围在 0~1 之间,输出值可以用来检测某一类样本的概率,达到多分类的目的。

$$\text{Softmax} = \frac{e^{z_i}}{\sum_{m=1}^M e^{z_m}} \quad (2)$$

式中 z_i 表示第 i 个节点的输出值, M 表示输出节点个数。

1.3 挤压和激励网络模块

受启发于挤压和激励网络在图像分类中优异

的表现^[14], 本文将结合挤压和激励模块来对恶意家族灰度图像进行分类。挤压和激励模块原理图如图4所示。

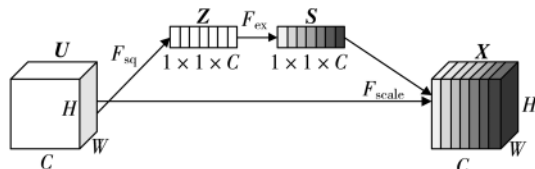


图4 挤压和激励模块原理图

挤压和激励模块直接连接在卷积神经网络模块之后, 输出的多通道特征图用 $U=[u_1, u_2, \dots, u_C]$ 来表示, 特征图的宽为 W , 高为 H , 通道数为 C 。通过挤压操作和激励操作, 最终得到加权特征图 $X=[x_1, x_2, \dots, x_C]$ 。下面详细介绍具体过程。

1.3.1 挤压块

为了得到多通道特征图之间的相关权重信息, 首先进行挤压操作, 对多通道特征图 U 进行全局平均池化(global average pool), 把 $H \times W \times C$ 的多通道特征图 U 挤压成 $1 \times 1 \times C$ 的向量 Z , 公式如下:

$$z_y = F_{sq}(u_y) = \frac{1}{H \times W} \sum_{i=1}^H \sum_{j=1}^W u_y(i, j) \quad (3)$$

其中 $z_y (y=1, 2, \dots, C)$ 表示向量 Z 中第 y 通道的特征向量, $u_y (y=1, 2, \dots, C)$ 表示 U 中第 y 通道的特征图。分别把 C 个通道 $H \times W \times 1$ 的特征图池化为 $1 \times 1 \times 1$ 的信息点, 经过 C 个通道的池化, 最后形成 $1 \times 1 \times C$ 的向量 Z 。因此经过挤压操作后的向量 Z 集合了整个特征图像信息。

1.3.2 激励块

激励块的目的是为了利用挤压操作后聚集的信息, 充分获取通道之间的关系, 激励有用的信息传递。为了限制模块的复杂性和提升通用性, 激励操作包含了两个全连接层, 即一个降维层和一个升维层。通过两个全连接层后得到 $1 \times 1 \times C$ 的向量 S , S 的公式如下:

$$S = F_{ex}(Z, W) = \sigma(g(Z, W)) = \sigma(W_2 \delta(W_1, Z)) \quad (4)$$

其中 σ 表示 Sigmoid 函数, δ 表示 ReLU 激活函数, $W=[W_1, W_2]$, W_1 代表维度缩减参数, W_2 代表维度增加参数, F_{ex} 为加权函数。Sigmoid 函数表示如式(5)所示:

$$S(x) = \frac{1}{1 + e^{-x}} \quad (5)$$

通过激励操作中降维层和 ReLU 激活函数以及升维层在进行非线性变换时建立通道之间的相关

性, 然后通过 Sigmoid 激活函数将权重信息归一化到 $[0, 1]$, 最终输出的向量 S 包含了各通道重要程度的信息。

将向量 S 中 C 通道的权重系数 $s_y (y=1, 2, \dots, C)$ 对 C 通道特征图 U 进行加权, 得到加权特征图 X , 即:

$$x_y = F_{scale}(u_y, s_y) = s_y u_y \quad (6)$$

式中 $x_y (y=1, 2, \dots, C)$ 为加权特征图 X 中第 y 通道特征图, 函数 $F_{scale}(u_y, s_y)$ 把特征图 u_y 和对应通道的权重系数相乘, 其结果为 x_y 。最终每个通道特征图都被赋值了不同的权重, 用来表示特征信息的重要程度, 从而实现激励有用信息, 抑制无用信息。

1.4 Dropout 层

深度学习的神经网络模型在训练时容易出现过拟合现象, Dropout 算法的提出, 在一定程度上可以防止模型过拟合^[15], 提高模型泛化性。图5展示了其中一种忽略部分隐层节点的情况。Dropout 算法的原理是在模型训练中随机忽略部分隐层节点, 在每批次的训练过程中, 由于每次忽略的隐层节点不同, 使得每次训练网络都有所不同, 每次训练都可以看作是一个“新”模型。此外, 由于节点都是以一定概率随机忽略, 使模型不会过分依赖于某些局部特征。

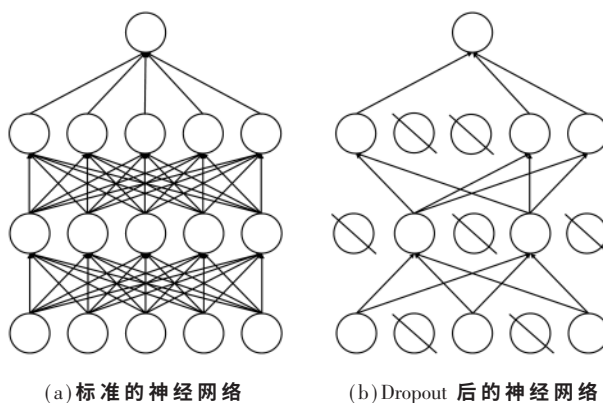


图5 使用 Dropout 的网络模型

Dropout 过程是一个有效的神经网络模型平均的方法, 通过训练大量不同的网络来平均预测结果, 这使得模型泛化性更强, 防止了模型过拟合。

Dropout 取值大小代表随机忽略的节点占总节点比率。Dropout 的取值会一定程度影响实验结果, Dropout 的值太小, 会导致防止过拟合的作用效果不佳, Dropout 值过大就会导致模型欠拟合。所以针对

Dropout 的取值本文做了对比实验。如图 6 所示,当 Dropout 取值为 0.6 时,准确率最高,所以本文选取 Dropout 为 0.6 作为实验参数设置。

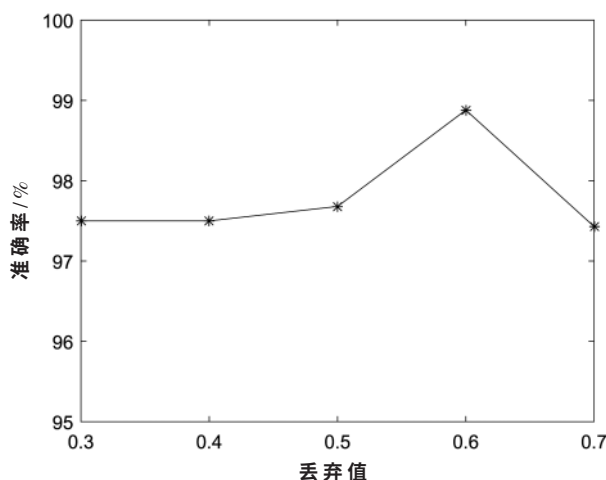


图 6 Dropout 取值对实验结果的影响

2 实验过程与结果分析

2.1 实验数据集和实验环境

本文使用的数据集是 Nataraj 等人公开的 Maling 数据集,该数据集包含了 9 339 个样本,共有 25 类。本文实验环境为 64 位 Windows 10 操作系统,处理器为 Intel Core i5-6300HQ,8 GB 内存,NVIDIA GeForce GTX 960 M 显卡以及 2.2.0 版本 Tensorflow-GPU。

2.2 评估指标

本文使用的评价指标分别为准确率(Accuracy)、精确率(Precision)、召回率(Recall)和 F1-Score,这些评价指标也被广泛地用在了最近研究中^[9,15]。此外本文还引入了方差评价指标,用来评估模型在 25 类家族准确率的稳定性。

Accuracy:是指分类正确的样本占总样本的比例,即用分类正确的正样本和分类正确的负样本之和除以总样本数。其公式如下:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (7)$$

Precision:是指分类为正确样本占真正正确样本的比例。其公式如下:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (8)$$

Recall:是指分类为正样本占总正样本的比例。其公式如下:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (9)$$

F1-Score:是指精确率和召回率的调和平均数,结合两者综合考虑的指标。其公式如下,其中 P 代表 Precision, R 代表 Recall。

$$\text{F1-Score} = 2 \times \frac{P \times R}{P + R} \quad (10)$$

方差:是每个样本值与全体样本值的平均数之差的平方值的平均数。其公式如下:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (X_i - u)^2 \quad (11)$$

式中 σ^2 表示总体方差, X_i 表示第 i 个变量, u 代表总体均值, N 表示总体个数。

上述评价指标参数说明如下:TP 为将正样本预测为正样本的数量;TN 为将负样本预测为负样本的数量;FP 为将负样本预测为正样本的数量;FN 为将正样本预测为负样本的数量。

2.3 不同大小图片对模型的影响

不同的图像尺寸对实验的效果也有一定的影响。针对输入图像尺寸问题,本文做了对比实验,分别将恶意代码图像缩放为 64×64 、 128×128 、 256×256 的尺寸作为输入。如表 2 所示, 256×256 尺寸的输入图像效果最好, 64×64 尺寸的输入图像相比于其他尺寸在准确率方面有一定下降,原因是尺寸过小会忽略一些原始图像边缘信息,导致分类效果稍微差一些。从表中还可以看出 256×256 尺寸的输入图像时间开销是 128×128 尺寸的输入图像的 10 倍左右,但是 128×128 尺寸的输入图像的效果和 256×256 尺寸的输入图像效果相差不大,时间开销却更少,所以本文最终选取输入图像的尺寸为 128×128 。

2.4 实验评估

为了验证本文提出模型的检测效果,将本文提出的方法和支持向量机、随机森林等传统的机器学习方法以及未结合 SE 模块的卷积神经网络方法做了对比,对数据集中 25 类家族的 Accuracy、Precision、Recall、F1-Score 和方差等指标进行了比较。最后本文还和最近文献使用了 Maling 数据集的机器学习方法和深度学习检测分类方法^[8-11]进行了比较。

在上述实验条件下,本文获得了 98.86% 的分类准确率。表 3 展示了 5 种方法总体的准确率、精确率、召回率和 F1-score。结果表明,本文提出的分类方法在分类精度方面优于其他所有分类方法,比未结合 SE 模块的卷积神经网络模型准确率提高了 0.32%。

表 2 不同尺寸图像的检测结果

输入图片	准确率/%	精确率/%	召回率/%	F1-Score/%	运行时间/s
64×64	97.79	97.76	97.79	97.75	6.19
128×128	98.86	98.88	98.86	98.86	19.13
256×256	99.11	99.10	99.11	99.09	192.53

表 3 模型检测结果对比

分类方法	准确率/%	精确率/%	召回率/%	F1-Score/%
随机森林	97.68	97.71	97.68	97.53
支持向量机	96.96	96.80	96.96	96.78
决策树	94.90	94.98	94.89	94.90
卷积神经网络	98.54	98.56	98.54	98.53
本文方法	98.86	98.88	98.86	98.86

图 7~图 10 展示了 5 种方法在 25 类恶意代码家族的准确率、精确率、召回率和 F1-score。实验结果表明,本文提出的 SE-CNN 算法(98.86%)在准确度上高于支持向量机(96.96%)、随机森林(97.68%)、决策树(94.90%)、卷积神经网络(98.54%)等算法。并且在 Swizzor.gen!E 和 Swizzor.gen!I 两个恶意家族十分相似的情况下,其他分类方法都不能准确地区分它们,但本文方法还保持了较高准确度。

图 11 表示了各类方法在 25 类恶意家族准确率的方差。可以看出本文的方法方差最小,一定程度上反映了本文方法在 25 类恶意家族分类准确率的稳定性相较于其他方法更优一些。

如表 4 所示,本文与现有的基于机器学习和深度学习技术的恶意软件分类方法进行了比较。以下方法均使用了公开的 Maling 数据集。

由表 4 可见,与深度学习算法 VGG16 和 ResNet 相比,本文方法参数量更少。表中一些算法在网络训练过程中,采用了不同的技术来平衡 Maling 数据集中恶意代码家族不平衡的问题,如 Cui 等人使用了 BAT^[8]和 NSGA-II^[9]等平衡数据集的技术。本文方法没有采用任何平衡数据的技术,并取得了较高分类准确率。

本文在相似的家族分类中也保持了较高的召回率,如表 5 所示,对 C2LOP.gen!g 和 C2LOP.P 家族,以及 Swizzor.gen!E 和 Swizzor.gen!I 家族的召回率进行了对比。

本文方法在相似家族的分类召回率上有较大提升,是因为本文方法结合了 CNN 和 SE 模块的优点,通过激励特征通道中有用信息并抑制无用信息,可以更好地区分相似恶意家族之间的细微变化。

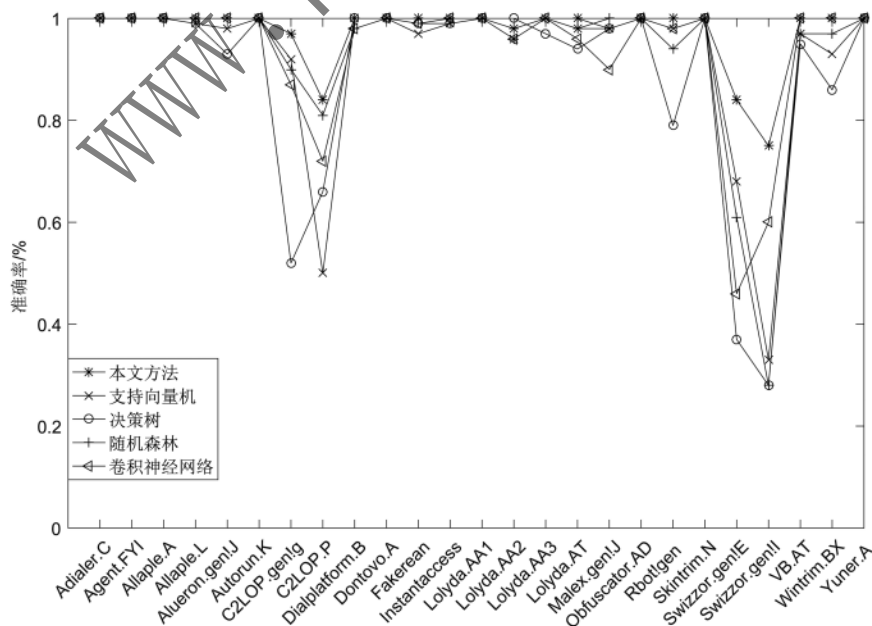


图 7 不同分类方法在 25 类恶意家族的准确率

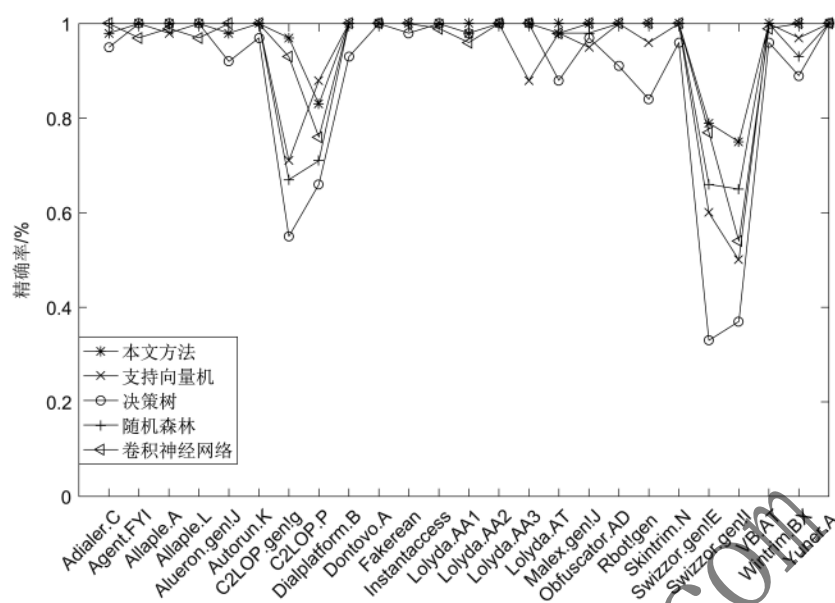


图 8 不同分类方法在 25 类恶意家族的精确率

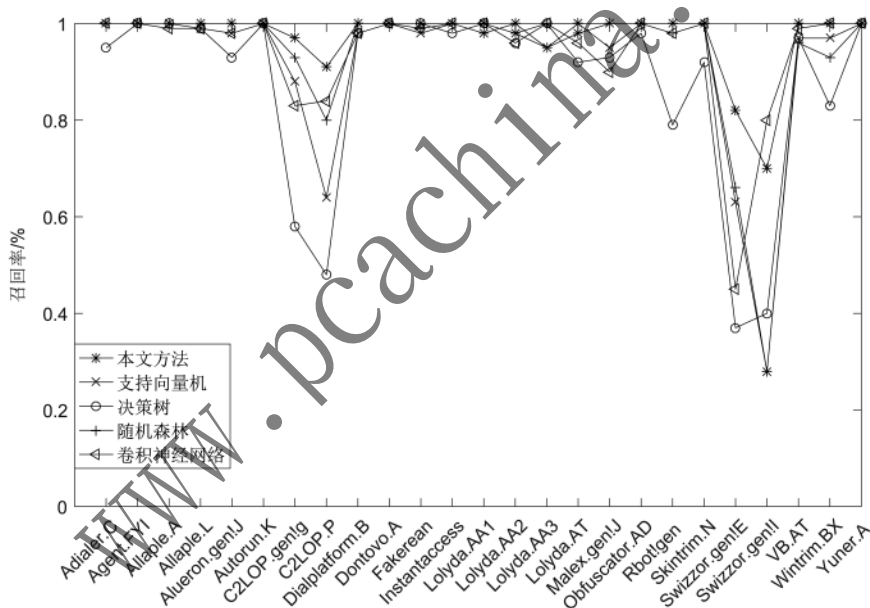


图 9 不同分类方法在 25 类恶意家族的召回率

3 结论

现有的反恶意代码的解决方案主要还是依赖于机器学习技术。虽然基于机器学习的方法已经被证明在检测新的恶意代码方面具有有效性,但同时也伴随着巨大的开发成本。机器学习所需要的有用特征需要花费大量时间以及一定的恶意代码分析专业知识。

在深度学习架构中,尤其是卷积神经网络,在检测恶意代码方面有出色的表现。因此,本文提出

了一种基于图像和卷积神经网络的恶意代码检测方法,把 SE 模块和卷积神经网络相结合,快速对恶意代码转换的灰度图像进行分类,准确识别恶意代码所属家族。实验结果证明,本文方法在大多数恶意样本下分类正确,即使在恶意代码家族相似的情况下也有较高的准确率;本文方法与传统的基于机器学习解决方案相比具有更高的准确率,同时避免了手动特征工程阶段,节省了大量时间。未来工作将恶意代码转换成彩色图像也是研究的一个方向。

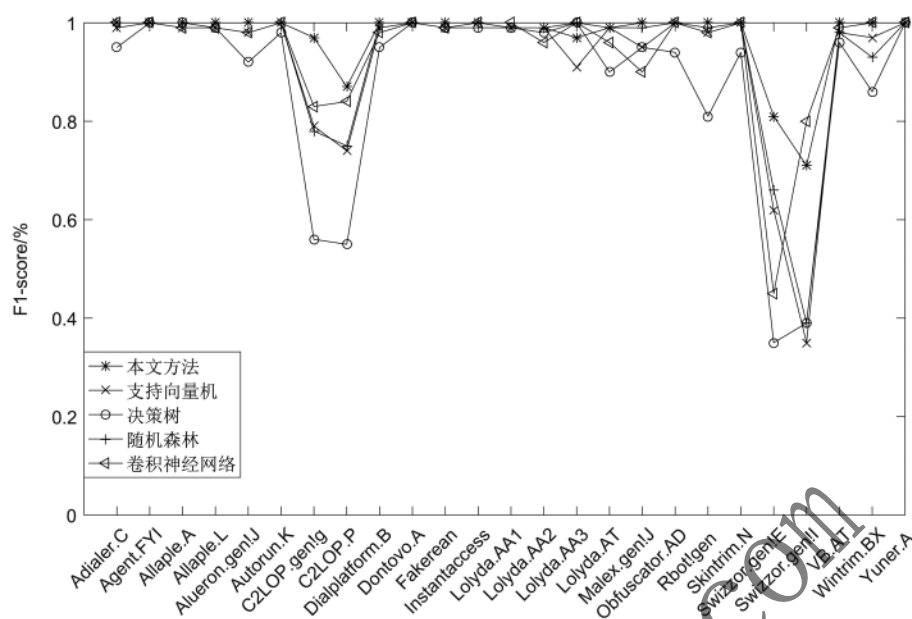


图 10 不同分类方法在 25 类恶意家族的 F1-score

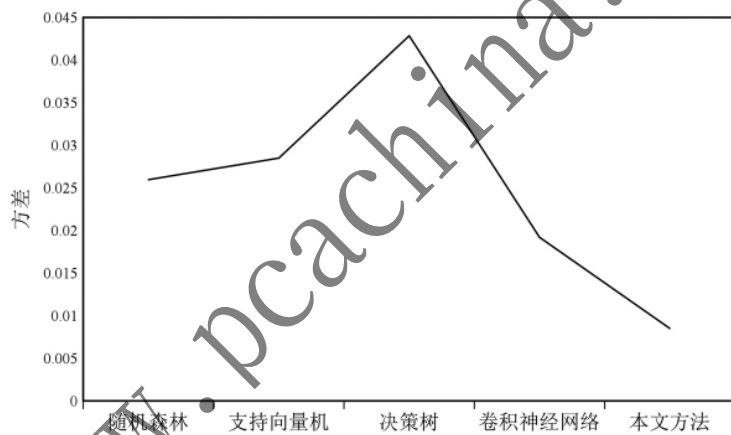


图 11 25 类恶意家族准确率的方差

参考文献

- [1] 北京瑞星网安技术股份有限公司. 瑞星 2020 年中国网络安全报告[J]. 信息安全研究, 2021, 7(2): 102-109.
- [2] 冀甜甜, 方滨兴, 崔翔, 等. 深度学习赋能的恶意代码攻防研究进展[J]. 计算机学报, 2021, 44(4): 669-695.
- [3] WANG K, SONG T, LIANG A L. Mmda: Metadata based malware detection on Android[C]//12th International Conference on Computational Intelligence and Security. IEEE, 2016: 598-602.
- [4] DAMODARAN A, DI TROIA F, VISAGGIO C A, et al. A comparison of static, dynamic, and hybrid analysis for malware detection[J]. Journal of Computer Virology and Hacking Techniques, 2017, 13: 1-12.
- [5] NATARAJ L, KARTHIKEYAN S, JACOB G, et al. Malware images: visualization and automatic classification[EB/OL]. (2011-07-20)[2021-07-10]. <https://doi.org/10.1145/2016904.2016908>.
- [6] CHOI S, JANG S, KIM Y, et al. Malware detection using malware image and deep learning[C]//International Conference on Information and Communication Technology Convergence. IEEE, 2017: 1193-1195.
- [7] SU J W, VASCONCELLOS D V, PRASAD S, et al. Lightweight classification of IoT malware based on image recognition[C]//42nd Annual Computer Software and Applications Conference. IEEE, 2018: 664-669.
- [8] CUI Z H, XUE F, CAI X J, et al. Detection of malicious

表 4 与其他文献方法比较

分类方法	准确率/%	精确率/%	召回率/%	参数量
VGG16+SVM(文献[10])	92.29	—	—	138 000 000
ResNet(文献[11])	98.62	—	—	23 534 592
CNN+BAT(文献[8])	94.50	94.60	94.50	—
GLCM+SVM(文献[8])	93.20	93.40	93.00	—
GIST+KNN(文献[8])	92.50	92.70	92.30	—
CNN+NSGA-II(文献[9])	97.60	—	88.40	—
本文方法	99.11	99.10	99.11	2 355 097

表 5 相似家族召回率比较

分类方法	C2LOP.gen!g	C2LOP.P	Swizzor.gen!E	Swizzor.gen!I
CNN+BAT	71	70	75	66
CNN+NSGA-II	72	43	40	52
本文方法	98	89	84	78

code variants based on deep learning[J].IEEE Transactions on Industrial Informatics, 2018, 14(7): 3187–3196.

- [9] CUI Z H, DU L, WANG P H, et al. Malicious code detection based on CNNs and multi-objective algorithm[J]. Journal of Parallel and Distributed Computing, 2019, 129: 50–58.
- [10] REZENDE E, RUPPERT G, CARVALHO T, et al. Malicious software classification using VGG16 deep neural network's bottleneck features[J]. Advances in Intelligent Systems and Computing, 2018, 738: 51–59.
- [11] REZENDE E, RUPPERT G, CARVALHO T, et al. Malicious software classification using transfer learning of RESNET-50 deep neural network[C]//16th IEEE International Conference on Machine Learning and Applications. IEEE, 2017: 1011–1014.
- [12] 周飞燕, 金林鹏, 董军. 卷积神经网络研究综述[J]. 计算机学报, 2017, 40(6): 1229–1251.
- [13] GAYATHRI S, GOPI V P, PALANISAMY P. A light-

weight CNN for diabetic retinopathy classification from fundus images[J]. Biomedical Signal Processing and Control, 2020, 62: 102–115.

- [14] HU J, SHEN L, SUN G. Squeeze-and-excitation networks[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2020, 42: 2011–2023.
- [15] CHEN Y Y, YI Z. Adaptive sparse dropout: Learning the certainty and uncertainty in deep neural networks[J]. Neurocomputing, 2021, 450: 354–361.

(收稿日期: 2022-03-03)

作者简介:

申高宁(1997–), 男, 硕士研究生, 主要研究方向: 恶意代码检测。

陈志翔(1982–), 通信作者, 男, 博士, 教授, 主要研究方向: 信息安全、图像处理。E-mail: zxchenphd@163.com。

王辉(1985–), 男, 博士, 副教授, 主要研究方向: 水声通信与网络、计算智能与应用。



版权声明

经作者授权，本论文版权和信息网络传播权归属于《信息技术与网络安全》杂志，凡未经本刊书面同意任何机构、组织和个人不得擅自复印、汇编、翻译和进行信息网络传播。未经本刊书面同意，禁止一切互联网论文资源平台非法上传、收录本论文。

截至目前，本论文已经授权被中国期刊全文数据库（CNKI）、万方数据知识服务平台、中文科技期刊数据库（维普网）、JST 日本科技技术振兴机构数据库等数据库全文收录。

对于违反上述禁止行为并违法使用本论文的机构、组织和个人，本刊将采取一切必要法律行动来维护正当权益。

特此声明！

《信息技术与网络安全》编辑部
中国电子信息产业集团有限公司第六研究所