

人工智能技术在等保测评领域的应用及重要性分析

武建双¹, 孙 宝², 田金玉²

¹合肥天帷信息技术有限公司总经办, 安徽 合肥

²合肥天帷信息技术有限公司数据安全、人工智能安全研究中心, 安徽 合肥

收稿日期: 2024年11月25日; 录用日期: 2024年12月23日; 发布日期: 2024年12月30日

摘 要

随着信息技术的快速发展, 等级保护(等保)测评成为确保信息系统安全的关键环节。作为网络安全保障的重要措施, 其测评工作至关重要。人工智能(AI)技术, 作为近年来迅速发展的前沿科技, 正在为等级保护测评工作提供新的方法和工具。本文深入探讨了人工智能(AI)技术在等保测评中的应用案例, 包括利用大模型出色的语义理解能力进行等保测评报告的风险分析自动生成、精确抽取复杂文本中的关键实体关系, 以及运用Bert模型进行等保测评报告中高风险项的判断。同时, 本文还探讨了将机器视觉算法与传统图像处理算法相结合, 实现了测评工具在特殊字符识别和用户状态识别方面的突破。通过案例研究, 本文分析了AI技术在提升测评效率、准确性和应对复杂性方面的重要性, 并展示了AI技术如何助力等保测评的自动化和智能化。同时, 本文也指出了当前面临的挑战和未来的发展方向, 为等保测评领域的进一步发展提供了宝贵的参考和启示。

关键词

人工智能, 等级保护测评, 信息安全

Application and Importance Analysis of Artificial Intelligence Technology in the Field of Equal Protection Evaluation

Jianshuang Wu¹, Bao Sun², Jinyu Tian²

¹Department of General Office, Hefei Tianwei Information Security Technology Co., Ltd., Hefei Anhui

²Department of Data Security and Artificial Intelligence Security Research Center, Hefei Tianwei Information Security Technology Co., Ltd., Hefei Anhui

Received: Nov. 25th, 2024; accepted: Dec. 23rd, 2024; published: Dec. 30th, 2024

文章引用: 武建双, 孙宝, 田金玉. 人工智能技术在等保测评领域的应用及重要性分析[J]. 计算机科学与应用, 2024, 14(12): 243-252. DOI: 10.12677/csa.2024.1412259

Abstract

With the rapid development of information technology, equal protection evaluation (equal protection) has become a key link in ensuring the security of information systems. As an important measure for network security assurance, its evaluation is crucial. Artificial Intelligence (AI) technology, as a rapidly developing cutting-edge technology in recent years, is providing new methods and tools for equal protection evaluation work. This paper discusses in depth the application cases of artificial intelligence (AI) technology in Equal Protection, including the use of the excellent semantic understanding ability of Large Language Models for the automatic generation of risk analysis of equal protection reports, the precise extraction of key entity relationships in complex texts, and the use of the Bert model to judge the high-risk items in reports. Meanwhile, this paper also discusses the combination of machine vision algorithms and traditional image processing algorithms to achieve breakthroughs in special character recognition and user status recognition of evaluation tools. Through case studies, this paper analyzes the importance of AI in improving the efficiency, accuracy and complexity of assessment, and shows how AI can help the automation and intelligence of equal protection. At the same time, this paper also points out the current challenges and future development directions, and provides valuable reference and inspiration for the further development of the field of Equal Protection.

Keywords

Artificial Intelligence, Equal Protection Evaluation, Information Security

Copyright © 2024 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

1.1. 等级保护(等保)测评的背景和目的

信息安全等级保护[1], 简称“等保”, 是我国为加强信息安全管理、确保国家信息系统安全而实施的一项重要制度。该制度的核心目标是通过信息系统进行分级保护, 确保不同等级的信息系统能够得到相应的安全防护。具体而言, 等级保护测评是对信息系统在网络安全层面的安全性、合规性方面进行评估, 以确保其符合国家规定的安全标准和要求。在全球网络安全形势日益复杂的背景下, 等保测评对于维护国家安全、保护公民隐私和保障企业运营具有重要意义。

1.2. 传统测评方法的局限性

传统的等保测评方法主要依赖于人工操作和专家经验, 虽然在一定程度上能够保障测评的全面性, 但在当前日益复杂的网络环境中, 传统方法的局限性愈发明显。其局限性具体如下:

1) 依赖人工: 传统的等保测评往往需要安全专家对信息系统进行详细的检查和分析。这种高度依赖人工的方式不仅耗时耗力, 而且容易受到人员专业知识和经验的限制, 导致测评结果的可靠性和一致性受到影响。

2) 效率低下: 随着信息系统规模的扩大和复杂性的增加, 传统测评方法的效率问题变得尤为突出。人工测评通常需要逐一检查系统中的各个安全控制点, 这不仅耗费大量时间, 还容易出现疏漏。此外,

在面对复杂的攻击手段和多样化的安全威胁时，人工测评难以实现快速响应。

3) 主观性强：由于传统测评方法依赖于安全专家的判断，不同的测评人员可能对相同系统产生不同的评价。这种主观性导致测评结果存在一定的不确定性，可能无法全面反映信息系统的真实安全状况。此外，人工操作中的疏忽或偏见也可能导致关键问题被忽视，进而影响系统的安全性。

因此，传统的测评方法已难以应对当今网络安全环境的挑战，亟需引入更加高效、客观和智能化的技术手段。而人工智能技术因其自动化和智能化的特点，能够提供更为精准的威胁检测和响应，逐渐成为维护网络空间安全的重要力量。

2. 人工智能技术在等保测评中的重要性分析

随着计算能力的提升和数据积累的增加，人工智能(AI) [2]技术近年来迅速发展，并在多个领域展现出强大的应用潜力，也为等保测评的自动化和智能化提供了新的技术手段。其主要优势如下：

1) 自动化能力[3]：人工智能技术能够自动化处理大量复杂数据，这一能力在等级保护测评中尤为关键。通过自动化数据分析和威胁检测，人工智能算法可以迅速识别潜在的安全漏洞和异常行为，减少对人工操作的依赖，并提高测评的速度和效率。

2) 智能化分析[4]：人工智能技术能够从大量历史数据中学习，进而对新出现的威胁进行预测和识别。与传统规则方法不同，AI 算法能够在不断变化的网络环境中自适应地调整，提升测评的智能化水平。这不仅提高了测评结果的准确性，还增强了系统应对新型攻击的能力。

3) 实时性与动态性：人工智能技术能够实现实时分析和动态调整，这对于应对当前复杂多变的网络威胁至关重要。在等级保护测评过程中，人工智能模型可以在网络流量和系统日志中持续监控异常行为，并及时发出警报，为安全管理者提供及时的决策支持。

可以看出，人工智能技术的引入为等级保护测评的自动化和智能化提供了强有力的支持，可显著提升了测评的效率、准确性和响应速度。

3. 人工智能技术在等保测评中的应用

3.1. 大模型(LLMs) [5]

LLMs (Large Language Models)是一种基于机器学习的技术，大型语言模型(LLMs)的发展演变经历了从早期的统计语言模型到如今的基于 Transformer 架构[6]的复杂模型。它通过在大量文本数据上进行预训练，来预测下一个词或下一段话的可能性，能够理解和处理人类语言，展现出深度的语言理解能力、类人类的文本生成能力、情境意识以及强大的问题解决技能。2023 年，如对英文较友好的 OpenAI 发布 ChatGPT 系列[7]、Meta 发布的 LLaMA 系列[8]，以及中文理解能力较强的阿里发布的 Qwen 系列[9]、智谱发布的 Glm 系列[10]等模型的发布，进一步推动了 LLMs 的普及和应用。随着技术的不断进步，LLMs 在理解和生成自然语言方面的能力越来越强，为各种语言处理应用提供了强大的支持，尤其是在自然语言处理直接相关的领域，如搜索引擎、客户支持和翻译。此外，它们还在代码生成、医疗、金融和教育等更广泛的场景中找到了应用，展示了适应性和简化语言相关任务的潜力。

等级保护测评中针对测评师结果记录生成问题描述、问题分析、危害分析、整改建议、关联威胁及风险等级，传统方法依赖于测评师经验人为书写，最终导出测评报告。此方式存在主观性强、效率低下等问题。大语言模型(LLMs)因其强大的自然语言理解和生成能力，为风险分析自动生成提供了新的思路。

等保测评领域已有不少相关应用实例。比如基于 qwen-7b 模型，对风险分析数据进行 sft 操作[11]，利用 LLM 的语义理解技术根据结果记录自动生成问题描述、问题分析、危害分析、整改建议、关联威胁及风险等级，这样不仅提高了准确率还大大提升了评测效率。首先，问题描述生成整改建议原始数据共

计 46,360 条，剔除空数据共计 35,000 条，并编写代码把原始数据转化成 json 格式，具体如下图 1 所示：

```
{
  "kind": "OpenQA",
  "input": "经核查，系统网络设备、服务器采用冗余部署，保证系统的高可用性，但存在部分关键设备未采用冗余设计。对上述文本进行问题描述的结果：",
  "target": "系统部分关键设备未冗余部署。",
  "kind": "OpenQA",
  "input": "未定期进行审查审批事项。对上述文本进行问题描述的结果：",
  "target": "未定期进行审查审批事项。",
  "kind": "OpenQA",
  "input": "未与被录用人签署",
  "target": "未与被录用人签署",
  "kind": "OpenQA",
  "input": "系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。对上述文本进行问题描述的结果：",
  "target": "系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。",
  "kind": "OpenQA",
  "input": "未制定详细的系统交付清单。对上述文本进行问题描述的结果：",
  "target": "未制定详细的系统交付清单。",
  "kind": "OpenQA",
  "input": "未对负责运行维护",
  "target": "未对负责运行维护",
  "kind": "OpenQA",
  "input": "经核查，",
  "target": "主机操作系统、网络设备、系统管理软件和安全设备分别采用SSH协议和HTTPS协议进行行",
  "kind": "OpenQA",
  "input": "系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。对上述文本进行问题描述的结果：",
  "target": "系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。",
  "kind": "OpenQA",
  "input": "用HTTPS进行远程管理，能够保证通信过程中鉴别数据的保密性。对上述文本进行问题描述的结果：",
  "target": "数据通信",
  "kind": "OpenQA",
  "input": "系统已针对互联网医院信息系统制定",
  "target": "主要内容包含项目基本情况及完成情况。文件包含",
  "kind": "OpenQA",
  "input": "网络设备和安全设备分别通过SSH协议和HTTPS协议进行通信传输，可保证通信过程中鉴别数据的完整性，但未采用校验技术或密码技术保证通信过程中审计数据和配置数据的完整性。",
  "target": "服务器通过RDP协议进行通信传输。",
  "kind": "OpenQA",
  "input": "经核查，",
  "target": "其中规定由网络与信息安全领导",
  "kind": "OpenQA",
  "input": "系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。对上述文本进行问题描述的结果：",
  "target": "系统通信设备未采取可信验证措施。",
  "kind": "OpenQA",
  "input": "所使用的软件为成熟软件，软件发售前进行了源代码检查和功能测试。对上述文本进行问题描述的结果：",
  "target": "未对开发单位提供的软件源代码做到全面审查，无法确",
  "kind": "OpenQA",
  "input": "未定期对系统进行漏洞扫描和其它安全性测试。",
  "target": "已记录和保存基本配置信息，包括网络拓扑结构、各个设备的参数信息、各个设备安装的软件组件、软件",
  "kind": "OpenQA",
  "input": "与关键岗位人员签署责任协议。",
  "target": "下级单位，各类制度均遵循",
  "kind": "OpenQA",
  "input": "未在软件交付前检测其中可能存在的恶意代码。对上述文本进行问题描述的结果：",
  "target": "未在软件交付前检测其中可能存在的恶",
  "kind": "OpenQA",
  "input": "暂无未经现场核查得知，系统通信网络相关设备未基于可信根对边界设备的系统引导程序、系统程序、重要配置参数和边界防护应用程序等进行可信验证。对上述文本进行问题描述的结果：",
  "target": "应用系统、安全设备采用SSH协议和RDP协议可以保证保证通信",
  "kind": "OpenQA",
  "input": "软件商未提供软件源代码。对上述文本进行问题描述的结果：",
  "target": "未提供软件代码的检测报告。",
  "kind": "OpenQA",
  "input": "经现场核查得知，",
  "target": "限公司未对系统日常运行情况和系统漏洞情况定期进行安全检测。对上述文本进行问题描述的结果：",
  "target": "未对系统日常运行情况。系统漏洞情",
  "kind": "OpenQA",
  "input": "系统的开发单位未提供系统的软件源代码，无法审查其中可能存在的后门及隐藏逻辑。",
  "target": "经现场核查得知。"
```

Figure 1. Rectification data processed into json format

图 1. 整改数据处理成 json 格式

而后，使用 lora 方式[12]对其进行微调，共计迭代 18,750 次，验证集上达到最优识别率 96.1%。在 ubuntu20.04、4 颗 Intel Xeon 6348H (24 C, 165 W, 2.3 GHz)、4080/16G*2 环境下，最佳模型对 1000 条数据进行测试，平均准确率为 96.3%，每条整改建议生成时间为 6 s，不仅避免了人为主观因素的影响，时间上还提升了近 20 倍，在资产只有 1 台情况下生成每份报告平均可以节省 1.58 小时。

测评师编写与审核测评报告过程中，既耗费了大量时间，又难以稳定保证准确率。此外，公司期望通过深入分析报告，特别是关键实体的准确抽取，来指导项目定位和投入方向，但面对海量的冗余数据，这项工作的推进显得尤为困难。为此，已有基于 baichuan-7B [13]大模型进行网络安全领域的实体抽取的成功应用。发明中指出，网络安全领域中的实体类型非常多样，包括威胁类型、漏洞、攻击手段等，这些实体的表现形式可能与大模型预训练数据有一定差异。通过大量的实践测试后发现，直接采用大模型底模，进行网络安全方面的实体抽取的成功率，最高只能达到 65%左右，远远不能满足实际应用的需求。为了提高识别率，发明中采用二次训练的策略，通过向 baichuan-7B 大模型提供大量的网络安全相关数据进行训练，来提高其在网络安全方向上实体识别的准确率。

```
{
  "kind": "NER",
  "input": "在安全计算环境层面，所有安全组件、数据库、应用系统、操作系统均提供身份标识和鉴别功能，身份标识具有唯一性，同时采用 HTTPS 或 SSH 加密协议进行远程管理，能够防止鉴别信息在网络传输过程中被窃听；所有安全组件、数据库、应用系统、操作系统均已启用安全审计功能，审计记录到每个用户，对重要的用户行为和重要安全事件进行审计，审计记录实时发送至阿里云控制台的云日志审计中；服务器遵循了最小安装的原则，仅安装了需要的组件和应用程序，未启用多余的系统服务和高危端口，所有服务器仅允许通过 远程管理设备，限制了终端的接入方式，通过云安全中心对服务器操作系统的入侵行为和恶意代码进行检测和防护；设备和系统采用 HTTPS 或 SSH 加密协议进行远程管理，能够保证通信过程中鉴别数据和配置数据的保密性与完整性；应用系统和数据库仅采集和保存了业务必需的用户个人信息，已禁止未授权访问和非授权使用用户个人信息。",
  "target": "不存在实体"
},
{
  "kind": "NER",
  "input": "在安全管理人员层面，人力资源部负责人员录用，《M306.216-人员信息安全 管理规定》规定对被录用人的身份、背景和专业资格等进行审查，人员录用 后与被录用人签署《保密协议》，与关键岗位人员签署《员工信息安全责任 承诺书》；明确及时终止离岗人员的所有访问权限，根据《员工离职申请》办 理严格的调离手续，并签署《离岗协议》才能离开；人力资源部和信息化安全管理部对各类人员进行安全意识和岗位技能培训，并告知相关的安全责任和 惩戒措施，制定针对不同岗位的不同培训计划，并要求员工通过 平台按时完成，内容涵盖安全基础知识、岗位操作规程等；已制定《M306.215-物理安全管理规定》，行政部建立有访客在公司内行动轨迹的监视机制，外部 人员接入受控网络访问系统前通过VPN 访问申请，批准后由部门主管开设账户、分配权限，《互联网访问权限申请》中明确外部人员离场后及时清除所有的访问权限，《项目保密协议》中明确了外部人员的保密义务。",
  "target": "文档：《M306.216-人员信息安全 管理规定》，《员工信息安全责任 承诺书》，《员工离职申请》，《离岗协议》，《M306.215-物理安全管理规定》，《互联网访问权限申请》，《项目保密协议》\n部门：人力资源部，信息安全管理部，行政部"
},
{
  "kind": "NER",
  "input": "整改建议：建议使用符合国家密码管理主管部门认证核准的密码技术或密码产品。",
  "target": "不存在实体"
},
{
  "kind": "NER",
  "input": "网络安全等级保护测评是依据国家网络安全等级保护制度，按照有关管理规范和标准，对已定级备案的非涉及国家秘密的网络（含信息系统、数据资源等）的安全保护状况进行检验评估的活动。",
  "target": "地点：不存在实体"
},
{
  "kind": "NER",
  "input": "抽取上下文中的实体：",
  "target": "对智慧教学平台开展网络安全等级测评，通过对目标网络系统的 安全技术状态及安全管理状况依据相应网络安全等级保护要求进行测试、评估与分析，并将测评结论作为委托方进一步完善系统安全制度策略及安全技术防护措施依据。",
  "target": "组织： 技术有限公司"
```

Figure 2. Processed data json format file

图 2. 处理后的数据 json 格式文件

为了实现对 baichuan-7B 的高效训练，需要大量收集网络安全领域的相关数据并进行全面清洗，以获得高质量数据集。普通的测评报告 PDF 文件中，通常含有大量停用词，同时数据的目标标签也并不清晰，这种未经处理的原始数据难以满足模型训练的基本要求。此外，模型训练通常需要的是结构化的 json 格

式文件。根据工程实践经验,要构建出成熟且具备高准确率的模型,往往需要数十乃至上百个这样的文本集作为数据支撑。因此,在利用 baichuan-7B 大模型进行训练之前,对原始测评报告进行预处理,包括去除停用词、明确目标标签以及格式转换等工作,以确保训练数据的准确性和有效性。训练的数据 json 文件,其中输入为 input,是要进行实体抽取的句子,目标是 target 模型的训练目标。用于二次训练的标准数据格式如图 2。

模型整个训练过程共计完成了 32,343 次迭代。最终,模型在验证集上展现出 89.4%最优识别率。经测试,在 ubuntu20.04、4 颗 Intel Xeon 6348H (24 C, 165 W, 2.3 GHz)、4080/16G*2 环境上,在 2000 条测试语句,以及目标实体上,模型平均准确率为 88.1%,与底模相比,关键实体抽取准确率显著提高了 23%。因底模的抽取能力较为关键,在后续算法优化过程中,可以通过引入更先进的预训练模型进行二次微调的方式,以提高模型整体的实体抽取准确率。

3.2. 自然语言处理(NLP)

BERT (Bidirectional Encoder Representations from Transformers) [14]自 2018 年提出以来,已经成为自然语言处理(NLP) [15]领域的一个重要里程碑。它的核心优势在于使用 Transformer 架构的双向编码器,能够同时考虑句子中所有单词的上下文信息。BERT 通过预训练和微调两个阶段来训练模型,预训练阶段在大量未标记的文本上进行,学习上下文嵌入,微调阶段则针对特定任务调整模型参数。BERT 的这种设计使得它在多项 NLP 任务上取得了显著的性能提升,并且具有高度的灵活性和可扩展性。其应用优势在于其强大的上下文理解能力,能够处理复杂的语言结构和语义。它在多项 NLP 任务中表现出色,如文本分类、问答和命名实体识别等。BERT 的预训练-微调范式几乎是一种“一刀切”的解决方案,可以轻松适应各种 NLP 任务,从而减少了从头开始训练模型的复杂性和计算成本。此外,BERT 的灵活性使其可以轻松适应多种任务和行业特定的需求。BERT 和其他基于注意力的模型提供了一定程度的解释性,通过分析注意力权重,可以了解模型在做决策时到底关注了哪些部分的输入。这些特性使得 BERT 在多语言应用中表现出强大的多功能性。

等级保护测评过程中,测评师需要基于评测条款,结合不同的系统等级、结果记录等进行风险等级评定,传统方法主要靠人工分析判断。此方式存在主观性强、效率低下等问题。Bert 利用其快速捕捉文本中的深层次语义关系,在文本分类任务上表现优异,在等保测评风险等级评定文本分类任务上得到成功应用。

```
{
  "result_record": "经检查,服务器未设置审计管理员,设置了合理的审计日志文件访问权限",
  "system_level": "三级",
  "label": 1
},
{
  "result_record": "经核查得知未添加可信芯片以及可信程序实现可信验证,未集于可信根",
  "system_level": "三级",
  "label": 0
},
{
  "result_record": "经检查,服务器未设置审计管理员,设置了合理的审计日志文件访问权限",
  "system_level": "三级",
  "label": 1
},
]
```

Figure 3. Bert risk classification json format training dataset

图 3. Bert 风险分类 json 格式训练数据集

在针对等保测评的风险等级评定中,可将不同的系统等级、结果记录作为模型训练的输入,所对应的风险等级作为标签,进行模型训练。其中,风险等级有低、中、高三等级,原始数据分别有 17 多万

条、80 多万条、15 多万条。任务需求是进行二分类，重点提高高风险的识别精度及召回率。通过对原始数据及任务需求分析，将低风险、中风险统一划分为非高，作为一类(如图 3 中的 label:0)；高风险作为单独一类(如图 3 的 label:1)。可以看出，非高数据集数量(97 多万条)与高风险数据数量相差较大，样本数量严重不平衡。为此，案例中首先按照 8:2 比例对非高、高风险数据进行训练集、验证集划分，之后只对训练集中的高风险数据进行文本增强，以扩充高风险样本数量，最后获得高风险训练集样本 70 多万条。用于模型训练的标准数据格式如图 3 所示。

因需同时保证高风险的推理准确率和召回率，模型训练及验证过程中，采用 F1-score [16]作为模型评价指标，如公式 1 所示：

$$F1=\frac{2*acc*recall}{acc+recall}$$

(公式 1)

可以看出 F1-score 对模型的识别准确率和召回率都进行兼顾。模型训练迭代至 70,000 次时，总体验证精度 F1 为 48%。经测试，模型在高风险测试集上准确率(acc)为 84%，在非高风险测试集上准确率(acc)为 88%。后续为了进一步提高模型的准确率，以及确保模型在实际应用中的有效性和可靠性，可从数据层面，采用重采样技术，通过采样少数类(非高数据)或欠采样多数类(高风险数据)，以减少偏差(见图 4)。

```
{
  "epoch": 2.492433683460922,
  "eval_f1": 0.48351005873129765,
  "eval_loss": 0.3554157316684723,
  "eval_runtime": 2204.0501,
  "eval_samples_per_second": 149.578,
  "eval_steps_per_second": 6.233,
  "step": 70000
},
```

Figure 4. Model training record
图 4. 模型训练记录

3.3. 机器视觉(CV)

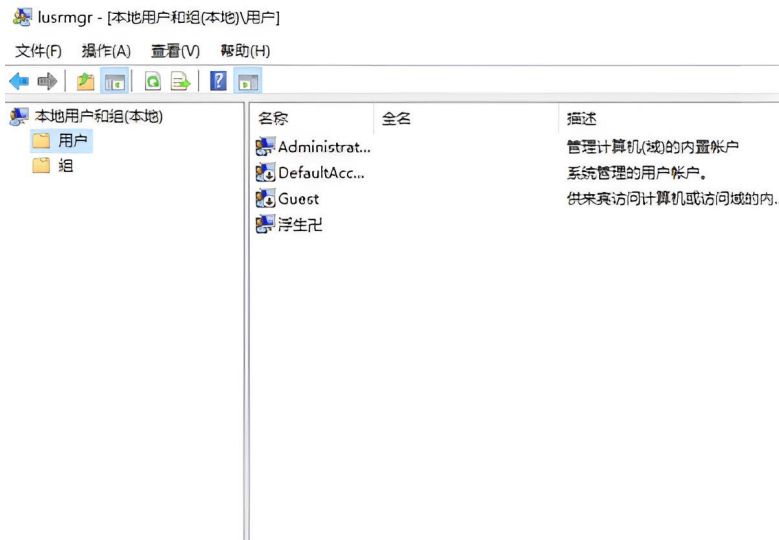


Figure 5. Special character diagram
图 5. 特殊字符图

机器视觉算法(Computer Vision, CV [17])一系列用于处理和分析图像数据的技术,它们在图像识别、分析和理解方面发挥着关键作用。在目标检测领域[18],算法能够识别图像中的特定对象,并确定它们的位置和大小。图像分类技术[19]则进一步将图像归入预定义类别中,这对于自动化处理数据至关重要。语义分割算法[20]不仅识别图像中的所有对象,还能区分每个对象的边界,从而提供更细致的图像理解。实例分割则在此基础上更进一步,能够区分同一类别中不同的个体,这对于复杂场景的分析尤为重要。目标跟踪算法在视频序列中跟踪特定对象的移动,这对于监控、安全和自动驾驶等领域至关重要。边缘检测技术通过识别对象轮廓来增强图像的边缘信息,这对于图像处理和特征提取非常有用。特征点检测与匹配技术,如尺度不变特征变换(SIFT) [21],能够检测图像中的关键点并进行匹配,这对于图像匹配和三维重建等领域具有重要意义。

在等保测评管理工具中,需要准确识别出用户名状态,在已有系统中,用户名及其状态以文字和图标形式存在,如图 5 所示。针对一般文本图像,光学字符识别(OCR) [22]可以把其中的文字识别出来,但是针对文本图像中的各种图标无法识别,用户状态(图标)无法获得。

为此,案例中提出了一种基于 PaddleOCR [23]和 PeleeNet [24]识别特殊字符的方法,首先使用 PaddleOCR 对特殊字符进行定位,再把特殊字符区域送入改进后的 PeleeNet 小网络中对其进行分类,从而解决了 PaddleOCR 不能识别特殊字符的问题。PaddleOCR 是一种基于深度学习的光学字符识别(OCR)算法。OCR 是一种将图像中的文字信息转换为可读文本的技术。PaddleOCR 使用 PaddlePaddle 作为其背后的深度学习框架,通过将图像输入到深度卷积神经网络(CNN) [25]中,对文字区域进行检测和识别,提供了具有高精度和鲁棒性的 OCR 解决方案,具有多场景支持、高性能和准确性、多语言支持、开放源代码等优点。PeleeNet 是一种轻量级的深度神经网络模型,专门设计用于目标检测任务。它采用了一系列创新的设计和优化策略,以在保持高准确率的同时,显著减小了模型的尺寸和计算量。它引入了双向密集层、根块、瓶颈中的动态通道数、转换层压缩以及常规的激活函数,以减少计算成本和提高速度。双向密集层有助于获得感受野的不同尺度,使更易于识别更大的目标。

但在实际使用 PeleeNet 对用户状态图标分类时存在一定难点:由于用户名状态图标像素点只有 23×23 左右,图片输入网络后经过卷积下采样后续的特征图只有 1×1 大小,通过把输入图像压缩至成 100×100 会导致原始图像信息很少,都会影响最终的识别率。因此,案例中提高了将改进后的 PeleeNet 小网络对用户状态进行识别的方法。首先使用 PaddleOCR 对文本图像进行字符检测,使用文字坐标来定位用户名状态图标,使用 opencv 把图标保存下来,而后对 Peleenet 骨干网络进行裁剪,把所有 stage2、stage3、stage4 的层及 stage1-tb/pool 删除,保留剩下的作为最终的 backbone,然后训练模型测试。模型改进前后网络结构如图 6、图 7 所示:



Figure 6. Network structure before improvement
图 6. 改进前的网络结构

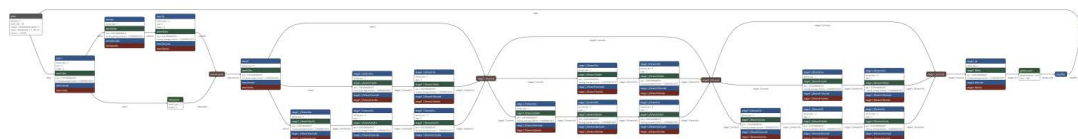


Figure 7. Network structure after improvement
图 7. 改进后的网络结构

在 ubuntu20.04、i5-11400、3070/8G 环境上、分别使用改进前及改进后的 PeleeNet 训练用户名状态数据,改进前模型针对 200 张用户启用状态图像识别率为 71.5%,消耗时间共 356 ms(图 8 左所示)、200 张用户禁用状态图像识别率为 78%,消耗时间共 359 ms(图 8 右所示)。改进后模型针对 200 张用户启用状态图像识别率为 95.5%,消耗时间共 155 ms(图 9 左所示)、200 张用户禁用状态图像识别率为 96%,消耗时间共 173 ms(图 9 右所示),不仅提高了识别率而且还缩短了识别时间。

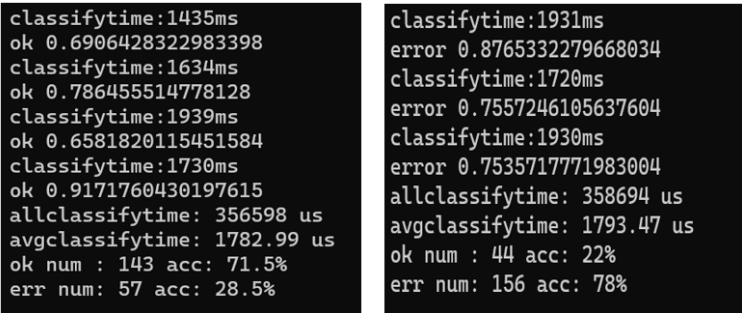


Figure 8. Test results of network model before improvement
图 8. 改进前网络模型测试结果

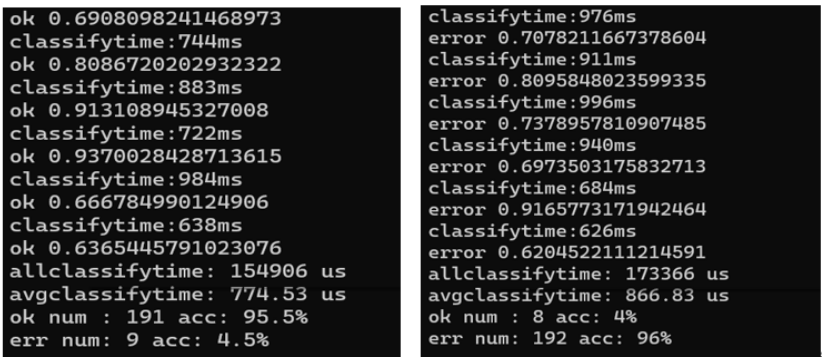


Figure 9. Test results of network model after improvement
图 9. 改进后网络模型测试结果

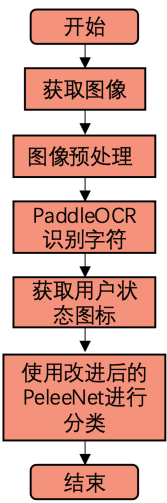


Figure 10. User status recognition flow chart
图 10. 用户状态识别流程图

该案例总体实现流程如图 10 所示。主要分为三大步骤：

1. 对用户状态图进行去噪、对比度增强操作；
2. PaddleOCR 对预处理后图像进行文字提取，定位状态图标区域；
3. 使用改进后的 PeleeNet 对用户状态图进行分类。

4. 总结及展望

本文深入探讨了人工智能技术在等级保护(等保)测评领域的应用及其重要性。通过分析传统测评方法的局限性，明确了引入人工智能技术的重要性。人工智能技术，特别是大语言模型(LLMs)、自然语言处理(NLP)和机器视觉算法(CV)，已经在等保测评的多个方面展现出显著优势。在案例中，通过自动化生成测评报告，提高了测评效率和准确性，减少了人为主观因素的影响；在风险等级评定中，利用深度学习模型捕捉文本中的深层次语义关系，提高了风险识别的精度和召回率；为解决特殊字符识别的问题，通过结合 PaddleOCR 和改进后的 PeleeNet，提升了用户状态识别的准确率和效率。这些技术的应用不仅提高了等保测评的自动化和智能化水平，还为测评师提供了更高效、更准确的工具，从而更好地应对日益复杂的网络安全挑战。

尽管 AI 技术在等保测评领域取得了显著进展，但仍面临一些挑战和未来的发展方向，具体如下：

- 1) 数据质量和隐私保护：AI 模型的训练需要大量高质量的数据。如何确保数据的质量和隐私保护，同时满足测评需求，是一个亟待解决的问题。
- 2) 模型的可解释性和透明度：AI 模型的决策过程往往被视为“黑箱”。提高模型的可解释性，使测评师能够理解和信任 AI 的决策，是提升 AI 技术在等保测评中应用的关键。
- 3) 跨领域融合：AI 技术在等保测评中的应用需要与其他领域，如云计算、物联网等，进行更深入的融合，以应对不断变化的安全威胁。
- 4) 持续的技术创新：随着 AI 技术的不断发展，新的算法和模型将不断涌现。持续的技术创新和应用研究将推动等保测评向更高水平发展。

随着算法和计算能力的进一步提高，人工智能在网络安全等级保护测评中的应用将更加广泛和深入。自主测评系统有望成为现实，能够独立完成从测试设计到结果分析的全流程工作，大幅提升测评效率和准确性。AI 驱动的智能威胁预测将更加精准，通过分析海量数据和复杂模式，提前预警潜在风险，使防护措施从被动响应转向主动预防。跨领域知识整合将更加深入，AI 系统将能够综合利用网络安全、法律法规、业务流程等多方面知识，提供更全面的安全评估和建议。人机协作模式将更加成熟，AI 系统将成为安全专家的得力助手，辅助决策并处理复杂任务。随着可解释 AI 技术的发展，测评结果的可信度和透明度将显著提高，有助于提升等级保护测评的权威性和说服力。这些发展将推动等级保护测评向更智能、更精准、更高效的方向演进，为构建安全可靠的网络空间提供有力支撑。

总之，AI 技术在等保测评领域的应用前景广阔，但同时也需要在技术、法规和伦理等多个层面进行深入研究和探索，以实现其在信息安全领域的最大化效益。

参考文献

- [1] 肖国煜. 信息系统等级保护测评实践[J]. 信息网络安全, 2011(7): 86-88.
- [2] 郑南宁. 人工智能新时代[J]. 智能科学与技术学报, 2019, 1(1): 1-3.
- [3] 郭天霞. AI 对工作流程和工作质量的影响: 自动化, 错误减少和优化[J]. 文化与社会心理学, 2023(12): 1-7.
- [4] 王志宏, 杨震. 人工智能技术研究及未来智能化信息服务体系的思考[J]. 电信科学, 2017. 33(5): 1-11.
- [5] Yao, Y., Duan, J., Xu, K., Cai, Y., Sun, Z. and Zhang, Y. (2024) A Survey on Large Language Model (LLM) Security and Privacy: The Good, the Bad, and the Ugly. *High-Confidence Computing*, 4, Article 100211.

- <https://doi.org/10.1016/j.hcc.2024.100211>
- [6] Vaswani, A. (2017) Attention Is All You Need. <https://arxiv.org/abs/1706.03762>
 - [7] An, J., Ding, W. and Lin, C. (2023) Chatgpt: Tackle the Growing Carbon Footprint of Generative AI. *Nature*, **615**, 586-586. <https://doi.org/10.1038/d41586-023-00843-2>
 - [8] Touvron, H., et al. (2023) Llama: Open and Efficient Foundation Language Models.
 - [9] Bai, J., et al. (2023) Qwen Technical Report.
 - [10] Glm, T., et al. (2024) ChatGLM: A Family of Large Language Models from GLM-130b to GLM-4 All Tools.
 - [11] Ding, N., Qin, Y., Yang, G., Wei, F., Yang, Z., Su, Y., et al. (2023) Parameter-Efficient Fine-Tuning of Large-Scale Pre-Trained Language Models. *Nature Machine Intelligence*, **5**, 220-235. <https://doi.org/10.1038/s42256-023-00626-4>
 - [12] Mao, Y., et al. (2024) A Survey on Lora of Large Language Models.
 - [13] Yang, A., et al. (2023) Baichuan 2: Open Large-Scale Language Models.
 - [14] Devlin, J. (2018) Bert: Pre-Training of Deep Bidirectional Transformers for Language Understanding.
 - [15] Patil, R., Boit, S., Gudivada, V. and Nandigam, J. (2023) A Survey of Text Representation and Embedding Techniques in NLP. *IEEE Access*, **11**, 36120-36146. <https://doi.org/10.1109/access.2023.3266377>
 - [16] Yacoub, R. and Axman, D. (2020) Probabilistic Extension of Precision, Recall, and F1 Score for More Thorough Evaluation of Classification Models. *Proceedings of the First Workshop on Evaluation and Comparison of NLP Systems*, Online, November 2020, 79-91. <https://doi.org/10.18653/v1/2020.eval4nlp-1.9>
 - [17] Voulodimos, A., Doulamis, N., Doulamis, A. and Protopapadakis, E. (2018) Deep Learning for Computer Vision: A Brief Review. *Computational Intelligence and Neuroscience*, **2018**, 1-13. <https://doi.org/10.1155/2018/7068349>
 - [18] Kaur, R. and Singh, S. (2023) A Comprehensive Review of Object Detection with Deep Learning. *Digital Signal Processing*, **132**, Article 103812. <https://doi.org/10.1016/j.dsp.2022.103812>
 - [19] Sharma, S. and Guleria, K. (2022) Deep Learning Models for Image Classification: Comparison and Applications. *2022 2nd International Conference on Advance Computing and Innovative Technologies in Engineering (ICACITE)*, Greater Noida, 28-29 April 2022, 1733-1738. <https://doi.org/10.1109/icacite53722.2022.9823516>
 - [20] Li, L., Zhou, T., Wang, W., Li, J. and Yang, Y. (2022) Deep Hierarchical Semantic Segmentation. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, New Orleans, 18-24 June 2022, 1236-1247. <https://doi.org/10.1109/cvpr52688.2022.00131>
 - [21] Tsourounis, D., Kastaniotis, D., Theoharatos, C., Kazantzidis, A. and Economou, G. (2022) SIFT-CNN: When Convolutional Neural Networks Meet Dense SIFT Descriptors for Image and Sequence Classification. *Journal of Imaging*, **8**, Article 256. <https://doi.org/10.3390/jimaging8100256>
 - [22] Mori, S., Suen, C.Y. and Yamamoto, K. (1992) Historical Review of OCR Research and Development. *Proceedings of the IEEE*, **80**, 1029-1058. <https://doi.org/10.1109/5.156468>
 - [23] Peng, Q. and Tu, L. (2023) Paddle-OCR-Based Real-Time Online Recognition System for Steel Plate Slab Spray Marking Characters. *Journal of Control, Automation and Electrical Systems*, **35**, 221-233. <https://doi.org/10.1007/s40313-023-01062-w>
 - [24] Wang, R.J., Li, X. and Ling, C.X. (2018) PELEE: A Real-Time Object Detection System on Mobile Devices. *Advances in Neural Information Processing Systems*, **2018**, Article No. 31.
 - [25] Lecun, Y., Bottou, L., Bengio, Y. and Haffner, P. (1998) Gradient-Based Learning Applied to Document Recognition. *Proceedings of the IEEE*, **86**, 2278-2324. <https://doi.org/10.1109/5.726791>