

致课题组未来的学生

如果你有兴趣加入我的课题组，成为其中一名学生（包括**硕士/博士研究生**或者**科研助理**），请认真阅读以下材料，然后按要求提交申请。

1. 自我介绍：

我是杭州电子科技大学网络空间安全学院副教授乔通，硕士生导师。更多信息可以从我的[学院主页](#)和[个人主页](#)查到。

2. 课题组简介：

我的课题组属于[浙江-法国数字媒体取证联合实验室](#)。每位老师负责一个课题，指导学生开展研究。我目前是实验室的负责人，目前围绕人工智能在媒体领域的安全应用，开展以下四个方向研究：

（1）媒体伪造鉴别

通过传统、新型检测技术，识别和定位伪造的数字媒体内容。

（2）媒体追踪溯源

利用设备、模型和身份指纹，追踪伪造媒体的来源，识别生成设备或模型。

（3）媒体智能合成

利用人工智能技术生成语音、图像和视频，支持人脸替换、表情重现等技术。

（4）媒体智能安全

通过水印嵌入等方法，加强对生成内容的监管，保障生成模型的安全合规应用。

具体研究方向细节查看课题组主页 <https://dmf.hdu.edu.cn/> 阅读发表的论文。

3. 课题组目前开放学生岗位：

（1）硕士/博士研究生

面向获得考研复试资格的网络空间安全、计算机科学与技术、电子信息或相关学科的**优秀本科生/硕士生/博士生**。

（2）科研助理

面向对课题组研究方向感兴趣的**优秀在校本科生和硕士生**。

4. 课题组收获:

我会一对一指导课题组每一位学生, 提供舒适的工作环境和丰厚的助研津贴, 并资助学生积极参与国内和国际交流, 在杭电的平台上为每位学生提供最大的助力。

5. 课题组要求申请学生必须具备的能力:

(1) 目标明确及热爱所学

明确自己为什么要加入我的课题组, 明确自己在课题组想要获得什么。要热爱自己的学科和专业, 对将来所学知识表现出极大的兴趣, 要让真正的兴趣和喜爱来内驱自己不断前进。

(2) 乐观开朗及坚韧不拔

懂得科研不是一帆风顺, 要坐得住板凳和耐得住寂寞, 要注意点滴积累。

(3) 勇于创新及乐于挑战

在探索未知事物中, 充满好奇心和求知欲, 富有怀疑的精神, 乐于尝试, 并以正确心态面对挫折。具有良好的团队合作精神, 乐于阐述自己的想法, 有效与导师和其他伙伴沟通。

(4) 数学扎实及编程熟练

具有扎实的数学功底和钻研精神, 熟悉 C 或 C++ 或 Python, 至少用过一门深度学习框架语言, 比如 PyTorch、Tensorflow、Keras, 并有视频图像处理等科研项目经验。

6. 申请加入课题组流程:

(1) 阅读报告: 在满足以上具备能力要求的基础上, 了解目前课题组的任意一个研究方向, 阅读至少一篇我近五年来发表的论文, 写一份不超过一页的阅读报告, 阐述自己的理解和想法。

(2) 个人简历: 将上述报告、个人简历、成绩单一起发到 tong.qiao@hdu.edu.cn。简历不一定特别漂亮, 但请排版工整, 内容清晰明了。邮件标题为【申请读研-学生姓名-学校名称】。

(3) 远程面试: 对于我认为合适的学生, 会在三天内回复, 并约定时间, 发出远程面试邀请。对于未按照规定命名邮件标题和未认真阅读本材料的邮件, 概不回复。

每年名额有限, 若感兴趣, 请早做准备提前联系! 期待我们一起在科研道路上共同奋战!

乔通
2025. 03. 04

中法数字媒体取证 联合实验室



成员招募 2024 DMF-LAB

团队简介

中法数字媒体取证联合实验室成立于2017年，由杭州电子科技大学与法国特鲁瓦科技大学联合创建，依托浙江省首批网络安全学院，数据安全治理省工程中心，信息化实验平台和法国国家科学研究中心（Centre National de la Recherche Scientifique, CNRS）的网络空间安全平台“Plateforme CyberSec”，由教育部黄大年式教师团队、海外优青、国家高端外专、浙江省高校领军人才等师资队伍共同支撑。

研究方向

● AIGC安全

● 来源取证

● 伪造检测

● 主动防御



联系方式



所在位置

一教南607



邮箱地址

tong.qiao@hdu.edu.cn

招生要求

面向网络空间安全、信息安全、人工智能、网络工程等计算机相关专业及方向

具有良好的数学基础、编程技能，熟悉图形学、图像处理、机器学习、深度学习、信息论与编码者优先

有快速学习新知识的能力、优秀的团队精神、良好的表达沟通和协调能力

热爱科研工作且具有较强的自我驱动能力

我们热忱欢迎对人工智能领域感兴趣的本科生、硕士及博士研究生通过邮件自荐，加入我们的学术团队。期待您的精彩加入！



研究方向及成果介绍

生成式人工智能安全

AIGC Security

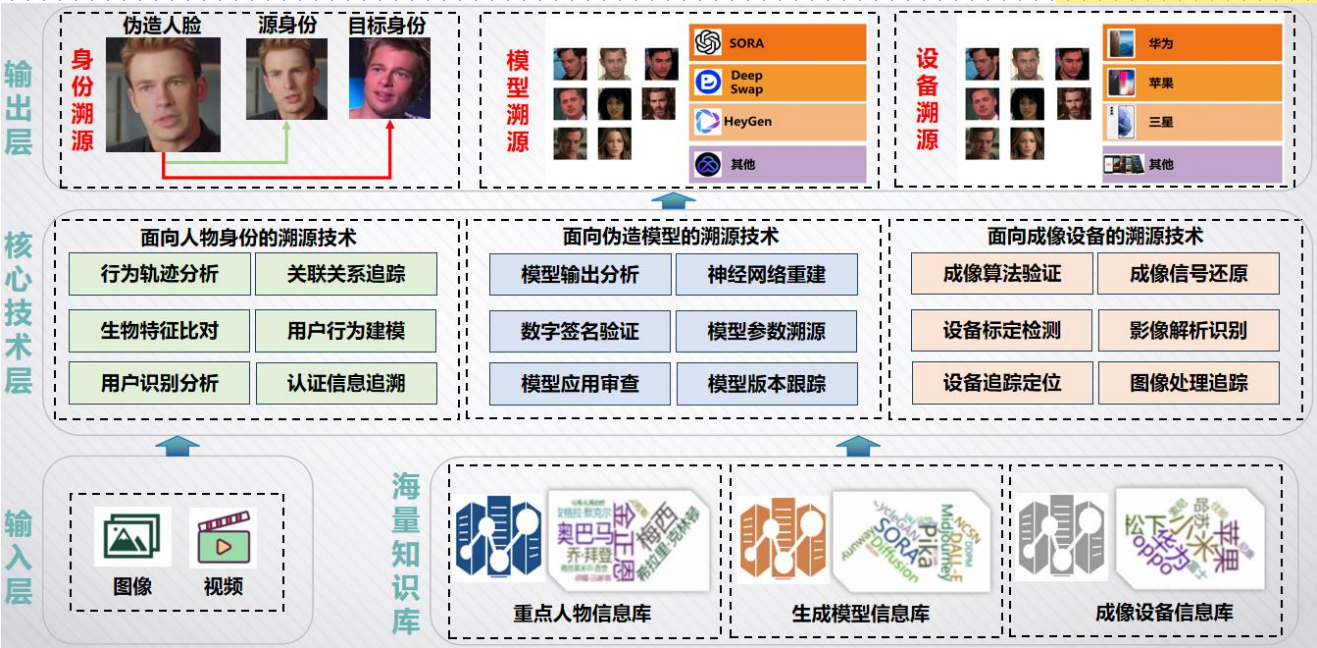


深度学习技术在生成式人工智能安全领域取得了显著的成功。其主要应用领域涵盖深伪取证、活体检测和模型水印等任务。深伪取证专注于对数字内容进行取证，旨在检测和识别深度伪造内容。借助深度神经网络，可以高效地进行深伪取证工作。活体检测的主要目标是辨别真实的活体目标。深度神经网络能够准确地捕捉真实活体目标的特征和行为，为活体检测提供可靠的技术支持。模型水印是在AI系统中嵌入隐藏信息，以验证系统的真实性和可靠性。深度学习技术在解决模型水印任务方面表现出色，能够在保持模型性能的同时，训练带有模型水印的模型。通过正确的水印标签，即可完成模型的验证。

研究方向及成果介绍

来源取证

Source Attribution

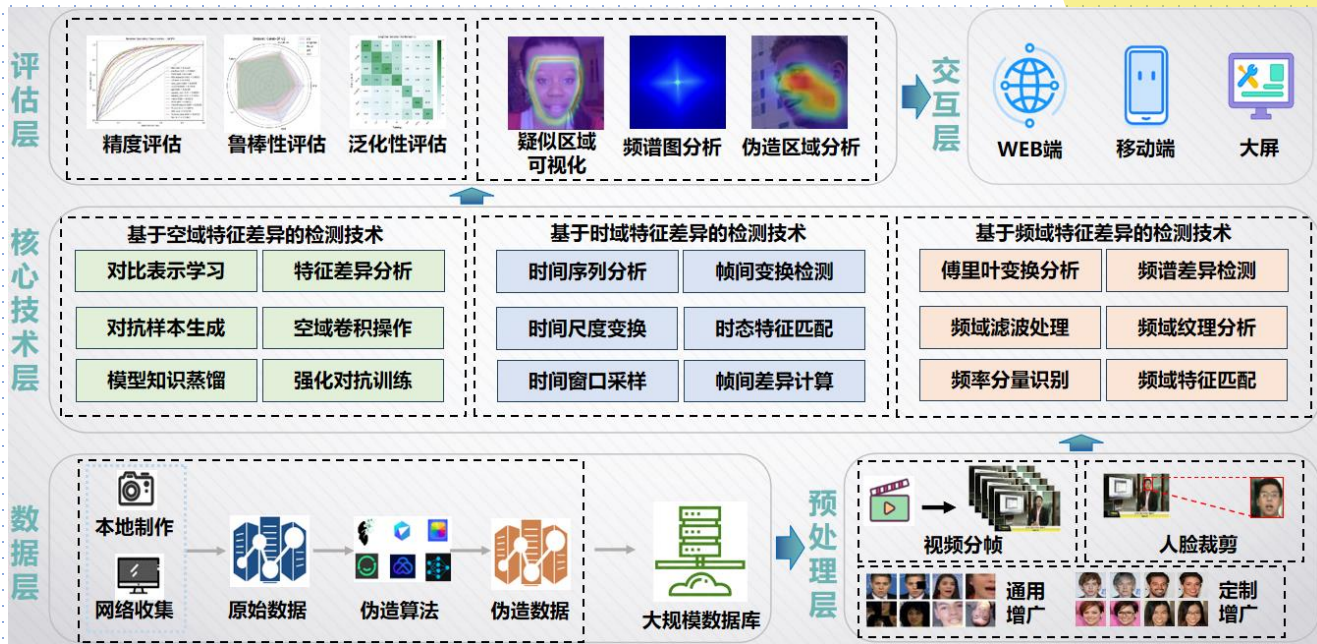


来源取证是计算机科学和法律领域中一项重要的研究方向，旨在通过结合技术手段和法律手段，明确数字证据的来源和可信度，从而应对计算机犯罪、网络欺诈、知识产权侵犯等问题。该领域的研究具有重要的实际应用价值。在数字时代，各类数字证据被广泛应用于法律案件、安全调查、商业纠纷等多个领域。通过来源取证的研究，我们能够更加有效地确认数字证据的真实性和可信度，为解决涉及数字证据的各种问题提供有力的支持。这对于保障司法公正、确保安全调查的准确性，以及防范商业纠纷的发生都具有重要意义。同时，来源取证的研究也有助于推动相关法律法规和标准的制定和完善，提高数字证据在法律程序中的法律效力和可信度。

研究方向及成果介绍

伪造检测

Tampering Detection

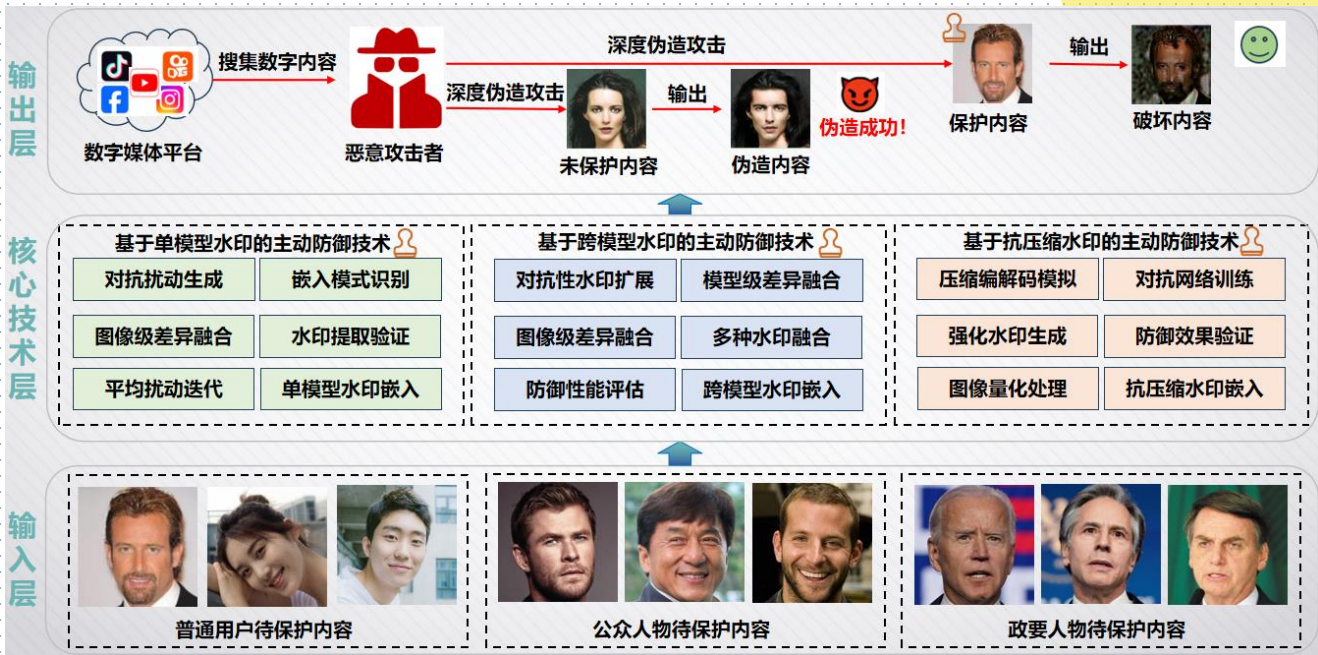


伪造检测是指利用计算机技术和图像处理算法，对图像进行自动分析，以识别和检测图像是否被篡改或修改过的过程。图像伪造可能包括图像拼接、复制粘贴、对象移除、色彩篡改等多种手段。图像伪造检测技术在许多领域都有应用，如数字取证、安全监控、社交媒体内容真实性验证等。图像伪造检测的关键在于提取图像的特征并比较它们之间的差异。这些特征可能包括像素值、纹理、颜色、形状等。通过比较原始图像和疑似伪造图像之间的特征差异，可以检测出图像是否被篡改。此外，基于深度学习的算法，研究人员可以通过训练模型来识别和检测图像中的伪造痕迹。

研究方向及成果介绍

主动防御

Active Defense



深度伪造主动防御通过在人脸图像或视频中加入特定的扰动或水印，从而达到破坏深度伪造模型的性能或对伪造图像溯源、认证的目的。深度伪造主动防御技术初步实现了深度伪造的“事前防御”，当用户希望自己上传到互联网的人脸图像或视频不被恶意用户用于伪造，则可以在上传之前对这些图像或视频嵌入特定的保护性扰动或水印。随着未来研究的深入，深度伪造主动防御技术将有望在实际的社交多媒体场景中使我们的人脸图像和视频免受深度伪造的威胁，构建更加安全和健康的互联网环境。



DMF-LAB

期待你的到来

顶级平台

中法数字媒体取证联合实验室作为浙江省级平台，拥有一流的科研平台和完善的配套设施。同时，实验室提供中法线上线下交流平台，与法国特鲁瓦科技大学共享学术资源，提供公费赴法交流学习的机会，以及进一步赴法留学深造的宝贵机会。



丰硕成果

联合实验室已在国内外知名刊物及会议录上公开发表了100余篇论文，涵盖多篇CCF A类等高质量期刊论文。

未来发展

联合实验室拥有优秀的师资力量、能够在不同的方向指导学生从事科研工作。现已经同法国特鲁瓦科技大学联合培育5名博士、10名硕士。学生在众多学术竞赛与项目申报中脱颖而出，屡获殊荣，包括多个重量级奖项与项目申报，毕业生先后拿到了浙江大学、阿里巴巴、百度、华为、字节跳动的录取通知。实验室大力支持团队成员进行赴法访学交流，拓展国际视野。