

数智化时代教育伦理问题的法律规制

史立民

广西师范大学教育学院，桂林

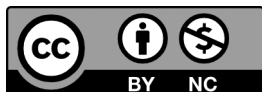
摘要 | 人工智能是一种模仿并扩展人类智能的技术，其应用领域广泛。它的发展不仅代表着科技的进步与突破，更是各行各业借助人工智能这一新质生产力带来益处的不断进行迭代、更新。在众多领域中，教育领域对其应用最为广泛，学生作为人工智能的使用者，常借助其进行辅助学习、开展虚拟实验，以及使用智能推荐学习系统等，进而通过跨界融合创新与发展自身的专业。然而，其背后存在的问题仍需进一步探索并进行相关规制。

关键词 | 人工智能；智能；辅助学习；创新发展

Copyright © 2024 by author (s) and SciScan Publishing Limited

This article is licensed under a [Creative Commons Attribution-NonCommercial 4.0 International License](https://creativecommons.org/licenses/by-nc/4.0/).

<https://creativecommons.org/licenses/by-nc/4.0/>



1 问题缘起

科技时代掀起数据技术热潮，教育智能体——人工智能以其内在的高效率、自动化及精准性等特征，通过多种样式的人工智能功能，如机器学习、深度学习、自然语言处理、计算机视觉、专家系统、知识图谱、情感分析等，助力学子畅游知识海洋^[1]。《中国教育现代化2035》《加快推进教育现代化实施方案（2018—2022年）》跨越信息化时代教育变革，致力于运用现代技术加快推动人才培养模式改革，实现规模化教育与个性化培养的有机结合。作为时代的产物，人工智能是科技进步的典型代表，亦逐渐演变为教育领域炙手可热的重点关注工具，其热度高居不下。教育迎来的必将是数智

化时代，而其带来的教育伦理问题日益突显，精准驾驭人工智能所带来的机遇与挑战已然成为当务之急。^[2]本文旨在阐述人工智能在教育领域中的道德伦理现状及问题，探究其背后缘由，并为法制处理手段提供相应建议。

2 数智化时代多重样式人工智能教育伦理问题折射阐述

数智化时代，泛指数字化和智能化相结合的时代，数据是重要的驱动力，其数字化技术涵盖云计算（cloud computing）、大数据分析（Hadoop）、人工智能和物联网等的发展与应用，已成为社会发展和人们生活的基石。智能化技术的广泛应用为人

作者简介：史立民，广西师范大学在读研究生，研究方向：教育法学。

文章引用：史立民：数智化时代教育伦理问题的法律规制[J]．教育研讨，2024，6（5）：1269-1276.

<https://doi.org/10.35534/es.0605168>

们提供了极大便利,但其多重样式的人工智能同步发展也为数智化时代人工智能的应用带来挑战。在此,多重样式又可分为狭义、通用、超级人工智能。这里主要是在微观的层面上研究其所带来的违背人自主性、价值观误区、知识产权、伦理问责等教育伦理问题。

2.1 单一式人工智能自动化决策背离人的自主性

“单一式人工智能”(Monolithic AI)这里指专门设计用于执行特定任务或解决特定问题的人工智能系统。其自动化决策受制于以“大数据,小定律”技术范式为核心的概率统计逻辑以及追求效率而非绝对精确、强调相关性而非因果性的效率主导逻辑影响,往往会出现人的主体性丧失和决策结果不公正等问题。^[3]本文探究重点是其在教育领域中所涉及的伦理现象。

虽然不可否认人工智能技术的应用使教学从单一讲授形式转变为研究型的多样化教学,摆脱了传统时空固有模式的局限,具有积极的教育变革意义。^[4]但其本质是从冗杂的庞大数字系统内容中挑选并为学生提供所需的自学方式,其中夹杂着大数据贯穿的不确定性以及风险控制的不完全性,由此造成的困扰触及教育伦理中最基础的原则——人的自主性。例如,现有小学生开始充当人工智能学习博主,三年级学生参加ChatGPT夏令营,并且在自主学习的过程中,以人工智能充当辅助完成本该独立完成任务。设计伦理学家特里斯坦·哈里斯断言人类的自控力和决策能力正受到计算和诱惑的干扰。^[5]此思想和行为的干预过程,使学生对问题的主动性思考和多元角度的探索逐渐减少,该群体更依赖已有数据技术的支持,而忽略了自身的创新思维,继而导致其创造力和想象力匮乏,禁锢在常规机械思维之中。^[6]

2.2 多样式人工智能引发价值观传播误区

多样式人工智能(Multi-modal Artificial Intelligence,简称Multi-modal AI)指代一种能够处理、理解并整合来自不同模态信息的人工智能技术。面对社会变迁和不确定性,价值观的探讨变得尤为关键。价值观是人们内在信念的体现,指引着我们的行为和决策。价值对齐是人工智能领域一个日益受到关注的研究议题。斯图尔特·罗素(Stuart

Russell)在《人工智能的神话》(The Myth of AI)一文中强调了价值对齐问题的重要性。他提出,我们的目标是开发能够与人类价值观保持一致的智能系统,而不仅仅是追求智能本身。对齐的核心在于确保机器的动机与人类的意图和价值观相匹配,而不仅是它们所掌握的知识或能力。^[7]

人工智能在教育领域充当辅助工具而非替代教师,其效用取决于使用者的意图和应用方式。学生在使用过程中可能无法形成独立的判断力和正确的价值观,或对事物看法的价值观念逐渐模糊或偏离社会主流价值观。故此,学生群体应通过系统搜集信息来构建知识网络,而非依赖零散答案,以培养批判性思维和逻辑推理能力。^[8]在学习过程中,智能系统的统一性使得学生在长期人机交互中耳濡目染,同质化严重。学生过度依赖智能体传授知识,可能会传播狭窄的价值观,或者形成智能化机器传播优于教师知识传授的观念。

2.3 生成式人工智能引发剽窃隐患、原作者知识产权受损

生成式人工智能(Generative Artificial Intelligence,简称GAI)是一种利用复杂的算法、模型和规则,从大规模数据集中学习以创造新的原创内容(AI-generated content,简称AIGC)的人工智能技术。AIGC系统的输出通常是基于大量的输入数据生成,因此很难确定其是否为原创内容,这可能导致原作者的作品被误认为是人工智能生成,从而进一步损害他们的声誉和权益。AIGC技术是分析学习现有文本生成新文本内容,现有材料来源广泛但不全是公开免费,倘若侵犯受版权保护的材料也会侵犯他人知识产权。^[9]另外,AIGC输出、传播阶段,演绎权侵权也是其一。演绎作品是通过改编、汇编或翻译等方式对现有作品进行创新性再创作而形成的新作品。当AIGC技术生成的内容具有独创性,同时又保留了原作的基本表达时,它可能被视为演绎作品。如果AIGC的开发者或服务提供者未获得原作品权利人的授权,这些内容可能同时具有侵权的法律风险。^[10]除AIGC生成作品本身可能侵犯原作者知识产权外,其本身也是知识产权的一部分,遭受他人的恶意篡改也会使原作者权益受损。例如,斯坦福某AI团队被曝抄袭面壁智能开源成果,有“套壳”嫌疑:模型结构和代码展现出人

人的相似度,经实证该项目使用的模型结构和代码与面壁智能不久前发布的MiniCPM-Llama3-V2.5呈现出高度的相似,仅修改了部分变量名。^[11]

2.4 管理式人工智能伦理问题问责主体探讨

管理式人工智能(Managing Artificial Intelligence)在教育领域是指将人工智能技术集成到教育管理体系中,以实现更高效、智能化的教育管理和决策过程。这种应用不仅限于教学层面,还涵盖教育机构的行政、资源配置、学生事务管理、教学质量监控等多个方面。然而,这类人工智能存在各类伦理道德问题。例如,在智能化辅导领域,ALEKS(Assessment and Learning in Knowledge Spaces)是一个基于知识空间理论(Knowledge Space Theory, KST)的系统。初始知识评估伊始,通过选择题测试确定学习者的知识水平,并据此创建个性化的学习路径。^[12]整个过程中需要收集学生的数据,学生的数据如果没有得到妥善保管或者受到妥善保管但未经授权被应用于营销活动或研究,仍然存在操控学生学习行为的可能性。另外,学生数据的准确性有待考量,其成绩、作业的真实性不一,教师的评判标准也不一致,这些都可能受到人工智能模型结果的误导。

3 科技缺陷是数智化时代人工智能教育伦理问题的根源所在

人工智能广泛应用,在见证其益处的同时,也需清晰地认识到其固有缺陷。而人工智能内在的缺陷问题是促使其在教育领域伦理问题显露的主要原因。本文重点阐述以下四点问题。

3.1 信息过滤及资源预测诱发个体信赖人工智能而忽略主观判断

信息过滤及资源预测是将信息传递给需要它的用户处理过程的总称。具体而言,信息过滤系统是一个针对非结构化或半结构化信息的系统,主要处理文本信息及巨大的数据量。首先进行数据的收集,接着进行数据预处理,然后过滤信息,通过设定关键词筛选相关内容,并对这些内容进行量化分析(如使用SPSS、R、Python的Pandas和NumPy等)、计算机辅助内容分析(Computer-Aided Content Analysis, CACA)、深度学习(如卷积神经网络(CNN)、循环神经网络(RNN)、短时记

忆网络(LSTM)等),以对未来教育资源的需求和供应情况进行合理过滤和预测。

在教育领域中,资源预测即指收集大量数据和信息,对教育资源的未来需求和供应情况进行合理推测和预测。这些教育资源包括但不限于学生、教师、教室、教学设备、教材、经费等。资源预测的目的是为了提前了解未来一段时间内教育资源的需求状况,以便教育管理者和决策者能够做出合理的规划和决策,确保教育资源的有效分配和利用,这里特指学生学习资源预测。虽然信息过滤能够最大限度地满足学习者对知识的需求速度,资源预测能够极致化地预判教育资源地所需所求,但背后也存在智能系统本身的弊端。在信息过滤的过程中,只能对现有的知识逻辑进行筛选,而不能对知识之间未来可能存在的关系作出合理的解读和预测,因而就会使得可能存在关系的知识体系被自动过滤,过度依赖智能系统的人也自然无法洞察这层关系,对无止境的知识探索缺少一份思辨。信息过滤过滤掉的只是关于知识的条件设置,而忽略了未知知识的关联性。

3.2 算法缺陷及数据偏见致使个体获取资源价值取向有失偏颇

所谓算法(Algorithm)是指依靠固有程序指令解决现存问题,并通过完整方案进行呈现。^[13]算法公平是人工智能发展中的关键伦理追求,在教育领域尤为关键。它涉及诸如算法偏见、歧视和不透明性等问题。人工智能算法的先进性为教育领域带来了科学和安全的解决方案,但它们的固定输入、规则和输出模式可能导致教育内容的表面化。同时,“黑箱”现象、算法偏见和算法推荐可能对教育产生负面影响,亟需我们给予足够的关注和审慎的处理。^[14]算法偏见通常源于研发者的价值观和偏见,这种偏见通过代码编写过程被嵌入算法中,可能导致不公平的政策决策和严重的社会后果。^[13]因此,算法通常设定的价值观呈现是部分而非全面覆盖,形成一种算法认知结构的固化。对于只传播具有正负方向价值观的人工智能,亦使得教育领域学生在学习过程中无法进行正反价值观的判断。

3.3 智能摘取及未经授权复制传播他人版权未进行编著署名

人工智能资料智能摘取(Intelligent

extraction), 或自动文摘 (automatic abstracting), 通过AI技术从文本中自动提取关键信息, 生成简洁摘要, 并可用于信息提取 (Information extraction)、文本总结 (Text summary)、关键词提取 (Keyword extraction)、内容分类 (Content classification)、自动化处理 (Automated processing) 和可视化呈现 (Visual presentation) 等功能。若智能摘要技术被用于非法获取或复制版权内容, 将侵犯原作者权益, 可能导致原创力受损、内容真实性受质疑、创作环境被扭曲以及传播方式发生变异。^[15]在提取和学习真实数据后, 虽可对内容进行评估和优化, 但即使数据完全真实, 也存在生成虚假信息风险。一旦这些虚假信息被传播, 就可能误导用户决策, 造成错误的判断和后果。^[16]

未经授权复制传播他人版权内容且未进行编著署名, 是对版权人合法权益的严重侵犯。这不仅损害原作者的合法权益, 而且破坏知识产权保护体系, 进而削弱创作者的创新动力。在教育领域, 该行为会误导公众对其作品的来源和真实性的认知, 削弱学术研究的严谨性和可信度, 影响学术生态的健康发展。更有甚者, 会破坏学术界的市场秩序公平竞争, 损害国家形象和权益。

3.4 问题显露无法追责于人工智能本身, 但又无法精准地归咎于某群体或个体

人工智能在道德责任归责方面确实面临特殊挑战。由于它们既不具有与人类相同的道德主体地位, 也不同于一般机器的无道德主体地位, 这导致了责任归属的模糊, 从而产生了所谓的责任鸿沟。^[17]因此, 当人工智能道德伦理问题浮现时, 需对具有主体性的个体或群体进行追责。例如, AI开发者和设计师负责AI系统的创造与道德原则的设置, 对其决策和行为负有责任。另外, 运维人员虽不参与设计, 但对系统运行中的伦理问题需承担相应责任。AI研发团队和监管机构, 如政府, 都应负责确保技术遵循道德标准, 监督并预防潜在的伦理风险。

4 人工智能教育伦理问题的法律规制

法律作为限制个体和群体违反道德伦理的最后

底线, 它的绝对权力是发展社会主义市场经济的最大保障, 也是遏制数字领域违纪行为的底牌。通过法律形式约束违德事件, 并在事发前和事发后最大限度地保障个人参与人工智能自主化决策、防止算法偏见和数据歧视、实现知识产权最优化与全面保护、明确责任和追溯机制。2023年4月28日, 中共中央政治局召开会议, 分析研究当前经济形势和经济工作, 会议指出“要重视通用人工智能发展, 营造创新生态, 重视防范风险”, 阐释了社会主义核心价值观引领生成式人工智能立法的双重目标。法规表明对于生成式人工智能这种新技术形态, 我国已有具体规则进行规制。然而, 从长期来看, 在立法上要形成完善的规制体系, 仍有必要探讨基础的规制原则。^[18]

4.1 自动化决策进行时, 主体的自主性参与的法律监管

人工智能在自动化决策中应融入自主性, 实现从“物我交互”到“自我意识独立”的转变。^[19]《中华人民共和国人工智能法(学者建议稿)》(以下简称《学者建议稿》)第三十五条“人工智能决策解释权与拒绝权”规定, 当人工智能个人或组织的合法权益受到影响时, 主体有权拒绝人工智能作出的决定, 并重新进行个体判断参与。是以, 主体执行人机交互时自主性遭受侵害的情况应予以规制。

4.1.1 决策可视化依法阐释

智能科技的决策框架依托于数据驱动的逻辑演绎, 以产生决策性结论。数据固有的复杂性促使决策制定趋向于抽象化, 决策的制定是数据收集、信息整理、知识建构和决策输出的系统化序列, 个体对产出的实质接收有“钝感”, 而对图形建构化的信息接收速度是普通数据的十倍。因此, 对于个体而言, 就需要用通俗易懂的语言或者简化数据、非专业语言形式进行可视化处理, 让用户直观地感受到自己的每个选择是如何得出的, 并且每个决策对接下来的行为会产生何种作用。数据可视化的交互式操作, 使用户能够实时分析和展示数据, 直观识别关键信息, 从而为决策提供可靠支持。^[20]

4.1.2 人工智能操纵性的依法降低

依法促使人工智能的操纵性降低, 转变为辅助引导的形式, 而非结果的呈现。人工智能扶持

主体、辅助疏导思绪,从繁冗的选择中挑选出适合自己的途径,循法介入阈值,调控人工智能参与性,赋予学生核心地位。法律界定人工智能提供者设置可调节智能系统参数,个体输入数据后,人工智能依次对应的选择,实现数据的选择偏差,允许学生自行填入多元化数据,在结果呈现中行使自主决策,发挥人工智能最大化的辅助性。涉及标准的事物参与需要人类进行知识的判定与决策。

4.1.3 主体决策权依法保障

主体所做出过的选择如同一个“记号”刻印在大数据库中,当主体再次进行决策时,这类“记号”会自动调取,以“记号”形式“寄存”于网络数据之中。当主体进行决策时,人工智能会再次根据之前的选择来有所作为,此时决策显然忽略主体生理和心理变化,有违主体的自主性决策。因此,在进行决策权的实施时,除了主体的参与,还要对主体所作出的选择授予删除权,即:在主体进行每次的决策之后,主体是否删除已选数据。该删除权不仅主动排除自己隐私暴露的可能,此外还能有效防止接收人工智能之前曾经做过决策的可能,有效地保护主体决策权。

4.2 文字性传播与人类价值观适配原则的律例界定

《学者建议稿》第十七条“算法模型创新”中规定,国家加强算法模型创新,促进依法应用、推广和流通。人工智能主管部门指导行业组织制定人工智能算法模型推荐目录与合作指引,完善算法模型流通中的利益分享机制。国家支持相关主体开展基础模型创新,发展面向通用人工智能的基础价值观适配理论体系。

4.2.1 内容真实性按律审核

生成式人工智能的底层算法与当前的推荐型算法并不相同。后者往往通过数据输入与分析实施“用户画像”等预测标签的行为,而生成式人工智能算法则强调利用数据进行迭代学习,并根据外部反馈生成全新的数据内容。^[21]数据的迭代更新会对旧的内容以全新的“包装”或者“拼接”重新呈现,但其真实性却有待商榷。“包装”使得知识的真实性被技术侵蚀,虽在某种程度上保持了知识的美感,但挟持了内在真实属性。基于《学者建议

稿》,开发者的训练人员在进数据库的搭建与构筑时应设立子系统,并在数据的收集和输出时经由该子系统自动筛选真实信息,标注其来源,阻止尚未确定的内容输出,保障输出内容的客观性,不得进行有偏好性的排放。该子系统也允许将关键词之间的异质性以多元形式进行输送。

4.2.2 价值观按律导向与审查

《生成式人工智能服务管理暂行办法》(以下简称《暂行办法》)明确界定,提供和使用生成式人工智能服务,应当遵守法律、行政法规,尊重社会公德和伦理道德,坚持社会主义核心价值观。不论是原文忠实复述还是创造虚构演示,赋予人工智能呈现内容的能力需要契合人类的意图、偏好、道德观标准等。因此,“价值观对齐”对人机交互非常有必要。按照律法对人工智能的基础性价值观进行“训练”“监管”,对具有偏失价值观的智能系统做出惩治。定期对人工智能决策实施监管并进行评估,对于有违伦理价值观的情况需尽快改进迭代。督导者对伦理参数的制定执行实施奖惩政策,该政策需依法执行。

4.2.3 隐私权和名誉权按律维护

《网络安全标准实践指南—人工智能伦理安全风险防范指引》对于人工智能伦理安全风险防范4.1基本要求指出,应尊重并保护个人基本权利,包括人身、隐私、财产等权利。《暂行办法》第一章总则第四条规定,尊重他人合法权益,不得危害他人身心健康,不得侵害他人肖像权、名誉权、荣誉权、隐私权和个人信息权益。为保障人工智能产品在维护用户隐私权与名誉权方面的合规性,迫切需要构建一个包含严格标准制定、独立审查流程、权威监管机构以及持续监管机制的准入许可体系,以此作为产品市场准入的先决条件。

4.3 文章生成引发知识产权受侵犯的立法约束

《学者建议稿》第十一条“知识产权保护”规定,国家依法保护人工智能领域的知识产权,探索创新人工智能科研成果转化机制,健全人工智能领域技术创新、知识产权保护和标准化的互动支撑机制。第二十三条“知识产权保护”指出,国家建立完善训练数据、算法、人工智能生成内容等知识产权保护规则。

4.3.1 著作权立法归属与保护

著作权乃人产物合理之根基，只有对他人的作品心存敬畏，才能尊重他人的思想劳动成果。并且，确定著作权的所有权归属能够在其个体权益受损后第一时间维护自身利益。因此，需通过立法进行保护。其一，对人工智能违反著作权的行为给予明确的法律侵权规定，任何违法行为都应受到法律的约束。其二，严格规范作品开放渠道，确保作品发布前不侵犯他人版权，依法建立版权审查机制。其三，当侵权行为发生时，法律应提供经济赔偿、精神补偿、公开致歉、法律禁令、刑事制裁以及加强版权教育和监管等综合措施，确保原作者在侵权行为发生时得到公正的补偿和权利的恢复。

4.3.2 合理引用立法规范

《著作权法》第二十二条详细规定了合理引用的情形，其中第二款特别提到了“适当引用”的概念，即为介绍、评论某一作品或者说明某一问题，在作品中适当引用他人已经发表的作品。但应指明作者姓名和作品名称，并不得影响作品的正常使用或不合理地损害著作权人的合法权益，同时对文章的出处进行标注，如MLA、APA等格式。创作者需要通过著作权获得经济利益和声誉，法律应促使知识授权人鼓励开放授权和知识共享，可采用开放授权协议，如知识共享协议（Creative Commons），允许人工智能在遵守特定条件下使用其作品，同时保障著作权人的利益。

4.3.3 技术保护措施立法实施

人工智能基于多样文章或者文献内容之间的可衔接性进行拼接，一般情况下，评价者或者审查者无法判断其正误或者是否为拼接而成。因此，要设立知名版权检测工具，例如：Turnitin、Grammarly、Copyleaks、Crossref Similarity Check等，对原作者知识产权的内容进行识别保护，并且使该自动化监测与识别技术程序持续更新，留存保护技术痕迹。广泛使用学术诚信检测工具，检测学生作业、论文等内容的原创性，并与全球的数据库进行比对，以识别抄袭行为。

4.4 责任归属与问责机制的法则限制

责任归属追溯需明确界定，问责机制建立的首要措施即对智能系统操作者、提供者予以法则限制。通过对人工智能的责任规制降低其滥用率，进

而在伦理矛盾产生时能够剖析缘由，使得智能系统涉及者明确责任分担。

4.4.1 责任方以法则界定与分类

涉及人工智能的责任方应以法则的形式进行界定并进行不同的分类，以便进行责任的归属。人工智能涉及研发者、持有者、数据供应者、所有者、使用者和监管者等多方主体。通过将经济利益与责任紧密结合，构建链式责任分配机制，确保各方合理承担相应责任，避免单一主体负担过重，以促进人工智能的可控发展。^[22]

4.4.2 责任分配结构化担责

实现责任分工的最佳方式是建立结构化的责任分担体系。通过明确各方关联，对人工智能进行有效管理：技术研发者负责研发失误，技术持有者负责运营责任；数据供应者负责数据传输错误，数据所有者负责数据问题；技术使用者负责数据使用后果，技术监管者负责监管过失。法律应调整人工智能产品生产者与使用者所承担的主要义务大小。生产者应承担更多义务，而使用者则应承担更少或不承担主要义务。

4.4.3 问责程序链接透明度以责限制

美国政府问责办公室（GAO）为联邦机构提供了一个包含治理、数据、性能和监测的人工智能问责框架，以确保人工智能的透明和负责任使用。中国也随之发布《新一代人工智能治理原则》，旨在指导负责任的智能系统发展。^[23]制定详细的操作规程和使用指南，确保人工智能系统按照既定的规则 and 标准运行。设立专门的监控机构或者机制，对人工智能系统的运行状态进行实时监控，及时发现并处理异常情况。按照责任的大小程度对问责分类分级监管，阐明问责程序的透明度，透明度未达标准的不能在市面上出现。

5 结语

在数字化与智能化的交汇点，人工智能技术在教育领域的渗透引发了深刻的变革，为学习者带来了多样化的学习工具，极大地扩展了教育的边界。然而，这种技术的进步也伴随着一系列伦理挑战，如个体自主性的削弱、价值观的误导、知识产权的侵权以及伦理责任的不明确。为了应对这些问题，需要采取法律规范、伦理教育以及明确技术责任等综合措施，以确保人工智能技术的应用不会侵

蚀教育的核心价值和伦理标准。通过持续的研究、法律的完善和伦理意识的提升,我们可以共同塑造一个既利用人工智能的优势,又维护教育伦理的新时代,从而保障教育的健康发展和学习者的全面成长。

参考文献

- [1] 毛楷仁,舒馨月.人工智能背景下对于教育管理的思考[J].现代商贸工业,2024,45(4):192-194.
- [2] 赵磊磊,张黎,代蕊华.教育人工智能伦理:基本向度与风险消解[J].现代远距离教育,2021(5):73-80.
- [3] Viktor Mayer-Schönberger, Kenneth Cukier. Big Data: A Revolution that Will Transform How We Live, Work and Think [M]. Boston: Houghton Mifflin Harcourt Publishing Company, 2013.
- [4] 辛继湘,李瑞.人是技术的尺度——智能教学中人的主体性危机与化解[J].中国电化教育,2023(7):23-28,42.
- [5] 新华社.玩智能手机越来越上瘾?因为你被设计了! [EB/OL]. (2017-02-05) [2024-04-28]. http://www.xinhuanet.com/world/2017-02/05/c_129466659.htm.
- [6] 于雪.自动还是自主——论人工智能时代的“自动化生存”[J].探索与争鸣,2024(2):153-161,180.
- [7] Gabriel I. Artificial Intelligence, Values, and Alignment [J]. Minds and Machines, 2020, 30(3):411-437.
- [8] 极目新闻.剑桥大学研究:过度依赖AI语音助手可能影响儿童智力和情商[EB/OL]. (2022-10-19) [2024-04-28]. <https://baijiahao.baidu.com/s?id=1747069909347930930&wfr=spider&for=pc>.
- [9] 李志锴,张骁.人工智能生成内容(AIGC)应用于学位论文写作的法律问题研究[J].学位与研究生教育,2024(4):84-93.
- [10] 刘云开.人工智能生成内容的著作权侵权风险与侵权责任分配[J/OL].西安交通大学学报(社会科学版):1-13[2024-06-14]. <http://kns.cnki.net/kcms/detail/61.1329.C.20240604.1353.004.html>.
- [11] 中国经营报.被斯坦福AI团队套壳 面壁智能公开回应[EB/OL]. (2024-06-04) [2024-07-28]. <https://baijiahao.baidu.com/s?id=1800924929751143598&wfr=spider&for=pc>.
- [12] Falmagne J C, Koppen M, Villano M, et al. Introduction to knowledge spaces: How to build, test, and search them [J]. Psychological Review, 1990, 97(2):201.
- [13] 汝绪华.算法政治:风险、发生逻辑与治理[J].厦门大学学报(哲学社会科学版),2018(6):27-38.
- [14] 冯永刚,屈玲.ChatGPT运用于教育的伦理风险及其防控[J].内蒙古社会科学,2023,44(4):34-42;213.
- [15] 娄超,申文君.数据智能参与知识生产的实践机理与法治机制——以ChatGPT为例[J/OL].情报资料工作:1-13[2024-06-17]. <http://kns.cnki.net/kcms/detail/11.1448.G3.20240528.1703.002.html>.
- [16] 吴宗宪,肖艳秋.生成式人工智能的应用、风险与法律回应——以ChatGPT为视角[J].天津师范大学学报(社会科学版),2024(3):12-20.
- [17] 赵子明.人工智能体道德责任归责问题探析[J/OL].中国矿业大学学报(社会科学版),1-20[2024-05-30]. <http://kns.cnki.net/kcms/detail/32.1593.C.20240422.2019.004.html>.
- [18] 《思想政治工作研究》杂志社.以社会主义核心价值观引领生成式人工智能的立法[EB/OL]. (2023-07-13) [2024-07-28]. <https://m.12371.gov.cn/app/template/displayTemplate/news/newsDetail/5452/446027.html>.
- [19] 于秋月,刘坚.“人工智能+教育”的伦理分析及主体重塑[J].学术探索,2024(1):148-156.
- [20] 刘欢,汤维中,任友群.数据可视化促进教育决策科学化:内涵、策略与挑战[J].

- 教育发展研究, 2018, 38 (5): 75-82. 119-130.
- [21] 韩世鹏. 生成式人工智能算法备案的法律属性与控制路径 [J]. 河南财经政法大学学报, 2024, 39 (2): 119-128.
- [22] 袁曾. 生成式人工智能责任规制的法律问题研究 [J]. 法学杂志, 2023, 44 (4):
- [23] 科技部. 发展负责任的人工智能: 新一代人工智能治理原则发布 [EB/OL]. (2019-06-17) [2024-07-28]. https://www.most.gov.cn/kjbgz/201906/t20190617_147107.html.

Legal Regulation of Educational Ethics in the Age of Digital Intelligence

Shi Limin

College of Education, Guangxi Normal University, Guilin

Abstract: Artificial intelligence is a technology that mimics and extends human intelligence and has a wide range of applications. Its development not only represents the progress and breakthrough of science and technology, but also all walks of life continue to iterate and update with the benefits of artificial intelligence, a new quality of productivity. In many fields, it is the most widely used in the field of education. Students, as users of artificial intelligence, often use it to assist learning, carry out virtual experiments, and use intelligent recommendation learning systems, etc., and then innovate and develop their own majors through cross-border integration. However, the problems behind it still need to be further explored and regulated.

Key words: Artificial intelligence; Intelligence; Assisted learning; Innovative development