

生成式 AI 情感引导适用范围与禁区——以医学院校思政课为例

马欧阳

昆明医科大学, 云南昆明, 650500;

摘要: 以 DeepSeek 为代表的生成式人工智能迅速兴起, 为医学高校思想政治理论课教学带来了新的机遇。生成式 AI 通过情境模拟、个性化答疑、教学反馈分析等功能, 有助于增强思政教学的情境真实感和医学生的情感共鸣, 提高教学的针对性和实效性。同时, 生成式 AI 介入课堂也潜藏师生主体弱化、算法偏见渗透、虚拟情感依赖等风险隐患。因此, 应明确生成式 AI 情感引导的适用范围, 充分发挥其情感赋能优势, 充分遵循教育伦理规范和意识形态安全原则, 通过教师监管和政策引导确保技术应用不偏离立德树人的根本任务。

关键词: 生成式 AI; 思政课; 情感引导

DOI: 10. 64216/3080-1516. 25. 11. 041

党的二十届三中全会指出, 教育、科技、人才是中国式现代化的基础性、战略性支撑。当前, 以 deepseek 为代表的生成式人工智能的科技出现, 掀起了新一轮数字技术革命浪潮, 人类教育领域也随之迎来深刻变革。在医学高等院校思想政治理论课教学中, 引入生成式人工智能具有重要意义。一方面, 医学专业思政课强调将人文关怀融入医学伦理、医患关系等情境, 培养医学生的社会责任感和道德情操。另一方面, 生成式 AI 赋能思政课也带来新的挑战。开放的技术既提供便利, 也可能冲击课堂主导权和内容严肃性。所以如何既发挥生成式 AI 的情感引导优势, 又严守教育伦理和意识形态安全底线, 已成为亟待研究的课题。

1 以 AI 情感引导推进思政教学革新

马克思主义人学立场是分析生成式 AI 教育应用的根本视角。马克思主义强调以人民为中心的价值取向, 认为人的自由而自觉的实践活动构成人的本质特性。技术进步是人类改造世界、延伸自身能力的成果, 但技术本身也有独立的发展规律, 必须处理好人与技术的关系才能实现符合人类利益的结果。这蕴含两层要求: 其一, 在教育领域应用人工智能必须坚持“技术为器、育人为道”的原则, 以技术手段服务育人目标, 人始终是教育的主体, 不能本末倒置。教师和学生的人文主体地位不容动摇, 生成式 AI 应作为辅助而不能取代人, 更不能让人过度依赖技术而丧失主观能动性。其二, 要防范技术异化对人的侵蚀。马克思批判过资本主义社会中技术统治人的现象, 指出技术可能异化为反过来支配人的异己力量。如果听任生成式 AI 不加监管地参与教学, 任

由算法逻辑左右课堂话语权, 就可能削弱甚至取代教师的教育主导权和学生的独立思考权。这实质上是一种新的异化: 人的主体性沦陷于技术的统治之下。因此, 必须以马克思主义人学立场审视 AI 应用边界, 确保人工智能始终为人所用而不喧宾夺主^[1]。

技术伦理原则为生成式 AI 融入教育提供了基本准则。2021 年出台的《新一代人工智能伦理规范》进一步强调人工智能活动应始终由人控制, 并坚持隐私保护、公平公正等原则。映射到思政课情感引导场景, 这意味着教师必须对 AI 输出内容严格把关, 有权随时中止不当的生成内容传播; AI 不得侵害学生人格尊严和隐私, 涉及学生个人情感数据时必须依法保护, 不可利用学生心理弱点进行情感操控牟利; 并确保 AI 引导的价值取向积极健康, 决不能为追求情感效果而传播低俗、偏激或违背公序良俗的内容。只有将 AI 的功能纳入道德法治的轨道, 才能真正服务于人的全面发展和教育公益。

意识形态安全是评估生成式 AI 教育应用的政治底线。习近平总书记强调“意识形态工作是党的一项极端重要的工作”^[2]。所以各领域都必须把维护意识形态安全作为重要任务。科学技术并非中立无害, 新兴技术同样蕴含意识形态属性, 需要用正确导向引领其发展。生成式 AI 凭借强大的话语生成能力, 正深度嵌入政治、经济、文化等社会领域, 成为影响舆论和观念的重要变量。基于此, 国家已着手构建生成式 AI 应用的意识形态安全边界——2023 年国家网信办等七部门联合发布《生成式人工智能服务管理暂行办法》, 明确规定生成式 AI“不得生成煽动颠覆国家政权、推翻社会主义制度,

危害国家安全和利益、损害国家形象”等内容，严禁生成任何可能扰乱社会稳定的有害信息^[3]。这些条款划定了AI应用的政治红线，确保AI技术不被利用来危害国家制度和安全。

2 以合理定位发挥生成式AI情感引导优势

在马克思主义人本理念和技术伦理规范指引下，应有条件地将生成式AI应用于医学院校思政课堂，使其发挥辅助育人的积极作用，同时避免越俎代庖或引发负效应。

2.1 提升思政案例的真实感

生成式AI可用作情境模拟工具，增强医学生的价值体验。传统思政课往往侧重说教，缺乏具体生动的场景，使学生难以产生情感共鸣。尤其在医学专业背景下，医学生对临床案例、医患故事更易产生兴趣与共情。因此，AI可以根据教学需要生成贴近医学实际的典型案例和虚拟情境。如利用大模型模拟紧急救治场景下的医患对话，让学生仿佛置身舍己救人的抢救现场，在虚拟体验中感受医者仁心的价值追求；又如生成历史上医学家投身革命或抗疫前线的故事，再现医学史重大事件中蕴含的爱国奉献精神。通过这些定制的数字情境，抽象的思想政治道理被赋予了鲜活具体的载体，学生以第一人称视角参与其中，易被情境中所蕴含的信念与情感所感染，激发内心深处的道德情感共鸣，促进思政课理论向信仰转化。研究表明，借助生成式AI进行临床情境模拟和伦理决策训练，可以有效缓解医学思政教学中情境真实感不足、价值引导抽象化等问题，使价值引领转化为可体验的实践过程^[4]。这种由AI辅助（而非替代）的情境教学模式，能够大幅提升思政课的亲和力和感染力。

2.2 促进学生自主参与学习过程

生成式AI可以充当智能助教，与学生开展个性化的思想交流。在大班教学中，教师往往难以及时照顾到每个学生的思想困惑和情感需求。一些医学生对思政理论存在疑问或心理障碍，但可能因为羞涩等原因未主动向老师求助。在这种情况下，配置一个生成式AI驱动的思政问答聊天助手，可为学生提供私密且持续的交流渠道。学生可以随时就课程内容、人生理想、价值困惑等向AI发问，AI即时回答并引导其思考。由于AI助手具备24小时在线、耐心不倦的优势，能够弥补教师精力的不足，对学生提出的重复、琐碎问题给予耐心回应，从而增加每位学生参与思政讨论的机会和深度。同时，

AI可以通过分析学生的提问模式和语言情感来判断其思想状态，并动态调整交流策略。针对情绪低落的学生，AI助手能够识别其言语中的消极情绪词汇，适当提供安慰和正向激励；对于表现出功利倾向的学生，则强调医者仁心和社会责任，引导其纠正观念偏差，使思政课真正做到春风化雨、润物无声。

2.3 强化教师情感引导能力

生成式AI也可成为思政教师的“智慧助手”，帮助改进教学策略和资源建设。一方面，AI能够协助教师进行课堂情感反馈分析。通过对课堂讨论、学生提问以及（若条件允许）学生面部表情的实时分析，AI可以判断某些知识点讲解时学生的情绪反应和注意力变化。如果AI检测到多数学生在讨论某医学伦理案例时表现出迷茫或抵触，说明该案例引发了情感困惑，教师可据此反思教学方法并在后续课程中作出调整。另一方面，生成式AI具备自动生成教学素材的能力。教师输入特定主题（如“新时代医学道德榜样事迹”）后，AI能够快速搜集整合相关资料，并以生动语言生成故事或短文供备课参考。尤其在情感引导环节，恰当的案例与语言至关重要。AI可以起草充满情感张力的教学情境剧本、设计富有感染力的PPT讲稿片段，帮助教师丰富课堂呈现形式。借助AI反馈，教师可以更清晰地了解自己的长处与不足，不断改进教学艺术，增强课堂对学生的吸引力和情感感召力。

3 加强风险管控划定应用禁区

在明确了生成式AI积极可为的应用领域后，更关键的是划定其不可逾越的禁区。这些禁区对应教育理念、技术伦理和意识形态安全三方面的底线，一旦突破将产生严重后果。

3.1 防止AI取代师生主体地位

要警惕生成式AI异化教育主体性。思政课的靈魂在于“以生为本、以师为导”，教师是灵魂塑造者、学生是成长主体，绝不能让AI喧宾夺主、取而代之。具体来说，要防止AI完全充当“思政教师”。目前有些学校尝试用AI讲授部分理论内容或参与课堂提问，但无论技术多先进，也不能让学生把AI当主要老师。思想政治理论的教学严肃而权威，必须由政治立场坚定、理论功底深厚的真人教师把关讲授；如果任由AI自动生成观点传输给学生，极易出现偏差甚至严重错误，导致学生误入歧途。这是由于AI缺乏真正的情感和价值

立场,无法承担教师“传道授业解惑”的使命。即使 AI 再聪明,也无法如真人教师般以人格魅力和真挚情感感染学生。

要警惕学生主体性的异化。学生是思政教育的实践主体,只有经过亲身思考领悟并付诸行动,才能真正成长。学生过度依赖 AI 完成学习任务和思想活动,其自主能动性将大大弱化。近日教育主管部门发布的《中小学生生成式人工智能使用指南(2025 年版)》明确禁止学生将 AI 生成内容直接作为作业或考试答案,并限制在创造性任务中滥用人工智能^[5]。这一举措旨在防止学生让 AI 代劳作业,保护其独立思考 and 创新能力。在高校思政课中也应严肃对待:坚决杜绝医学生利用 deepseek、豆包等程序代写思想汇报或心得体会等作业,否则不仅违背诚信原则,更让其失去思想锤炼的机会。思政课的价值在于心灵塑造,无法由机器替学生完成。因此,无论教学还是学习环节,都必须坚守师生主体地位不动摇。AI 可以做助手,但绝不能成为主人——这是不可逾越的红线。

3.2 严控算法偏见和错误思潮渗透

价值立场偏差是生成式 AI 应用中的致命风险区,必须筑起防火墙予以隔离。正如前文所述,AI 的大模型训练数据可能隐藏各种意识形态倾向,若缺乏有效甄别,输出内容就可能带有不良价值导向。必须警惕算法偏见导致价值观偏移。大模型从网络学到功利、消费主义等扭曲观念时,生成的回答看似中立却可能暗含偏向,潜移默化地影响学生三观。西方垄断的技术平台也可能夹带文化优越论,使 AI 潜意识里高估西方制度的合理性而贬低中国制度成就。因此必须从技术上校准 AI 的价值取向,并对输出内容进行严格监控,防止偏见信息流入思政课堂。

更危险的是敌对思潮的渗透。近年来西方反华势力不断利用互联网传播历史虚无主义、新自由主义等错误思潮,企图动摇我国主流意识形态。生成式 AI 如果缺乏防护,很可能成为这些思潮的载体和放大器。因此,在高校思政课情境中,必须严格禁止 AI 生成任何违反我国宪法和“四项基本原则”的言论,包括否定党的领导、质疑社会主义制度、丑化英雄人物、吹捧西方政治模式等内容。一经发现,必须立即叫停并向学生澄清影响。学校还应建立 AI 内容审核机制和应急处置预案,确保意识形态安全万无一失。

3.3 杜绝 AI 逾越人类情感伦理

生成式 AI 在情感引导中还有一个必须严肃对待的禁区,即防止情感和伦理的越界。首先要杜绝利用 AI 进行情感操控或令学生沉迷于虚拟情感。AI 能够分析学生的情绪数据并作出回应,这本是积极的一面,但如果使用不当也可能走向反面:一味迎合学生情绪、过度安慰甚至刻意煽情以提高“参与度”,会让学生对 AI 产生危险的心理依赖。思想政治教育应引导学生热爱现实生活、关心现实中的他人,而不是沉溺虚拟情感世界。因此 AI 情感引导必须把握尺度,不鼓励学生对 AI 寄托真情,更不允许 AI 假扮真人与学生发展出超越教学关系的密切联系。教师应密切关注学生使用 AI 的动机,一旦察觉苗头应及时制止。

要严防 AI 突破伦理禁忌生成不当内容。思政课讨论的很多是伦理道德议题,AI 在生成情景或回答问题时可能无意间越界。例如在医学伦理案例模拟中,若 AI 提出了违背医疗伦理的极端做法(即便只是作为反例呈现),都可能对学生造成不良刺激。如果这些内容未经教师引导就展现给学生,容易引起价值混淆甚至道德麻木。同样,AI 若生成涉及宗教信仰或民族敏感问题的不当言论,可能触及社会伦理高压线,后果不堪设想。因此学校必须为 AI 设定明确的伦理禁区,任何高风险敏感话题一律不准 AI 涉及。如教学确需讨论,也必须由教师亲自掌控、谨慎处理,不能交由 AI 自由发挥。

参考文献

- [1] 代金平,覃杨杨. ChatGPT 等生成式人工智能的意识形态风险及其应对[J]. 重庆大学学报(社会科学版), 2023(5): 101-110.
- [2] 《习近平著作选读》第一卷,北京:人民出版社 2023 年版,第 147 页.
- [3] 国家互联网信息办公室等. 生成式人工智能服务管理暂行办法[Z]. 2023-07-15.
- [4] 刘成,莫生叶. 生成式人工智能应用中的意识形态安全边界厘清、风险监测及规制理路[J]. 新疆师范大学学报(哲学社会科学版), 2025(3): 45-53.
- [5] 教育部基础教育教学指导委员会. 中小学生生成式人工智能使用指南(2025 年版)[Z]. 2025-05-12.

作者简介:马欧阳(1990.1-)男,壮族,河南商丘人,讲师,博士,研究方向:思想政治教育。

本文是“昆明医科大学习近平新时代中国特色社会主义思想研究中心”党的二十届三中全会精神研究专项一般项目的阶段性成果。项目号:KYDJZX25262。