

· 标准与规范探讨 ·

中国眼科规范应用生成式人工智能大模型专家共识(2026 年)

国际电气电子工程师学会技术与工程管理分会智能医疗数字化技术委员会智能眼科分部(中国)

通信作者:陈有信,中国医学科学院 北京协和医学院 北京协和医院眼科,北京 100730,
Email:chenyouxinpumch@163.com

【摘要】 随着人工智能(AI)生成技术和 AI 大模型快速发展,其在眼科领域医疗流程的诸多方面有望发挥重要作用。然而,技术快速普及也带来应用规范、伦理约束及法律合规等多重挑战。为确保生成式 AI 大模型在眼科领域安全、有效及合规应用,国际电气电子工程师学会技术与工程管理分会智能医疗数字化技术委员会智能眼科分部(中国)组织国内相关专家,参考中国医疗行为规范、数据安全及隐私保护法律等多项法规,结合最新研究进展,在全国范围内经过多次线上及线下研讨会议,并基于推荐分级的评估、制订与评价系统以及良好实践主张、德尔菲法,针对生成式 AI 大模型在眼科应用中的伦理、技术和法律问题,达成可随技术进步持续更新的共识性意见,为临床开展工作提供依据和支持。

【关键词】 生成式人工智能; 大语言模型; 伦理学,医学; 眼科学; 多数赞同

基金项目: 中央高水平医院临床科研业务费(2025-PUMCH-A-189, 2022-PUMCH-B-101); 国家自然科学基金青年项目(82301239); 北京市科学技术委员会-AI+ 健康协同创新培育项目(Z211100003521020); 首都卫生发展科研专项项目(首发 2022-2-4015)

Chinese expert consensus on the standardized application of generative artificial intelligence large models in ophthalmology (2026)

Intelligent Ophthalmology Subcommittee (China) of Technical Committee on Smart Medical Digitalization of IEEE Technology and Engineering Management Society

Corresponding author: Chen Youxin, Department of Ophthalmology, Peking Union Medical College Hospital, Chinese Academy of Medical Sciences & Peking Union Medical College, Beijing 100730, China, Email: chenyouxinpumch@163.com

【Abstract】 With the rapid development of generative artificial intelligence (AI) technologies and large AI models, these technologies are expected to play an important role in multiple aspects of medical workflows in ophthalmology. However, their rapid and widespread adoption has also introduced multiple challenges regarding application specifications, ethical constraints, and legal compliance. To ensure the safe, effective, and compliant application of generative AI large models in ophthalmology, the Intelligent Ophthalmology Subcommittee (China) of Technical Committee on Smart Medical Digitalization of IEEE Technology and Engineering Management Society convened domestic experts in the field. Drawing upon relevant Chinese laws and regulations governing medical practice, data security, and privacy protection, and incorporating the latest research advances, the expert panel held multiple online and offline seminars nationwide.

DOI: 10.3760/cma.j.cn112142-20250721-00317

收稿日期 2025-07-21 本文编辑 黄翊彬

引用本文: 国际电气电子工程师学会技术与工程管理分会智能医疗数字化技术委员会智能眼科分部(中国). 中国眼科规范应用生成式人工智能大模型专家共识(2026 年)[J]. 中华眼科杂志, 2026, 62(5): 329-345. DOI: 10.3760/cma.j.cn112142-20250721-00317.



Using the Grading of Recommendations Assessment, Development and Evaluation (GRADE) system, Good Practice Statements, and the Delphi method, the panel developed a consensus framework addressing the ethical, technical, and legal issues associated with the application of generative AI large models in ophthalmology. Designed to be continuously updated as technology advances, this framework aims to provide evidence-based support for clinical practice.

【Key words】 Generative artificial intelligence; Large language models; Ethics, medical; Ophthalmology; Consensus

Fund program: National High Level Hospital Clinical Research Funding (2025-PUMCH-A-189, 2022-PUMCH-B-101); National Natural Science Foundation of China Youth Program (82301239); Beijing Municipal Science and Technology Commission-AI+Health Collaborative Innovation Cultivation Project (Z211100003521020); Capital's Funds for Health Improvement and Research (CFH2022-2-4015)

一、前言

(一)制订背景和目的

2023 年 4 月中国公布了《生成式人工智能服务管理暂行办法》,2024 年 1 月世界卫生组织发布了《卫生领域人工智能的伦理与治理:多模态大模型指南》。生成式人工智能(artificial intelligence, AI)及 AI 大模型逐步成为医学 AI 的研究热点。眼科作为医学 AI 研究的前沿领域,尚缺少研究和应用的相关共识、指南或规定。生成式 AI 及 AI 大模型在眼科领域的应用属于新兴用途^[1-2],需要明确存在的潜在益处和风险,以确保此类 AI 模型的设计、研发、部署、使用等流程符合伦理规范、法律法规、人权和安全准则等。

制订本共识旨在确保眼科领域应用生成式 AI 及 AI 大模型符合法律法规及伦理规范,并提供明确的技术要求和操作标准,提升研究方法和应用的透明度,提高眼科医务人员及研究人员对生成式 AI 及 AI 大模型技术的理解和应用能力,推动该技术进步和广泛合作。

(二)概念和定义

生成式 AI 与 AI 大模型不是同一概念,随着 AI 技术发展,二者逐步融合。生成式 AI 是一种 AI 技术,这种深度学习模型在经过训练后,可以生成文本、算法代码、多媒体内容、特定格式的文件或其他专业领域的内容。生成式 AI 已用于医学、设计、模拟、计算机等诸多领域。眼科曾研究使用生成式 AI 生成眼底病变治疗后图像,是典型的应用范例。

AI 大模型作为基础模型的核心实现形态,通常是指基于海量多源数据及至少千卡级图形处理器集群,经超大规模算力训练的参数量达亿级以上的复杂 AI 模型。其核心特性包括规模效应、能力范式及功能输出。规模效应是指通过规模化法则实现参数量、数据量、计算量的三重拓展。能力范

式主要包括涌现能力及强泛化能力。涌现能力是指 AI 大模型在训练数据和参数量超越临界规模后,表现出在较小模型或较少训练数据下未曾具备的功能或行为,尤其复杂推理或创造能力,如针对具体问题的逻辑推理、自行编写计算机代码、跨语言翻译等。强泛化能力是指 AI 大模型在分布外数据及零样本任务中保持高鲁棒性,即在训练数据之外的未知数据方面仍可保持良好的表现能力,反映了模型将从有限的训练样本中学习到普遍规律或知识应用于新情境的能力,如处理新语言、认识新的异常病灶等。功能输出是指 AI 大模型支持生成式任务(文本、图像、代码等生成)及非确定性推理(如多步逻辑推演、反事实分析)。AI 大模型的典型代表为大语言模型,这是一种特殊的单模态 AI 大模型,如 ChatGPT 的早期版本。目前多模态 AI 大模型是学界的研究热点,近期的 DeepSeek 等 AI 大模型均属于多模态 AI 大模型,可以处理多媒体信息。

将 AI 大模型用于内容生成和输出,从而为用户提供服务,称为生成式 AI 大模型。生成式 AI 大模型在眼科领域的应用具备重大潜力,有望在辅助诊断、影像分析、医学文献综述、患者沟通、医疗管理等多个方面发挥作用。然而,随着该技术快速发展和普及,其在应用、伦理和法律方面逐渐面临挑战。为此,制订必要的规范以确保生成式 AI 大模型在眼科领域的应用尤为重要。

(三)共识制订方法

基于目前生成式 AI 大模型在眼科应用的有关问题,国际电气电子工程师学会技术与工程管理分会智能医疗数字化技术委员会智能眼科分部(中国)于 2025 年 4 月成立《中国眼科规范应用生成式 AI 大模型专家共识》制订专家组,于 2025 年 5 月 1 日开始针对生成式 AI 大模型在中国眼科医疗中



的应用进行讨论,整理相关领域涉及的医学伦理问题以及相关技术临床应用面临的挑战及其对策。

由于生成式 AI 大模型在眼科应用尚未形成统一的可遵守的规范,专家组基于相关文献,包括法律法规、伦理要求、研究进展,召开线下和线上会议,针对收集的生成式 AI 大模型在眼科应用中的相关规范观点进行充分讨论和论证,形成 47 条核心陈述。证据质量评价方法采用推荐分级的评估、制订与评价(grading of recommendations, assessment, development and evaluation, GRADE)系统^[3-6]以及良好实践主张(good practice statement, GPS)。

采用德尔菲法,以电子问卷征询形式,根据李克特 5 分量表法,针对 47 条核心陈述内容,专家组成员独立选择回答“强烈同意、同意、不确定、不同意、非常不同意”^[7],以 75% 以上专家选择回答“强烈同意”和“同意”为该陈述内容达成共识^[8]。根据专家组成员的回答形成共识性意见初稿,提交每名专家审阅后完成共识性意见终稿(共识性意见要点

框架见图 1)。

二、眼科应用生成式 AI 大模型需要遵守相关法律法规并符合医学伦理要求

说明:这是关于眼科应用生成式 AI 大模型须遵循的总体原则共识。

1. 符合相关法律法规(GPS;非常同意:78.79%,同意:21.21%)。中国已出台关于生成式 AI 的法规,并在不断完善中^[9]。在眼科领域开发、测试及使用生成式 AI 大模型,应遵循中国医疗行为规范、数据安全保护、患者隐私保护等相关法律法规,包括《生成式人工智能服务管理暂行办法》《人工智能法示范法 1.0(专家建议稿)》《网络数据安全条例》等。此外,中国的《网络安全法》《数据安全法》《个人信息保护法》等,也是生成式 AI 大模型应用中应当遵守的基础性法律框架内容。应依法依规在国家互联网信息办公室备案大模型。当生成式 AI 大模型用于眼科医疗实践时,如基于眼底影像诊断疾病,应将 AI 的软件和硬件视为医疗设备,因此应获得国家药品监督管理局批准,确

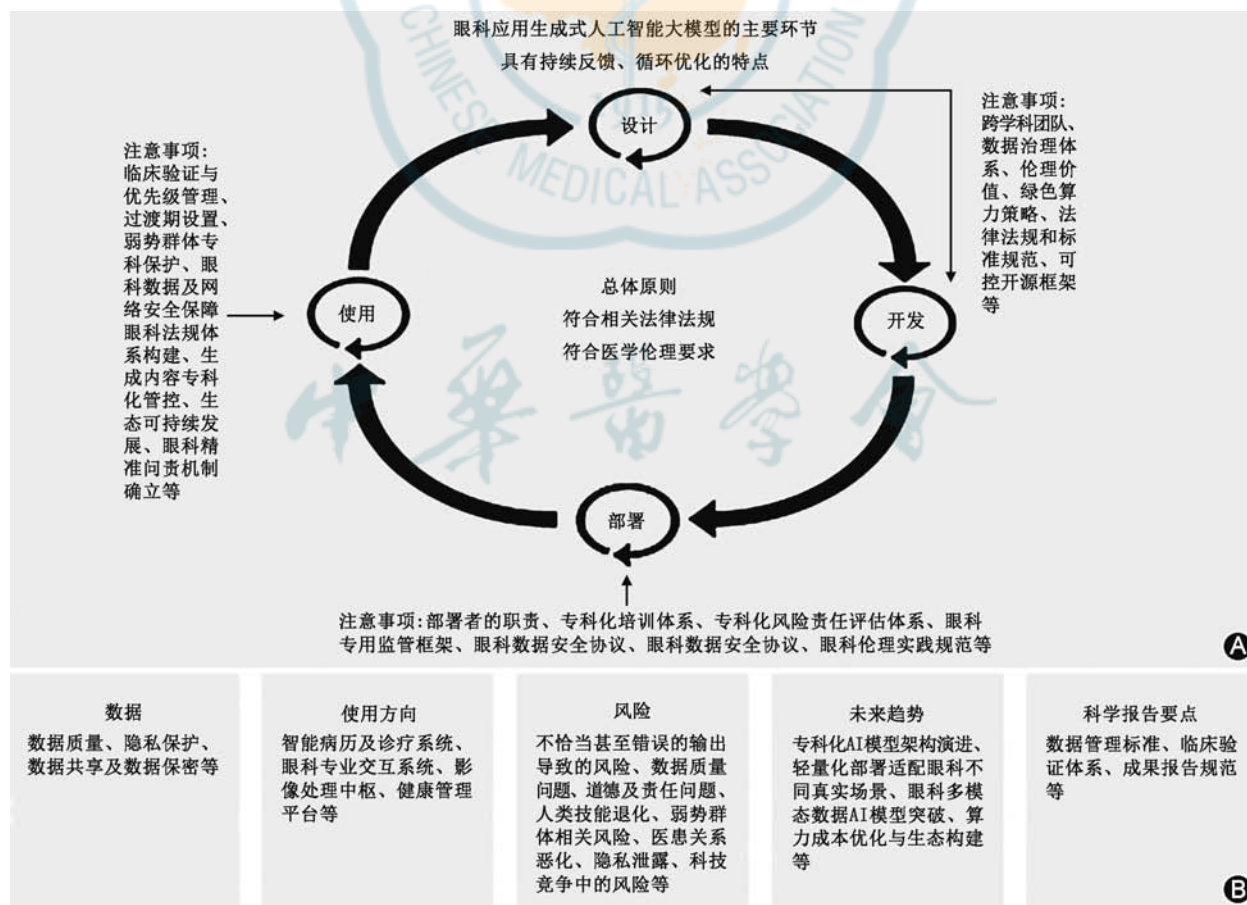


图1 中国眼科规范应用生成式人工智能(AI)大模型专家共识性意见要点框架示意 A 示在总体原则指导下进行的生成式 AI 大模型设计、开发、部署及使用四大环节具有持续反馈、循环优化特点,且各有注意事项;B 示应用生成式 AI 大模型的 5 个需要关注的方面,即数据、使用方向、风险、未来趋势及科学报告要点

保技术的有效性和安全性。

2. 符合医学伦理要求 (GPS; 非常同意: 93.94%, 同意: 6.06%)。在眼科领域开发、使用生成式 AI 大模型, 应遵循中国发布的《新一代人工智能伦理规范》及联合国教科文组织发布的《人工智能伦理建议书》, 以人机对齐为前提, 强调增进人类福祉、促进公平公正、保护隐私安全、确保可控可信、强化责任担当、提升伦理素养, 同时应注重促进环境与生态系统发展, 保证多样性和包容性^[10]。例如在眼科诊疗过程中可产生大量包含影像、数值、文字等不同类型的数据, 必须确保数据安全并保护患者隐私; 对于存在的可能风险, 须获得患者知情同意, 并保证患者有权拒绝参与或随时终止参与^[11]。生成式 AI 大模型的目标及输出应当与人类价值观相符, 从而保证大模型的能力及行为与人类意图、价值观、伦理要求、道德规范等一致。生成式 AI 大模型的训练数据和输出结果可能存在偏倚, 导致针对弱势群体出现不公平医疗结果, 应采取多种方法尽可能避免这些偏倚, 防止加剧健康不平等。例如建立主动的偏见审核与缓解机制, 在大模型部署后持续监控其不同亚组人群中的表现差异。医学伦理属于科技伦理, 目前科技伦理强调健康和可持续发展, 并发展中国特色的科技伦理, 因此生成式 AI 大模型也应当注意符合科技伦理的宏观框架。

三、眼科应用生成式 AI 大模型需要至少涉及设计、开发、部署和使用四大环节

说明: 这是关于生成式 AI 大模型在研发至应用于眼科全流程中, 应包括哪些环节的共识。

眼科应用生成式 AI 大模型需要至少涉及设计、开发、部署和使用四大环节 (GPS; 非常同意: 87.88%, 同意: 12.12%)。

1. 设计。生成式 AI 大模型的初期设计至关重要, 需要全面考量模型架构、规模、参数、应用目标、数据集质量、资源消耗 (包括数据、人力、硬件、碳足迹、水足迹等)、使用者角色、生成内容、维护策略、医学伦理及政府监管要求等。眼科应用的生成式 AI 大模型应由眼科医务人员、相关学科医师、AI 专家、信息安全专家、眼科患者、大模型用户及法律专家等共同参与设计^[12]。设计阶段的任何微调均可能对后续环节产生深远影响, 故设计过程应兼具包容性和审慎性。

2. 开发。生成式 AI 大模型的研究和开发主体通常为大型科技企业。然而, 眼科作为特定应用领

域, 其生成式 AI 大模型的开发机构亦可为高等院校、初创科技公司、政府机构或其他资源丰富的实体, 个人开发者通常较难胜任。开发途径包括全新构建、微调通用 AI 大模型、集成至眼科软件系统, 或者建立眼科知识库与 AI 大模型协同工作机制。开发由提供者执行, 通常负责将通用 AI 大模型适配至特定用途, 可能为原始通用 AI 大模型研发人员和使用者之外的第三方。开发须忠实于设计初衷, 因设计与开发人员可能存在专业差异, 建议设立专职管理人员以建立开发质量管控体系, 使其符合法律法规和医学伦理要求; 建议引入医疗 AI 产品经理角色, 通过临床需求文档到技术规格文档的标准化转换流程, 保障临床与技术需求对齐, 验证生成式 AI 大模型的工作效能和输出准确性 (如系统性眼底图像阅片能力评测)^[13], 确保设计意图贯彻。数据获取、标注及管理也是开发过程中的重要环节。

3. 部署。部署环节即将开发完成的生成式 AI 大模型以产品或服务形式提供给用户, 使其具备使用能力。此环节为应用前的准备、安装与适配阶段, 保证可与医疗机构环境适配, 通过眼科临床验证^[14-15], 建设动态监测机制捕捉异常输出。部署者身份可为开发者、提供者、政府医药卫生管理部门、医疗机构、保健机构、制药企业、医疗器械企业、研究机构等企业等; 当多个主体同时为部署者时, 应当选定其中一个主体为主要部署者, 承担主要责任和义务。眼科应用生成式 AI 大模型, 输出须确保对患者无害, 故部署时应充分验证大模型在即将使用场景中的真实效能 (如诊断效能)^[16]; 还应根据使用者需求调整和优化, 甚至允许对特定功能提出修改要求。部署阶段可同时对多个生成式 AI 大模型选型比对, 根据准确性、灵敏度等具体指标, 确定最适合部署的具体模型。

4. 使用。眼科生成式 AI 大模型的主要使用者为眼科医师和医疗机构, 政府卫生部门、制药企业及个人亦可使用。本共识聚焦高度专业的眼科应用场景, 患者所需科普内容并非重点。理想的生成式 AI 大模型应具备直观、易用且可解释的交互界面, 支持文本、语音、图像等多模态输入, 实时响应, 并在必要时提供推理依据查看功能, 为临床决策提供人机协同支持。该自然友好的体验是生成式 AI 大模型广受接纳的关键之一。然而, 使用者须认识到, 大模型的输出系算法生成, 可能存在偏差或不完全符合医学伦理要求、法律法规, 需要人工复核,



故 AI 无法完全替代医务工作者。负责监测生成式 AI 大模型应用及复核输出结果的人员应当接受系统性培训并完成考核,方可进入眼科医疗环境工作,培训内容包括但不限于如何操作大模型,批判性解读大模型输出结果,识别幻觉(指大模型编造其认为真实存在,甚至看似合理或可信,但与现实世界事实或用户输入不一致信息现象)和局限性,并建立 AI 建议与临床判断冲突时的明确处置流程。

明确上述四大环节的主要任务和责任主体,有助于厘清生成式 AI 大模型的责任划分^[17-19]。生成式 AI 大模型具有迭代本质,这四大环节内部多个步骤具有持续反馈、循环优化过程,如更新数据集和知识库,优化算法架构,改进应用方式及场景等;这四大环节之间亦具有反馈机制,如应用环节中的临床反馈,应反馈至大模型的设计、开发人员,指导大模型的设计优化和开发升级。建议建立标准化的部署、使用后评估框架,包括临床效能监测、用户满意度调查、安全性追踪等指标;评估结果应形成结构化报告,定期反馈至开发团队,作为大模型迭代优化的依据;定期进行较为全面的评估,对发现的重大问题实施即时反馈机制。但是,为了避免部署后大模型在持续学习过程中出现灾难性遗忘、效能下降、责任归属模糊等问题,禁止在临床环境中对已注册的生成式 AI 大模型进行无监督持续学习,任何模型效能相关参数更新必须视为新版本,须重新进行医疗器械变更注册流程。

四、眼科应用生成式 AI 大模型应当注意数据质量、隐私保护、数据共享及数据保密等方面

说明:这是关于眼科生成式 AI 大模型所用数据的共识。合规性管理数据是确保眼科应用生成式 AI 大模型临床有效性和伦理安全的核心基础。鉴于当前普遍存在数据透明度缺失问题(包括数据来源隐匿性、潜在偏见不可溯性及采集合法性存疑等),建议建立以下规范化框架。

1. 数据质量(GPS;非常同意:87.88%,同意:12.12%)。生成式 AI 大模型在眼科领域的应用广泛,因此应使用高质量数据训练大模型^[20]。数据质量包括准确性、完整性、一致性、时效性、代表性和多样性等^[21]。构建眼科专科多模态数据质量的标准和规范,有助于统一数据收集和处理流程^[22]。例如,彩色眼底图像数据应当具有高分辨率,以确保细节清晰;屈光度数、眼轴长度等测量数据应准确,避免错误或偏倚,以防止大模型生成误导性结果。

收集的眼科数据应当具有代表性和多样性,覆盖不同人群、病况和场景,以保证大模型结果的普适性、有效性和可靠性。

2. 隐私保护(GPS;非常同意:93.94%,同意:6.06%)。眼科应用生成式 AI 大模型,数据隐私保护是确保大模型伦理安全和临床应用可信度的核心因素。除了姓名、编号等数据外,虹膜、视网膜等组织亦具备生物身份识别特征。生成式 AI 大模型涉及的数据须严格遵守隐私保护和数据保密规定^[23],对患者数据进行去标识化处理,确保个人信息的安全^[24]。数据存储、传输、管理过程应采取参数加密、差分隐私、分布式部署(如联邦学习、区块链)、访问控制等技术方法,防止数据泄露^[25-26]。

3. 数据共享(GPS;非常同意:84.85%,同意:15.15%)。眼科应用生成式 AI 大模型,应当注意数据共享的法律法规和医学伦理要求,确保通过正式协议和合规程序进行数据共享,以保护患者权益和促进医疗创新。尤其进行跨机构或跨国数据共享时,须签订数据共享协议,明确数据使用范围和保护措施,确定各环节的责任人,确保各方遵守相应的法律法规^[27],如《中华人民共和国网络安全法》《中华人民共和国数据安全法》《中华人民共和国个人信息保护法》《中华人民共和国生物安全法》《人类遗传资源管理条例》等。

4. 数据保密(GPS;非常同意:90.91%,同意:9.09%)。国家法律法规规定的保密数据、国家或特定机构的专有数据及其他敏感数据不能输入至生成式 AI 大模型中,涉及人类遗传学的资源和非公开的研究数据亦应在监管下合理使用。暴露或有可能披露患者个人信息的数据应当在脱敏处理后使用^[28]。对于必须使用的受限数据,应建立受控训练环境和实时数据焚毁机制,应用匿名化技术及数字水印技术等实现数据保护目的。应避免直接输入训练系统的数据包括但不限于涉及国家战略安全的眼科流行病学监测数据,尚未公开发表的多中心临床试验原始数据集,以及可逆向还原个体身份的高分辨率虹膜和视网膜特征数据。应积极防范生成式 AI 大模型的记忆效应,避免在输出结果中泄露训练数据中的特定敏感数据。例如,在开发阶段使用差分隐私等技术,向训练数据注入受控噪声,从而保证大模型输出不会泄露任何单个训练样本的敏感信息。

五、眼科应用生成式 AI 大模型可实现建立智



能病历及诊疗系统、眼科专业交互系统、影像处理中枢、健康管理平台等功能,未来也可能用于医疗流程的其他环节

说明:这是关于眼科应用生成式 AI 大模型可能实现功能的共识。基于当前技术发展水平和医学伦理要求,眼科应用生成式 AI 大模型可实现诸多临床功能,部分功能尚待验证。

1. 智能病历及诊疗系统(1C;非常同意:87.88%,同意:9.09%)。集成电子病历或电子健康记录及医学知识图谱,构建智能病历系统^[28]。包括结构化病历生成^[29],如根据医学数字成像与通信标准影像和健康信息七层级标准数据自动生成病程记录、手术记录等医疗文书,采集门诊医患对话生成门诊病历;多模态报告辅助生成,如分析相干光层析成像术(optical coherence tomography, OCT)图像并生成报告,输出符合国际疾病分类第 11 次修订本编码的标准化诊断建议;鉴别诊断支持,如通过比对罕见病特征矩阵,降低漏诊率并提供参考文献。但是,大模型生成病历等资料的准确性可能因幻觉等局限性而受影响,生成内容应以草稿或建议等形式输出,而非医疗文书的最终版本,并由人类医师实名审核、验证、修改、确认等,形成可追溯的责任链。

2. 眼科专业交互系统(1C;非常同意:90.91%,同意:9.09%)。利用生成式 AI 大模型成熟的自然语言交互功能,构建分角色知识服务体系^[30]。包括面向临床医师、护师、技术人员、行政管理人员等不同角色的系统终端,基于循证医学数据库、病历资料等,提供实时诊疗方案优化建议^[31];患者教育端,结合患者居家、随身、便携设备收集的全天候健康数据,生成个性化健康指导,并针对眼科患者可能视力欠佳的特点,遵循无障碍设计的原则,提供语音说明、高对比度模式或大字体显示等^[32];教学科研端,生成医学生教育和继续教育资料以及虚拟患者,提供科研指导^[33-34]。

3. 影像处理中枢(1C;非常同意:84.85%,同意:15.15%)。眼科诊疗可能高度依赖影像学等多媒体资料,广泛使用 OCT、血管造影、眼底照相、裂隙灯显微镜照相、手术录像等多媒体数据。利用生成式 AI 大模型可在病灶分析、病程预测和多模态融合等方面实现多种功能^[35],包括病灶定量分析,如高精度解剖结构分割、视网膜层间水肿体积分析等;病程预测模拟,如生成糖尿病视网膜病变长期进展预测影像;多模态影像融合和推理,如将眼底血管

造影和相干光层析血管成像术结果转化为血流动力学参数可视化报告,或根据眼底影像推测全身健康指标。注意大模型的效能应通过多中心、多种族、多年龄段、不同性别等覆盖多元人口学特征且具备场景与任务多样性的数据集验证;若大模型的效能在不同人群间存在差异,须在应用过程中明确警示或限制大模型的特定应用人群。

4. 健康管理平台。生成式 AI 大模型具备出色的多维数据处理能力(GPS;非常同意:78.79%,同意:15.15%)。眼健康护理与日常营养、饮食、运动、身体指标及围治疗期护理密切相关,通过对接可穿戴设备数据流,有望提供全生命周期的眼健康管理^[36]。包括个性化干预方案,如根据 24 h 眼压曲线调整青光眼患者用药计划,并进行慢病管理;远程监护系统,如通过增强现实眼镜实时指导低视力患者进行视觉康复训练;低视力辅助和预警系统,如利用混合现实结合 AI 技术,理解患者视觉需求并生成适合低视力患者的高亮度放大影像,对生活中危险情境进行听觉触觉预警^[37];临床试验匹配方案,如基于中国临床试验注册中心和药物临床试验登记与信息公示平台数据库,筛选适用研究项目。

此外,部分待验证的拓展功能有望基于生成式 AI 大模型技术实现(GPS;非常同意:87.88%,同意:6.06%)。包括诊疗流程优化,如应用数字孪生技术构建虚拟诊疗单元进行路径模拟,协助管理复杂病例,优化医疗流程并提供改进意见^[38];医疗质量控制审计^[39],如基于电子病历数据检测诊疗方案偏离情况,审查眼科诊疗并尝试发现潜在临床错误,监控管理医疗经费及医疗保险;辅助药物和医疗软硬件设备研发^[40]。

六、设计和开发眼科应用的生成式 AI 大模型须注意跨学科团队、数据治理体系、伦理价值、绿色算力策略、法律法规和标准规范、可控开源框架等方面

说明:这是关于设计和开发眼科应用的生成式 AI 大模型注意事项的共识。

1. 跨学科团队(GPS;非常同意:84.85%,同意:15.15%)。研发团队需确立“临床-技术-法律”三元协同机制,建议构建临床团队、技术团队、监管法律团队^[41]。临床团队应包括眼科医师、眼科影像专家等医务工作者、科研人员、患者。技术团队应包括具备 AI 项目开发经验的编程人员、硬件工程师,尤其熟悉处理眼科图像数据的技术人员。监管法律团队应包括熟悉《医疗器械监督管理条例》等国内



法律法规和《通用数据保护条例》《加利福尼亚州消费者隐私法》等国际隐私法规的法律工作者。此外,建议设置具有眼科医学背景的项目经理岗位,负责统筹和管理开发工作。同时,建议政府监管机构制定并定时更新行业 AI 大模型开发的基本原则或标准,并监督相关工作的开展情况,以提升医疗产品和服务的质量和可靠性。

2. 数据治理体系(GPS;非常同意:90.91%,同意:9.09%)。重视数据治理体系有助于确保生成式 AI 大模型的有效性、安全性和伦理合规性等。构建符合国家标准的眼科数据立方体,如国际标准化组织/国际电工委员会 25012 等标准。在数据采集规范性方面,眼科数据(如眼部图像、患者病历、检查报告)应真实、准确、客观,满足大模型的多样性目标。由专业眼科医师对数据进行标注,所有标注人员应通过相应考核,从而获得标注资质,并应明确标注的来源、使用方式(如标注工具、场所)及处理方法。在标注流程中应设置质量控制环节,如多人交叉验证,标注一致性评估,设立质量控制委员会进行标注内容抽查等。在更新机制方面,建立动态数据湖(数据湖是用于存储、管理和分析大规模、多样化数据的数据存储架构),不断更新数据集、思考链、知识库,以满足不断发展的医学需求。在质量控制方面,应由具备高级资质的专家组成质量控制委员会,制订相关标准操作流程,对数据管理全流程和质量控制工作人员进行监管;应部署信任、认证机制及自动化数据清洗程序,避免不可靠、低质量或不可信第三方来源数据,避免过时、存在偏倚或缺陷的数据,确保数据可靠、专业、合法且可持续更新,提高医疗产品和服务的质量^[42]。对于需要人工审查的数据,须先审核审查人员的资质,并应保证工作人员的身心健康得到足够的支持和福利。在隐私保护方面,应确保使用隐私保护技术,并平衡隐私与模型效能,如 k-匿名化、差分隐私等技术,并严格执行知情同意流程,同时避免使用过多小型数据集,从而降低个人身份被识别的风险。

3. 伦理价值(GPS;非常同意:93.94%,同意:6.06%)。保护患者利益是医疗实践的核心价值^[43-44]。鉴于生成式 AI 大模型功能强大且服务广泛,应践行普世价值观,如维护人类尊严、自由、平等与团结。设计团队中应纳入代表患者利益的监督人员,甚至可成立官方监督组织^[45]。在设计过程中,应明确使用者可从中获益的具体方式,并预判潜在不良后果,开展影响评估。可考虑采用价值敏

感设计框架,训练阶段嵌入公平性约束,开发偏见检测模块,并部署反思日志系统^[46]。

4. 绿色算力策略(GPS;非常同意:84.85%,同意:15.15%)。生成式 AI 大模型的训练和运行可能消耗大量能源和水资源^[47-48]。设计时应尽早考虑碳足迹和水足迹,探索减小模型和数据体量的方法,以保护环境并促进可持续发展^[48-49]。中国已建立的中国算力网通过采用可持续运营模式,为环境保护提供了重要支持^[50]。

5. 法律法规和标准规范(GPS;非常同意:93.94%,同意:6.06%)。政府和行业协会在制订眼科应用的生成式 AI 大模型相关规范时,应考虑眼科领域的特殊性。法律法规和标准规范应包括以下内容:(1)要求大模型具备明确目标,确保在眼科疾病诊断和治疗中的准确性和可靠性(如达到符合诊疗实践要求或与金标准相当的灵敏度和特异度),并公开大模型的关键指标数据;(2)要求设定设计标准和指标,如可预测性、可解释性、可追溯性、安全性及网络安全性,并特别关注应用 AI 技术的眼科医疗器械须符合的医疗器械认证标准(如国家药品监督管理局认证)^[14];(3)要求强制保护眼科数据的隐私,鼓励识别并规避伦理和法律风险^[46];(4)要求对大模型全生命周期进行监管,重点审查开发者、大模型本身及其下游眼科医疗产品和服务;(5)要求披露环境影响,并注明大模型输出为算法生成而非人类创作;(6)鼓励早期注册,并公开负面结果(如误诊率),确保使用者遵守法律法规和道德原则^[14]。

6. 可控开源框架(GPS;非常同意:93.94%,同意:6.06%)。眼科应用的生成式 AI 大模型需要多大程度满足开放源代码促进会发布的开源条件,以及开源的意义、优势和风险,需要在谨慎评估后进一步明确^[51]。开源可提升透明度和公众参与度,促进眼科 AI 研究发展。例如共享视网膜疾病检测模型以加速创新,但须符合国家安全和公共利益的要求,尤其保护眼科患者的数据隐私和安全^[52]。开源并不意味免除监管责任,无论大模型是否开源,其应用于眼科医疗相关环境时,提供者和部署者均具有监管义务。使用政府资源开发的眼科生成式 AI 大模型应优先考虑开源,并接受公众监督,在授权后允许广泛访问,同时制订相应政策。例如要求对眼科图像数据进行去标识化处理,以确保数据安全和患者隐私。建议采用分级开放策略,对于包含算法架构、基础模型的核心层,需向监管机构提供技



术白盒验证,但可保留商业机密;对于包含输入输出规范、性能指标的接口层,可完全开放,确保可验证性和互操作性;对于包含临床决策支持系统的应用层,可开源基础版本,保留高级功能的商业授权。

七、部署眼科应用的生成式 AI 大模型,需要注意部署者的职责、专科化培训体系、专科化风险责任评估体系、眼科专用监管框架、眼科数据安全协议、眼科伦理实践规范等方面

说明:这是关于部署眼科应用的生成式 AI 大模型注意事项的共识。部署眼科应用的生成式 AI 大模型,需要进行多维度审慎考量,以确保其在眼科领域的安全性、实用性和合规性。以下是需要注意的 6 个关键事项。

1. 部署者的职责(GPS;非常同意:93.94%,同意:6.06%):部署者需要具备足够的专业能力、从业经验及资源,以持续承担责任,明确生成式 AI 大模型的使用场景和目的^[53],并向使用者充分告知潜在风险^[54]。眼科患者可能存在视功能损伤情况,部署者应考虑为视力残障者等提供无障碍支持,避免不当使用而引发风险。

2. 专科化培训体系(GPS;非常同意:81.82%,同意:15.15%)。培训应在部署阶段启动,使用者(如眼科医师或技术人员)需要熟悉相关法律法规,接受网络安全培训以保护数据,理解生成式 AI 大模型的原理及其局限性(如处理低质量图像能力不足)^[55-56],避免过度依赖。医务工作者应向患者解释 AI 在诊疗中的作用和风险,促进 AI 技术的社会接受度和长期可持续性发展,并能够应对大模型可能导致的医患纠纷。

3. 专科化风险责任评估体系(GPS;非常同意:96.97%,同意:3.03%)。在部署生成式 AI 大模型前,政府监管机构或由国家药品监督管理局等机构认定的独立第三方,需要根据中国《人工智能医疗器械质量要求和评价》及世界卫生组织与国际医疗器械监管机构论坛推荐的 AI 医疗器械风险分层,对大模型进行全面风险评估,使结果公开透明,以增强信任度^[12, 57-58]。责任主体包括开发者、提供者、眼科医师及医疗机构。评估内容除了大模型本身外,还应包括输出结果的准确性和透明性以及人群适配、法律法规和医学伦理要求遵守、隐私保护、数据保护及处理等情况,确保使用结果可靠并明确责任归属。

4. 眼科专用监管框架(GPS;非常同意:90.91%,同意:6.06%)。中国应建立由专职人员负

责的眼科 AI 监管体系,覆盖器械分类、数据监控和临床场景。若生成式 AI 大模型介入诊疗(如生成诊断或治疗方案),须视为医疗器械,通过国家认证(如国家药品监督管理局审批),确保使用符合预期并防范非预期场景的风险^[59]。与大模型相关的 AI 应用或下游程序,如眼科健康管理应用程序或在线咨询平台的 AI 工具,均应在被监管范围。鉴于使用者可能偏离开发者预期(如将大模型用于非预期场景),部署者须在监管下确保使用方法与预期一致,且下游产品或服务与大模型无明显差异^[46]。透明度是确保眼科应用的生成式 AI 大模型安全性和可解释性的关键。部署时须向监管机构公开大模型的代码和数据处理方法、效能指标(如在眼底病变检测中的准确率、灵敏度和特异度)、训练结果及资源消耗(如能源和水资源使用情况)。

5. 眼科数据安全协议(GPS;非常同意:93.94%,同意:6.06%)。在部署生成式 AI 大模型时须特别关注数据保护。眼科数据(如眼底图像、虹膜数据等)因含有生物身份识别特征,难以完全去标识化^[60-62]。大模型的所有数据应接受国家审查和监管^[23],即使允许用户删除的数据也须符合法律法规;采用隐私保护技术(如联邦学习、加密计算)处理眼科数据^[26],以及采用分级访问控制,确保患者隐私安全。

6. 眼科伦理实践规范(GPS;非常同意:93.94%,同意:6.06%)。在部署生成式 AI 大模型时,须确保大模型满足个人尊严、自主权、知情权及隐私保护等伦理要求,并持续评估 AI 大模型从部署到使用阶段的伦理合规性^[63]。应尤其注意对患者隐私进行保护;避免虚假输出,确保大模型不出现生成错误或误导性诊断信息;避免自动化偏见,眼科医师不应完全依赖 AI 输出结果;对弱势群体进行保护,对视力残障等边缘人群,大模型需要提供无障碍访问(如语音输出),并避免造成伤害^[64];在交互设计中,避免 AI 伪装为人类主体,如大模型使用“我”等拟人化称谓和表述,或使用人类头像,并避免用户对其产生错误的信任感和情感依赖;所有大模型生成内容或基于这些内容的产出,均应明确、显著标记为 AI 辅助生成,避免对患者造成伤害。

八、眼科运营、使用生成式 AI 大模型应注意临床验证与优先级管理、过渡期设置、弱势群体专科保护、眼科数据及网络安全保障、眼科法规体系构建、生成内容专科化管控、生态可持续发展、眼科精



准问责机制确立等方面

说明:这是关于眼科运营和使用生成式 AI 大模型注意事项的共识。

1. 临床验证与优先级管理(GPS; 非常同意: 90.91%, 同意: 9.09%)。在眼科应用的生成式 AI 大模型被充分测试并证明其有效性和安全性前, 不应将大模型的使用优先级置于已使用的 AI 技术之前, 也不应置于已证明更具经济效益的非 AI 或非计算机辅助的解决方案之前, 尤其对于经济效益更高的传统眼科诊疗方法^[14, 65]。例如, 在白内障筛查中, 若传统的裂隙灯检查法已被证明准确且成本效益高, 应优先使用, 直至大模型在诊断精度和安全性方面超越现有方法。

2. 过渡期设置(GPS; 非常同意: 84.85%, 同意: 15.15%)。眼科应用生成式 AI 大模型可能改变医疗模式、劳动岗位及政策法规, 为此需要设立过渡期, 确保相关人员和系统平稳适应^[66], 使眼科医师和技术人员逐步掌握 AI 工具的操作技能, 取得相应资质认证^[67], 并为因 AI 技术进步而面临职业转型的眼科工作人员(如传统影像分析师)提供再培训机会。政府和行业协会需要借此调整法规和行业标准, 如制订眼科 AI 设备的认证标准, 避免技术推广导致的混乱。

3. 弱势群体专科保护(GPS; 非常同意: 93.94%, 同意: 6.06%)。眼科应用生成式 AI 大模型应特别关注弱势群体的需求, 如儿童(儿童保护)^[68]、视力残障人士(无障碍服务)和低收入人群(普惠医疗)^[69]。例如, 大模型应支持无障碍功能, 如语音导航或大字体界面, 方便视力残障患者使用 AI 驱动的眼科筛查工具^[70]。同时, 为避免加剧社会不平等, 可通过政策支持提供免费或低成本的 AI 眼科筛查服务(如针对糖尿病视网膜病变的早期检测), 缩小富裕与贫困群体在眼科医疗资源获取方面的差距。此外, 参与大模型开发的眼科医师、数据标注员等群体, 应享有同等的职业保护和福利, 避免因技术进步导致劳动权益受损。

4. 眼科数据及网络安全保障(GPS; 非常同意: 93.94%, 同意: 6.06%)。眼科医疗系统因涉及敏感数据(如患者眼底图像和个人健康信息)而成为网络攻击的高风险目标, 生成式 AI 大模型的使用可能进一步放大该风险, 其使用的数据可能被操纵、劫持^[68]。例如, 眼底图像若被恶意篡改或加入肉眼不可见噪声, 则可能导致误诊, 如将正常视网膜误判为黄斑变性^[71], 应用于眼科临床的大模型必须通

过对抗攻击测试。为此, 须增强大模型的安全管控^[72], 严格控制大模型的访问权限, 禁止将患者身份信息直接输入大模型, 并采用加密技术和定期安全审计方法, 确保数据传输和存储的安全性^[73]。

5. 眼科法规体系构建(GPS; 非常同意: 84.85%, 同意: 12.12%)。眼科开发和利用生成式 AI 大模型应注意审批标准、数据合规、患者权益、消费者权益等^[74]。例如, 训练模型使用的眼底图像数据集可能未经患者明确同意, 或未提供数据删除选项, 这违反了数据保护原则。为此, 中国政府应牵头制定应用生成式 AI 大模型的审查和监管制度, 明确数据采集的合法性要求, 设定未成年人数据使用的年龄门槛, 并建立患者隐私保护机制^[46]。同时, 需要完善行业标准, 确保 AI 输出结果的准确性, 并为付费用户提供消费者权益保障。

6. 生成内容专科化管控(GPS; 非常同意: 87.88%, 同意: 9.09%)。生成式 AI 大模型的生成内容包括文字、图像、视频等信息, 可能存在误导性或错误内容, 如不准确的诊断或治疗建议。例如, 大模型可能将正常眼底图像误诊为视网膜脱离, 引发患者恐慌或不必要的干预。为此, 运营期间须对大模型输出进行严格审查, 确保诊断建议准确、客观, 并避免生成歧视性信息(如基于性别或种族的偏见性诊断)。生成式 AI 大模型在运营期间应受到严格审查、监管, 确保其有助于人民群众健康发展, 可引入眼科专家监督机制, 对 AI 生成的高风险诊断进行人工复核, 保护患者权益^[75]。目前许多患者直接向大模型询问眼病相关问题, 应禁止将未经医疗器械注册的生成式 AI 大模型以软件等形式, 直接面向患者提供诊断或治疗建议; 已上线的大模型应在显著位置标注“本产品仅供健康教育, 不作为诊疗依据”等提示, 并限制对高风险问题(如我是不是要失明)进行回答。

7. 生态可持续发展(GPS; 非常同意: 93.94%, 同意: 6.06%)。眼科长期运营生成式 AI 大模型将消耗大量能源(尤其电能)和水资源^[48-49], 尤其训练和运行大规模眼科图像和视频分析模型时。例如, 处理高分辨率 OCT 图像的大模型, 可能显著增加能量消耗。为此, 应鼓励使用可再生能源或无碳能源驱动的 AI 系统^[76], 并追踪大模型的碳足迹和水足迹, 探索节能技术, 如 AI 模型压缩或高效计算架构, 可减少环境影响, 实现眼科 AI 的可持续发展。

8. 眼科精准问责机制确立(GPS; 非常同意: 81.82%, 同意: 12.12%)。为应对眼科应用生成式



AI 大模型可能引发的错误或滥用,甚至造成的伤害,需要建立明确的问责制度,确保受害者获得赔偿^[77]。法律法规应界定 AI 在过错中的参与程度,明确开发者、提供者、部署者、使用者等不同身份的责任,避免责任推诿,并设立报告制度及补救机制,确保赔偿足以弥补患者的损失(如手术费用或生活质量下降)^[78]。可设立独立的第三方评估机制,对争议事件进行技术和责任评估。例如,若 AI 误诊青光眼导致患者视力损失,需要明确开发者、医疗机构和监管机构等在责任链中的角色。算法错误、训练数据偏倚应由开发者、数据提供方、部署机构共同承担责任,部署后性能下降应由开发者、部署机构及使用机构共同承担责任,操作不当应由操作人员、培训机构、系统设计者共同承担责任,监管机构监管不力也应承担相应责任,但具体责任划分应当根据具体情况确定,不可一概而论。同时,需要建立案例数据库,分析误诊原因,持续优化模型性能。可考虑开发眼科 AI 责任保险产品,设立专项风险准备金。

九、眼科应用生成式 AI 大模型可能存在不恰当甚至错误的输出及其导致的风险、数据质量问题、道德及责任问题,以及人类技能退化、弱势群体相关风险、医患关系恶化、隐私泄露、科技竞争风险等问题

说明:这是关于眼科应用生成式 AI 大模型潜在风险的共识。

1. 不恰当甚至错误的输出导致的风险(GPS; 非常同意:93.94%, 同意:6.06%)。生成式 AI 大模型通常基于海量数据训练,若数据质量存在问题,可能导致大模型输出错误结果,甚至是难以判断的、与真实情况相似的虚假回答,有时大模型可能编造并输出虚假信息,即幻觉^[35, 75]。生成式 AI 大模型不产生知识,也可能不理解输出的答案,没有道德约束和场景推理,其输出的错误内容可能被其他 AI 模型误学习,导致错误知识传播,甚至造成眼科医师水平下降^[74]。大模型常可用于多种场景,若其对应应用场景判断错误,可能给出不恰当的输出。有时对大模型的问题或提示稍作修改,其便可能输出完全不同的结果,甚至完全不合适的结果。大模型可执行多种任务,有时难以对其可执行的所有任务进行完善的效能测试,因此应当避免高估、夸大模型的效能而低估风险。对于此类风险,可强制所有存在高风险的输出结果附加“结果为 AI 生成”的说明;在大模型中部署不确定性模块,并设定强

制转为人工处理的阈值;建立全国眼科 AI 幻觉案例库,定期更新、发布以推动行业发展。

2. 数据质量问题(GPS; 非常同意:93.94%, 同意:6.06%)。目前鲜有商用 AI 大模型公布训练数据集,数据质量难以保证^[74]。同时,医疗 AI 大模型的训练数据通常来源于特定人群的医学信息,难以保证在种族、国家、性别、血统、年龄、收入方面不存在偏倚。因此,为使生成式 AI 大模型能覆盖所有使用人群,在模型设计阶段便应注意提高数据的质量^[35],尽可能避免优势者偏倚,如高收入国家的 AI 使用价格反而更为低廉^[79]。

3. 道德及责任问题(GPS; 非常同意:87.88%, 同意:12.12%)。目前尚无获得国家批准的医疗用生成式 AI 大模型。生成式 AI 大模型应用于临床后,存在医务工作者过度依赖大模型的可能,相信大模型输出的不符合道德、医学伦理要求甚至法律法规的结果,从而面临伦理甚至道德困境^[80]。生成式 AI 大模型若由不恰当的群体负责,如盈利组织,则对大模型的监管和审查将存在困难,同时由于利益驱动,可能使大模型输出不恰当的结果,甚至使使用者受到不恰当的引导,从而发生某些风险,包括但不限于经济损失、自我伤害等^[19, 54]。

4. 人类技能退化(GPS; 非常同意:84.85%, 同意:9.09%)。生成式 AI 大模型的高性能可能使医务工作者将更多的日常工作 and 责任转移给 AI,可能使人类的医务工作专业水平下降;而专业水平下降将使人类更难以发现大模型输出结果中的问题或错误^[81]。若生成式 AI 大模型用于医学教育,可能增加医学生放弃自己的判断,转而听从大模型意见的可能性,进而影响医学教育质量^[82]。当无法应用大模型时,或不具备相应硬件、网络时,医务工作者如何继续提供医疗服务,应当具有后备方案^[74]。对于住院医师培训,应确保其在不使用 AI 辅助的情况下,仍能达到对核心疾病进行独立诊疗的要求。

5. 弱势群体相关风险(GPS; 非常同意:87.88%, 同意:6.06%)。儿童是应首要关注的弱势群体。将生成式 AI 大模型应用于儿童,可能存在潜在的弊端和风险^[46]。不仅眼科,整个医学领域的 AI 模型将如何影响儿童的身心健康,尚无明确结论。训练和测试数据是否包含儿童资料,生成式 AI 大模型对儿童群体是否有普适性、是否可能对儿童造成伤害,其测试中是否有针对儿童群体的内容等,均应在设计阶段进行考虑。与此类似,部分



眼科患者视力不可逆下降,甚至达到视力残疾标准,是使用生成式 AI 大模型的弱势群体,应避免弱势群体数据被忽视,同时应考虑低视力群体能否获得使用大模型的资源^[54]。对于缺乏支持和资金的群体,应增强其对生成式 AI 大模型的可及性。此外,开发、运营大模型需要支付费用,可能仅有较为富裕的群体才能接触到大模型,因此需要警惕大模型有可能加剧贫富人群间的健康差距,进一步扩大健康不平等^[63]。

6. 医患关系恶化(GPS;非常同意:84.85%,同意:12.12%)。生成式 AI 大模型介入医患之间,医患沟通可能减少,双方更多依赖大模型的输出结果^[83],如治疗方案或病情预测方面,对医务人员专业判断和经验的依赖减少。目前阶段的大模型尚无法取代医务人员,而大模型介入医患之间,可能导致医患信任度下降,影响共情和互动,降低透明度^[54, 63],最终造成医患关系发生根本改变。

7. 隐私泄露(GPS;非常同意:90.91%,同意:9.09%)。生成式 AI 大模型使用的海量数据,可能被大模型意外泄露,或基于使用者的要求而泄露^[74]。此外,存在多种尝试攻破 AI 模型的算法,其对隐私保护构成威胁^[46, 62]。眼科数据的特殊性之一在于虹膜、视网膜等组织本身具有生物身份识别特征,即使将患者的姓名等资料脱敏,仍然存在数据泄露风险。部分使用合成数据的研究者认为,合成数据可降低隐私泄露风险,但目前仍存在逆向还原风险和病理特征失真可能,使用前需要经过真实数据验证大模型的效能无明显下降。

8. 科技竞争中的风险(GPS;非常同意:84.85%,同意:9.09%)。在生成式 AI 大模型的研发、部署过程中,大型科技企业常占据主导和中心地位,而科技企业可能在科技竞争中为了避免政府、高等院校、初创企业对自身利益的不利影响,不能坚持承担医学伦理方面的社会责任;即使自愿承诺遵守伦理和法律法规的要求,但仍缺少约束力^[84]。同时,大型科技企业可能没有医疗产品研发经验,无法满足医药卫生领域的需求以及人民群众健康的需求^[11]。此外,科技企业应当保证生成式 AI 大模型对监管、审查机构保持充分的透明度,不应为了商业利益而急于将不成熟的大模型推向市场,令人民群众承担不良后果^[85]。应当充分鼓励高等院校、科研院所、企业等进行生成式 AI 大模型的研发。支持政府探索医疗专病专科领域生成式 AI 大模型监管和备案机制。

十、短期内眼科应用的生成式 AI 大模型可能具有专科化 AI 模型架构演进、轻量化部署适配眼科不同真实场景、眼科多模态数据 AI 模型突破、算力成本优化与生态构建等发展趋势

说明:这是关于眼科应用的生成式 AI 大模型短期内发展趋势的共识。

1. 专科化 AI 模型架构演进(GPS;非常同意:87.88%,同意:9.09%)。随着通用 AI 大模型发展,高品质训练数据和突破性架构的短缺限制了其在眼科等领域的应用^[14, 86-87]。未来,AI 大模型有望趋于专业化,如基于混合专家 AI 模型架构等的眼科专用 AI 大模型,将提升诊断和治疗建议的精准性,尤其在眼科不同专业领域将表现突出。

2. 轻量化部署适配眼科不同真实场景(GPS;非常同意:87.88%,同意:6.06%)。眼科生成式 AI 大模型可能向轻量化部署以适配不同真实场景的趋势发展。通过采用知识蒸馏等技术,AI 大模型的知识可迁移至小型 AI 模型,保留核心推理和生成能力^[88]。轻量化模型部署更灵活,适用于移动设备和远程医疗场景,如便携式眼科筛查工具,既降低资源消耗,又能满足基层医疗需求^[86-87, 89]。

3. 眼科多模态数据 AI 模型突破(GPS;非常同意:87.88%,同意:9.09%)。生成式 AI 大模型更加注重高知识密度、高质量数据。高知识密度数据对提升 AI 大模型的推理能力至关重要。眼科领域需要依赖多模态标准化数据集^[74, 90],如专家标注的眼科数据和知识密度极高的数据,以提高眼科 AI 大模型的思维链质量和数量。目前数据透明度不足,未来需要建立开放共享机制,以推动 AI 大模型性能优化^[74]。

4. 算力成本优化与生态构建(GPS;非常同意:90.91%,同意:6.06%)。随着算力基础设施完善和政策支持,AI 大模型的训练和使用成本将逐步下降^[50, 91],这为眼科 AI 大模型研发提供了便利,使基层医疗机构和科研机构更易采用 AI 技术,加速眼科诊疗智能化进程。在中国政府的大力支持下,中国在建立开放共享数据集、促进多学科合作以及推动 AI 在眼科实际应用等方面已具备优势。

十一、撰写眼科应用的生成式 AI 大模型相关文章,应注意数据管理标准、临床验证体系、研究成果报告规范等方面

说明:这是关于撰写眼科应用的生成式 AI 大模型文章注意事项的共识。尽管生成式 AI 大模型应用于眼科临床诊疗尚需时日,但国际技术进展迅

速,中国需要加速研究步伐。以下是撰写相关文章时需要关注的要点^[92-95],以确保内容规范、透明且具有眼科针对性。

1. 数据管理标准(GPS;非常同意:90.91%,同意:9.09%)。(1)数据质量:训练数据须高质量^[90]。例如眼底图像数据应当具有高分辨率,屈光度数、眼轴长度等测量数据应准确,避免错误或偏倚,以防止生成式 AI 大模型生成误导性结果;收集的眼科数据应当具有代表性和多样性,以保证大模型结果的广泛适用性和可靠性。(2)隐私保护:在数据存储、传输、管理过程中,应确保患者隐私安全,防止数据泄露^[62]。除了姓名、编号等数据外,虹膜、视网膜等具有生物身份识别特征数据需要严格去标识化,采用加密和分布式部署等技术,有助于确保个人身份信息的安全。试点隐私计算技术,在数据不出域前提下实现多方协作建立大模型。(3)数据共享:进行跨机构或跨国数据共享时,需要签订数据共享协议,明确数据使用范围和保护措施,确定各环节责任人,确保各方遵守相应的法律法规。建立数据使用分级授权机制,基于研究目的、机构资质分配不同访问权限。(4)数据保密:国家法律法规规定的保密数据、国家或特定机构专有数据及其他敏感数据,不能输入至大模型中;涉及人类遗传学的资源和非公开研究数据,亦应在监管下合理使用;暴露或有可能披露患者个人信息的数据,应当在脱敏处理后使用。可采用数据信托模式,由政府批准的独立第三方管理敏感数据,在保障隐私的前提下支持研究。

2. 临床验证体系(GPS;非常同意:87.88%,同意:12.12%)。(1)验证测试:明确应用场景(如视网膜病变筛查、青光眼风险评估、诊疗方案推荐),基于循证医学研究等,逐一验证生成式 AI 大模型各指标的性能,确保外部测试集独立于训练集,评估结果稳定可靠,临床安全有效,最终形成高等级的循证医学证据^[96]。(2)结果解读:生成式 AI 大模型的输出结果在应用于临床实践前,应由眼科专业人士进行解读及审查,确保结果的合理性及可行性^[65]。结果报告须涵盖正面和负面发现,避免选择性披露。多模态大模型在输出结果的同时,强烈建议提供可视化注意力热图或关键病灶定位框等增强可解释性的内容。(3)风险管理:建议成立专家委员会,纳入眼科医师、计算机专家、数据采集人员、数据审查及分析人员等各方代表,识别并控制潜在风险(如输出不确定性、误判),制订应对流程,确保

大模型安全应用于医疗领域^[97]。

3. 研究成果报告规范(GPS;非常同意:90.91%,同意:9.09%)。(1)使用声明:明确披露使用的 AI 模型、版本及训练数据来源,确保使用情况清晰。(2)方法披露:向监管机构提供训练、验证和测试方法,包括数据预处理和参数设置,确保透明性和可重复性^[65]。(3)利益冲突:说明资金来源及商业合作关系,保证利益透明,避免潜在冲突。(4)标准规范:遵循业界报告格式和评价框架^[96, 98],提供详细附加资料(如代码结构、数据集描述)以供参考。(5)透明度:公开训练和验证数据集,鼓励分享代码和算法,接受独立专家审查,提升生成式 AI 大模型的可靠性和客观性。[扫描本文首页下方二维码,可查看眼科医学文献推荐生成式 AI 大模型报告的格式和内容(附录)]

十二、结语

通过制订使用规范和相关要求,可以确保生成式 AI 大模型在眼科实践应用中的透明性、规范性和可问责性,最大限度发挥其潜力,同时保障患者安全和数据隐私等利益。随着生成式 AI 大模型应用推广,本《共识》将在充分讨论的基础上不断更新。

形成共识意见的专家组成员:

陈有信 中国医学科学院 北京协和医学院 北京协和医院眼科[国际电气电子工程师学会技术与工程分会智能医疗数字化技术委员会智能眼科分部(中国)执行长]

[以下智能眼科分部(中国)顾问按姓氏拼音排序]

黄天荫 清华大学医学院
刘江 南方科技大学计算机科学与工程系
戚扬 北京汉唐自远技术股份有限公司
孙兴怀 复旦大学附属眼耳鼻喉科医院眼科
王宁利 首都医科大学附属北京同仁医院北京同仁眼科中心

[以下智能眼科分部(中国)常务委员按姓氏拼音排序]

陈蔚 温州医科大学附属眼视光医院
陈吉利 上海市市北医院眼科
陈新建 苏州大学电子信息学院
迟玮 深圳市眼科医院
丁晓伟 上海交通大学苏州体素信息科技有限公司
冯天宜 中国信息通信研究院
李冰 视微影像(河南)科技有限公司
李程 厦门大学医学院
李光 广东唯仁医疗科技有限公司



李超宏 苏州微清医疗器械有限公司
 李锡荣 中国人民大学信息学院
 林浩添 中山大学中山眼科中心
 刘庆淮 江苏省人民医院眼科
 邵毅 上海交通大学附属第一人民医院眼科
 孙鹏 天津鹏升科技有限公司
 孙宇辉 北京致远慧图科技有限公司
 王雁 天津市眼科医院
 许言午 华南理工大学未来技术学院
 杨卫华 深圳市眼科医院
 叶娟 浙江大学医学院附属第二医院眼科中心
 袁进 首都医科大学附属北京同仁医院北京同仁眼科中心
 张明 四川大学华西医院眼科
 张成文 北京邮电大学计算机学院
 张大磊 北京鹰瞳科技发展股份有限公司
 张冬冬 北京至真健康科技有限公司
 张铭志 汕头大学·香港中文大学联合汕头国际眼科中心
 张秀兰 中山大学中山眼科中心
 周少华 中国科学技术大学生物医学工程学院
 周行涛 复旦大学附属耳鼻喉科医院眼科
 [以下智能眼科分部(中国)委员按姓氏拼音排序]
 柴勇 江西省儿童医院眼科
 陈景尧 昆明市第一人民医院眼科
 戴琦 温州医科大学附属眼视光医院
 董贺 大连市第三人民医院眼科
 范雯 江苏省人民医院眼科
 郭竞敏 华中科技大学同济医学院附属同济医院眼科
 何鲜桂 上海市眼病防治中心
 胡丽丹 浙江大学医学院附属儿童医院眼科
 黄旭 杭州华夏眼科医院
 黄锦海 复旦大学附属耳鼻喉科医院眼科
 黄明海 广西医科大学第二附属医院眼科
 纪振宇 武汉大学人民医院眼科
 李飞 中山大学中山眼科中心
 李笠 福州大学附属省立医院眼科
 李凯彦 中国科学技术大学苏州高等研究院
 林嘉雯 福州大学计算机与大数据学院
 刘红玲 哈尔滨医科大学附属第一医院眼科
 马翔 大连医科大学附属第一医院眼科
 马韶东 中国科学院宁波材料技术与工程研究所
 苏婷 武汉大学人民医院眼科
 孙大卫 哈尔滨医科大学附属第二医院眼科
 孙旭芳 华中科技大学同济医学院附属同济医院眼科
 谭钢 南华大学附属第一医院眼科
 齐虹 北京大学第三医院眼科
 邱坤良 汕头大学·香港中文大学联合汕头国际眼科中心
 王烽 梅州市人民医院眼科

王少攀 厦门大学附属第一医院眼科
 吴振凯 常德市第一人民医院眼科
 谢华桃 华中科技大学同济医学院附属协和医院眼科
 杨阳 岳阳市中心医院眼科
 杨景元 中国医学科学院北京协和医学院北京协和医院(执笔)
 杨于力 陆军军医大学西南医院眼科
 张冰 杭州市儿童医院眼科
 张含 中国医科大学附属第一医院眼科
 张国明 深圳市眼科医院
 赵欣宇 中国医学科学院北京协和医学院北京协和医院
 赵一天 中国科学院宁波材料技术与工程研究所
 冯时 中国医学科学院北京协和医学院北京协和医院眼科,非委员,整理资料
 程世瑜 中国医学科学院北京协和医学院北京协和医院眼科,非委员,整理资料
 古兴旺 中国医学科学院北京协和医学院北京协和医院眼科,非委员,整理资料
 (参与讨论的其他专家按姓氏拼音排序)
 张纯 北京清华长庚医院眼科
 邹海东 上海交通大学医学院附属第一人民医院上海市眼病防治中心

声明 本共识制订过程遵循科学性、实用性及一致性原则,旨在为临床医疗服务提供指导,不是在各种情况下均必须遵循的医疗标准,也不是为个别特殊个体提供的保健措施。本共识可能存在未尽之处,与其不一致的实践方法不应被视为错误或不当。所有形成共识意见的专家组成员声明不存在利益冲突。本共识制订未接受赞助

参 考 文 献

- [1] Feng X, Xu K, Luo MJ, et al. Latest developments of generative artificial intelligence and applications in ophthalmology[J]. Asia Pac J Ophthalmol (Phila), 2024, 13(4): 100090. DOI: 10.1016/j.apjo.2024.100090.
- [2] Waisberg E, Ong J, Kamran SA, et al. Generative artificial intelligence in ophthalmology[J]. Surv Ophthalmol, 2025, 70(1): 1-11. DOI: 10.1016/j.survophthal.2024.04.009.
- [3] Alonso-Coello P, Oxman AD, Moberg J, et al. GRADE evidence to decision (EtD) frameworks: a systematic and transparent approach to making well informed healthcare choices. 2: clinical practice guidelines[J]. BMJ, 2016, 353: i2089. DOI: 10.1136/bmj.i2089.
- [4] Alonso-Coello P, Schünemann HJ, Moberg J, et al. GRADE evidence to decision (EtD) frameworks: a systematic and transparent approach to making well informed healthcare choices. 1: introduction[J]. BMJ, 2016, 353: i2016. DOI: 10.1136/bmj.i2016.
- [5] Schünemann HJ, Mustafa R, Brozek J, et al. GRADE guidelines: 16. GRADE evidence to decision frameworks for tests in clinical practice and public health[J]. J Clin Epidemiol, 2016, 76: 89-98. DOI: 10.1016/j.jclinepi.2016.01.032.
- [6] Guyatt GH, Oxman AD, Kunz R, et al. GRADE guidelines: 7.



- rating the quality of evidence—inconsistency[J]. *J Clin Epidemiol*, 2011, 64(12): 1294-1302. DOI: 10.1016/j.jclinepi.2011.03.017.
- [7] Tsuchiya Y, Hiramoto N. Measuring consensus and dissensus: a generalized index of disagreement using conditional probability[J]. *Information Sciences*, 2018, 439-440: 50-60. DOI: 10.1016/j.ins.2018.02.003.
- [8] Diamond IR, Grant RC, Feldman BM, et al. Defining consensus: a systematic review recommends methodologic criteria for reporting of Delphi studies[J]. *J Clin Epidemiol*, 2014, 67(4): 401-409. DOI: 10.1016/j.jclinepi.2013.12.002.
- [9] Han Y, Ceross A, Bergmann J. Regulatory frameworks for AI-enabled medical device software in China: comparative analysis and review of implications for global manufacturer[J]. *JMIR AI*, 2024, 3: e46871. DOI: 10.2196/46871.
- [10] World Health Organization. WHO calls for safe and ethical AI for health[EB/OL]. (2023-05-16) [2025-06-17]. <https://www.who.int/news/item/16-05-2023-who-calls-for-safe-and-ethical-ai-for-health>.
- [11] Farhud DD, Zokaei S. Ethical issues of artificial intelligence in medicine and healthcare[J]. *Iran J Public Health*, 2021, 50(11): i-v. DOI: 10.18502/ijph.v50i11.7600.
- [12] Li Z, Wang L, Wu X, et al. Artificial intelligence in ophthalmology: the path to the real-world clinic[J]. *Cell Rep Med*, 2023, 4(7): 101095. DOI: 10.1016/j.xcrm.2023.101095.
- [13] Wei Q, Qian K, Li X. FunBench: benchmarking fundus reading skills of MLLMs[J]. *arXiv*, 2025: 2503.00901. DOI: 10.48550/arXiv.2503.00901.
- [14] Artiaga JCM, Guevarra MCB, Sosuan GMN, et al. Large language models in ophthalmology: a scoping review on their utility for clinicians, researchers, patients, and educators[J]. *Eye (Lond)*, 2025, 39(15): 2752-2761. DOI: 10.1038/s41433-025-03935-7.
- [15] Domalpally A, Channa R. Real-world validation of artificial intelligence algorithms for ophthalmic imaging[J]. *Lancet Digit Health*, 2021, 3(8): e463-e464. DOI: 10.1016/S2589-7500(21)00140-0.
- [16] de Hond A, Leeuwenberg T, Bartels R, et al. From text to treatment: the crucial role of validation for generative large language models in health care[J]. *Lancet Digit Health*, 2024, 6(7): e441-e443. DOI: 10.1016/S2589-7500(24)00111-0.
- [17] World Health Organization. Ethics and governance of artificial intelligence for health: Guidance on large multi-modal models[EB/OL]. (2025-03-25) [2025-06-17]. <https://www.who.int/publications/i/item/9789240084759>.
- [18] Naik N, Hameed B, Shetty DK, et al. Legal and ethical consideration in artificial intelligence in healthcare: who takes responsibility?[J]. *Front Surg*, 2022, 9: 862322. DOI: 10.3389/fsurg.2022.862322.
- [19] Schwabe D, Becker K, Seyferth M, et al. The METRIC-framework for assessing data quality for trustworthy AI in medicine: a systematic review[J]. *NPJ Digit Med*, 2024, 7(1): 203. DOI: 10.1038/s41746-024-01196-4.
- [20] Luca AR, Ursuleanu TF, Gheorghe L, et al. Impact of quality, type and volume of data used by deep learning models in the analysis of medical images[J]. *Inform Med Unlocked*, 2022, 29: 100911. DOI: 10.1016/j.imu.2022.100911.
- [21] Guillen-Aguinaga M, Aguinaga-Ontoso E, Guillen-Aguinaga L, et al. Data quality in the age of AI: a review of governance, ethics, and the FAIR principles[J]. *Data*, 2025, 10(12): 201. DOI: 10.3390/data10120201.
- [22] Jin K, Yu T, Grzybowski A. Multimodal artificial intelligence in ophthalmology: applications, challenges, and future directions [J]. *Surv Ophthalmol*, 2026, 71(1): 158-167. DOI: 10.1016/j.survophthal.2025.07.003.
- [23] Wang C, Zhang J, Lassi N, et al. Privacy protection in using artificial intelligence for healthcare: Chinese regulation in comparative perspective[J]. *Healthcare(Basel)*, 2022, 10(10): 1878. DOI: 10.3390/healthcare10101878.
- [24] Martin KD, Zimmermann J. Artificial intelligence and its implications for data privacy[J]. *Curr Opin Psychol*, 2024, 58: 101829. DOI: 10.1016/j.copsyc.2024.101829.
- [25] Lim JS, Hong M, Lam W, et al. Novel technical and privacy-preserving technology for artificial intelligence in ophthalmology[J]. *Curr Opin Ophthalmol*, 2022, 33(3): 174-187. DOI: 10.1097/ICU.0000000000000846.
- [26] Teo ZL, Zhang X, Yang Y, et al. Privacy-preserving technology using federated learning and blockchain in protecting against adversarial attacks for retinal imaging [J]. *Ophthalmology*, 2025, 132(4): 484-494. DOI: 10.1016/j.ophtha.2024.10.017.
- [27] Li X, Cong Y. Exploring barriers and ethical challenges to medical data sharing: perspectives from Chinese researchers[J]. *BMC Med Ethics*, 2024, 25(1): 132. DOI: 10.1186/s12910-024-01135-8.
- [28] Miotto R, Wang F, Wang S, et al. Deep learning for healthcare: review, opportunities and challenges[J]. *Brief Bioinform*, 2018, 19(6): 1236-1246. DOI: 10.1093/bib/bbx044.
- [29] Van Veen D, Van Uden C, Blankemeier L, et al. Adapted large language models can outperform medical experts in clinical text summarization[J]. *Nat Med*, 2024, 30(4): 1134-1142. DOI: 10.1038/s41591-024-02855-5.
- [30] Bedi S, Liu Y, Orr-Ewing L, et al. Testing and evaluation of health care applications of large language models: a systematic review[J]. *JAMA*, 2025, 333(4): 319-328. DOI: 10.1001/jama.2024.21700.
- [31] Thirunavukarasu AJ, Mahmood S, Malem A, et al. Large language models approach expert-level clinical knowledge and reasoning in ophthalmology: a head-to-head cross-sectional study[J]. *PLOS Digit Health*, 2024, 3(4): e0000341. DOI: 10.1371/journal.pdig.0000341.
- [32] Clusmann J, Kolbinger FR, Muti HS, et al. The future landscape of large language models in medicine[J]. *Commun Med (Lond)*, 2023, 3(1): 141. DOI: 10.1038/s43856-023-00370-1.
- [33] Safranek CW, Sidamon-Eristoff AE, Gilson A, et al. The role of large language models in medical education: applications and implications[J]. *JMIR Med Educ*, 2023, 9: e50945. DOI: 10.2196/50945.
- [34] Lucas HC, Upperman JS, Robinson JR. A systematic review of large language models and their implications in medical education[J]. *Med Educ*, 2024, 58(11): 1276-1285. DOI: 10.1111/medu.15402.
- [35] Feng X, Xu K, Luo MJ, et al. Latest developments of generative artificial intelligence and applications in ophthalmology[J]. *Asia Pac J Ophthalmol (Phila)*, 2024,



- 13(4): 100090. DOI: 10.1016/j.apjo.2024.100090.
- [36] Swaminathan U, Daigavane S. Unveiling the potential: a comprehensive review of artificial intelligence applications in ophthalmology and future prospects[J]. *Cureus*, 2024, 16(6): e61826. DOI: 10.7759/cureus.61826.
- [37] RaviChandran N, Teo ZL, Ting D. Artificial intelligence enabled smart digital eye wearables[J]. *Curr Opin Ophthalmol*, 2023, 34(5): 414-421. DOI: 10.1097/ICU.0000000000000985.
- [38] Iliuță M-E, Moisesescu M-A, Caramihai S-I, et al. Digital twin models for personalised and predictive medicine in ophthalmology[J]. *Technologies*, 2024, 12(4): 55. DOI: 10.3390/technologies12040055.
- [39] Oala L, Murchison AG, Balachandran P, et al. Machine learning for health: algorithm auditing & quality control [J]. *J Med Syst*, 2021, 45(12): 105. DOI: 10.1007/s10916-021-01783-y.
- [40] Gandhi S, Balas M. Artificial intelligence for ophthalmic drug discovery and development: capabilities, applications, and challenges[J]. *Artif Intel Health*, 2024, 1(3): 26-30. DOI: 10.36922/aih.3341.
- [41] Bertagnolli MM. Advancing health through artificial intelligence/machine learning: the critical importance of multidisciplinary collaboration[EB/OL]. (2023-12-19) [2025-06-17]. <https://academic.oup.com/pnasnexus/article/2/12/pgad356/7477226>.
- [42] Yang X, Chen A, PourNejatian N, et al. A large language model for electronic health records[J]. *NPJ Digit Med*, 2022, 5(1): 194. DOI: 10.1038/s41746-022-00742-2.
- [43] Ning Y, Teixayavong S, Shang Y, et al. Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist[J]. *Lancet Digit Health*, 2024, 6(11): e848-e856. DOI: 10.1016/S2589-7500(24)00143-2.
- [44] Abdullah YI, Schuman JS, Shabsigh R, et al. Ethics of artificial intelligence in medicine and ophthalmology[J]. *Asia Pac J Ophthalmol (Phila)*, 2021, 10(3): 289-298. DOI: 10.1097/APO.0000000000000397.
- [45] Haltaufderheide J, Ranisch R. The ethics of ChatGPT in medicine and healthcare: a systematic review on Large language models (LLMs) [J]. *NPJ Digit Med*, 2024, 7(1): 183. DOI: 10.1038/s41746-024-01157-x.
- [46] Wang Y, Liu C, Zhou K, et al. Towards regulatory generative AI in ophthalmology healthcare: a security and privacy perspective[J]. *Br J Ophthalmol*, 2024, 108(10): 1349-1353. DOI: 10.1136/bjo-2024-325167.
- [47] Karakaş Ü, Özdemir V. Artificial intelligence and environmental impact: moving beyond humanizing vocabulary and anthropocentrism[J]. *OMICS*, 2025, 29(1): 2-4. DOI: 10.1089/omi.2024.0197.
- [48] Berreby D. As use of A.I. soars, so does the energy and water it requires[EB/OL]. (2024-02-06) [2025-06-17]. <https://e360.yale.edu/features/artificial-intelligence-climate-energy-emissions>.
- [49] Yu Y, Wang J, Liu Y, et al. Revisit the environmental impact of artificial intelligence: the overlooked carbon emission source? [J]. *Front Environ Sci Eng*, 2024, 18: 158. DOI: 10.1007/s11783-024-1918-y.
- [50] Alzoubi YI, Mishra A. Green artificial intelligence initiatives: potentials and challenges[J]. *J Clean Prod*, 2024, 468: 143090. DOI: 10.1016/j.jclepro.2024.143090.
- [51] Eiras F, Petrov A, Vidgen B, et al. Risks and opportunities of open-source generative AI[EB/OL]. (2024-05-14) [2025-06-17]. <https://arxiv.org/abs/2405.08597>.
- [52] Bildirici F. Open-source AI: an approach to responsible artificial intelligence development[J]. *Reflektif J Soc Sci*, 2024, 5(1): 73-81. DOI: 10.47613/reflektif.2024.145.
- [53] Beede E, Baylor E, Hersch F, et al. A human-centered evaluation of a deep learning system deployed in clinics for the detection of diabetic retinopathy[G]//*Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, Honolulu HI, 2020. New York: Association for Computing Machinery, 2020: 1-12. DOI: 10.1145/3313831.3376718.
- [54] Evans NG, Wenner DM, Cohen IG, et al. Emerging ethical considerations for the use of artificial intelligence in ophthalmology[J]. *Ophthalmol Sci*, 2022, 2(2): 100141. DOI: 10.1016/j.xops.2022.100141.
- [55] Heinke A, Radgoudarzi N, Huang BB, et al. A review of ophthalmology education in the era of generative artificial intelligence[J]. *Asia Pac J Ophthalmol (Phila)*, 2024, 13(4): 100089. DOI: 10.1016/j.apjo.2024.100089.
- [56] Ahmed MI, Spooner B, Isherwood J, et al. A systematic review of the barriers to the implementation of artificial intelligence in healthcare[J]. *Cureus*, 2023, 15(10): e46454. DOI: 10.7759/cureus.46454.
- [57] Reddy S, Rogers W, Makinen VP, et al. Evaluation framework to guide implementation of AI systems into healthcare settings[J]. *BMJ Health Care Inform*, 2021, 28: e100444. DOI: 10.1136/bmjhci-2021-100444.
- [58] Chen D, Geevarghese A, Lee S, et al. Transparency in artificial intelligence reporting in ophthalmology-a scoping review[J]. *Ophthalmol Sci*, 2024, 4(4): 100471. DOI: 10.1016/j.xops.2024.100471.
- [59] 国家药品监督管理局医疗器械技术审评中心. 人工智能医疗器械注册审查指导原则 (2022 年第 8 号) [EB/OL]. [2025-06-17]. <https://www.cmde.org.cn/flfg/zdyz/zdyzwbk/20220309091014461.html>.
- [60] Kim M, Kim YN, Jang M, et al. Synthesizing realistic high-resolution retina image by style-based generative adversarial network and its utilization[J]. *Sci Rep*, 2022, 12(1): 17307. DOI: 10.1038/s41598-022-20698-3.
- [61] Yang Y, Chen X, Lin H. Privacy preserving technology in ophthalmology[J]. *Curr Opin Ophthalmol*, 2024, 35(6): 431-437. DOI: 10.1097/ICU.0000000000001087.
- [62] Tom E, Keane PA, Blazes M, et al. Protecting data privacy in the age of AI-enabled ophthalmology[J]. *Transl Vis Sci Technol*, 2020, 9(2): 36. DOI: 10.1167/tvst.9.2.36.
- [63] Liu T, Wu JH. The ethical and societal considerations for the rise of artificial intelligence and big data in ophthalmology[J]. *Front Med (Lausanne)*, 2022, 9: 845522. DOI: 10.3389/fmed.2022.845522.
- [64] Simonini B. Research project in artificial intelligence ethics an ethical approach to trustworthy AI: framing the AI fairness principles for vulnerable and misrepresented groups[EB/OL]. (2021-11) [2025-06-17]. https://www.researchgate.net/publication/355446049_Research_project_in_Artificial_Intelligence_Ethics_An_Ethical_Approach_to_Trustworthy_AI_Framing_the_AI_Fairness_Principles_for_Vulnerable_and_Misrepresented_Groups.
- [65] Hashemian H, Peto T, Ambrósio R, et al. Application of artificial intelligence in ophthalmology: an updated comprehensive review[J]. *J Ophthalmic Vis Res*, 2024,



- 19(3): 354-367. DOI: 10.18502/jovrv.19i3.15893.
- [66] Bajwa J, Munir U, Nori A, et al. Artificial intelligence in healthcare: transforming the practice of medicine[J]. *Future Healthc J*, 2021, 8(2): e188-e194. DOI: 10.7861/fhj.2021-0095.
- [67] Alowais SA, Alghamdi SS, Alsuhebany N, et al. Revolutionizing healthcare: the role of artificial intelligence in clinical practice[J]. *BMC Med Educ*, 2023, 23(1): 689. DOI: 10.1186/s12909-023-04698-z.
- [68] Piotr R, Robert R, Marek N, et al. Artificial intelligence enhanced ophthalmological screening in children: insights from a cohort study in Lubelskie Voivodeship[J]. *Sci Rep*, 2024, 14(1): 254. DOI: 10.1038/s41598-023-50665-5.
- [69] Wang J, Wang S, Zhang Y. Artificial intelligence for visually impaired[J]. *Displays*, 2023, 77: 102391. DOI: 10.1016/j.displa.2023.102391.
- [70] Brilli DD, Georgaras E, Tsilivaki S, et al. AIris: an ai-powered wearable assistive device for the visually impaired[C]//2024 10th IEEE RAS/EMBS International Conference for Biomedical Robotics and Biomechatronics (BioRob), Heidelberg, 2024. New York: IEEE, 2024: 1236-1241. DOI: 10.1109/BioRob60516.2024.10719976.
- [71] Zbrzezny AM, Grzybowski AE. Deceptive tricks in artificial intelligence: adversarial attacks in ophthalmology[J]. *J Clin Med*, 2023, 12(9): 3266. DOI: 10.3390/jcm12093266.
- [72] Murdoch B. Privacy and artificial intelligence: challenges for protecting health information in a new era[J]. *BMC Med Ethics*, 2021, 22(1): 122. DOI: 10.1186/s12910-021-00687-3.
- [73] Balas M, Wong DT, Arshinoff SA. Artificial intelligence, adversarial attacks, and ocular warfare [J]. *AJO Int*, 2024, 1(3): 100062. DOI: 10.1016/j.ajoint.2024.100062.
- [74] Sonmez SC, Sevgi M, Antaki F, et al. Generative artificial intelligence in ophthalmology: current innovations, future applications and challenges[J]. *Br J Ophthalmol*, 2024, 108(10): 1335-1340. DOI: 10.1136/bjo-2024-325458.
- [75] Phipps B, Hadoux X, Sheng B, et al. AI image generation technology in ophthalmology: use, misuse and future applications[J]. *Prog Retin Eye Res*, 2025, 106: 101353. DOI: 10.1016/j.preteyeres.2025.101353.
- [76] Truhn D, Müller-Franzes G, Kather JN. The ecological footprint of medical AI[J]. *Eur Radiol*, 2024, 34(2): 1176-1178. DOI: 10.1007/s00330-023-10123-2.
- [77] González-Gonzalo C, Thee EF, Klaver C, et al. Trustworthy AI: closing the gap between development and integration of AI systems in ophthalmic practice[J]. *Prog Retin Eye Res*, 2022, 90: 101034. DOI: 10.1016/j.preteyeres.2021.101034.
- [78] Palaniappan K, Lin E, Vogel S, et al. Gaps in the global regulatory frameworks for the use of artificial intelligence (AI) in the healthcare services sector and key recommendations[J]. *Healthcare (Basel)*, 2024, 12(17): 1730. DOI: 10.3390/healthcare12171730.
- [79] Rashidisabet H, Sethi A, Jindarak P, et al. Validating the generalizability of ophthalmic artificial intelligence models on real-world clinical data[J]. *Transl Vis Sci Technol*, 2023, 12(11): 8. DOI: 10.1167/tvst.12.11.8.
- [80] Veritti D, Rubinato L, Sarao V, et al. Behind the mask: a critical perspective on the ethical, moral, and legal implications of AI in ophthalmology[J]. *Graefes Arch Clin Exp Ophthalmol*, 2024, 262(3): 975-982. DOI: 10.1007/s00417-023-06245-4.
- [81] Macnamara BN, Berber I, Çavuşoğlu MC, et al. Does using artificial intelligence assistance accelerate skill decay and hinder skill development without performers' awareness? [J]. *Cogn Res Princ Implic*, 2024, 9(1): 46. DOI: 10.1186/s41235-024-00572-8.
- [82] Civaner MM, Uncu Y, Bulut F, et al. Artificial intelligence in medical education: a cross-sectional needs assessment[J]. *BMC Med Educ*, 2022, 22(1): 772. DOI: 10.1186/s12909-022-03852-3.
- [83] Sauerbrei A, Kerasidou A, Lucivero F, et al. The impact of artificial intelligence on the person-centred, doctor-patient relationship: some problems and solutions [J]. *BMC Med Inform Decis Mak*, 2023, 23(1): 73. DOI: 10.1186/s12911-023-02162-y.
- [84] Tabuchi H. Understanding required to consider AI applications to the field of ophthalmology[J]. *Taiwan J Ophthalmol*, 2022, 12(2): 123-129. DOI: 10.4103/tjo.tjo_8_22.
- [85] Gerke S, Minssen T, Cohen G. Ethical and legal challenges of artificial intelligence-driven health care[J]. *SSRN Electron J*, 2020: 295-336. DOI: 10.2139/ssrn.3570129.
- [86] Ting D, Pasquale LR, Peng L, et al. Artificial intelligence and deep learning in ophthalmology[J]. *Br J Ophthalmol*, 2019, 103(2): 167-175. DOI: 10.1136/bjophthalmol-2018-313173.
- [87] Jin K, Ye J. Artificial intelligence and deep learning in ophthalmology: current status and future perspectives[J]. *Adv Ophthalmol Pract Res*, 2022, 2(3): 100078. DOI: 10.1016/j.aopr.2022.100078.
- [88] Meng H, Lin Z, Yang F, et al. Knowledge distillation in medical data mining: a survey[EB/OL]. (2022-03-01) [2025-0618]. <https://dl.acm.org/doi/10.1145/3503181.3503211>.
- [89] Jin K, Li Y, Wu H, et al. Integration of smartphone technology and artificial intelligence for advanced ophthalmic care: a systematic review[J]. *Adv Ophthalmol Pract Res*, 2024, 4(3): 120-127. DOI: 10.1016/j.aopr.2024.03.003.
- [90] Wang S, He X, Jian Z, et al. Advances and prospects of multi-modal ophthalmic artificial intelligence based on deep learning: a review[J]. *Eye Vis (Lond)*, 2024, 11(1): 38. DOI: 10.1186/s40662-024-00405-1.
- [91] Wagner SK, Hughes F, Cortina-Borja M, et al. AlzEye: longitudinal record-level linkage of ophthalmic imaging and hospital admissions of 353157 patients in London, UK[J]. *BMJ Open*, 2022, 12(3): e058552. DOI: 10.1136/bmjopen-2021-058552.
- [92] Kolbinger FR, Veldhuizen GP, Zhu J, et al. Reporting guidelines in medical artificial intelligence: a systematic review and meta-analysis[J]. *Commun Med (Lond)*, 2024, 4(1): 71. DOI: 10.1038/s43856-024-00492-0.
- [93] Hernandez-Boussard T, Bozkurt S, Ioannidis J, et al. MINIMAR (MINimum information for medical AI reporting): developing reporting standards for artificial intelligence in health care[J]. *J Am Med Inform Assoc*, 2020, 27(12): 2011-2015. DOI: 10.1093/jamia/ocaa088.
- [94] Shelmerdine SC, Arthurs OJ, Denniston A, et al. Review of study reporting guidelines for clinical studies using artificial intelligence in healthcare[J]. *BMJ Health Care Inform*, 2021, 28(1): e100385. DOI: 10.1136/bmjhci-2021-100385.
- [95] Vasey B, Nagendran M, Campbell B, et al. Publisher correction: reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial



- intelligence: DECIDE-AI[J]. Nat Med, 2022, 28(10): 2218. DOI: 10.1038/s41591-022-01951-8.
- [96] Yang WH, Shao Y, Xu YW, et al. Guidelines on clinical research evaluation of artificial intelligence in ophthalmology (2023) [J]. Int J Ophthalmol, 2023, 16(9): 1361-1372. DOI: 10.18240/ijo.2023.09.02.
- [97] Jeon S, Liu Y, Li JO, et al. AI papers in ophthalmology made simple[J]. Eye (Lond), 2020, 34(11): 1947-1949. DOI: 10.1038/s41433-020-0929-6.
- [98] Campbell JP, Lee AY, Abramoff M, et al. Reporting guidelines for artificial intelligence in medical research[J]. Ophthalmology, 2020, 127(12): 1596-1599. DOI: 10.1016/j.ophtha.2020.09.009.

· 图像精粹 ·

青光眼相关性视盘旁视网膜劈裂

雷春燕 张美霞

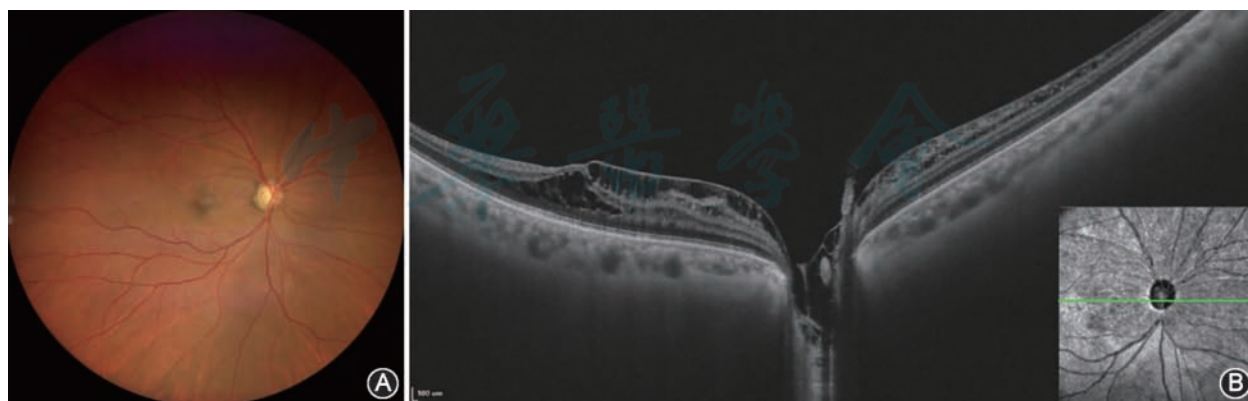
四川大学华西医院眼科, 成都 610041

通信作者: 张美霞, Email: zhangmeixia@scu.edu.cn

患者女性, 54 岁。因右眼视力下降 1 年余, 就诊于四川大学华西医院眼科。既往双眼原发性闭角型青光眼病史 10 年, 双眼曾行周边虹膜激光切除治疗; 1 年前在外院诊断为右眼视网膜静脉阻塞继发黄斑囊样水肿, 眼内注射抗血管内皮生长因子制剂 6 次。眼部检查: 右眼裸眼视力为 0.4, 眼压为 21.3 mmHg (1 mmHg=0.133 kPa), 眼轴长度为 22.05 mm; 角膜透明, 前房浅, 虹膜周边可见切口, 晶状体混

浊; 眼底视杯与视盘直径比约为 0.7, 黄斑区色素紊乱 (精粹图片 1 中 A)。扫频光源相干光层析成像术的红外像显示视盘旁视网膜呈低反射改变, 黄斑区呈放射状高反射改变; B 扫描显示视盘旁及黄斑区视网膜劈裂 (精粹图片 1 中 B)。临床诊断: 双眼原发性闭角型青光眼; 右眼青光眼相关性视盘旁视网膜劈裂。

利益冲突 所有作者声明不存在利益冲突



精粹图片 1 右眼青光眼相关性视盘旁视网膜劈裂患者右眼底检查图像 A 示彩色眼底图像; B 示扫频光源相干光层析成像术图像 (右下小图中绿线为扫描方位)

DOI: 10.3760/cma.j.cn112142-20250802-00328

收稿日期 2025-08-02 本文编辑 黄翊彬

引用本文: 雷春燕, 张美霞. 青光眼相关性视盘旁视网膜劈裂[J]. 中华眼科杂志, 2026, 62(5): 345. DOI: 10.3760/cma.j.cn112142-20250802-00328.

中华医学杂志社
Chinese Medical Association Publishing House

版权所有 违者必究

