

关于欧洲药品管理局发布的监管科学和药品监管工作中大语言模型的使用指导原则的介绍

Introduction to the guiding principles on the use of large language models in regulatory science and for medicines regulatory activities issued by the European Medicines Agency

潘建红^{1,2}, 曹 尚³, 赵 骏^{1,2}

(1. 国家药品监督管理局药品审评中心, 北京 100076; 2. 药品监管科学全国重点实验室, 北京 102629; 3. 国家药品监督管理局药品审评检查长三角分中心, 上海 201210)

PAN Jian - hong^{1,2}, CAO Shang³,
ZHAO Jun^{1,2}

(1. Center for Drug Evaluation, National Medical Products Administration, Beijing 100076, China; 2. State Key Laboratory of Drug Regulatory Science, Beijing 102629, China; 3. Yangtze River Delta Center for Drug Evaluation and Inspection of National Medical Products Administration, Shanghai 201210, China)

基金项目: 药品监管科学全国重点实验室课题
基金资助项目(2024SKLDRS0230)

作者简介: 潘建红(1986-), 男, 副研究员, 主要从事药品的技术审评工作

通信作者: 潘建红

Tel: (010) 80995822

E-mail: panjh@cde.org.cn

摘要:随着人工智能发展, 欧洲药品管理局于2024年发布《监管科学和药品监管工作中大语言模型的使用指导原则》, 强调安全、合规、有效使用大语言模型的重要性。该指导原则涵盖一般伦理考量、用户原则及组织管理机制, 明确指出大语言模型在辅助文本处理、数据挖掘等方面的潜力, 同时警示其在幻觉、数据隐私、偏见输出等方面的风险, 提出需通过持续学习、风险监控和机制建设等予以应对。我国虽尚未发布类似指南, 但大语言模型在审评资料辅助处理、指导原则制修订、风险识别等方面具备潜在的应用前景, 同时还面临决策可解释性、数据偏见、法律问责、人才储备等挑战。建议我国监管机构参考国际经验, 尽早构建符合国情的应用框架, 推动人工智能技术与药品监管领域的融合。

关键词: 大语言模型; 人工智能; 监管机构

DOI: 10.13699/j.cnki.1001-6821.2026.01.023

中图分类号: R-1

文献标志码: C

文章编号: 1001-6821(2026)01-0142-06

Abstract: With the advancement of artificial intelligence, the European Medicines Agency released "Guiding principles on the use of large language models in regulatory science and for medicines regulatory activities" in 2024, emphasizing the importance of the safe and responsible use of large language models. This guidance covers general ethical considerations, user principles and organizational principles, clearly outlining the potential of large language models in areas such as text processing assistance and data mining, while also warning of risks such as hallucinations, data privacy issues, and biased outputs. It proposes measures including continuous learning, risk monitoring and mechanism - building to address these challenges. Although China has not yet issued similar guidelines, large language models hold potential application prospects in areas such as assisting in the processing of review materials, formulating and revising guidance principles and identifying risks. At the same time, challenges related to decision interpretability, data bias, legal accountability and talent reserves remain. It is recommended that China's regulatory authorities draw on international experience to construct an application framework tailored to national conditions as soon as possible, thereby promoting the integration of artificial intelligence technology with the field of drug regulation.

Key words: large language model; artificial intelligence; regulatory authorities

药品监管科学是一门旨在开发和评估新工具、新标准、新方法,以优化药品全生命周期监管决策的科学体系。它并非单一学科,而是融合了医学、药学、毒理学、统计学、数据科学、工程学、社会学、法学等多学科前沿知识的交叉科学领域。其核心目标是弥合“科学研究前沿”与“监管决策实践”之间的鸿沟,确保监管体系能够基于最先进的科学证据与方法,高效、可靠地评价药品的安全性、有效性和质量。药品监管工作通常是指国家药品监督管理部门及其相关机构,依据法律法规和科学原则,对药品的研发、注册、生产、流通、使用及上市后监测等全生命周期各环节所进行的监督、管理、控制、评价和指导等一系列行政和技术行为的总称。药品监管工作通过设定标准、实施许可、持续监督、纠正违规和应对风险,在患者、产业和社会之间构建起至关重要的信任纽带。随着科技发展,现代药品监管工作正日益与监管科学深度融合,运用新工具、新方法不断提升其监管活动效能。

欧洲药品管理局(European Medicines Agency, EMA)与欧盟成员国药监局局长机构(Heads of Medicines Agencies, HMA)于2024年8月联合发布了《监管科学和药品监管工作中大语言模型的使用指导原则》,总共包括7个章节,分别是范围、术语表、引言、一般伦理考虑、用户原则、组织原则和结论,旨在为监管决策过程中应用大语言模型(large language models, LLM)提供遵循的原则与建议,以下内容(1~7部分)是对该指导原则的一般介绍。第8部分是我国药品监管应用的思考建议。

1 范围

指导原则描述了在监管科学和监管活动中如何使用通用大语言模型(large language models, LLM),并为监管机构提供了高级别的建议,以促进安全、负责任和有效的使用LLM。

2 术语表

该章节对人工智能(artificial intelligence, AI)、聊天机器人、基础模型、幻觉、通用人工智能、生成式人工智能等给出了详细定义。比如:人工智能指的是一种基于机器的系统,旨在以不同程度的自主性运行,在部署后可能表现出适应性,并且对于明确或隐含的目标,可以从接收到的输入推断出如何产生输出,如预测、内容、建议或决定,这些输出可以影响物理或虚拟环境。基础模型是指基于广泛数据进行训练的可以适应广泛下游任务的模型,例如生成式预训练变换

器(generative pre-trained transformer, GPT)^[1]。幻觉是指大语言模型的输出虽然看似合理,但偏离了用户输入、先前生成的上下文或事实^[2]。通用人工智能是指一个不是故意和专门设计的可以使用和适应广泛的应用程序的人工智能系统^[3]。

3 引言

3.1 什么是大语言模型?

LLM是一类生成式AI或基础模型,通过大量文本数据进行训练,使其能够对输入或请求(提示)生成自然语言响应。简而言之,LLM在2个阶段进行训练,在预训练期间,使用大量文本数据,将词元(文本单元或单词大小)之间的统计相关性(或权重)转换为数值形式。在微调阶段,预训练模型通过使用带标注示例的监督学习来更新权重并改进模型以适应特定的任务。尽管用户经常使用聊天机器人与LLM互动,但LLM和聊天机器人并不相同。聊天机器人可以由LLM或其他技术提供支持,LLM也可以在其他聊天机器人程序中使用。例如,GPT4是LLM,ChatGPT是聊天机器人。

3.2 大语言模型的使用

LLM通常用于通用或定制聊天机器人。通用聊天机器人允许用户询问任何问题,因为计算机程序可以作为通用聊天机器人处理各种各样的问题。通用聊天机器人如ChatGPT或Bing AI,使用开放的在线LLM。另一方面,定制聊天机器人是为特定领域的特定问题定制的。LLM也常用于信息自动化处理或知识/数据挖掘,开发者会向LLM提供一个标准的“锁定提示”并附上原始数据。用户可以获得处理后的数据,或者导航结果,但不会更改提示,只有原始数据会发生变化。例如,LLM可以创建一个应用程序,该应用程序以特定格式总结会议记录。提示始终相同,改变的是提供的会议记录。LLM还可以用于AI虚拟助手,并为用户完成某些任务,比如安排会议。

LLM根据其可用性和可访问性分为2大主要类别,开源模型(免费向公众开放)和闭源模型(作为商业产品开发,通常需要许可证或订阅才能使用)。LLM的访问通常由第三方提供,其中可能包括基于LLM技术开发或与LLM技术集成的专业工具的组织。用户与LLM互动的方式可能会影响互动过程中的灵活性、控制权、资源需求和集成度。因此,在为特定工作选择模型时,必须考虑这些因素。

3.3 LLM的安全合规使用

LLM具有执行多种任务的强大功能,包括文本生成、翻译、代码支持和知识检索。这些功能中的许多功能也可用于支持药品监管系统内的各种任务和流程。与此同时,LLM的使用也面临挑战和风险。LLM在一些看似简单的任务中出现低级错误,返回无关或不准确的响应—即所谓的“幻觉”。许多由监管机构执行的任务是创新性的(例如涉及新的活性物质),LLM可能未曾接触到去回答这些科学或监管问题所需要的信息,从而导致这些技术的验证具有非常大的挑战。

LLM可以处理从简短的段落到整个文档的各种用户输入。对于提示的存储和进一步处理,存在涉及机密信息、数据保护和隐私方面的风险。LLM基于大量的文本源进行训练,其中可能包括公开来源的数据,公开的网络爬取数据集可能包含不准确或虚假的(个人)数据。因此,在没有适当控制措施的情况下,应对数据保护风险非常具有挑战。考虑到这些LLM以数十亿或万亿个权重的形式存储数据,因此纠正、删除甚至请求访问LLM学习的个人数据实际上是不可能的。

如果没有适当的安全措施,LLM的输出可能会泄露训练数据集中包含的敏感或私人信息,从而导致潜在或实际的数据泄露。最后,输出的内容可能侵犯其他合法权利(例如版权),或具有伦理影响(例如影响人们的选择自由或道德价值观)。

因此,监管机构应该授权员工只能在安全合规的使用LLM。

4 一般伦理考虑

2018年,欧洲科学与新技术伦理小组发布了1份关于使用AI、机器人和自动化系统的伦理声明。该小组是欧盟委员会在伦理事务方面的顾问,他们提出了1套基于《欧盟条约》和《欧盟基本权利宪章》基本价值观的原则。这些原则包括人的尊严、自主性、责任、正义、公平、民主、法治与问责、保障、安全、身心完整、数据保护和隐私、可持续性^[4]。随后,欧盟委员会设立的独立高级专家小组(High-Level Expert Group, HLEG),于2019年发布了人工智能道德准则《Ethics Guidelines for Trustworthy AI》(可信人工智能伦理指南)。HLEG倡导以基于基本权利为基础的方法,强调在AI设计、开发和使用过程中应考虑4个伦理原则:自主、避免伤害、公平和可解释性^[5]。在这些原则的基础上,HLEG提出了7项要求:问责制、人力代理和监督、技术稳健性和安全性、隐私

和数据治理、透明性、多样性、非歧视和公平、社会和环境福祉^[5]。

上述原则和要求适用于所有AI系统。WEIDINGER等^[6]于2021年发表了1篇关于伦理和社会伤害风险的重要论文,从开发者和用户的角度探讨了LLM的使用和交互。论文指出以下风险:①歧视、排斥和有害内容—LLM生成歧视性或排斥性文本,包括助长刻板印象、抵触规范和恶毒语言等内容,从而带来社会危害;②信息危害—可能的危害,包括数据误用、不安全使用、个人或其他敏感信息泄露,以及不安全的数据存储和管理实践;③错误信息危害—LLM传播虚假或误导性信息、或内容鼓励不道德甚至非法行为;④恶意使用—刻意使用LLM造成危害的风险,包括欺诈、大规模虚假信息传播、生成用于网络攻击的代码;⑤人机交互危害—风险可能来源于LLM应用程序(如聊天机器人)的设计和交互中存在的不安全应用,包括对LLM输出的过度依赖(自动化偏见)、透明度和问责制的转变,以及刻板印象的延续,例如为烹饪聊天机器人命名为“女性”名字;⑥自动化、获取和环境风险—风险可能源于环境危害、社会不平等、人类劳动力的替代、硬件和数字素养的不平等获取。

基于这些伦理原则,个人数据保护也值得特别关注。虽然指导原则主要关注LLM的使用,但在LLM的开发、实施和使用过程中,可能会在无意识的情况下处理个人数据。因此,必须确保在LLM整个生命周期的各个阶段遵守适用的数据保护规则和其他法律要求。

5 用户原则

5.1 采取适当措施确保数据输入的安全性

与LLM的交互通常从用户提出的1个提示、任务或查询开始,让LLM解答。尽管提示可能只是1个简单的问题,例如“用100个字解释药物警戒”,但提示往往使用上下文信息,或者要求LLM使用用户提供的数据,例如“总结这段文字:……”。了解这种细微差别很重要,因为对于一些LLM交互,LLM仅使用提示中的信息,而对于其他交互,LLM使用其他地方生成的信息。

与LLM的第1次交互对于确保其被合法、安全、合规的使用至关重要。用户必须清楚了解LLM的功能和局限性,以便他们能够有效地利用其潜力,同时警惕潜在的风险^[7]。这种意识将使他们在与这些工具交互时做出明智的决策。在确保可以安全使用LLM聊天机器人的提示,或者在LLM信息提取自动

化和知识挖掘工具中能安全使用原始数据,有2个关键要素需要注意,①了解 LLM 的部署方式:即它们是被本地部署、由组织完全控制,还是开放的在线模型;②谨慎的提示工程(或称为提示设计):这对于避免暴露敏感数据给 LLM,确保更好的输出,并减少由偏见提示导致的偏见输出风险。

另外,需应用批判性思维对输出进行交叉核对。LLM 在生成内容方面表现出色,因此人们很容易直接使用其生成的内容,但也会带来重大的风险。因为 LLM 使用大量文本建立词元之间的统计相关性。这些统计相关性可以用来生成新的内容,意味着 LLM 可能会输出看似连贯但实际上无关、不真实、有偏见或直接从源文档中复制的文本。因此,用户应谨慎对待输出文本,并运用自己的判断来决定文本是否值得信赖/正确,是否可靠(模型是否遵循了指示,输出是否相关?),是否公平(输出是否有偏见?),以及最后是否合法(例如,是否未侵犯版权或剽窃?)。潜在危害越大,审查的程度就应越高^[8]。

5.2 持续学习如何有效使用 LLM

要做到合规地使用 LLM,需要熟悉这项技术并了解如何与它们进行交互。在药品监管中使用 AI 的主要目的是为了公众健康,其中,有效的提示是通过在财务和环境方面以较低的成本产生最可靠的产出,从而可以最大限度地发挥使用 LLM 的好处。此外,LLM 正在快速发展,上述识别出的风险可能会发生变化,持续学习有助于防止 LLM 的风险轮廓发生变化。LLM 用户应不断学习,以确保 LLM 能够在监管科学和药品监管工作中安全、有效和合规地使用。

5.3 知道在面对疑虑和报告问题时应该向谁咨询

从 LLM 用户的角度来看,他们可能面临的主要问题大致分为3类:①用于训练 LLM 的数据本身存在的问题,例如用于改进模型的数据包含个人数据或知识产权;②在提示中可以分享哪些信息,包括哪些 LLM 可以被用来操纵敏感信息;③输出的可靠性、真实性、公平性和合法性。知道应该联系谁是非常重要的,这可以确保:①如果有必要,进行安全风险分析、威胁分析或数据保护影响评估。将使组织能够识别并实施额外的安全和隐私增强措施;②调查并解决在不应该分享敏感信息的情况下分享此类信息的情况;③迅速解决严重偏差或错误的输出,确保其不会对药品监管机构构成风险。

由于 LLM 在不断更新和改进以提高性能或防止恶意使用,因此持续关注和解决相关问题显得尤为突出。用户的反馈在此过程中起着至关重要的作用。

持续监控,包括监控个人数据泄露或严重偏差或错误的输出,对于确保遵守“4”中描述的一般伦理原则至关重要^[9]。严重偏差或错误的输出包括决策中产生负面影响的生成文本,或者严重影响某些基本权利等。例如,关于合同纠纷的会议纪要更改了关键文本,比如从批评转为支持。

6 组织原则

6.1 制定有助于用户安全、合规地使用人工智能的管理机制

药品监管机构可以通过使用 LLM 获得很多好处。同时,工作人员对这些模型的使用也会直接影响到该机构或组织,因此,应努力确保工作人员安全、合规地使用 LLM。建立 LLM 使用的管理机制有助于用户规避这些模型的风险。首先,对于不同类型 LLM 的使用、可以接受的 LLM 风险概况、可以获得的 LLM 以及使用这些 LLM 可能遇到的风险,通过管理机制可以提供限制措施。还可以考虑是否需要针对整个机构或组织、或基于使用案例的风险,总结 LLM 引发的错误、问题的信息,对安全、合规使用 LLM 进行强制性的培训。管理机制的另一个关键方面与 LLM 的模型开发有关,而不仅仅是它们的使用。虽然一些机构或组织可能不会使用自己的数据来微调或提高 LLM 在特定任务上的性能,但那些这样做的机构或组织应考虑数据保护和信息安全的需求。

6.2 帮助用户从 LLM 中获得最大价值

可以通过提供或促进培训来帮助用户最大化 LLM 的使用价值。机构或组织应该认识到 LLM 不断变化的性质,并考虑允许用户从 LLM 中获得最大价值的变更管理活动。通常包括大规模的培训,最好是按需培训。还可能包括建立1个支持团队,用户可以联系他们来帮助优化关键提示,或者构建支持工具,例如由机构或组织完全控制的工具,在提示/输入用于 LLM 之前筛除个人或机密信息。

6.3 合作与分享经验

鉴于人工智能快速变化的特征,将经验分享进行网络化是非常重要的。它减少了不确定性,促进了更快的共识。依托欧洲专家社区等专业论坛平台,特别是人工智能兴趣小组,是1种通过网络分享经验的有效方式。另一种方式是通过欧盟网络培训中心向工作人员提供相关的课程培训。

7 小结

在监管科学和药品监管工作中使用大语言模型的指导原则至关重要,这将确保所有工作人员能够安全、合乎伦理和有效地使用这些技术。在这个快速

发展的领域,提供一种结构化的方法将先进的人工智能工具整合到监管过程中,以保障公众健康和维护对药品监管机构的信任,已成为当务之急。

8 我国应用建议

大语言模型因其强大的自然语言处理能力,是提升药品监管效率的重要工具。对于我国来说,结合相关的文献调研情况,大语言模型在药品监管环节可能的应用包括但不限于:①自动检查申报资料的完整性、格式合规性,比对与相关指导原则的一致性,标记潜在的数据矛盾或缺失项,大幅缩短形式审查时间^[10]。②快速解析冗长的毒理学研究报告、临床试验总结摘要,自动提取关键疗效终点,如总生存期(overall survival, OS)、无进展生存期(progress free survival, PFS)、安全性事件(serious adverse event, SAE)发生率,并与同类产品进行初步的横向对比分析^[11]。③针对创新靶点或罕见病药物,自动检索、摘要、整合全球相关文献,生成证据图谱,辅助审评员评估立项依据与临床价值^[12]。④根据结构化审评要点和数据输入,自动生成审评报告的初稿或部分章节(如药学概览、非临床综述),保证报告结构的规范性与术语的准确性,提升工作效率^[13]。⑤构建基于LLM的监管知识助手,7 d×24 h 回应企业关于法规、指南、申报流程的具体问题,提供精准的条文引用和解释,降低沟通成本,提升监管透明度^[14]。⑥持续扫描学术前沿,自动摘要新兴技术(如AI制药、类器官芯片)的监管挑战,辅助监管科学家确定研究优先级。⑦分析历次药品生产质量管理规范(good manufacturing practice, GMP)、药物临床试验质量管理规范(good clinical practice, GCP)检查报告,识别企业常见缺陷模式和高风险环节,为制定基于风险的检查计划提供数据支持。⑧跟踪全球监管动态和最新科学共识,辅助分析其对现有指导原则的影响,为新指南的起草提供文献和案例支持。

同时在监管应用过程中,可能需要注意的事项包括:①在监管决策中,由于模型的决策过程不透明,如何对其输出进行验证与追责是重大挑战^[15]。②模型的性能严重依赖于训练数据。监管数据中可能存在报告偏倚、不均衡性(如某些人群数据少)等问题,导致模型输出带有偏见^[16]。③AI作为辅助或潜在决策参与者,其法律地位、认证标准、验证流程、出现错误时的责任划分均处于空白^[17]。④成功应用LLMs需要既懂监管业务又懂AI技术的复合型人才。新技术的引入可能改变既有工作流程和权力结构,面临来自组织内部的文化与适应性挑战^[18]。⑤所有关键决

策必须由人类监管者在充分理解AI辅助分析的基础上作出。这要求培养监管人员的“数字素养”,使其具备指挥、审查、质疑AI输出的能力^[19]。⑥审评决策等高风险核心环节,在可预见的未来应以人类为主,AI仅提供参考信息等。

9 讨论

通过对EMA监管科学和药品监管工作中大语言模型的使用指导原则的介绍可以看出,EMA不仅认识到大语言模型在药品监管领域的潜力,同时强调在使用过程中必须遵循安全、合规和有效的原则,为EMA工作人员提供了高层次的原则和建议。尽管我国尚未发布类似或相关的指导原则,但大语言模型已在国内部分医药企业和医疗机构中得到应用实践。我国药监部门也充分的认识到了大语言模型对于药品监管的重要性,如陈锋等^[20]对于大语言模型在药品监管领域的应用实践的基础考量进行了探讨,同时,国家药监局将人工智能列为监管科学课题。本文对EMA大语言模型指导原则的介绍,可以为我国药品监管科学与监管工作中大语言模型的应用提供参考和借鉴。

参考文献:

- [1] BOMMASANI R, HUDSON D A, ADELI E, *et al.* On the opportunities and risks of foundation models [EB/OL]. 2022-07-12 [2025-09-23]. <https://arxiv.org/abs/2108.07258?ref=nural.cc>.
- [2] ZHANG Y, LI Y, CUI L, *et al.* Siren's song in the AI ocean: A survey on hallucination in large language models [EB/OL]. 2025-09-14 [2025-09-23]. <https://arxiv.org/abs/2309.01219>.
- [3] EUROPEAN PARLIAMENT. Artificial intelligence act: Amendment 169 (Article 5, Paragraph 1, Point ba (new)) [EB/OL]. 2023-06-14 [2025-09-23]. https://www.europarl.europa.eu/doceo/document/TA-9-2023-0236_EN.html.
- [4] EUROPEAN GROUP ON ETHICS IN SCIENCE AND NEW TECHNOLOGIES. Statement on artificial intelligence, robotics and 'autonomous' systems [EB/OL]. 2018-03-09 [2025-09-23]. <https://data.europa.eu/doi/10.2777/531856>.
- [5] AI HLEG. Ethics guidelines for trustworthy AI [EB/OL]. 2019-04-08 [2025-09-23]. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- [6] WEIDINGER L, MELLOR J, RAUH M, *et al.* Ethical and social risks of harm from language models [EB/OL]. 2021-12-08 [2025-09-23]. https://arxiv.org/abs/2112.04359?utm_source=openai.
- [7] CEDPO. Generative AI: The data protection implications [EB/OL]. 2023-10-26 [2025-09-23]. <https://cedpo.eu/generative-ai-the-data-protection-implications/>.

- [8] VRIJE UNIVERSITEIT BRUSSEL. Responsible use of artificial intelligence for research purposes [EB/OL]. 2023 - 01 - 01 [2025 - 09 - 23]. <https://www.vub.be/en/research/research-data-management/responsible-ai>.
- [9] YAN B, LI K, XU M, *et al.* On protecting the data privacy of large language models (LLMs): A literature review [J/OL]. *High - Confidence Computing*, 2025, 5 (2): 100300. 2025 - 06 - 30 [2025 - 09 - 23]. <https://www.sciencedirect.com/science/article/pii/S2667295225000042>.
- [10] LIU Y, HAN X, LIU Z, *et al.* AI in regulatory submissions: opportunities and challenges for automated document review [J]. *Nat Rev Drug Discov*, 2023, 22 (5): 345—358.
- [11] XU J, WANG S, LEE M, *et al.* Leveraging natural language processing for clinical trial data abstraction: A systematic review [J/OL]. *J Biomed Inform*, 2022, 135: 10428—10428. 2021 - 06 - 30 [2025 - 09 - 23]. <https://www.sciencedirect.com/science/article/abs/pii/S2210844021000411>.
- [12] WANG L L, LO K. Text mining for regulatory science: A review [J]. 2021, 109 (2): 322—331.
- [13] FDA. Artificial intelligence and machine learning in drug development: Discussion paper [EB/OL]. 2023 - 05 - 10 [2026 - 01 - 26]. <https://www.fda.gov/media/167973/download>.
- [14] 陈晓,王华. 人工智能在药品智慧监管中的应用场景与法律问题研究 [J]. 中国药学杂志, 2023, 58 (10): 865—870.
- [15] JI Z, LEE N, FRIESKE R, *et al.* Survey of hallucination in natural language generation [J]. *ACM Comput Surv*, 2023, 55 (12): 1—38.
- [16] OBERMEYER Z, POWERS B, VOGELI C, *et al.* Dissecting racial bias in an algorithm used to manage the health of populations [J]. *Science*, 2019, 366 (6464): 447—453.
- [17] COHEN I G, MELLO M M, STUDDERT D M, *et al.* The FDA and AI - driven software: A regulatory paradigm in flux [J]. *NEJM AI*, 2022, 1 (1): 1056—1056.
- [18] 张伟,李静. 药品监管数字化转型中的人才能力框架构建研究 [J]. 中国药事, 2024, 38 (1): 1—8.
- [19] 国家药品监督管理局. 关于实施中国药品监管科学行动计划第二批重点项目的通知 [EB/OL]. 2021 - 06 - 28 [2026 - 01 - 26]. <https://www.nmpa.gov.cn/xxgk/fgwj/gzwp/gzwjyp/20210628172854126.html>.
- [20] 陈锋,吴欣然. 大语言模型在药品监管中的应用实践与思考 [J]. 中国食品药品监管, 2025, 3 (254): 4—13.

(本文编辑:马凌云 收稿日期 2025 - 07 - 01)

· 科学文摘 ·

一项 Adagrasib 对比多西他赛治疗 KRAS^{G12C} 突变 非小细胞肺癌的随机、开放、Ⅲ期临床试验

引自:BARLESI F, YAO W, DURUISSEAU M, *et al.* Adagrasib versus docetaxel in KRAS^{G12C} - mutated non - small - cell lung cancer (KRYSTAL - 12): A randomised, open - label, phase 3 trial [J]. *Lancet*, 2025, 406 (10503): 615—626.

本研究旨在比较 Adagrasib 与多西他赛在既往接受过化疗和免疫治疗的 Kirsten 大鼠肉瘤病毒癌基因同源物 (KRAS) G12C 突变型晚期晚期非小细胞肺癌 (NSCLC) 患者中的疗效和安全性。本研究是一项随机、多中心、开放标签的Ⅲ期试验,在 22 个国家的 230 个中心进行。纳入患有 KRAS^{G12C} 突变的局部晚期或转移性非小细胞肺癌,且之前接受过基于铂的化疗和抗程序性细胞死亡蛋白 1 或抗程序性死亡配体 1 治疗的患者,将患者以 2:1 的比例随机分配,接受 Adagrasib 600 mg (每天口服 2 次) 或 75 mg · m⁻² 多西他赛 (每 3 周静脉注射 1 次)。治疗持续至发生疾病进展、不可接受的毒性、研究人员或患者的决定停止治疗或死亡。主要终点是所有随机患者 (意向治疗人群) 的无进展生存期。对所有接受治疗的患者进行了安全性评价。本研究共纳入 453 例患者,将其随机分配至 Adagrasib 组 301 例 (66%) 和多西他赛组 152 例 (34%)。Adagrasib 组中有 298 例患者 (99%) 接受了治疗,多西他赛组中有 140 例患者 (92%) 接受了治疗。意向治疗人群的中位随访时间为 7.2 个月 (95% CI 5.8 ~ 8.7), Adagrasib 组的中位无进展生存期为 5.5 个月 (95% CI 4.5 ~ 6.7), 多西他赛组的中位生存期为 3.8 个月 (95% CI 2.7 ~ 4.7), 风险比为 0.58 [95% CI 0.45 ~ 0.76]; *P* < 0.0001]。298 例接受 Adagrasib 治疗的患者中有 140 例 (47%) 发生了 3 级及以上治疗相关不良事件,140 例接受多西他赛治疗的患者有 64 例 (46%) 发生了此类事件。Adagrasib 组有 4 例 (1%) 与治疗相关的死亡,多西他赛组有 1 例 (1%)。与多西他赛相比,Adagrasib 在先前治疗过 KRAS^{G12C} 突变 NSCLC 的患者中,无进展生存率有统计学上的显著改善,但没有新的安全信号。