

2025

中国网民AI认知 调研报告

对外经济贸易大学数字经济与法律创新中心

中国人民大学数字经济研究中心

蚂蚁集团研究院

2026/05

研究团队介绍 >>>>>

课题负责人

许可 对外经贸大学数字经济与法律创新中心主任

程华 中国人民大学数字经济研究中心执行主任

课题组成员

王尉泽、刘雯菲、李子恬、金宛怡、万承志、卢侃、李昭涵、李勇国



对外经济贸易大学数字经济与法律创新研究中心，旨在以前瞻视野和探索精神，跨学科、多维度地研究数字经济法律规律，凝聚各界共识，推动制度变革，从而为我国数字经济发展和数字中国的建立贡献力量。“数字经济与法律创新研究中心”由多家知名企事业单位担任理事单位，中国人民大学信息法中心主任张新宝教授和中国法学会网络和信息法学研究会常务副会长周汉华研究员担任学术委员会联席主席，汇聚经济学、法学、金融学等多学科的研究人员，共同造就国内外一流的数字经济法律研究中心。



中国人民大学数字经济研究中心 (Center for Digital Economy Research, Renmin University of China) 是中国人民大学经济学院主办的研究机构。宗旨是跟踪、研究数字经济理论和实践的发展前沿。中心致力于通过原创研究、深度调研、论坛讲座、案例研究、企业与国际合作等一系列活动，来打造一流的数字经济学术交流平台。

Ant Group Research
蚂蚁集团研究院

蚂蚁集团研究院是蚂蚁集团的专业研究部门，从事公司战略、监管政策和学术理论研究，覆盖AI产业、数字支付、数字金融、数智健康、数据与平台治理等研究领域。通过开展客观、扎实和深入的研究工作，为公司战略决策提供支持，为政府部门公共政策提供参考和研究服务。研究院于2018年创建了国内第一个数字经济研究开放平台和数据开放研究实验室，成为推动国内数字经济理论和学术研究的最前沿平台，为学术研究和政策研究提供数据、案例等支持。目前实验室升级为AI与数据开放实验室，旨在携手行业机构与研究机构，探索医疗健康等领域AI for Science新模式。



摘要

>>>>>

随着生成式人工智能技术的迅猛发展，AI 大模型已深度融入公众日常生活与工作场景。为全面把握社会大众对这一新兴技术的认知水平、使用态度、风险感知及治理期待，相关研究机构于 2025 年开展了全国范围的问卷调研。本报告基于大规模有效样本，系统梳理了当前我国公众与 AI 大模型互动的基本图景，并通过纵向与横向比较，揭示其认知演进趋势与本土特征。总体而言，中国公众对 AI 大模型在认知上不断深入，在态度上整体积极，在展望上日趋理性，在治理上期待切实有效的规则形成。

一、公众对 AI 大模型普遍知晓，但理解深度有限，使用集中在日常生活工作场景

调研发现，AI 大模型在公众中的知晓度和使用率已达到较高水平，绝大多数人至少有过接触或尝试。然而，这种高渗透并未同步转化为深度认知。多数用户对技术原理、能力边界及潜在局限缺乏系统了解，其使用行为主要建立在直观体验和功能试错之上，而非理论掌握。技能获取多依赖自发探索，缺乏正规教育或培训支持。

TOGETHER

在应用场景方面，公众普遍将 AI 大模型视为提升效率的工具，优先用于学习辅助、职场任务处理和日常生活问题解决。相较初期以信息查询和娱乐为主的使用模式，当前应用更趋务实，尤其在专业写作、数据分析、创意生成等复杂任务中展现出更强的融合趋势。这反映出用户需求正从浅层体验向深度赋能演进，AI 大模型逐渐成为日常生产力体系的一部分。

二、公众对 AI 大模型整体持积极态度，但在严谨领域持审慎态度，对风险的关注转向深层社会影响

公众普遍认可 AI 大模型在便捷性、创新性和效率提升方面的价值，但对其信任具有差异化特征：在低风险、非关键决策场景中充分信赖 AI 输出，而在涉及医疗、财务、法律等严谨领域则保持高度审慎。这种“情境化信任”体现了公众具备理性评估风险与收益的能力，也反映出在 AI 在严谨领域应用需要搭建信任机制。

公众对 AI 大模型风险的担忧正在发生结构性转变。早期聚焦于隐私泄露、数据滥用或就业岗位替代等技术外部性问题，如今更多转向对技术内生性影响的反思。例如，人们开始忧虑过度依赖 AI 大模型可能导致人类自身认知能力退化、批判性思维弱化，甚至产生情感依附；同时，对 AI 大模型生成虚假或误导性内容的警惕也显著增强。这些新关切表明，公众的思考已超越工具安全层面，深入到人机关系、主体性维护及社会公平等更根本的伦理维度。

三、公众对 AI 大模型的未来趋于理性，具有中国特色的“人机共生”观念开始浮现

公众对 AI 大模型的态度趋于更加冷静和务实，初期的狂热追捧或过度恐慌情绪有所消退，对技术潜力的认可依然稳固，但对其局限与风险的认识也更加全面，更侧重于个人生活层面的安全感与实用性，体现出鲜明的本土文化与社会结构特征。

“人机共生”而非“人机替代”的理念逐渐成为主流认知，公众更希望 AI 作为增强人类能力的伙伴，而非取代人类的对手。

四、公众对 AI 大模型“有效治理”的诉求强烈，但权利意识与政策认知明显滞后

面对 AI 带来的复杂挑战，公众普遍认为有必要加强规范与引导，一方面是倾向于完善立法司法制度，不断提升治理规则的明晰化；另一方面是结合企业自律、社会监督、第三方核查等社会共治方式，构建稳定透明、可预期、可落实的多元治理框架。

尽管治理呼声高涨，但公众对自身在算法社会中的权利边界、维权途径及相关政策内容知之甚少。多数人未曾主动了解国家出台的 AI 治理法规，也不清楚在遭遇算法不公或数据侵权时应如何主张权益。这种对于治理规则和工具“高诉求、低认知”的状态，暴露出当前数字素养教育与公众参与机制的不足，可能制约治理体系的有效落地。

综上所述，当前公众与 AI 大模型的关系正处于从“工具采纳”向“生态共建”过渡的关键阶段。一方面，AI 已成为不可或缺的日常助手；另一方面，围绕其伦理、治理与长远影响的公共讨论日益深入。未来，推动 AI 健康发展不仅需要技术创新与制度完善，更需着力提升全民数字素养，弥合认知鸿沟，激发公众的主体意识与参与能力。唯有如此，才能构建一个技术可信、权利可及、发展可持续的人工智能社会。

目录 >>>>>

一、前言	1
1. 背景	1
2. 问卷设计理念	4
二、样本的基本特征	6
1. 调查对象的性别、年龄和区域分布	6
2. 调查对象的学历、职业和收入分布	7
三、消费者对 AI 大模型的认知和态度	8
1. 用户对 AI 大模型的认知与使用现状	8
2. 用户对 AI 大模型的使用体验与信任态度	10
3. 用户对 AI 大模型的问题遭遇与风险认知	13
4. 用户对 AI 大模型的规范诉求与权益保护行为	15
四、与国际社会的对比	21
1. 对比报告简介	21
2. 与国际社会的比较——对 AI 的认知与场景信任比较	21



五、消费者对 AI 大模型认知的变化	27
1. 关于 AI 大模型认知与使用基础水平	27
2. AI 大模型使用场景与体验	29
3. 个人使用 AI 大模型遇到的问题	30
4. 对 AI 大模型的整体评价	32
5. AI 大模型发展期待与监管认知	36
六、交叉分组分析	39
1. 不同年龄对 AI 大模型的感知差异	39
2. 不同学历对 AI 大模型的感知差异	43
3. 不同性别对 AI 大模型的感知差异	45
七、人工智能健康发展倡议	50
附件：AI 大模型调查问卷	52

一、前言

在数字技术全面渗透社会肌理的今天，人工智能（AI）已从实验室的前沿探索，演变为支撑现代社会运转的核心基础设施。尤其是 2022 年 ChatGPT 的问世，彻底打破了 AI 技术与普通用户之间的壁垒，推动 AI 大模型从专业领域走向大众生活，成为信息搜索、内容创作、学习辅助、工作协作等日常场景中不可或缺的工具。这种技术普及速度之快、影响范围之广，在人类科技发展史上实属罕见。截至 2025 年 10 月 31 日，中国已发布约 1509 个大模型，位居全球首位，形成了世界领先的研发与应用规模；同时，个人用户注册总数超 31 亿，API 调用用户突破 1.59 亿，这组数据不仅彰显了我国在 AI 领域的技术实力，更直观反映出 AI 大模型已深度融入个人生活、产业服务与数字经济的各个层面，重塑着人们的生产方式与生活形态。

1. 背景

人工智能作为计算机科学的核心分支，其本质是通过程序与机器模拟人类智能行为，赋予计算机感知、理解、判断、推理、学习、生成和交流等类人能力，从而实现多样化任务的自主执行。近年来，随着大模型技术的迭代、多模态学习的突破以及自主规划与工具调用能力的提升，AI 的应用边界持续拓展，从单一的内容生成向多步骤任务执行延伸，在医疗、金融、工业制造、交通等关键领域发挥着越来越重要的作用。在医疗领域，AI 通过分析医学影像与患者数据，为疾病诊断提供辅助支持，帮助医生提升诊断准确率与治疗效率；在金融领域，AI 技术广泛应用于风险评估、投资分析与欺诈识别，推动金融服务向精准化、高效化转型；在工业制造领域，AI 借助数据分析与自动化手段优化生产流程，有效提升了生产效率与产品质量。这些应用场景的落地，不仅展现了 AI 技术的巨大潜力，

更印证了其作为新型生产力对社会发展的推动作用。

然而，技术的飞速发展往往伴随着新的风险与挑战。与早期以信息匹配为主的 AI 应用不同，生成式 AI 的广泛普及带来了更为复杂的问题：错误信息生成的门槛降低，可能引发认知误导；模型训练数据中隐含的偏见，可能导致歧视性结果；责任主体界定模糊，使得 AI 应用引发的纠纷难以追责。这些风险若得不到有效管控，不仅会损害消费者的合法权益，还可能对社会秩序与公共利益造成潜在威胁。更为重要的是，随着 AI 大模型从单纯的工具向“智能体”“具身智能”等形态演进，其自主性与复杂性不断提升，人和机器的关系正从“人用工具”的单向互动，转向人机协作甚至部分场景下的机器自主运行模式。这种关系变革不仅对技术治理提出了更高要求，也引发了人们对技术伦理、社会公平等深层问题的思考。

2023 年 8 月 15 日，国家有关部门联合发布的《生成式人工智能服务管理暂行办法》正式实施，标志着我国开始对生成式人工智能进行系统性、规范化监管，为 AI 技术的健康发展奠定了制度基础。但技术治理是一项系统工程，仅靠政府监管远远不够，还需要企业履行主体责任、行业加强自律、公众积极参与。其中，消费者作为 AI 技术的直接使用者与体验者，其认知水平、使用态度与行为选择，不仅影响着个人在数字时代的福利水平，更在一定程度上决定着 AI 技术商业化应用的速度与方向，进而对我国数字经济发展与数字中国建设产生深远影响。

本次调研着重从消费者的认知层面与应用感知维度展开，主要基于以下三方面的考量：

第一，用户对 AI 大模型的认知水平与使用能力正日益成为影响个人福利的重要因素。随着 AI 大模型逐渐融入学习、工作和生活的各个场景，用户在理解能力、使用熟练度、风险识别与自主判断方面的差异不断显现。一方面，部分用户能够熟练运用 AI 工具提升效率、拓展能力边界；另一方面，也有部分用户因缺乏相关知识与技能，面临“智能鸿沟”带来的不便，甚至可能因不当使用或过度依赖 AI 而遭遇权益损害。因此，有必要通过调研全面了解消费者在使用 AI 大模型过程中的能力现状，识别其在享受技术便利的同时面临的挑战，如是否能够有效保护个人隐私、避免过度依赖导致的能力退化等。这不仅有助于评估公众的人工智能素养水平，也能为后续开展消费者教育、提升全民数字素养提供针对性的对策建议。

第二，AI 大模型的不当使用可能对消费者权益造成新的侵害形式。除了传统的数据过度采集、隐私泄露等问题外，生成式 AI 还可能引发知识产权争议、内容失真、模型偏见等新型权益侵害。消费者在实际使用过程中，对自身权益是否受到侵害以及侵害程度具有最直接的感知。通过问卷调查，我们能够系统呈现用户在现实应用场景中的真实体验与判断，识别当前 AI 大模型应用中存在的突出问题，为权益保护机制的完善提供参考。

第三，消费者对 AI 大模型的态度可能存在持续的内在不一致性。已有研究表明，用户在态度表达上往往对人工智能的风险保持警惕，关注隐私侵犯、伦理问题与潜在社会影响；但在实际行为中，为了获得更高效率与更优体验，又往往愿意持续使用 AI 大模型，并在一定程度上容忍其不完美性。随着人工智能从辅助工具向深度参与决策和执行的方向发展，这种态度与行为之间的张力可能更加突出。例如，部分用户可能在口头表达中担忧 AI 对就业的冲击，但在工作中仍会积极使用 AI 工具提升竞争力；又如，一些用户虽关注数据隐私问题，但为了使用便捷性可能会随意授权 AI 产品获取个人信息。通过在问卷中设置不同使用场景与问题，我们能够精准识别这种不一致性的具体表现形式，从而更为准确地判断 AI 大模型对消费者福利的实际影响，为技术治理与市场规范提供更具针对性的思路。

此外，从宏观层面来看，了解消费者对 AI 大模型的认知与态度，对于推动 AI 技术的健康发展、促进数字经济高质量发展具有重要意义。消费者的需求是技术创新的重要导向，其对 AI 大模型功能、体验、安全等方面的诉求，能够为企业的技术研发与产品优化提供明确方向；而消费者对风险的担忧与治理的期待，也能为政府完善监管政策、构建多元治理体系提供参考。同时，通过对比不同群体（如不同年龄、学历、区域、收入水平）对 AI 大模型的认知差异，还能够识别出技术普及过程中存在的不平衡问题，为促进数字包容、缩小数字鸿沟提供实证支持。

综上所述，本次调研旨在通过全面、系统的问卷调查，深入了解我国消费者对 AI 大模型的认知现状、使用态度、风险感知与治理诉求，揭示技术发展与用户反馈之间的互动关系。调研结果不仅能够为政府部门制定相关政策、加强行业监管提供数据支撑，也能为企业优化产品设计、履行社会责任提供参考，同时还能帮助公众更好地认识 AI 技术、提升使用能力与风险防范意识，最终推动 AI 技术在规范中创新发展，在发展中保障权益，为我国数字经济高质量发展与数字中

国建设注入强劲动力。

2. 问卷设计理念

问卷共设计了 33 个问题，包含五个部分，第一部分是关于 AI 大模型的认知与使用情况的 6 个问题，第二部分是关于 AI 大模型的信任、体验与风险感知的 10 个问题，第三部分是关于“智能体”与“具身智能”的 5 个问题，第四部分是关于 AI 大模型的治理、权利意识以及伦理原则方面的 9 个问题，第五部分是关于消费者个人基本信息的 3 个问题。问卷具体内容请参看附录。其中一些问题与本课题组 2024 年的调研是相同的，我们试图发现消费者认知及态度是否有变化。

在第一部分，问卷重点考察用户对 AI 大模型的认知基础与使用实践。通过设置对 AI 大模型了解程度、使用频率、主要使用场景、自评熟练度、是否接受过相关培训以及学习方式等问题，全面刻画用户在接触、学习和使用 AI 大模型过程中的能力结构差异。这一设计有助于区分不同认知水平和使用经验群体，识别潜在的人工智能素养差异。

第二部分围绕 AI 大模型的信任、使用体验与风险感知展开，重点关注消费者在实际使用过程中的主观评价与心理反应。通过设置正向体验、整体效用感知、能力进步判断、任务信任边界以及使用过程中遇到的问题等问题，问卷力图揭示消费者如何在效率提升与潜在风险之间进行权衡，尤其是在不同风险等级任务中对 AI 大模型的信任程度。这一部分的设计体现了当前生成式人工智能研究中“信任校准”与“风险分级”的前沿关注。

第三部分聚焦“智能体”与“具身智能”，反映出本次调查对人工智能新形态的重点关注。“智能体”与“具身智能”被视为 AI 大模型从内容生成走向自主规划和执行的重要发展方向，其更高的自主性与复杂性可能带来新的机遇与风险。通过测量受访者对“智能体”的了解程度、使用情况、期待与担忧，本部分旨在评估公众对这一前沿技术形态的认知基础和社会接受度，为后续治理与政策讨论提供经验依据。

第四部分聚焦 AI 大模型的治理、权利意识与伦理原则，回应人工智能广泛

应用所引发的规范性问题。问卷不仅考察受访者对 AI 大模型整体信任程度、治理紧迫性的判断，还进一步测量其对相关法律法规的了解情况、对算法解释和拒绝自动化决策等权利的认知，以及在遭遇不公正或歧视性结果时的潜在维权选择。同时，通过对公平、透明、隐私保护、安全可靠、人类监督等伦理原则关注度的调查，力图呈现消费者在价值层面对 AI 治理的核心诉求。

第五部分为人口学变量，包括学历、职业和收入水平等信息，旨在支持不同群体之间的比较分析，进一步揭示人工智能认知、使用和态度差异背后的社会结构因素。

二、样本的基本特征

本次调查共回收了 7061 份有效问卷，调查对象的性别、年龄、地域分布如下。

1. 调查对象的性别、年龄和区域分布

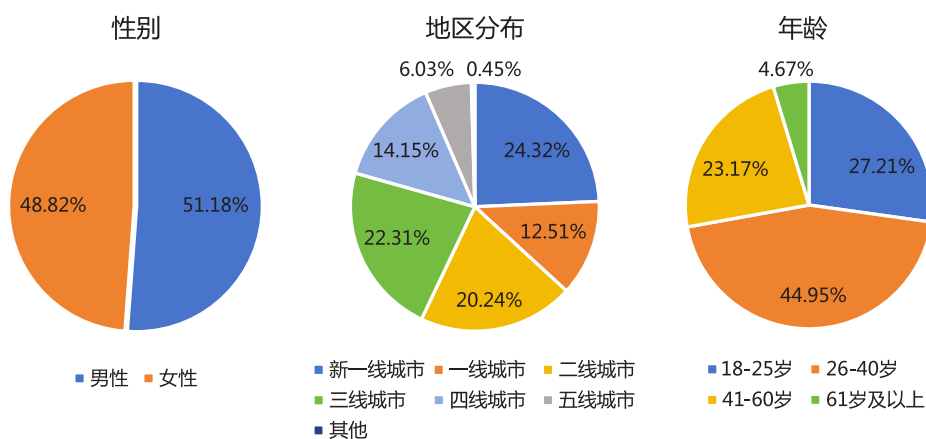


图 1 调查对象的性别、年龄、区域分布

- 从性别结构看，男性占比 48.82%，女性占比为 51.18%，占比相对均匀。
- 从年龄结构看，21 至 40 岁占比最高，60 岁以上占比最低，其余年龄段分布较为平均，大致符合中国互联网用户年龄分布情况。
- 从区域结构看，五线城市占比为 6%，相对较低，其他等级城市占比较为均匀。

2. 调查对象的学历、职业和收入分布

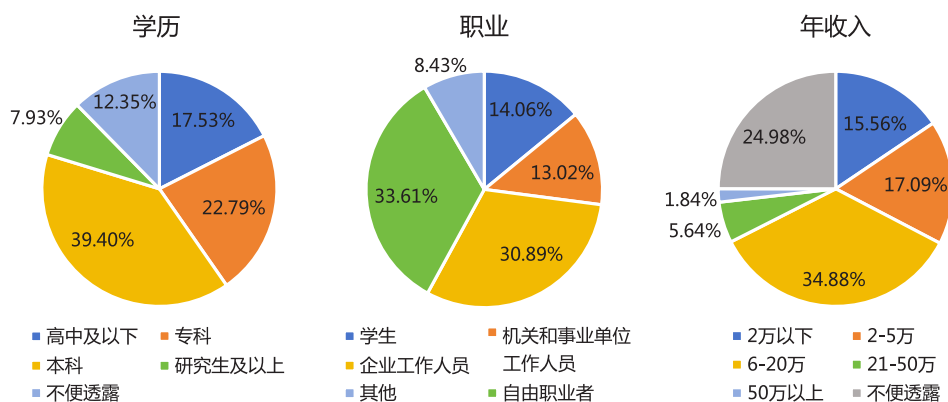


图 2 调查对象的学历、职业和收入分布

- 从学历构成看，半数以上为本、专科学历，研究生及以上高学历群体比重仅为 8%。
- 从职业分布看，企业工作人员、机关事业单位工作人员的比重为 30.89% 和 13.02%。自由职业者比重为 33.61%，占比最高，这可能与自由职业者多在网上工作，回答问卷的概率较高有关。
- 从收入分布看，调查对象 70% 年收入在 20 万元以下，其中年收入在 6 万-20 万的群体占比 34.88%，占比最高。

三、消费者对 AI 大模型的认知和态度

1. 用户对 AI 大模型的认知与使用现状

第一，虽然使用渗透率超过 50%，但用户整体对 AI 大模型的了解和使用情况差异依然明显。样本整体存在明显的双重差距。第一重是使用过与未使用过群体的差距，约 54.27% 的受访者使用过 AI 大模型，但仍有 45.73% 的受访者未实际使用。第二重是“会用”与“精通”的差距，使用过 AI 大模型的受访者中，72.94% 仅“部分了解”相关知识，“非常了解”者仅占 12.57%，自我评估使用熟练度时，62.21% 的用户认为自己“一般熟练”，“比较熟练”者仅 17.94%，这揭示当前用户多能通过基础操作解决明确问题，但对技术原理及高级功能缺乏系统认知。

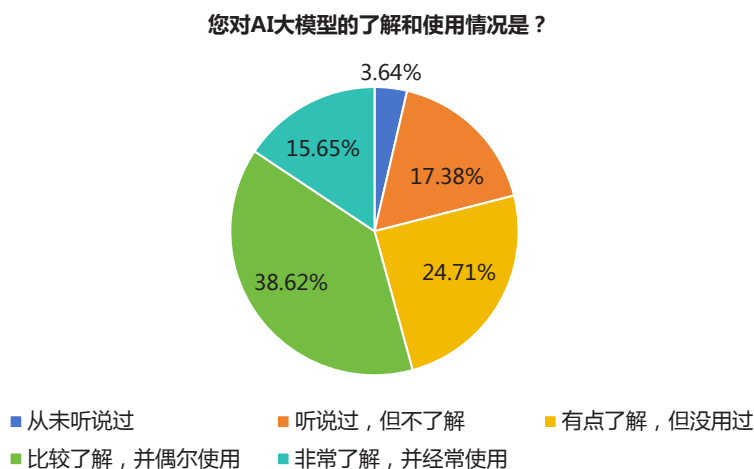


图 3 调查对象对 AI 大模型的了解和使用情况

您对AI大模型相关知识的了解程度是？

您觉得自己在使用AI大模型方面的熟练程度如何？

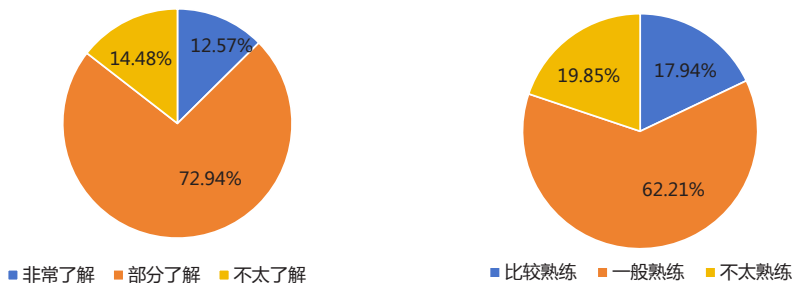


图 4 调查对象对 AI 大模型的了解及熟练程度

第二，AI 大模型的使用场景高度聚焦实用类需求，学习与工作是核心应用场景。对使用过 AI 大模型的受访者进行场景分布调研发现，用户需求高度聚焦实用类场景。其中，学习用占比 56.23%，工作用占比 53.30%，生活用占比 48.87%，三类场景占比均显著高于娱乐用（36.67%），这一数据表明，AI 大模型主要被定位于解决学习与工作难题的生产力工具。

您常在什么方面使用AI大模型？

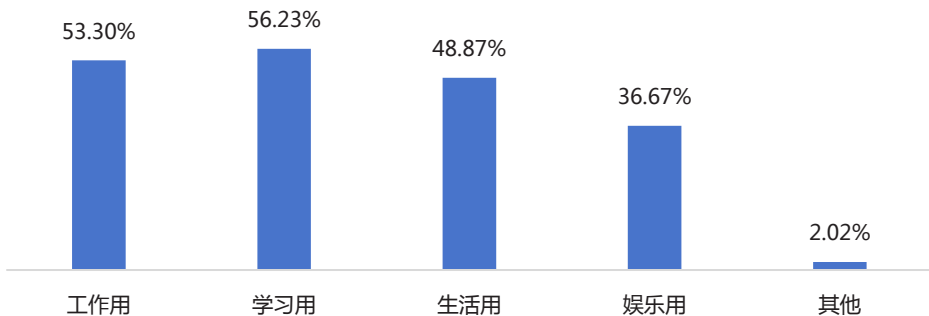


图 5 调查对象使用 AI 大模型的场景分布

第三，用户 AI 技能的学习路径呈现自主化、碎片化特征，“被动使用”是主流学习方式，系统性培训覆盖不足。从培训覆盖情况看，超七成受访者（72.71%）未接受过任何 AI 大模型相关培训，仅 4.87%的受访者接受过课程、工作培训等正式培训，22.42%的受访者通过自学教程、他人指导完成非正式培训，外部结构化的学习支持普遍缺失。在学习方式的选择上，“被动使用”的占比最高，达 40.28%；“主动探索”的受访者占 29.05%，这表明当前用户主要依靠用中学和自我驱动的模式提升 AI 技能，缺乏体系化的外部学习引导。

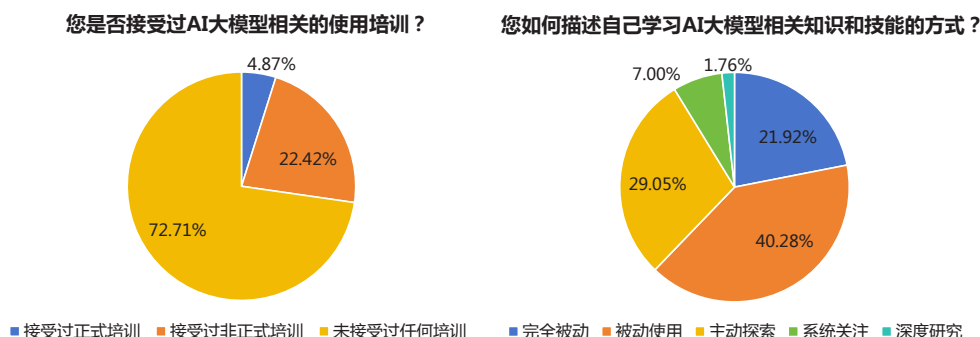


图 6 调查对象对 AI 大模型的培训接受情况及知识技能学习方式

第四，用户对智能体、具身智能等 AI 新形态的认知尚处于初级阶段，深度认知群体极少。智能体方面，仅 6.95% 的受访者表示“非常了解”，28.83% 的受访者表示“完全不了解”；具身智能方面，仅 5.98% 的受访者表示“非常了解”，“完全不了解”的受访者则高达 31.57%。这一现状表明，智能体、具身智能等更复杂的 AI 技术仍处于市场教育初期，用户对其功能、应用场景的认知深度不足，后续需通过场景化科普、应用示范等方式，逐步缩小新形态 AI 的认知缺口。

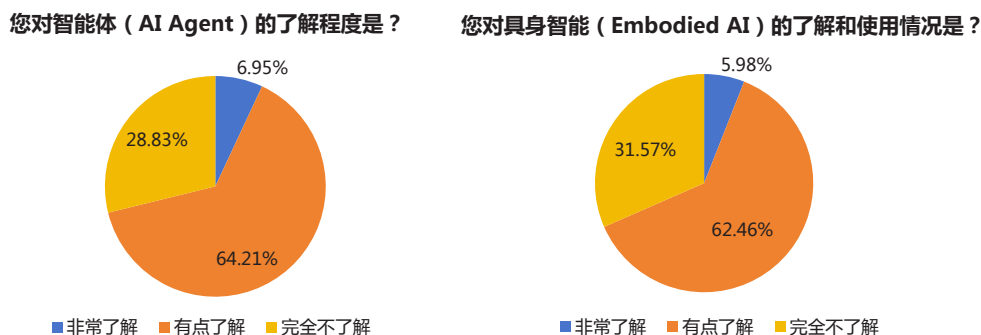


图 7 调查对象对智能体与具身智能的了解程度

2. 用户对 AI 大模型的使用体验与信任态度

第一，受访者对 AI 大模型的态度以积极为主，且体验整体正向，能力进步感知显著。提到 AI 大模型时，“高效”“便利”等正向联想词的占比均超 64%，

而“威胁”“失控”等负向词占比不足 14%，这说明多数用户对 AI 大模型的主观印象偏向积极。在已经使用过 AI 大模型的群体中，分别有 73.60%、71.69% 和 46.18% 的受访者表示其在提高工作/学习效率、解决生活问题和激发创意方面得到了正向体验，证实了 AI 大模型在这些领域具有较高实用性；结合体验程度与能力感知，57.75% 的受访者认为正向体验“非常显著”，58.85% 的受访者感知到 AI 能力“有明显进步”，进一步说明 AI 大模型的实际效用与技术迭代已获得使用群体的普遍认可。

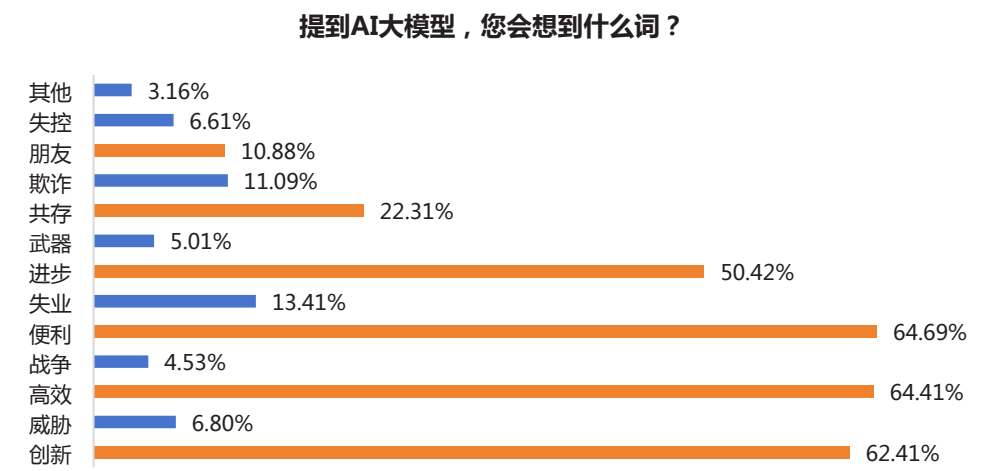


图 8 调查对象对 AI 大模型的联想词汇分布

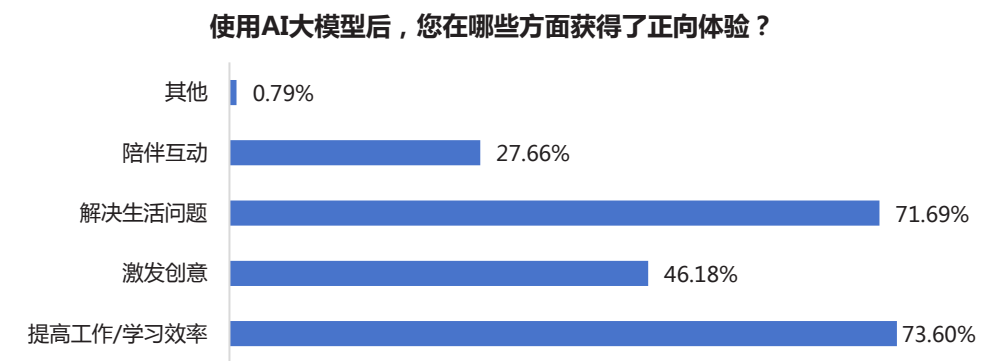


图 9 调查对象使用 AI 大模型后的正向体验分布

第二，受访者对 AI 大模型的信任呈明显分层特征，更倾向于信任低复杂度、低风险的任务与场景。对于基础信息查询、辅助创作、简单事务处理这类轻量、明确任务，愿意信任 AI 大模型的受访者占比分别达 59.51%、63.43%、61.45%，

均接近或超六成；而对专业问题参考、复杂事务处理这类任务，愿意信任 AI 大模型的受访者占比则明显下降，仅为 46.69%和 29.00%，可见对 AI 的信任会随任务自主性与责任要求提高而递减。场景层面，对于日常生活（68.51%）、教育学习（67.47%）等低风险场景，愿意信任 AI 的受访者占比均超六成；而对金融服务（18.42%）、司法领域（17.64%）等高风险场景，愿意信任 AI 的受访者占比均不足两成，反映出用户在涉及核心利益、专业决策时的谨慎态度。整体来看，用户对 AI 大模型的整体信任态度以谨慎接纳为主，60.71%的用户处于“半信半疑”状态，明确信任与不信任占比偏低。

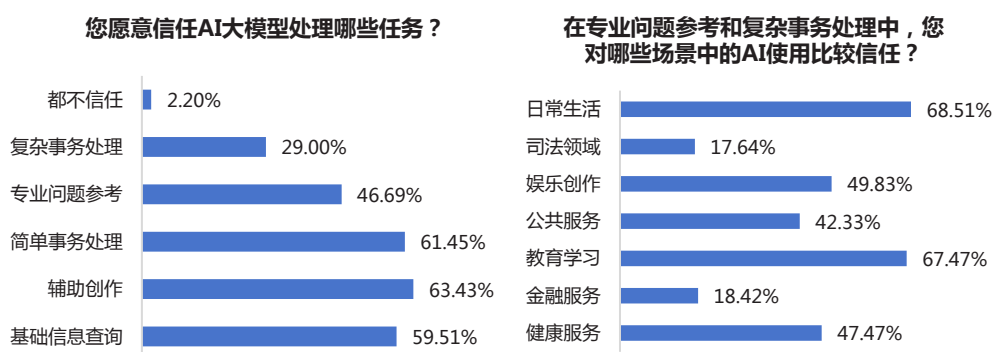


图 10 调查对象对 AI 大模型的任务信任及专业场景信任分布

整体上，您对AI大模型的信任程度是？

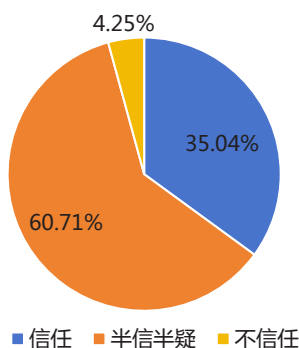


图 11 调查对象对 AI 大模型的整体信任程度分布

第三，受访者对 AI 新形态的认知局限在生活场景，对其未来发展以谨慎观望为主。对于“智能体”等人工智能新形态，受访者的核心期待在于提升生活便利度（69.04%）和激发个人创造力（63.93%），但对工作场景中的替代人类从事高风险活动和减少重复劳动的期待较低。关于智能体和具身智能是否要大力

发展，受访者普遍持观望态度（54.58%），仅 34.27%认为需要大力发展，反映出受访者既看到 AI 新形态的潜力，也对其潜在风险存在顾虑。

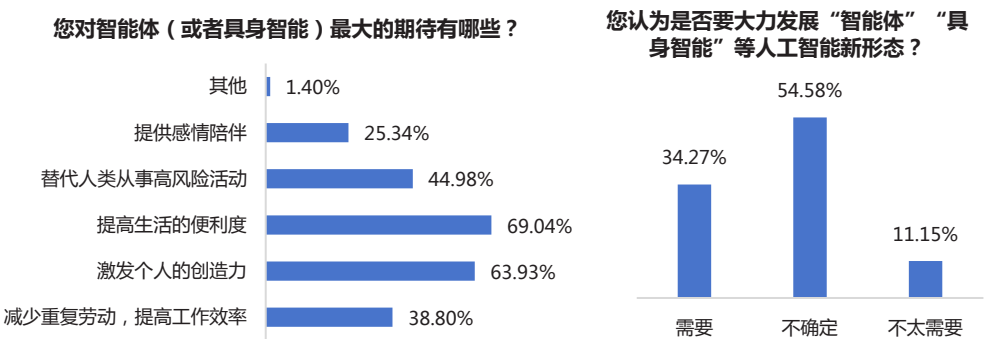


图 12 调查对象对 AI 新形态的期待及发展态度分布

3. 用户对 AI 大模型的问题遭遇与风险认知

第一,用户使用 AI 大模型的核心痛点相对集中在隐私安全与结果质量方面。在使用问题反馈中,“数据隐私和安全性问题”占比最高,达到 39.61%。“生成的结果不令人满意”占比为 37.19%,紧随其后。“个人创造性下降”、“错误和无用信息”等问题的反馈占比也相对较高,反映出用户对 AI 生成内容实用性和自身信息安全的关注。

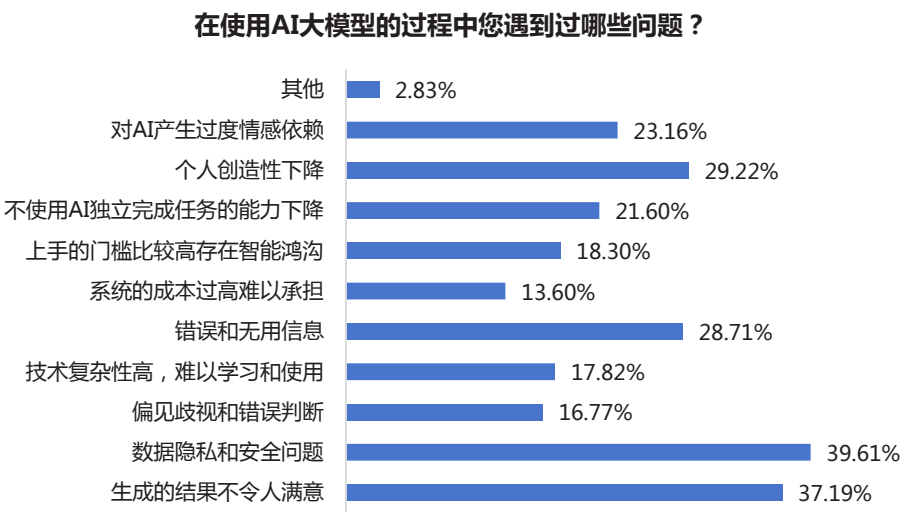


图 13 调查对象对 AI 大模型的问题遭遇

第二，用户对 AI 大模型的风险担忧强度依次是能力退化、隐私问题、工作岗位替代和 AI 幻觉，40%左右的用户都表现出上述四个方面的担忧。在风险认知调研中，“人类过度使用 AI 工具，能力退化”的认同占比最高，达到 45.79%。“不受控制地搜集和处理个人信息，隐私问题”占比为 41.17%。“工作岗位被 AI 替代”的担忧占比为 39.44%，体现出用户对过度依赖 AI 的负面影响及个人权益的顾虑。

您认为AI大模型对人类的风险在于？

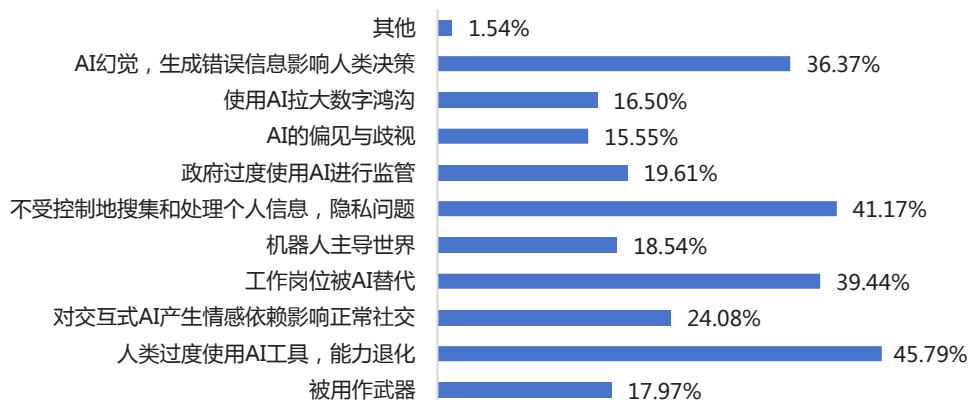


图 14 调查对象对 AI 大模型的具体风险认知

第三，用户对 AI 新形态的顾虑主要表现为失业风险和技术失控。针对智能体、具身智能等 AI 新形态，“导致大规模的失业”是用户最担心的问题，占比达到 46.42%。“放大单个智能体的固有缺陷，形成系统性风险”占比为 42.69%。“技术脱离人类的掌控”、“事故责任难以界定”等问题的关注占比也均超三成，反映出用户对 AI 新形态发展潜在隐患的谨慎态度。

您对“智能体”或“具身智能”主要的担心有哪些？

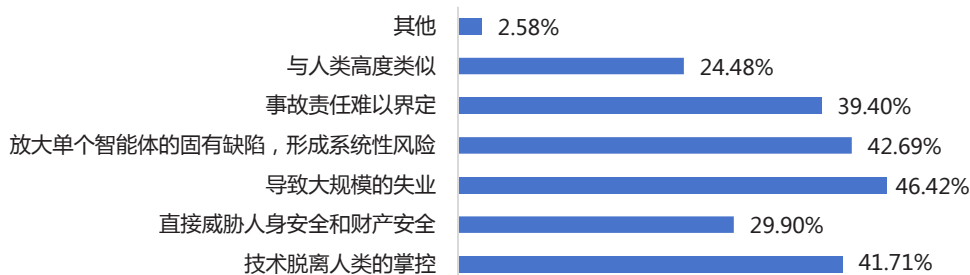


图 15 调查对象对 AI 新形态的未来隐患评估

第四，用户面对 AI 广泛应用的情绪呈**复杂矛盾特征**，**不确定感和新鲜感最为突出**。在情绪反馈中，“不确定感（无法判断 AI 结果的可靠性）”占比最高，达到 35.99%。“新鲜感（赶上现在潮流）”占比为 35.73%。同时“焦虑（担心依赖 AI 后自身能力退化）”、“恐慌（担心被 AI 取代工作）”等负面情绪也有一定占比，说明用户在体验技术便利的同时，也伴随着对技术可靠性和自身发展的担忧。

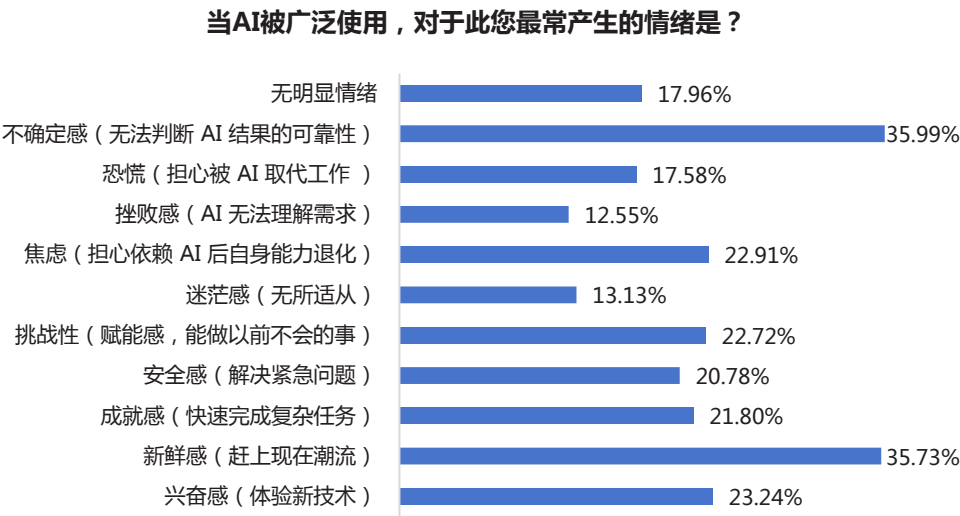


图 16 调查对象对 AI 应用的感性情绪

4. 用户对 AI 大模型的规范诉求与权益保护行为

第一，关于 AI 大模型的未来发展，法律法规的完善最受用户关注，其次是技术的创新和突破。在未来发展关注问题的调研中，“法律法规的完善”占比最高，达到 50.67%。“技术的创新和突破”占比为 47.08%。“服务应用的拓展”、“伦理道德的讨论”、“社会接受度的提高”的关注占比也均超三成。

您认为AI大模型在未来的发展中最应该关注什么问题？

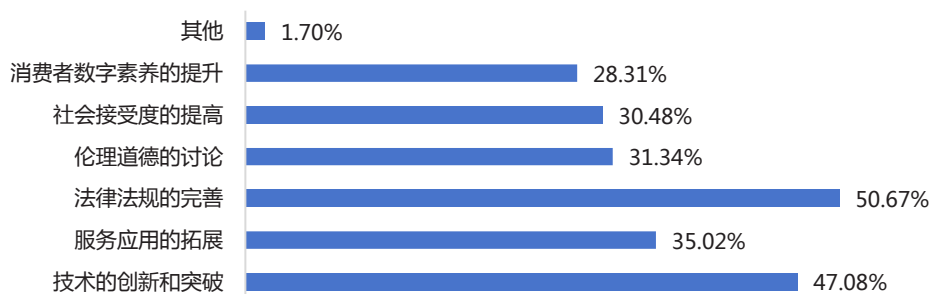


图 17 调查对象对 AI 未来发展的关注问题

第二，在伦理原则关注上，“隐私保护与数据治理”的占比遥遥领先，达到61.12%。“透明性与可解释性”、“安全与可靠”的关注占比也均超四成，体现出用户对 AI 发展规则约束和安全底线的重视。

您更关注以下哪些人工智能领域的伦理原则？

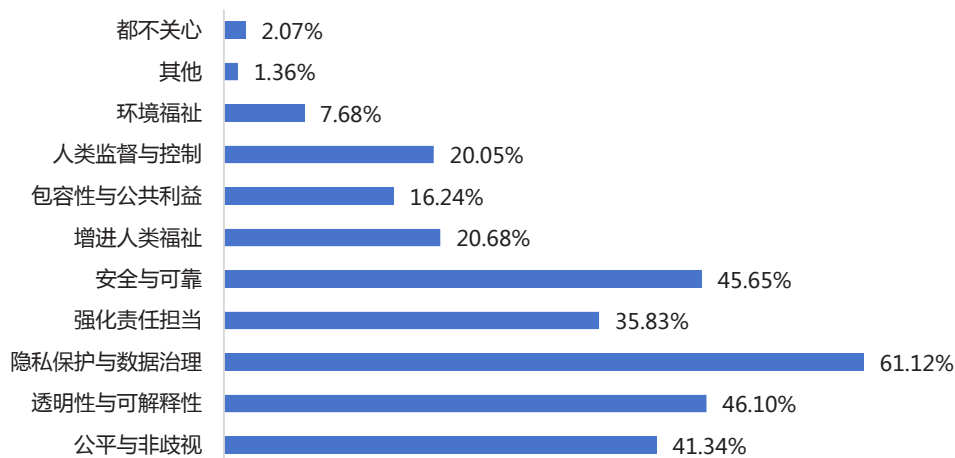


图 18 调查对象对 AI 领域的伦理原则关注

第三，为应对 AI 大模型带来的风险，在个人及社会层面，用户都主张积极主动作为。在个人层面，用户的权益保护行为以主动防护与知识学习为主。在个人规避不利影响的措施中，“注意保护个人隐私和数据安全”占比最高，达到66.82%。“主动学习基础知识，增强对 AI 大模型的认知和理解”占比为 60.61%。“加强对自身能力的培养”占比为 42.97%，表明用户倾向于通过自我提升和主

动防护的方式，降低 AI 使用带来的风险。在社会层面，用户主张通过**强化监管与多元治理降低 AI 风险**。在风险应对措施的调研中，“加强 AI 监管方面的立法和司法实践”占比最高，达到 54.89%。“构建全社会开放、合作、协商的多元治理机制”占比为 37.60%。“鼓励技术创新以解决可能的风险”、“加大对企业的监管和处罚力度”等措施也得到较多用户认同，反映出用户希望通过顶层设计和多方协作的方式规范 AI 发展。整体而言，超七成用户（72.11%）认为需要**加快 AI 治理进程**，其中“比较紧迫”占比 54.31%，“非常紧迫”占比 17.80%。

为避免AI大模型应用带来的不利影响，您通常会？

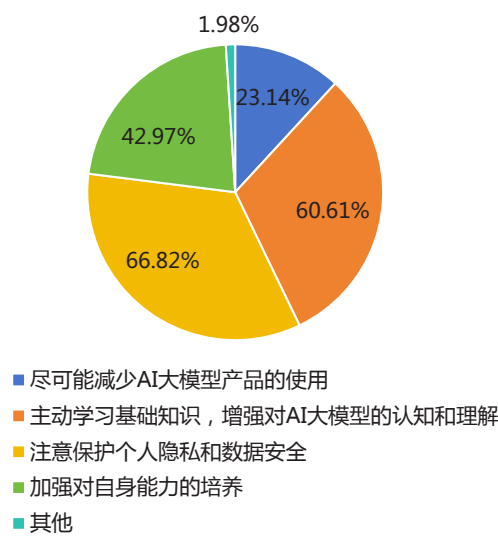


图 19 调查对象为应对 AI 风险的个人行为选择

您认为应该采取哪些措施以降低AI大模型可能带来的风险？

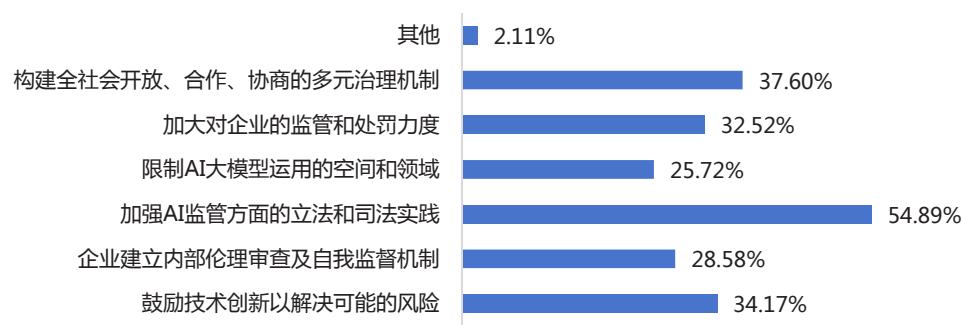


图 20 调查对象为应对 AI 风险的宏观措施选择

总体而言，您认为当前对AI发展进行规范和治理的紧迫程度如何？

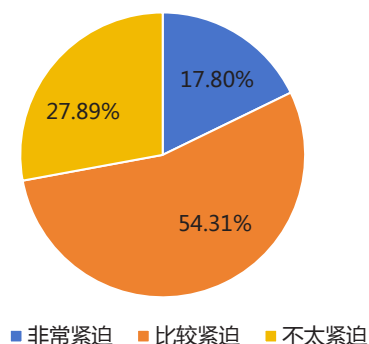


图 21 调查对象对 AI 治理的紧迫程度评估

第四，用户对自身在使用 AI 产品时拥有的合法权利的认知程度普遍偏低。

仅 5.74%的用户亲自阅读过国家有关算法安全、AI 大模型安全的相关规定及自身合法权利的相关内容，35.01%的用户听说过部分内容，25.70%的用户只听说过名字，高达 33.55%的用户表示完全不了解。在 AI 产品算法知情权和信息收集拒绝权的认知上，仅 5.58%的用户非常了解，25.42%的用户大致了解，41.35%的用户听说过但不清楚具体内容，还有 27.64%的用户完全不了解。以上数据反映出用户对 AI 相关权益的认知存在明显短板，多数人并不清楚自身在使用 AI 产品时拥有的合法权利，也缺乏对相关政策法规的了解。

您是否了解国家有关算法安全、AI大模型安全的相关规定和公民拥有的合法权利？

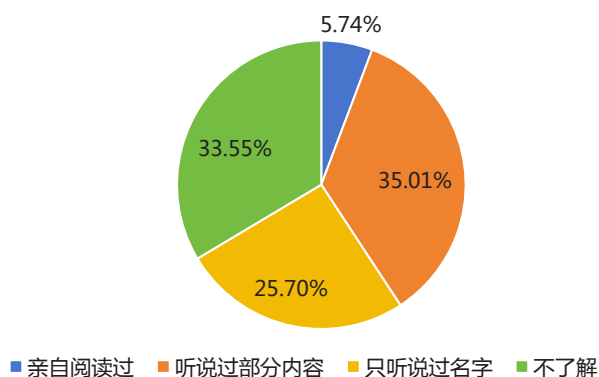


图 22 调查对象对 AI 领域安全的规定与自身权利的了解程度

您是否了解，在使用AI产品时您依法享有要求企业进行算法解释
并在特定情况下拒绝自动信息收集的权利？

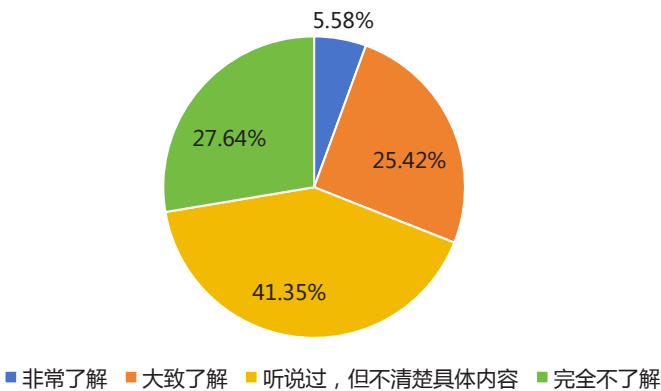


图 23 调查对象对 AI 算法知情权和信息收集拒绝权的了解程度

第五，用户遭遇 AI 使用带来的不利影响时的维权行为偏向理性合规。当用户认为 AI 大模型提供的答案存在严重偏见或歧视并对自身造成不利影响时，43.21%的用户会优先选择向提供该服务的公司投诉和反馈，32.64%的用户会向网信办、消协等监管机构举报，21.57%的用户选择在社交媒体上曝光以提醒他人，30.59%的用户反应消极，表示会不再使用该产品但不主动维权。仅 12.73%的用户表示不知道能采取什么行动。这一结果表明，大部分用户在权益受损时，更倾向于通过官方投诉渠道和监管举报途径解决问题，维权行为较为积极和理性，消极性和无知性用户占比不高。

如果您认为AI大模型对您造成了不利影响，您会采取哪些行动？

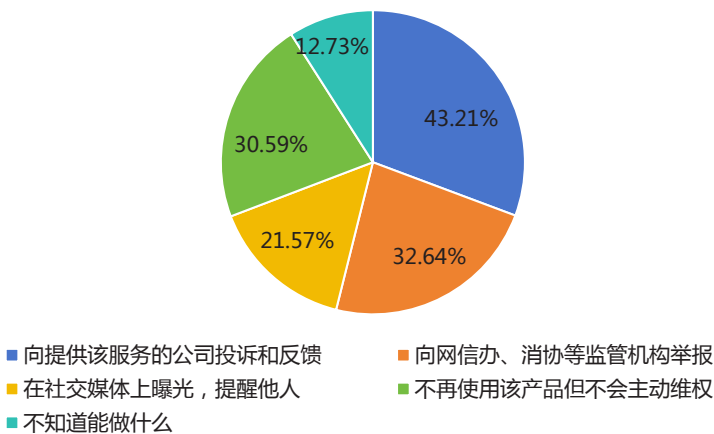


图 24 调查对象在 AI 使用过程中的维权行为选择

第六，用户对 AI 生成内容标识的态度以理性核实为主，极端排斥态度占比极低。在面对带有“AI 生成”标识的内容时，多数用户（57.44%）会“理性谨慎，核实后使用”。“产生戒备，减少使用”、“基本无视、习惯照旧”也有一定占比。仅有 8.45% 的用户会“强烈排斥，避免使用”。这表明，多数用户能够理性看待 AI 生成内容的标识，既不会盲目排斥，也不会完全忽视，而是倾向于通过核实来规避风险。

您看到 AI 生成内容上的“AI 生成”标识时的反应是什么？

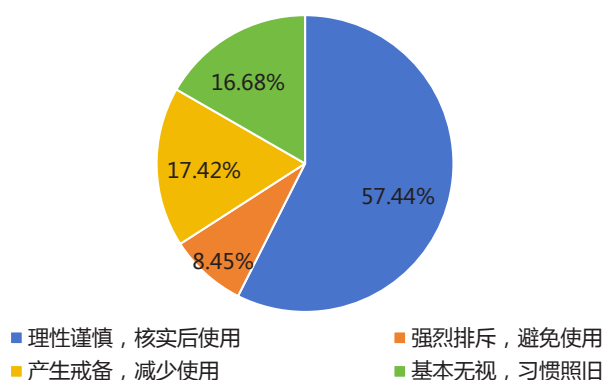


图 25 调查对象对“AI 生成”标识的反应

四、与国际社会的对比

1. 对比报告简介

牛津大学路透新闻研究所于 2024 年 5 月发布《What Does the Public in Six Countries Think of Generative AI in News?》^[1]，基于阿根廷、丹麦、法国、日本、英国、美国六国 12217 份消费者样本（每国约 2000 份）的在线调研（2024 年 3-4 月），核心聚焦六国消费者对生成式 AI 的认知普及、使用行为、场景偏好、风险信任及治理诉求，为全球消费者 AI 认知研究提供了重要国际参照，从消费者视角呈现了国际社会对生成式 AI（尤其是新闻领域）的认知图景，其核心维度、关键数据与国内消费者 AI 大模型认知调研呼应，为跨区域对比分析提供了直接参照。

2. 与国际社会的比较——对 AI 的认知与场景信任比较

第一，AI 认知率普遍较高，且中国 AI 认知普及率更高。中国“从未听说过 AI”的仅 3.64%，远低于六国平均 22.5%，AI 渗透更充分；从认知深度与使用看，中国不仅 96.36%听说过 AI，更有超 54%（38.62%+15.65%）达到“了解+使用”层级，而六国 ChatGPT 认知率最高仅 61%（丹麦），且多为“仅听说”，国内 AI 认知的深度与使用程度均更突出。

^[1]FLETCHERR,NIELSEN RK.WhatDoes the Public in Six Countries Think of Generative AI in News?[R].Oxford:Reuters Institute for the Study of Journalism,University of Oxford,2024.<https://doi.org/10.60625/risj-4zb8-cg87>.

您对AI大模型的了解和使用情况是？——中国

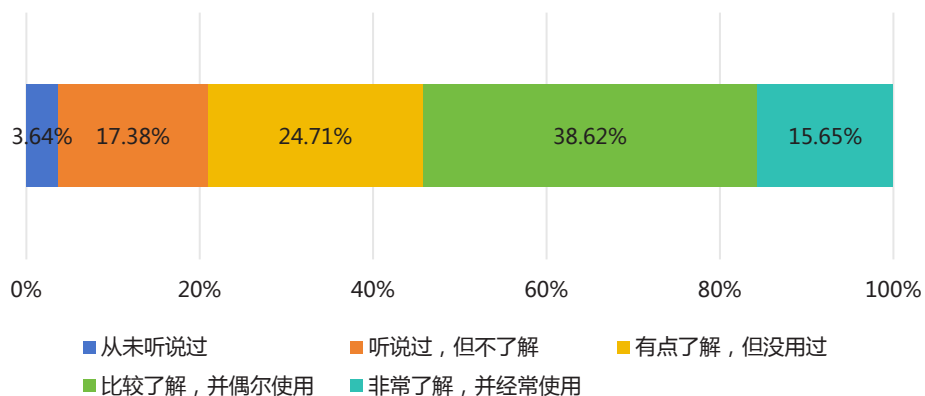


图 26 中国消费者对 AI 大模型的了解程度

Proportion who have heard of each generative AI tool

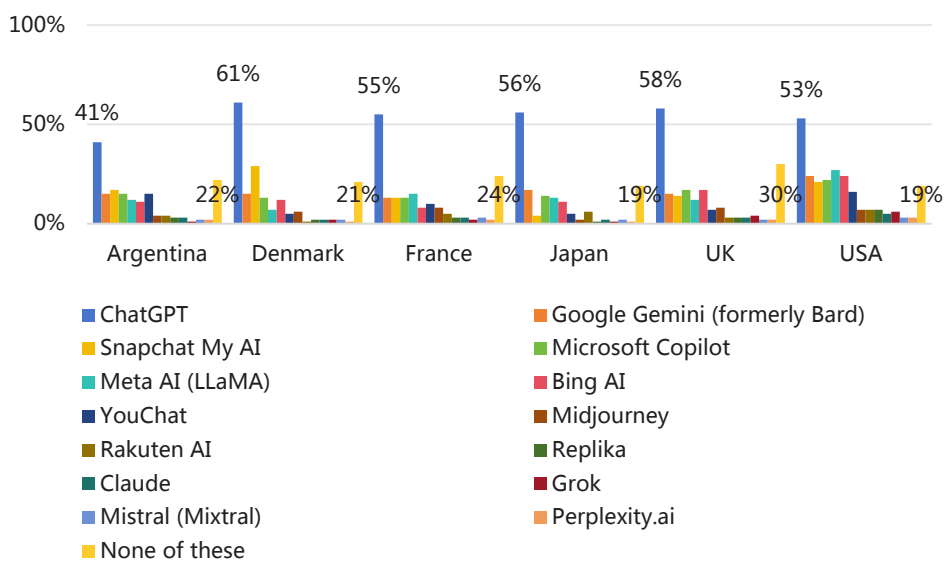


图 27 六国消费者对 AI 大模型的知晓程度

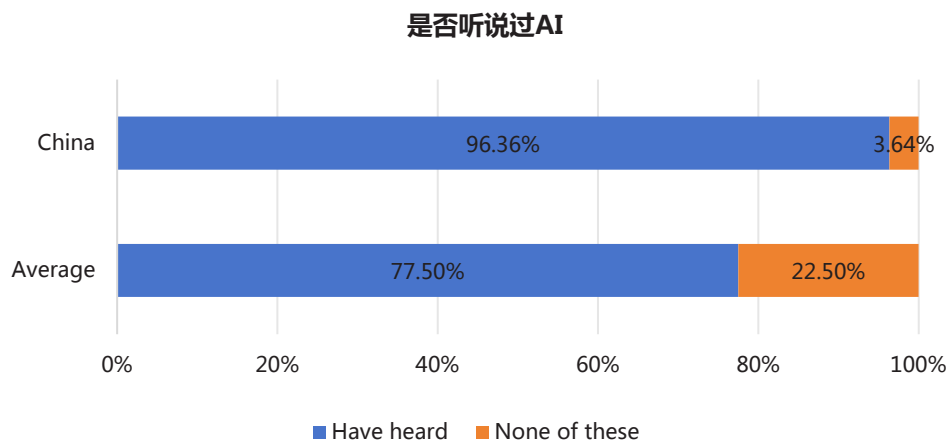


图 28 消费者对 AI 大模型的知晓程度——中国与国际对比

第二，从使用场景看，中国 AI 应用渗透度显著高于六国。工作学习场景中国使用率达 44%，远超六国平均 21%，其中美国最高仅 28%；生活场景中国使用率 42%，同样高于六国平均 27%，其中美国最高 35%。可见中国消费者已将 AI 更深度融入日常工作学习与生活，场景使用的普及性优于国际水平。

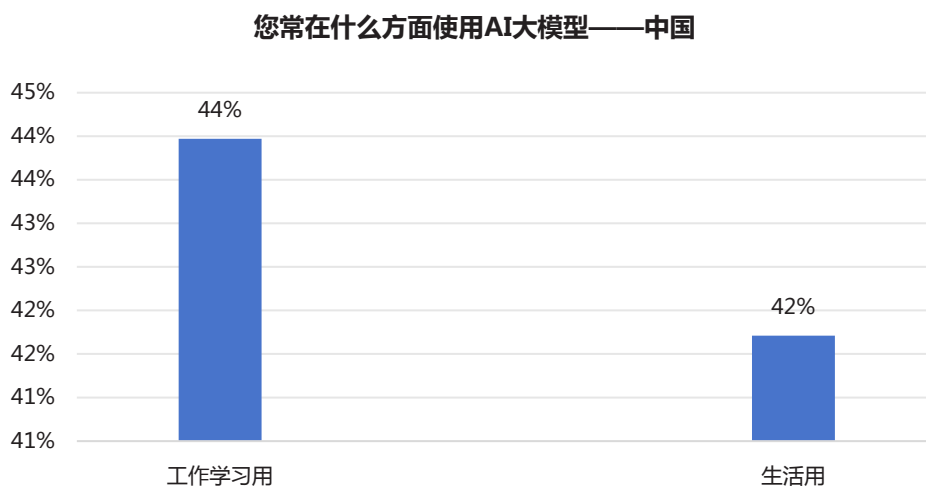


图 29 AI 大模型在中国的使用场景

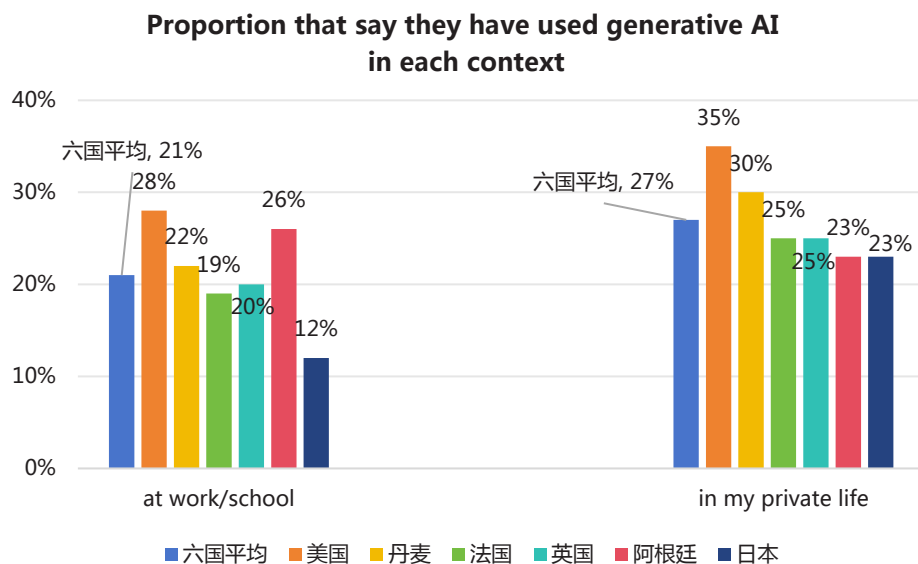


图 30 AI 大模型在六国的使用场景

第三，从任务信任度看，中国与国际存在一定差异。中国消费者对 AI 处理各类任务的信任，覆盖基础信息查询、辅助创作、事务处理等多类任务，信任度分布相对均衡，未呈现明确的信任优先级区分，整体对 AI 的信任场景较为多元。国际消费者对 AI 任务的信任呈清晰排序：“获取最新新闻”信任度最低，其次是“编程”，而“事实性问题回答、创意生成、图像制作、邮件写作”等任务的信任度相对更高，信任梯度区分明确。

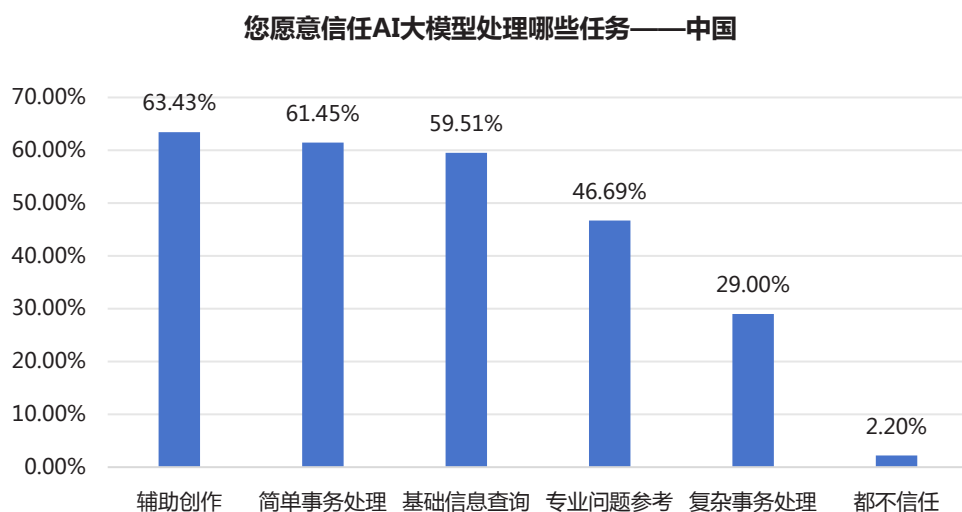


图 31 中国消费者对 AI 大模型的任务信任度

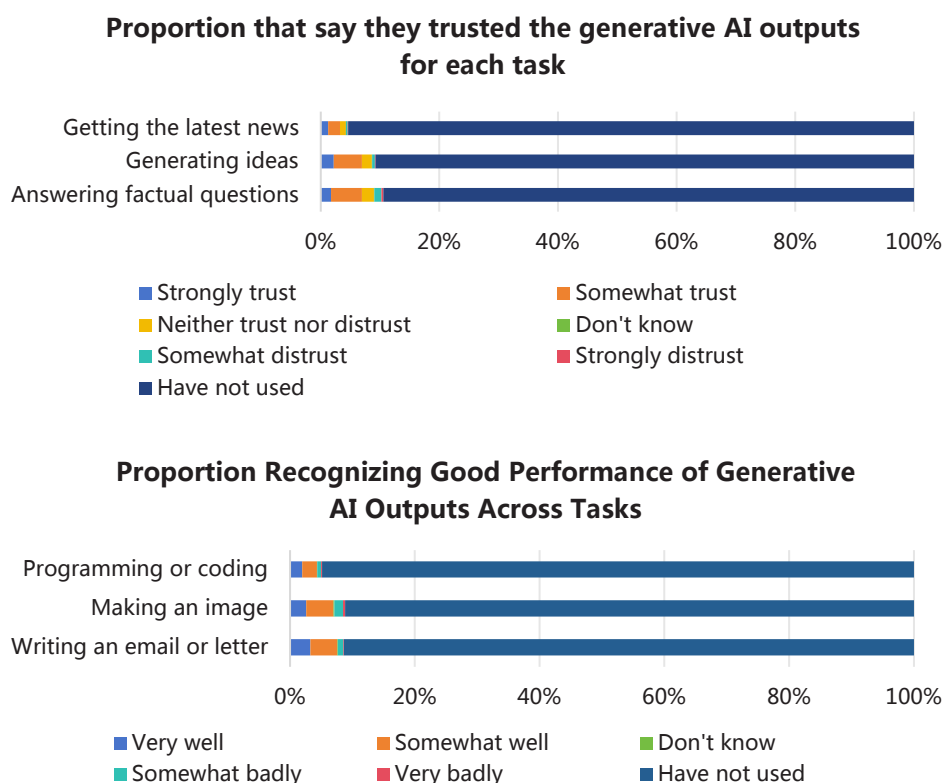


图 32 国际消费者对 AI 大模型的任务信任度

第四，尽管国内外调研维度存在差异（国内聚焦 AI 在当下场景中的使用信任度，国际侧重对 AI 未来影响的预判），但仍可清晰观察到公众对 AI 应用场景的关注存在显著分野。中国消费者对 AI 的信任高度集中于个人高频使用场景，日常生活（66.81%）、教育学习（62.69%）的信任度显著领先，娱乐创作、健康服务等场景也获得近五成信任，而对金融服务（18.76%）、司法领域（17.43%）等专业领域的信任度则明显偏低。国际公众则更前瞻地审视 AI 对社会结构的变革影响，普遍预判 AI 将对医疗从业人员（59.10%）、金融机构（58.83%）、政府（53.36%）及执法部门（49.61%）产生显著的影响，对普通民众个体层面的影响预期（48.00%）虽不低，但仍排在专业机构之后。这一差异表明，国内公众更倾向于将 AI 视为服务于个人生活的实用工具，信任度与个人使用场景的适配性高度绑定；而国际公众则更关注 AI 在专业领域和社会系统层面的变革性作用，认为其核心价值将首先在专业机构中释放，二者对 AI 发挥核心价值的场景认知存在明显分野。

您对以下场景中的AI使用比较信任

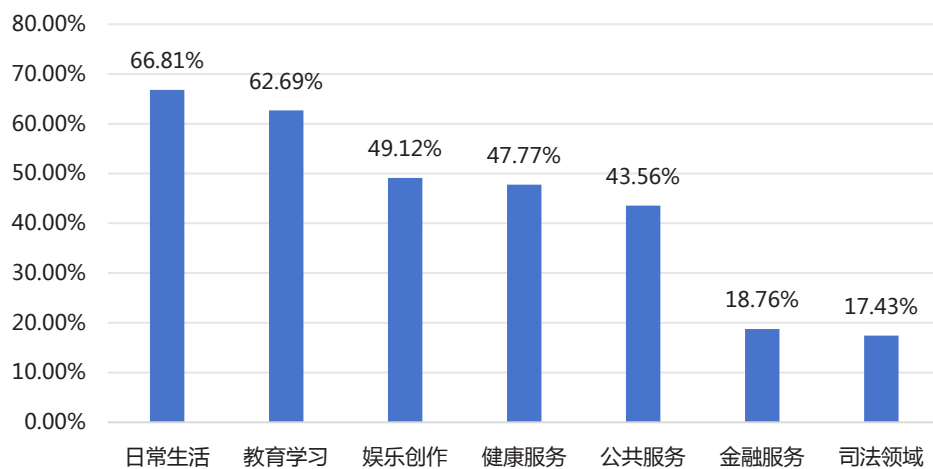


图 33 中国消费者的 AI 大模型信任场景

Proportion that think generative AI will have a large impact upon each

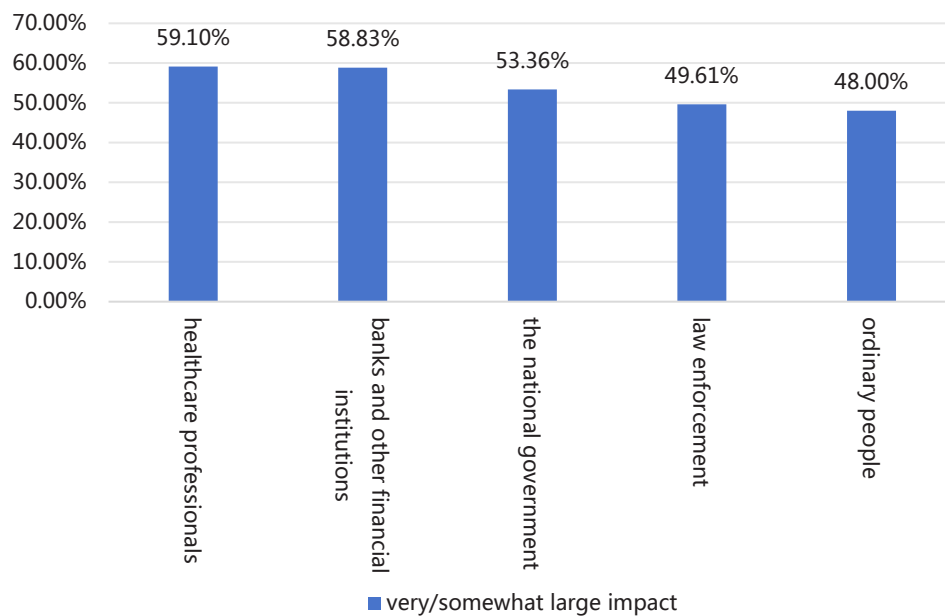


图 34 国际消费者对 AI 大模型的预期影响

五、消费者对 AI 大模型认知的变化

2025 年的调研中，部分关于 AI 大模型的消费者认知的题目与 2024 年的研究《算法与 AI 大模型的用户认知调研报告》^[2]是相同的。因此，本报告展开动态对比分析。整体来看，这些变化主要体现在五个方面：首先，消费者对 AI 大模型的认知普及度显著提升，显示出从简单的知晓向深度理解迈进的趋势。其次，尽管中间群体对 AI 大模型的理解有所加深，但深度认知群体的比例却略有收缩，这反映了普通消费者在面对技术快速迭代时可能遇到了知识更新的挑战。第三，AI 大模型的应用场景重心发生了转移，工作和娱乐场景中的应用增长尤为显著，而学习场景的使用则有所回落，显示了消费者需求的变化。第四，在体验感知上，虽然仍保持正面评价，但边际效益呈现递减趋势，反映出随着 AI 工具的常态化使用，其带来的新鲜感和独特价值正在逐渐减弱。最后，消费者对于 AI 大模型的整体评价和期待也变得更加务实和理性，关注点从单纯的隐私安全和技术复杂性转向更深层次的人机关系、伦理道德以及社会影响等方面，体现了公众对于 AI 技术发展更加成熟的看法。这些变化共同描绘了一个动态发展的图景，展示了消费者对 AI 大模型认识的深化和演变。

1. 关于 AI 大模型认知与使用基础水平

第一，消费者对 AI 大模型的认知普及度显著提升。2025 年“从未听说过”AI 大模型的受访者比例仅为 3.64%，较 2024 年的 10.72%下降超 7 个百分点，降幅达 66%；而 2025 年“比较了解，并偶尔使用”“非常了解，并经常使用”的受访者合计占比达 54.27%，远超 2024 年“偶尔使用”与“经常使用”合计

^[2]对外经济贸易大学数字经济与法律创新中心，中国人民大学数字经济研究中心，蚂蚁集团研究院.算法与 AI 大模型的用户认知调研报告[R].北京，2024.

的 49.21%，表明消费者对 AI 大模型有了更多的认识和理解，从“知晓”向“深度理解”层面快速迈进。

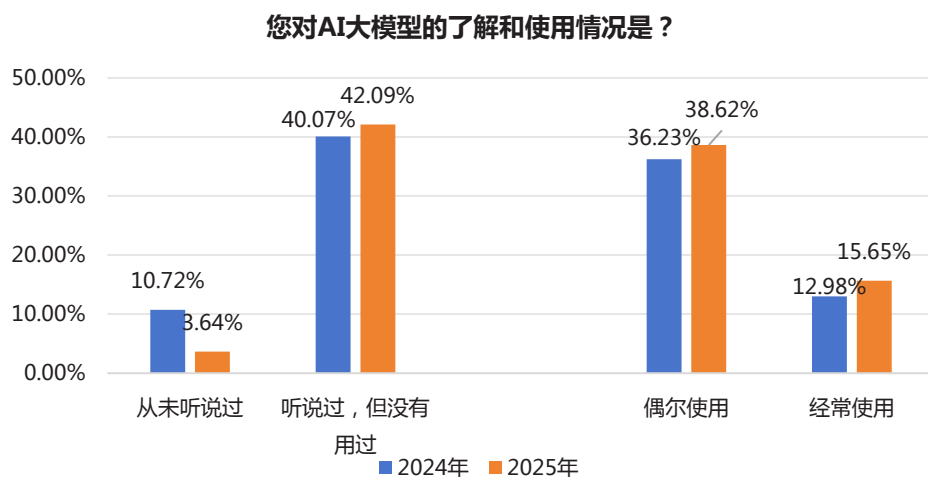


图 35 消费者对 AI 大模型的了解和使用情况变化

第二，消费者关于知识理解程度略有分化。2025 年的调研数据显示，60.86% 的受访者对 AI 大模型的功能特性、技术类型及运行原理处于“部分了解”状态，较 2024 年的 57.21% 提升 3.65 个百分点。值得注意的是，深度认知群体出现小幅收缩——“非常了解”的比例从 2024 年的 10.19% 回落至 2025 年的 8.19%；同时，“不太了解”的受访者占比从 27.99% 攀升至 30.96%。这一数据变化勾勒出“中间群体扩张、深度认知群体收窄”的认知结构演变特征，折射出在 AI 技术高速迭代的背景下，普通消费者可能面临知识更新滞后的现实困境。

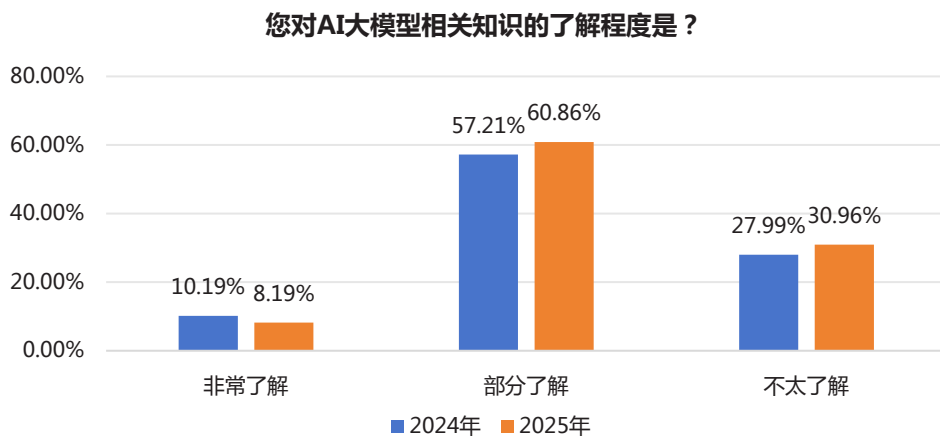


图 36 消费者对 AI 大模型相关知识的了解程度变化

2. AI 大模型使用场景与体验

本次调研发现，2024 年-2025 年消费者使用 AI 大模型的场景重心有所转移。

第一，工作场景渗透深化。2025 年，将 AI 大模型应用于“报告撰写、数据分析”等办公场景的受访者占比达 42.63%，较 2024 年在工作场景应用的 35.32%，提升 7.31 个百分点。这一数据变化表明，AI 大模型正加速向日常办公场景下沉，成为职场人士的得力助手。

第二，娱乐场景需求激增。在内容创作领域，AI 大模型的实用性获得广泛认可，其在娱乐场景的应用需求呈现爆发式增长。2025 年，用于“社交媒体文案与视频生成”的受访者占比飙升至 39.39%，相较 2024 年“文字生成图片”等应用的 19.69%实现翻倍，成为增长最为显著的应用场景。

第三，学习场景占比回调。2025 年，AI 大模型在“论文辅助、解答疑惑”等学习场景的应用占比为 45.31%，较 2024 年“ChatGPT、文心一言”等工具应用的 57.41%下降 12.1 个百分点。这一回落或归因于 2024 年学习场景需求的集中释放后趋于理性。

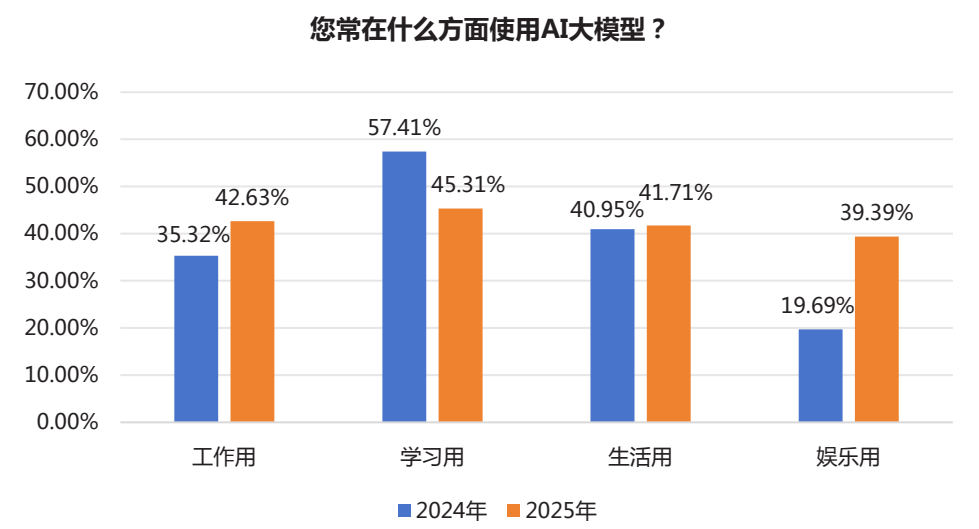


图 37 消费者 AI 大模型的主要使用场景

第四，消费者对 AI 大模型的体验感知呈现出边际效益递减趋势。2025 年，明确表示感受到 AI 大模型带来“非常显著”正向体验的受访者占比为 45.52%，相较于 2024 年“很有帮助”（42.32%）与“不可或缺”（11.42%）的合计占比略有回落。与此同时，认为 AI 大模型正向体验“一般”的消费者占比，从 2024 年的 40.56% 攀升至 2025 年的 48.52%。这一变化或可归因于消费者对 AI 技术的常态化使用，使得 AI 大模型所能带来的边际效益逐步降低。

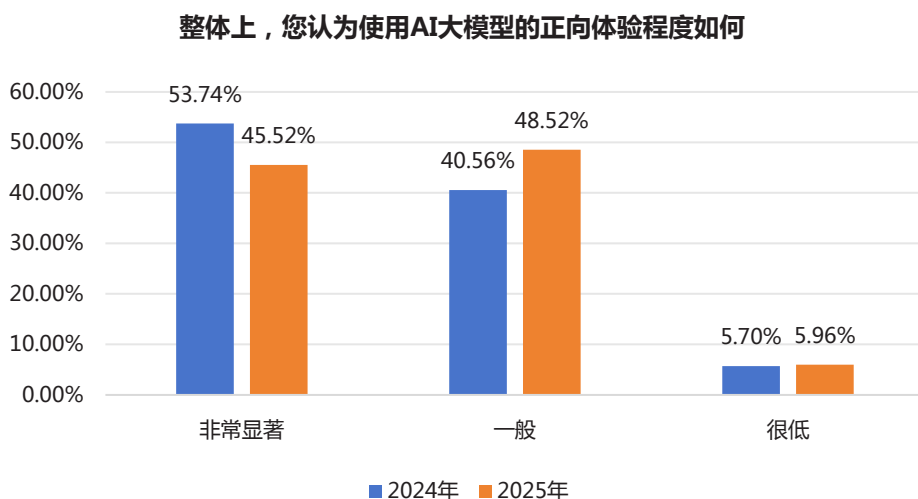


图 38 消费者对 AI 大模型的正向体验程度变化

3. 个人使用 AI 大模型遇到的问题

第一，隐私安全问题担忧但仍居首位。2025 年认为“数据隐私和安全隐患”是使用中遇到的问题的受访者占比 39.61%，较 2024 年的 52.15% 大幅下降 12.54 个百分点，但该问题仍为 2025 年消费者最关注的风险点，表明隐私安全仍是 AI 大模型推广的核心障碍。

第二，成本与门槛压力减轻。2025 年“系统的成本过高难以承担”的比例为 13.60%，较 2024 年的 26.69% 下降 13.09 个百分点；“技术复杂性高，难以学习和使用”的比例从 30.55% 降至 17.82%，表明 AI 大模型的“平民化”进程加快，使用门槛与成本大幅降低。

第三，偏见歧视与错误判断的顾虑有所降低。2025 年认为“偏见歧视和错误判断”是使用中遇到的问题的受访者占比 16.77%，较 2024 年的 24.79%下降 8.02 个百分点，这一变化表明 AI 大模型在算法公平性优化、内容错误校验机制等方面的改进已被消费者感知，但也意味着算法的客观中立性仍需持续打磨，以进一步降低此类问题的影响。

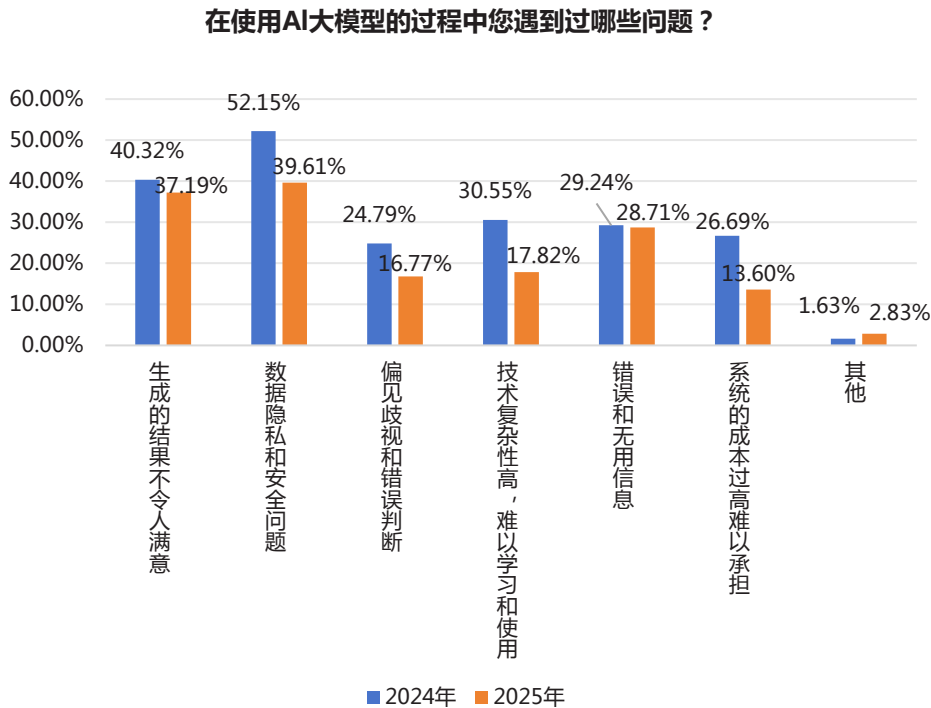


图 39 消费者使用 AI 大模型所遇问题变化

第四，自身主体性弱化焦虑凸显。2025 年新增“个人创造性下降”（29.22%）、“对 AI 产生过度情感依赖”（23.16%）、“不使用 AI 独立完成任务的能力下降”（21.60%）、“上手的门槛比较高存在智能鸿沟”（18.30%）四项风险认知，反映出消费者在依赖 AI 提升效率的同时，开始担忧自身主体性弱化：既怕依赖 AI 导致创造力、独立能力退化，也顾虑过度情感依赖稀释现实人际联结，还担心智能鸿沟加剧群体能力差距。这意味着消费者对 AI 的审视，已从“工具是否实用”转向“工具如何影响人的价值与社会关系”，这是 2024 年调研中未体现的新趋势。

在使用AI大模型的过程中您遇到过哪些问题？
(2025新增问题)

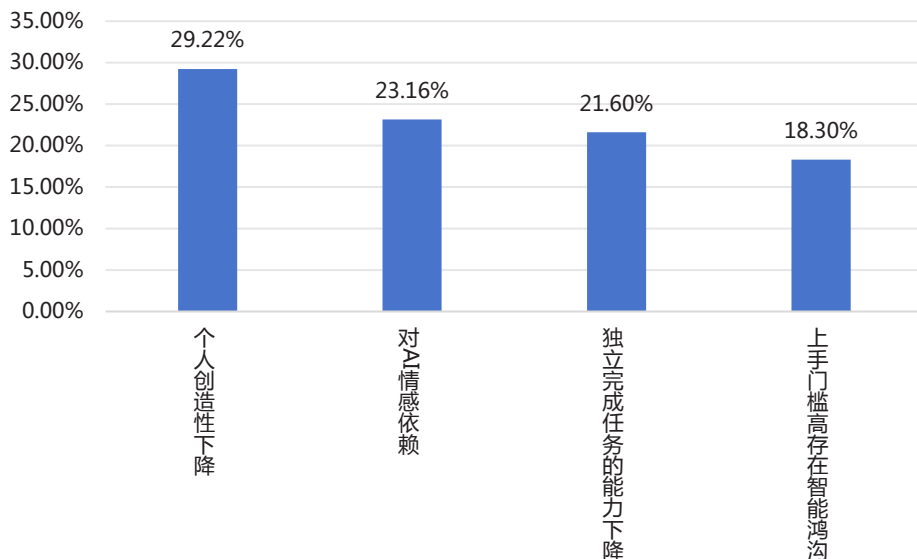


图 40 2025 年消费者在使用 AI 大模型时遇到的新问题

4. 对 AI 大模型的整体评价

第一，AI 大模型对人类的风险预测变化。2024 年消费者最关注的两项宏观风险为“个人信息不受控搜集（隐私风险）”（68.36%）与“工作岗位被 AI 替代”（60.50%），二者是当年占比均超 60% 的核心风险项。到 2025 年，这两项风险的关注占比均出现大幅回落：“个人信息不受控搜集（隐私风险）”降至 41.17%，“工作岗位被 AI 替代”降至 39.44%——尽管仍是 2025 年占比较高的两项风险，但担忧幅度已明显收缩。其中隐私风险虽然依旧是人们核心关注的风险项，但其占比的显著下降。究其原因，或源于近年来行业在数据安全与隐私保护领域的持续发力和积极探索，推动相关问题得到了一定程度的缓解与改善。对比可见，消费者对“隐私风险”与“岗位替代风险”的担忧均显著缓解，其中岗位替代风险的认知变化，或源于 2025 年“人机协同”的工作模式逐渐被消费者接受，AI 与就业市场的融合更趋成熟。而对 AI 大模型“被用作武器”的担忧由 30.72% 降至 17.97%，“机器人主导世界”的风险项由 30.14% 降至 18.54%，担

心“政府过度使用 AI 进行监管”的比例由 35.72%回落到 19.61%，2025 年认为“AI 的偏见与歧视”对人类存在风险的消费者占比为 15.55%，比 2024 年的 28.14%下降了 12.59 个百分点。

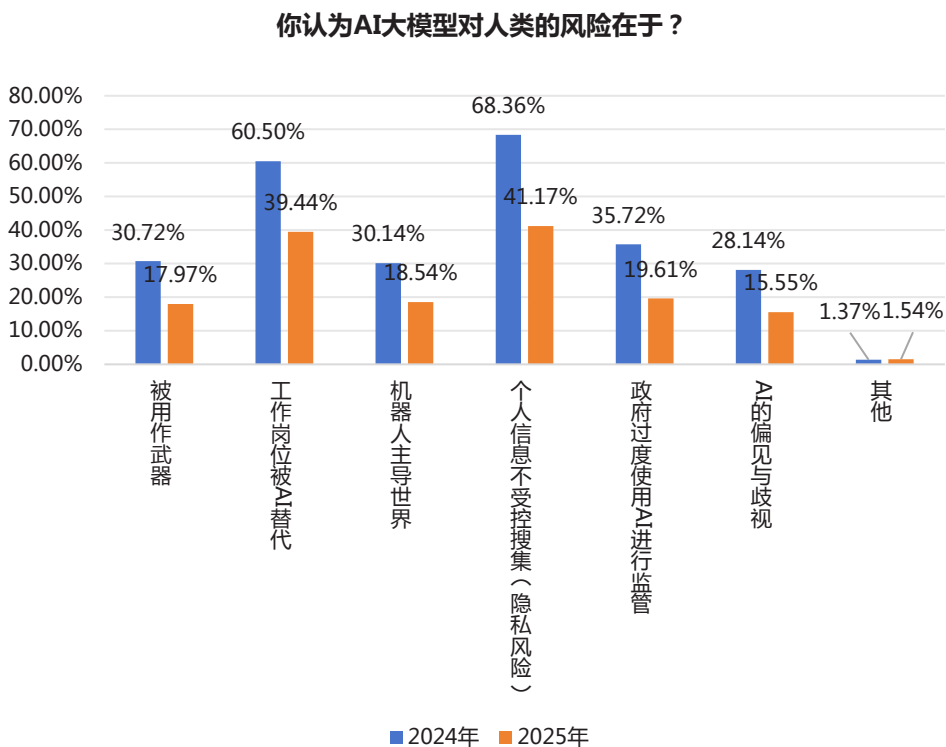


图 41 消费者关于 AI 大模型对人类风险预期的变化

与此同时，2025 年人们预测 AI 大模型对人类有新的风险。其中，“过度使用 AI，能力退化”以 45.79%的占比位居新增风险首位，随着 AI 在工作、生活中渗透加深，人们开始关注“依赖 AI 工具后，自身独立完成任务、自主思考的能力弱化”这一长期内生性风险。“AI 幻觉，生成错误信息”的占比达 36.37%，是第二受关注的新增风险。这一风险指向 AI 技术本身的局限性，反映消费者对“AI”的“内容可信度”抱有谨慎态度，对 AI“实用性”的认知加深。“对交互式 AI 产生情感依赖”占比 24.08%，体现人机交互场景普及后，消费者开始关注“对 AI 形成情感依附”的心理风险。“使用 AI 拉大数字鸿沟”占比 16.50%，指向“AI 使用门槛导致不同群体的权益/能力差距扩大”的社会分层风险。

你认为AI大模型对人类的风险在于？（2025新增风险点）

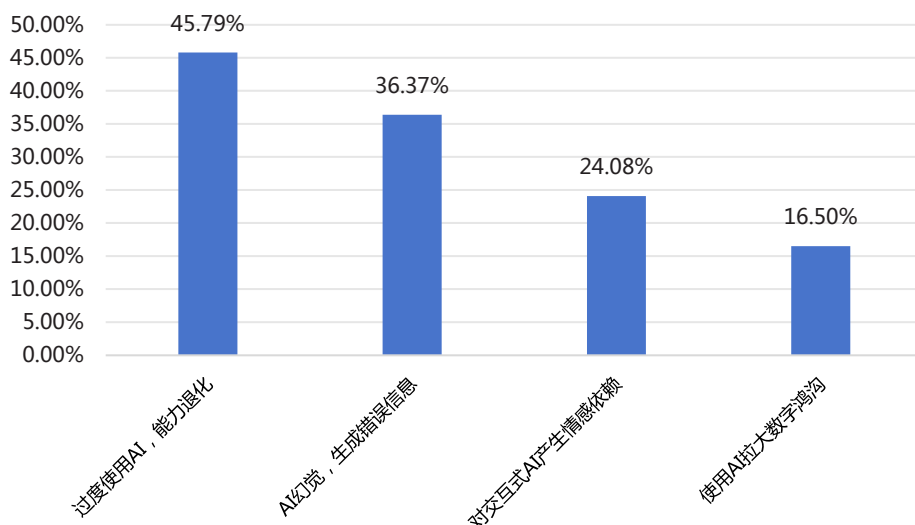


图 42 2025 消费者关于 AI 大模型对人类风险的新预期

第二，对 AI 大模型的第一印象基调仍积极。乐观态度仍是主流但占比微降：“创新”“高效”“便利”“进步”等正面词的 2024 年占比均超 60%，其中“高效”最高，达 76.36%。2025 年虽仍为认知主流，但占比普遍回落，如“高效”从 76.36%降至 64.41%，反映消费者对 AI 的积极认知仍稳固，但热情略有降温。负面关联感知普遍减弱：“威胁”“失业”“武器”等负面词的 2025 年占比均低于 2024 年，如“威胁”从 12.00%降至 6.80%，2024 年悲观总体态度占比 15.43%，2025 年降至 14.57%，说明消费者对 AI 的风险类负面联想在减少。中性关联认知小幅回落：“共存”这类中性词的认知占比从 2024 年的 38.36%降至 2025 年的 22.31%，而“朋友”等情感类关联词占比变化不大，中性总体态度占比从 2024 年的 12.51%小幅回落至 2025 年的 10.19%。整体看消费者对 AI 的情感联结类认知更趋理性。整体而言，消费者对 AI 大模型的认知仍以积极为主，只是从初期的高热情转向更冷静、更务实的关联判断。整体乐观态度的占比从 2024 年的 71.61%微升至 2025 年的 74.27%，积极认知的稳固性说明整体乐观基调仍在。

当提到AI大模型，您会想到什么词？

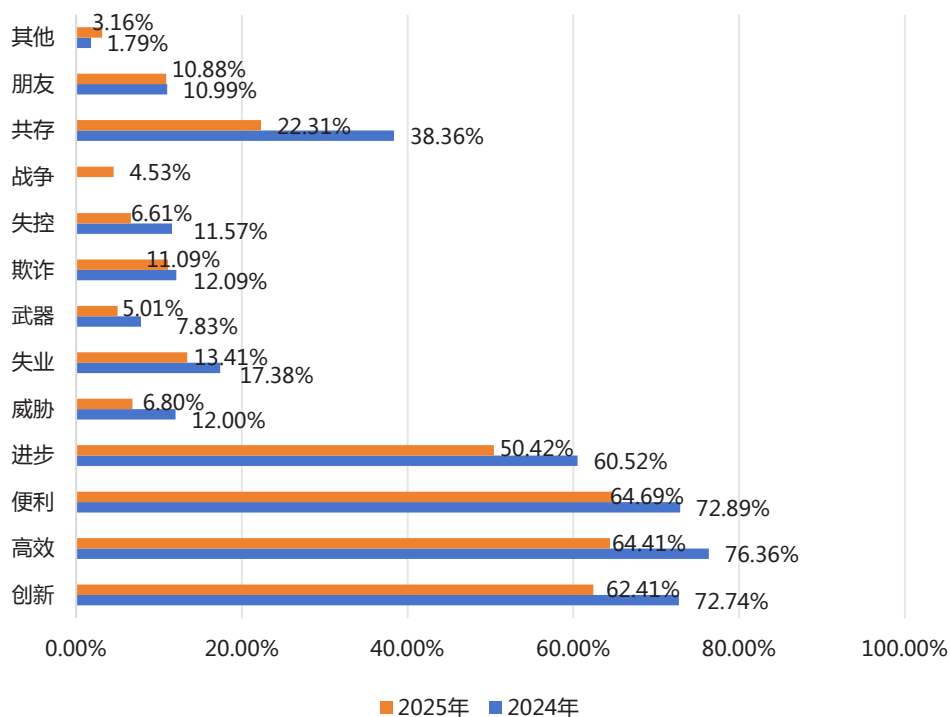


图 43 消费者对 AI 大模型的态度变化

当提到AI大模型，您会想到什么词

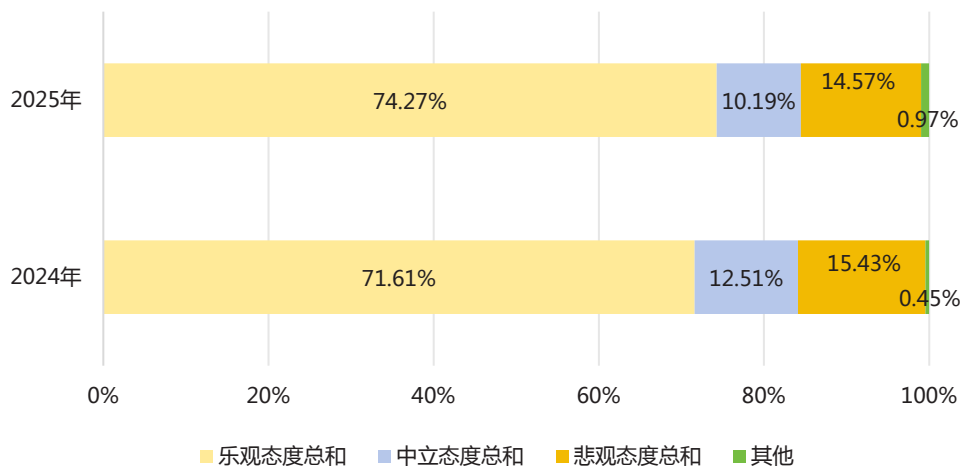


图 44 消费者对 AI 大模型的汇总态度变化

5. AI 大模型发展期待与监管认知

第一，发展关注重点转移。2025 年消费者认为 AI 大模型发展最应关注的问题是“法律法规的完善”（50.67%）和“技术的创新和突破”（47.08%），而 2024 年的重点是“法律法规的完善”（70.28%）和“商业应用场景的拓展”（51.62%）。其中，“法律法规的完善”占比下降 19.61 个百分点，“消费者数字素养的提升”占比从 15.43% 升至 28.31%，表明消费者从“依赖外部监管”向“监管与自身素养提升并重”转变，同时对“商业应用”的关注让位于“技术创新”，反映出对 AI 技术本质发展的期待增强。“伦理道德的讨论”由 2024 年的 37.31% 回落至 2025 年的 31.34%，“社会接受度的提高”由 2024 年的 25.90% 上升至 2025 年的 30.48%。这表明消费者对 AI 伦理的基础讨论需求有所降温，同时对 AI 融入社会的实际落地进程关注度提升。

您认为AI大模型在未来的发展中最应该关注的问题是什么？

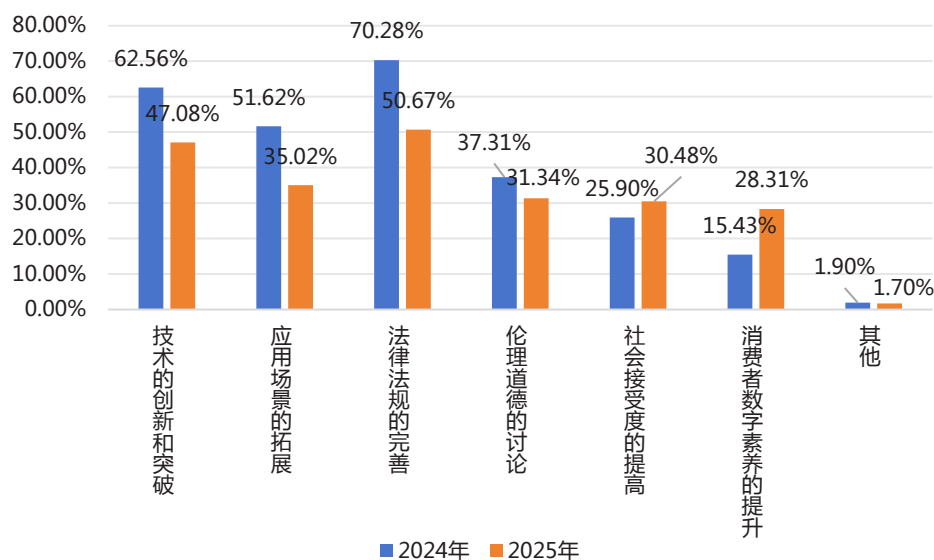


图 45 消费者认为 AI 大模型未来发展应关注的问题

第二，风险应对倾向变化。在“人工智能健康发展立法”已经进入我国立法进程的背景下，“加强 AI 监管方面的立法和司法实践”成为消费者最认可的风险应对措施（54.89%）。与之前相比，“鼓励技术创新以解决可能的风险”的认

可度从 43.34%降至 34.17%，消费者不再单纯依赖技术优化控风险；“企业建立内部伦理审查及自我监督机制”占比也从 54.69%滑落至 28.58%，反映在 AI 产业迅速发展和不断深化的同时，消费者对外部监督机制更加重视。但这不能解释为对 AI “严刑峻罚”的偏好，事实上，“加大对企业的监管和处罚力度”的占比从 44.67%降至 32.52%，“限制 AI 大模型运用的空间和领域”的比例从 45.40%降至 25.72%，这些均反映出消费者关切的是风险治理的有效性，而不是严厉性。因此，如何通过引入外部监督机制和第三方伦理审查，提升多元治理的效果，已经成为未来重要的发展方向。

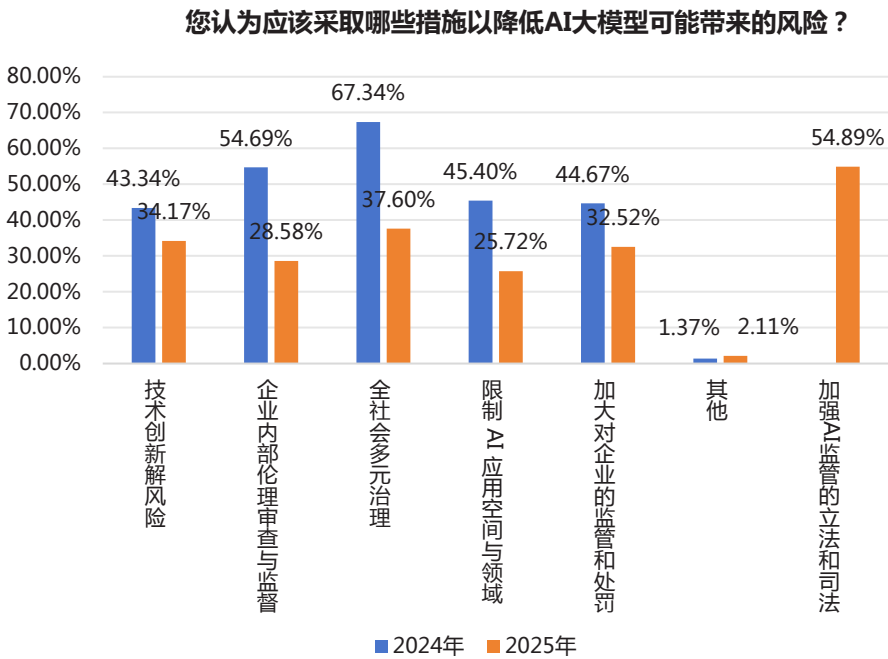


图 46 消费者认为降低 AI 大模型风险的措施

第三，政策认知水平下滑。2025 年“亲自阅读过”国家有关 AI 大模型安全规定的受访者仅占 5.74%，较 2024 年的 17.77%大幅下降 12.03 个百分点；“不了解”的比例从 17.04%升至 33.55%，翻倍有余。这一变化与 AI 认知普及度提升的趋势相悖。2025 年相关政策更新频率加快，但政策宣传渠道未有效触达普通消费者，导致消费者对政策的认知滞后于技术发展。

您是否了解国家有关算法安全、AI大模型安全的相关规定和公民自身所拥有的合法权利？

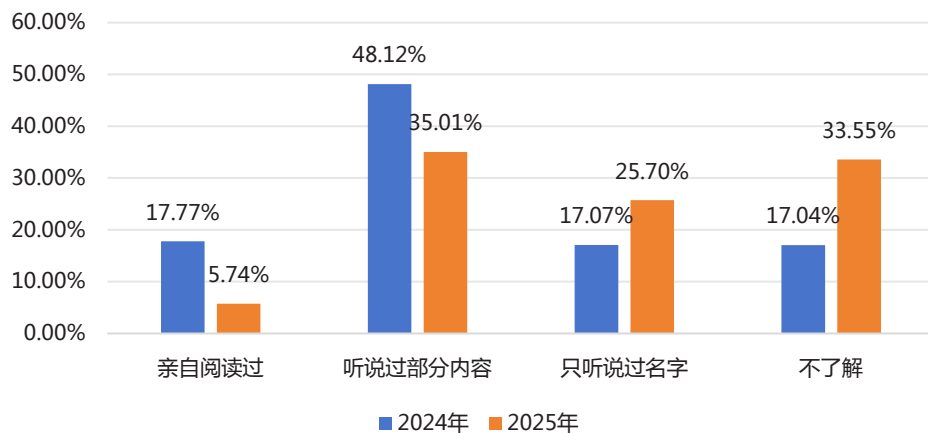


图 47 消费者对相关法规和合法权利的了解程度变化

六、交叉分组分析

课题组基于被调查者的年龄、学历、性别等基本信息，对调查结果进行了交叉分组分析，得到了一些有趣的差异性结论。

1. 不同年龄对 AI 大模型的感知差异

第一，年轻人对 AI 大模型的了解和使用程度超过了老年人。随着年龄的增长，各年龄段受访者对 AI 大模型的了解和使用程度呈现负相关关系。18 至 25 岁的年轻受访者使用过 AI 大模型的比例为 66.06%，而老年受访者中这一比例不超过 50%。这一调查结果相较此前发生了较大变化，2024 年研究结果显示老年人对大模型的认知度高于年轻人。

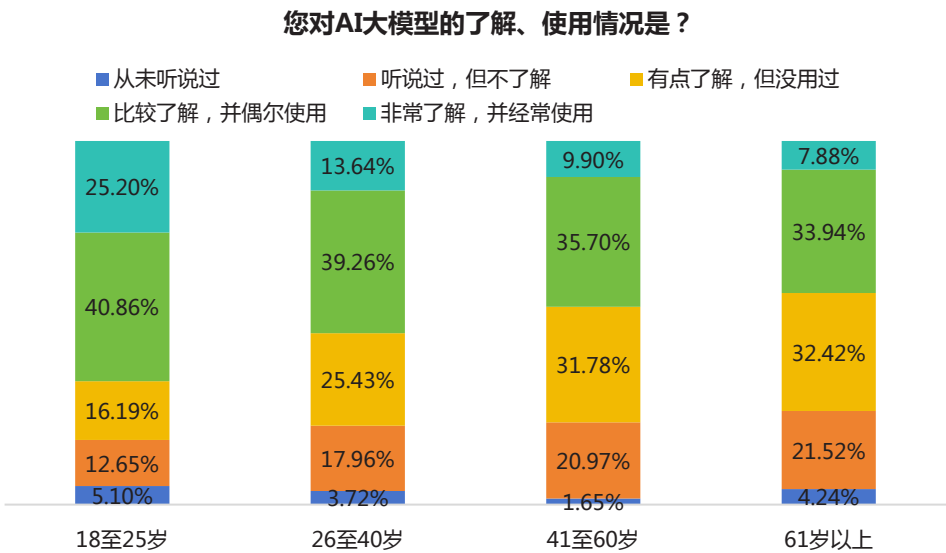


图 48 对 AI 大模型的了解和使用程度

第二，老年受访者更多在生活领域使用大模型，并对 AI 提供的健康服务信息展现更高信任。各个年龄段使用 AI 大模型的场景结构基本相似，18 至 25 岁受访者学习用、61 岁以上的老年受访者生活用 AI 大模型的比例明显高于其他年龄段。在涉及专业问题参考和复杂事务处理的场景中，受访者对 AI 提供的健康服务信息的信任程度随年龄增长而上升。

您常在什么方面使用AI大模型？（多选）

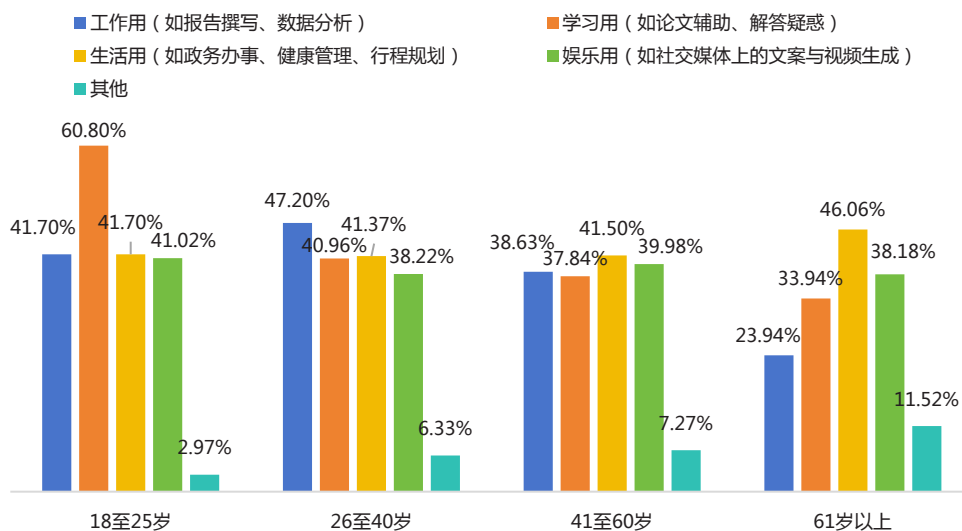


图 49 使用 AI 大模型的领域

您对哪些场景中的AI使用比较信任？（多选）

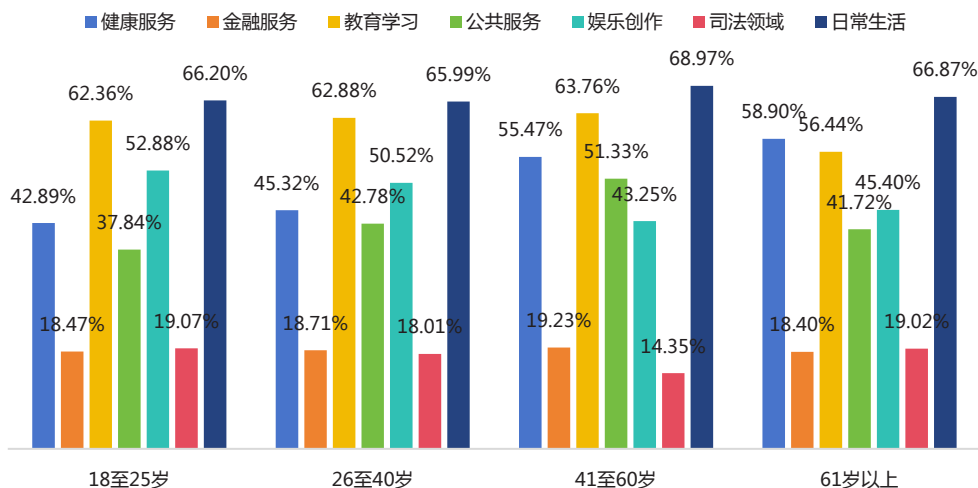


图 50 对 AI 大模型应用场景的信任程度

第三，不同年龄段对 AI 大模型的印象存在差异。在正向认知上，各年龄段均普遍认可 AI 的“创新、高效、进步”属性，而 60 岁以下群体整体认可度高于 60 岁以上受访者。在负面认知上，18 至 25 岁受访者对 AI 的“风险感知”最突出——其选择“失业”（16.81%）、“威胁”（8.85%）等负面关键词的占比，明显高于其他年龄段，且随年龄增长，这类风险类选项的占比逐步降低，61 岁以上群体选择“失业”的比例仅为 9.39%，选择“威胁”的仅为 3.03%。

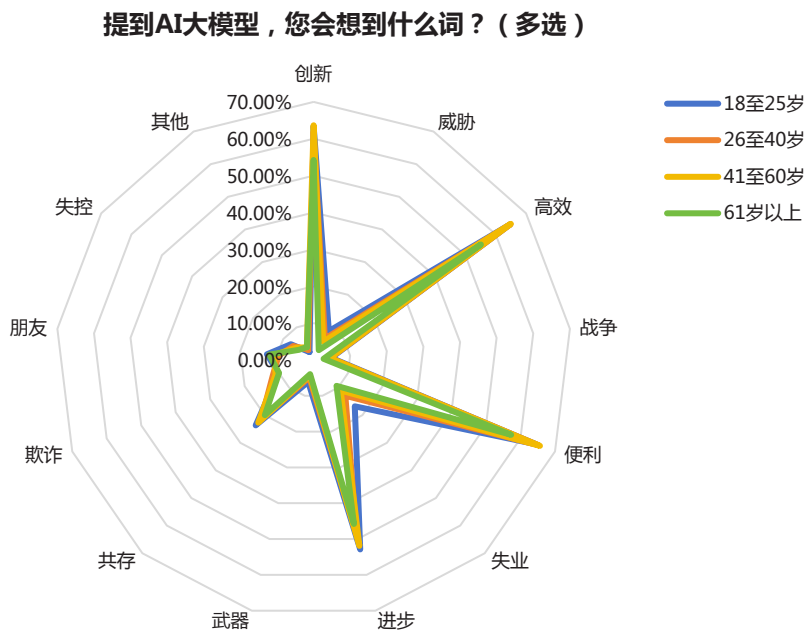


图 51 对 AI 大模型的印象

第四，不同年龄段对 AI 大模型的问题和风险的感知不同。18 至 25 岁的年轻人比其他年龄段的受访者在 AI 大模型生成结果的质量方面遇到更多问题，其他年龄段受访者则更关注数据隐私和安全问题。在对 AI 大模型的风险的感知上，60 岁以下的受访者对 AI 带来能力退化和工作岗位替代的担忧与年龄呈现显著的正相关。相较其他群体，60 岁以上受访者更担心 AI 大模型被用作武器。

在使用AI大模型的过程中您遇到过哪些问题？（多选）

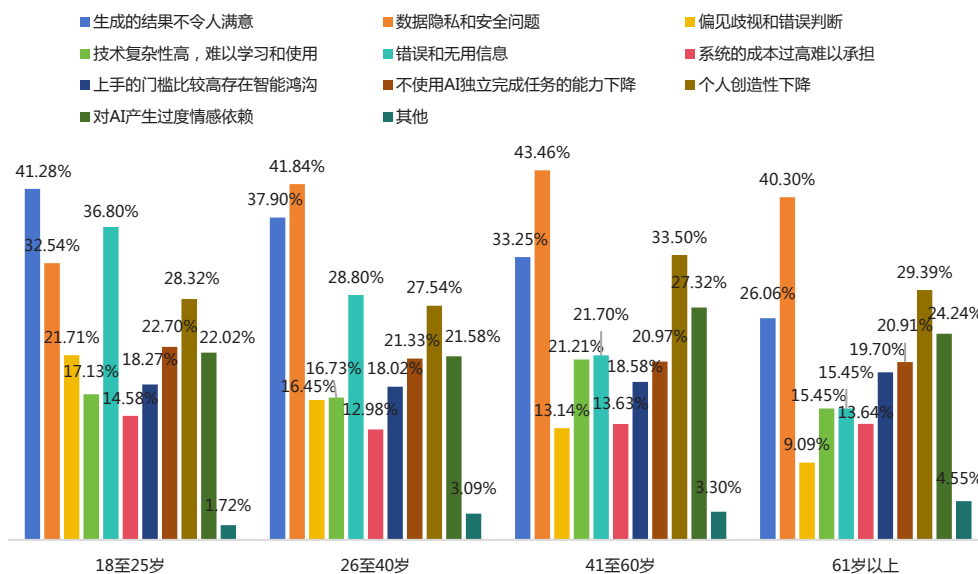


图 52 使用 AI 大模型遇到的问题

您认为AI大模型对人类的风险在于？（多选）

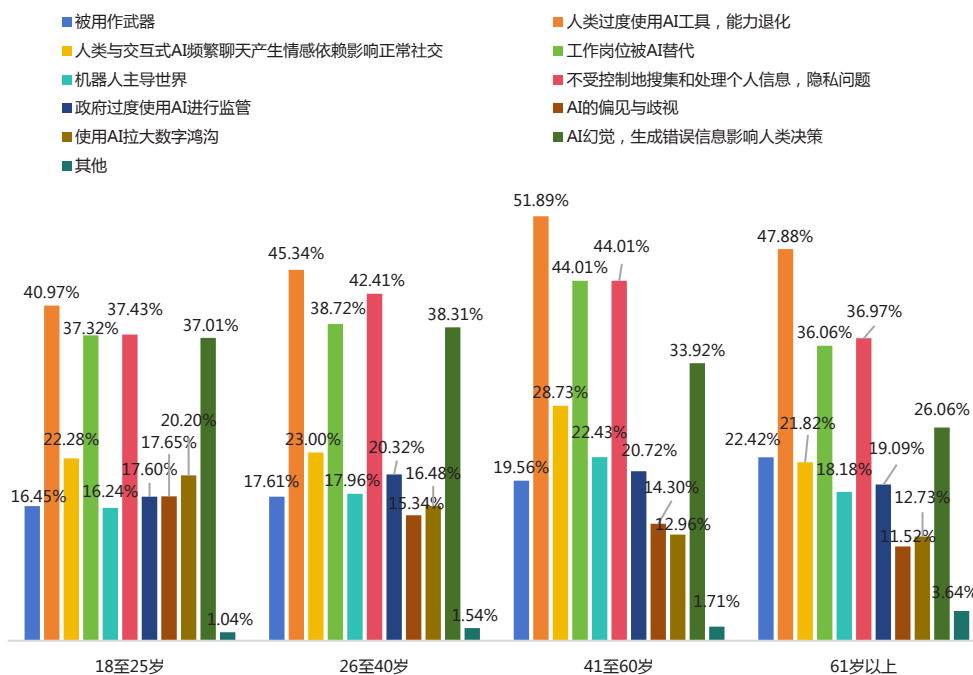


图 53 关于 AI 大模型的风险感知

2. 不同学历对 AI 大模型的感知差异

第一，不同学历受访者对 AI 大模型实际体验状况存在差异。在正向体验方面，学历越高对 AI 大模型正向评价越高，研究生及以上学历受访者认为正向体验非常显著的比例为 57.17%，高中及以下受访者则为 40.23%。在负向体验方面，高学历受访者更关注生成结果的准确度和质量，“生成的结果不令人满意”“偏见歧视和错误判断”“错误和无用信息”的反馈比例随着学历的提升而提高；低学历受访者更集中于 AI 大模型使用的门槛，专科及以下学历认为“上手的门槛比较高存在智能鸿沟”“技术复杂性高，难以学习和使用”的比例分别为 20.40%、20.34%，高于本科及以上学历的 16.87%、16.37%。此外，对 AI 产生过度情感依赖的情况也与学历呈现负相关。

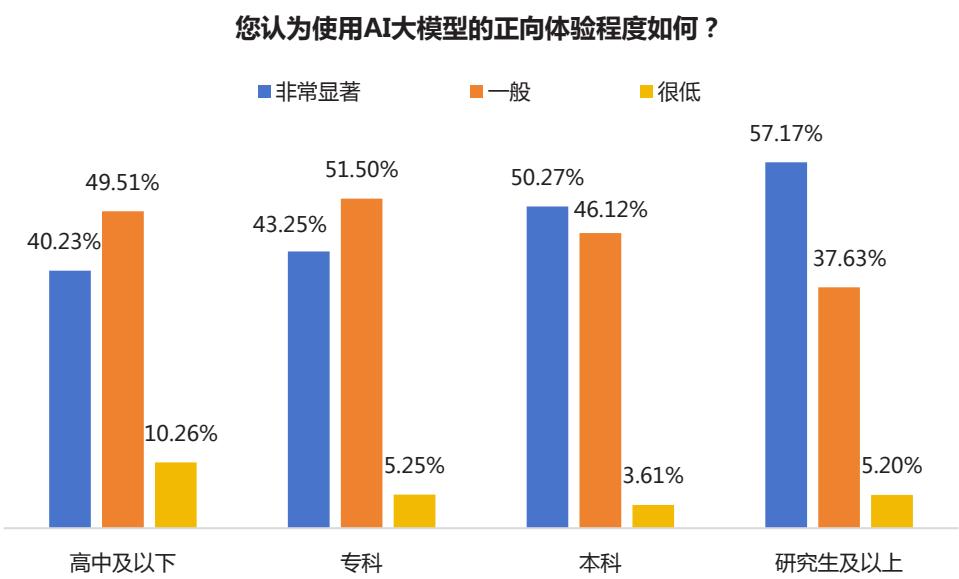


图 54 对 AI 大模型的正向体验认知

在使用AI大模型的过程中您遇到过哪些问题？（多选）

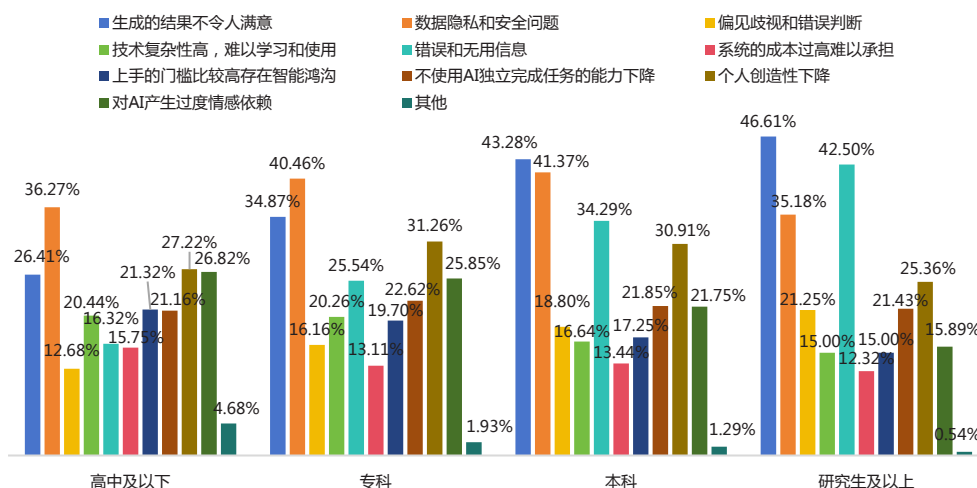


图 55 对 AI 大模型的负向体验认知

第二，不同学历对 AI 大模型的印象存在差异。在正向认知上，学历越高，对“高效”“便利”的认可度越强，本科及以上受访者比例达 70.88%和 68.54%，高于专科及以下受访者的 59.30%和 61.56%。在风险认知上，对“失业”“威胁”的选择占比随学历提升而上升，研究生及以上“失业”16.31%、“威胁”8.78%，均高于高中及以下组。

提到AI大模型，您会想到什么词？（多选）

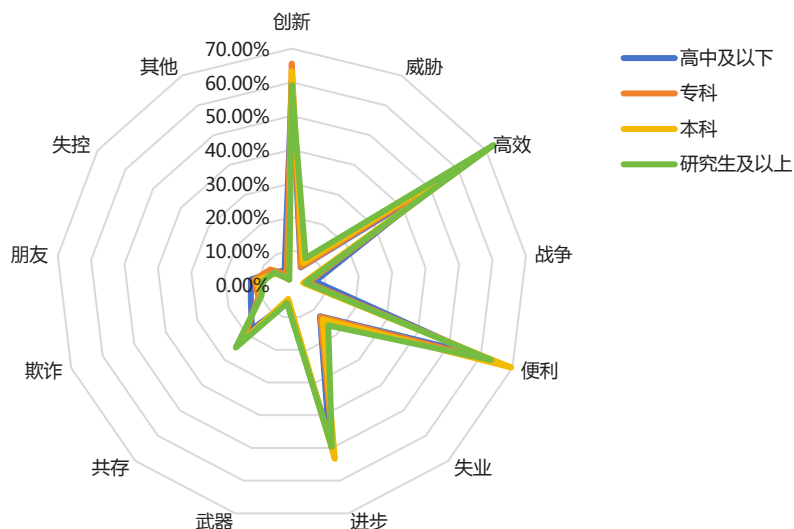


图 56 对 AI 大模型的印象

第三，不同学历对 AI 伦理原则的关注重点存在差异。低学历受访者更关注 AI 伦理的公平性，专科及以下学历对公平与非歧视和包容性与公共利益关注达 45.23%和 16.79%。高学历受访者更关注 AI 的可理解性与可控性，研究生及以上学历对透明性与可解释性的关注度达到 49.29%，对人类监督与控制的关注度为 21.43%。

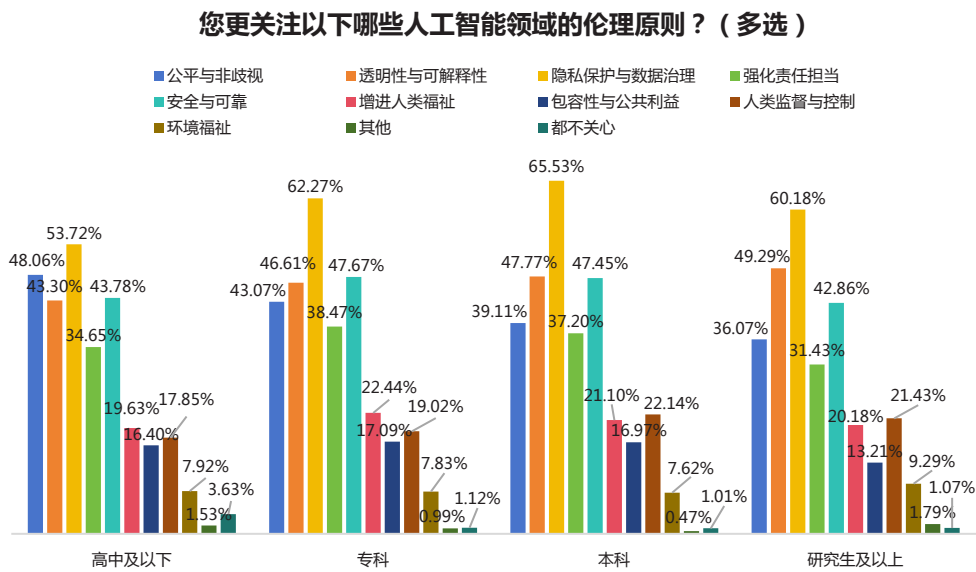


图 57 对 AI 伦理原则的关注

3. 不同性别对 AI 大模型的感知差异

第一，不同性别受访者使用 AI 大模型的场景存在差异。56.11%的女性受访者表示使用过 AI 大模型，高于男性受访者的 52.60%。在具体使用场景方面，女性使用最多的场景是学习（48.87%），男性使用最多的场景是生活（43.02%）。在娱乐场景，女性使用 AI 大模型的比例略高于男性，达到 40.02%。在工作场景，男女 AI 大模型使用差异不显著。

您常在什么方面使用AI大模型？（多选）

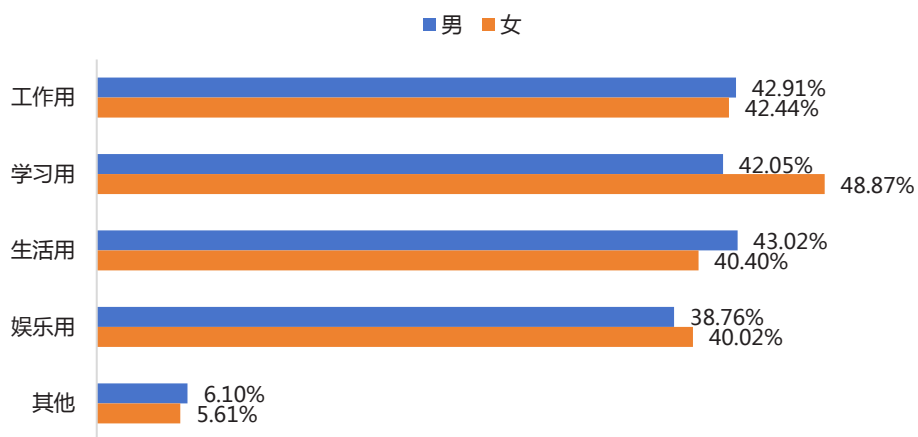


图 58 AI 大模型应用场景

第二，男女受访者使用 AI 大模型的感受和信任存在差异。女性对 AI 大模型的使用整体评价较高，47.53%的女性受访者认为 AI 大模型的正向体验非常显著。但是女性受访者对 AI 大模型的信任程度为 33.97%，低于男性整体水平（36.05%）。女性受访者相较男性更担心数据隐私和安全性问题，更担忧生成信息的准确性。

您认为使用AI大模型的正向体验程度如何？

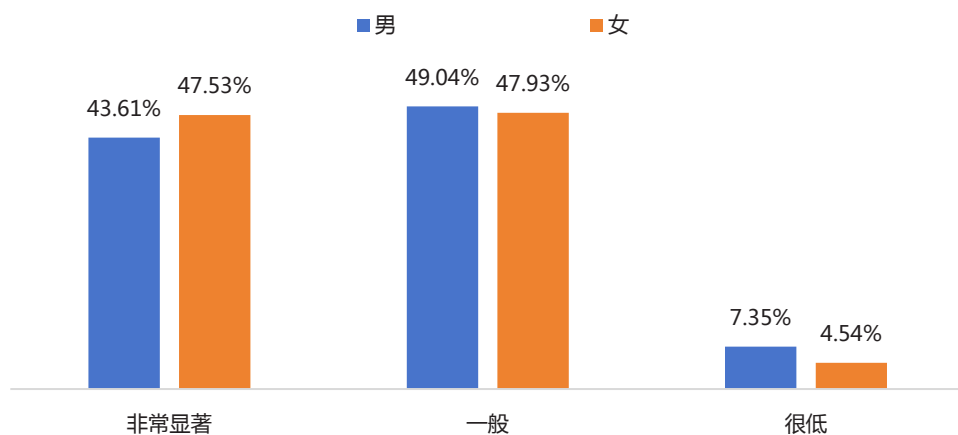


图 59 AI 大模型正向体验程度

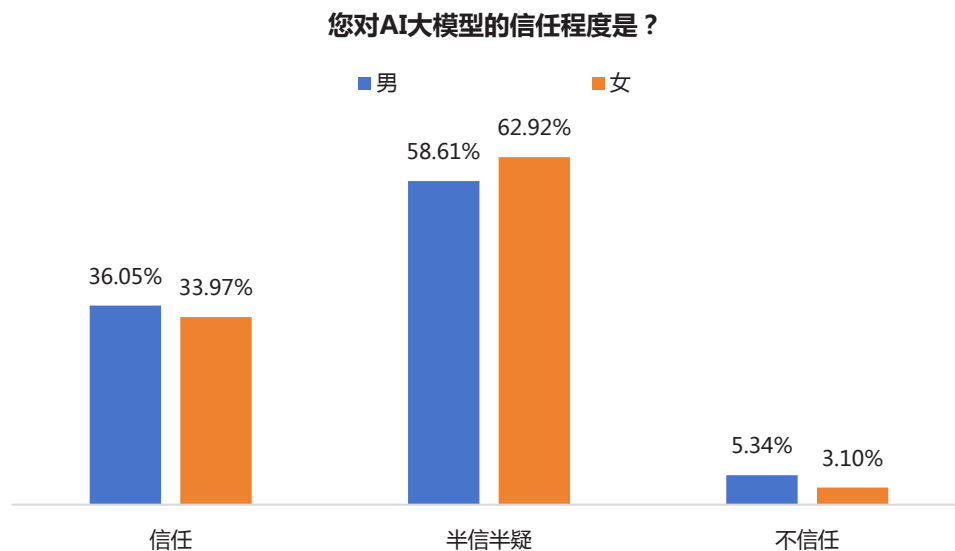


图 60 AI 大模型信任程度

第三，男女受访者对 AI 大模型的印象与伦理关注存在差异。在正向认识上，无论男女选择比例最高的前三项均为“高效”“便利”“创新”，且女性在这三项上的比例均略高于男性，尤其在“便利”（女性 67.92%，男性 61.51%）和“高效”（女性 66.84%，男性 62.16%）上更为明显。在负面认识上，男性在“战争”“武器”等具有攻击性的选项上比例高于女性；女性在“失业”“欺诈”等社会经济与安全风险上的担忧略高于男性。在伦理原则上，女性更关注隐私、安全与监督，对 AI 的风险防控类伦理诉求更突出；男性更关注公平与透明，对 AI 的规则公平性诉求更强。

提到AI大模型，您会想到什么词？（多选）

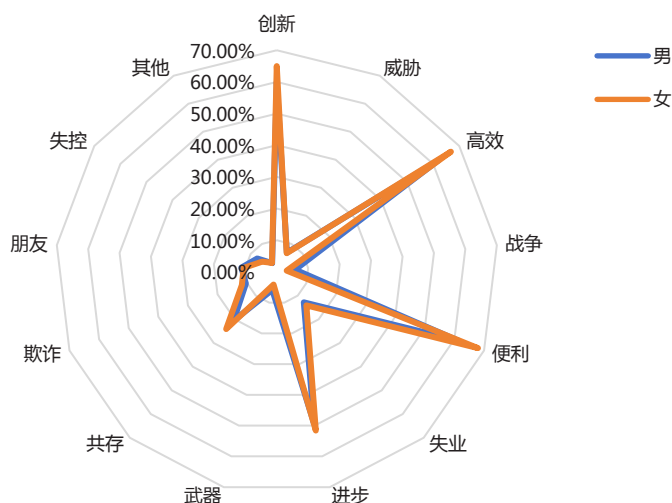


图 61 对 AI 大模型的印象

您更关注以下哪些人工智能领域的伦理原则？（多选）

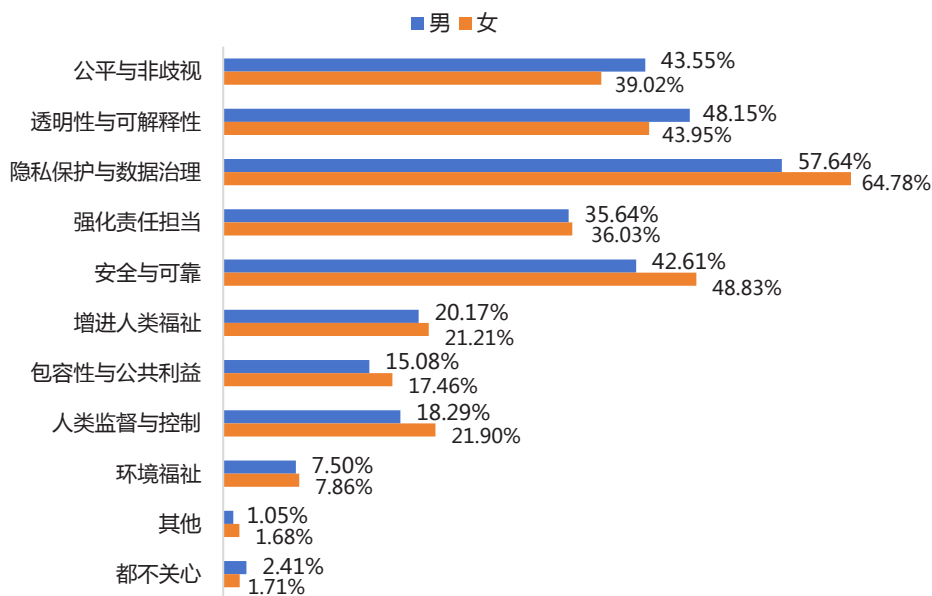


图 62 对 AI 伦理原则的关注

第四，不同性别受访者对“AI 生成”标识的反应不同。女性对 AI 生成内容谨慎程度更高，选择“理性谨慎，核实后使用”的占比（61.13%）显著高于男性

(53.93%)。男性选择“基本无视，习惯照旧”的占比（19.31%）高于女性（13.93%）。

面对“AI生成”标识，您的第一反应是什么？

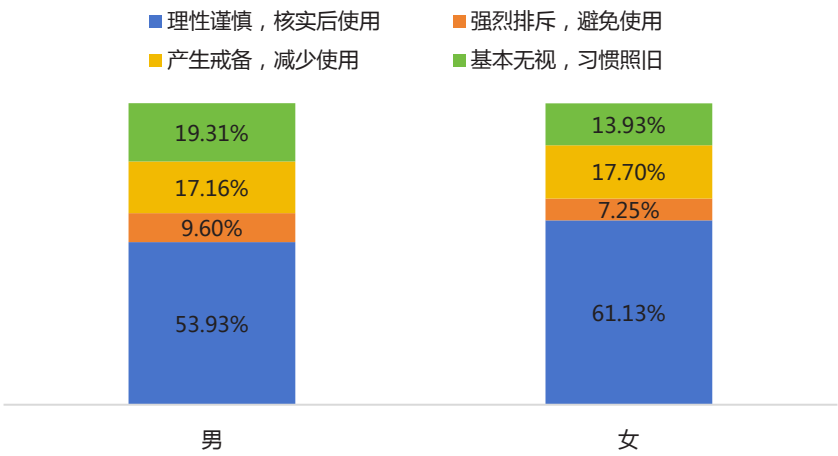


图 63 对“AI 生成”标识的反应

七、人工智能健康发展倡议

结合本文调研结果可见，当前公众对 AI 大模型已形成较高接受度，应用场景持续拓展，在工作、学习和生活中的实用价值不断增强。但与此同时，公众对 AI 大模型的深度理解仍较为有限，信任呈现明显的情境化特征，对隐私安全、结果质量、能力退化、AI 幻觉以及新形态人工智能潜在风险的关注亦在持续上升。公众对完善立法和司法实践、提升治理效能普遍具有较强诉求，但对相关政策规定、自身合法权利及维权路径的认知仍显不足。因此推动人工智能健康发展，需要在技术创新、制度完善、行业自律与公众参与之间形成更加协调的发展格局，推动公众与 AI 大模型的关系从“工具采纳”进一步走向“生态共建”。基于上述发现，本文提出以下倡议：

第一，面向监管层面，应坚持规范与发展并重，进一步夯实人工智能健康发展的制度基础。围绕公众高度关注的隐私保护与数据治理、安全与可靠、透明性与可解释性等重点议题，有必要持续完善法律法规及相关配套规则，强化立法与司法实践，不断提升制度供给的明确性、稳定性与可预期性。在此基础上，还应根据不同应用场景、风险等级和技术形态，推进更为精准的分类分级治理，尤其是在专业决策、高严谨性场景以及智能体、具身智能等新形态发展过程中，提前完善风险识别、责任界定和救济机制。此外，还需加强政策宣传与权利普及，推动相关规定、合法权利和维权渠道更加有效地触达普通用户，以缓解当前“高诉求、低认知”的现实矛盾。

第二，面向行业和企业层面，应坚持技术创新与责任落实并重，切实提升 AI 大模型产品和服务的可信度与可用性。调研显示，公众对 AI 大模型的认可主要建立在效率提升、问题解决和创意激发等正向体验之上，但对结果不准确、信息失真、隐私泄露和主体性弱化等问题仍保持较高警惕。基于此等现象，行业主体应把提升内容质量、优化结果准确性、降低 AI 幻觉和偏差风险作为持续改进的重要方向，并将隐私保护与数据治理嵌入产品设计、训练、部署和运营全过程，

强化安全底线和责任担当。另外还应健全内部伦理审查、风险评估、投诉反馈和纠错响应机制，完善“AI 生成”标识、算法说明、用户提示等必要安排，帮助用户在使用中形成更为合理的信任边界。对于智能体、具身智能等前沿应用，则应坚持审慎推进、场景验证与循序拓展，避免脱离实际需求和安全约束的盲目扩张。

第三，面向用户层面，应持续提升数字素养、权利意识与风险防范能力，以更加理性、谨慎地方式参与人工智能应用。调研结果表明，多数用户已将 AI 大模型视为提升效率的重要工具，但在技术原理、能力边界、合法权利和风险识别等方面仍存在明显短板。对此，用户既应主动学习基础知识，增强对 AI 大模型的认知和理解，也应更加重视个人隐私和数据安全保护，合理把握使用频率与使用边界，避免因过度依赖而导致独立完成任务能力、创造性和自主判断能力下降。在面对 AI 生成内容时也应持续保持理性谨慎、核实后使用的态度；在遭遇偏见、歧视、侵权或其他不利影响时，要积极通过平台反馈、投诉举报等正式渠道依法维护自身权益。只有不断提升主体意识和参与能力，公众才能更好地在人工智能时代实现技术使用、风险防范与权利保护的统一。

总体而言，人工智能健康发展是一个需要监管、行业与用户共同参与的系统性过程。未来应在规范中推动创新，在发展中保障权益，在协同中增进信任，逐步构建一个技术可信、权利可及、发展可持续的人工智能社会。



附件：AI 大模型调查问卷

AI 大模型调查问卷

尊敬的先生/女士：

您好！

AI 大模型（如通义千问、DeepSeek、豆包等）已经在我们的生活和工作中得到广泛应用。从撰写文案、编写代码、答疑解惑，到激发创意、辅助决策，AI 的潜力正被人类源源不断地挖掘和利用。

我们课题组试图了解用户对于 AI 大模型的认识和看法，希望得到您的配合。本问卷搜集的数据将只用于课题研究，无任何商业目的或信息泄露的可能，非常感谢您的参与。

对外经济贸易大学数字经济与法律创新研究中心
中国人民大学数字经济研究中心
联合课题组

认知&使用

1. 您对 AI 大模型（如文字生成类、图像生成类、生活助手类等）的了解和使用情况是？【单选题】

- A. 从未听说过
- B. 听说过，但不了解
- C. 有点了解，但没用过
- D. 比较了解，并偶尔使用
- E. 非常了解，并经常使用

2. 您常在什么方面使用 AI 大模型？【最多选 3 项】

- A. 工作用（如报告撰写、数据分析）
- B. 学习用（如论文辅助、解答疑惑）
- C. 生活用（如政务办事、健康管理、行程规划）
- D. 娱乐用（如社交媒体上的文案与视频生成）
- E. 其他

3. 您对 AI 大模型相关知识（如功能、类型、工作原理、发展状况）的了解程度是？【单选题】

- A. 非常了解
- B. 部分了解
- C. 不太了解

4. 您觉得自己在使用 AI 大模型方面的熟练程度如何？【单选题】

- A. 比较熟练
- B. 一般熟练
- C. 不太熟练

5. 您是否接受过 AI 大模型相关的使用培训？【单选题】

- A. 接受过正式培训（如课程、工作培训）
- B. 接受过非正式培训（如自学教程、他人指导）
- C. 未接受过任何培训

6. 您如何描述自己学习 AI 大模型相关知识和技能的方式？【单选题】

- A. 完全被动（很少主动学习，仅参与必要培训或从信息推送中偶然获取）
- B. 被动使用（仅在遇到具体问题时，搜索和学习相关知识以完成任务）
- C. 主动探索（会主动关注科技资讯，并积极尝试体验新的 AI 工具和功能）
- D. 系统关注（通过固定渠道定期关注技术动态，系统更新自己的知识库）
- E. 深入研究（通过阅读技术报告、论文或参与专业课程等，进行专业研究）

信任&体验&风险

7. 当提到 AI 大模型，您会想到什么词？【最多选 5 项】

- A. 创新
- B. 威胁
- C. 高效
- D. 战争
- E. 便利
- F. 失业
- G. 进步
- H. 武器
- I. 共存
- J. 欺诈
- K. 朋友
- L. 失控
- M. 其他

8. 使用 AI 大模型后，您在哪些方面获得了正向体验【最多选 3 项】？

- A. 提高工作/学习效率（如快速整理资料、写报告）
- B. 激发创意（如生成文案、设计思路）
- C. 解决生活问题（如查询攻略、解答常识）
- D. 陪伴互动（如聊天、趣味问答）
- E. 其他

9. 整体上，你认为使用 AI 大模型的正向体验程度如何？【单选题】

- A. 非常显著
- B. 一般
- C. 很低

10. 与最初接触 AI 大模型时相比，您认为当前 AI 大模型的能力是否有进步？

【单选题】

- A. 有明显进步
- B. 有轻微进步
- C. 没进步

11. 您愿意信任 AI 大模型处理哪些任务【最多选 5 项】？

- A. 基础信息查询（如天气、常识）
- B. 辅助创作（如写初稿、列提纲）
- C. 简单事务处理（如制定日程、翻译）
- D. 专业问题参考（如工作建议、健康诊断、投资建议）
- E. 复杂事务处理（如制定差旅方案、并预定酒店、机票）
- F. 都不信任

如果没有选择 D、E，跳转 13 题

12. 在专业问题参考和复杂事务处理中，您对哪些场景中的 AI 使用比较信任【最多选 4 项】

- A. 健康服务（如疾病咨询、诊断辅助、健康管理）
- B. 金融服务（如理财咨询、风险评估、信贷审核）
- C. 教育学习（如作业辅助、技能培训、知识科普）
- D. 公共服务（如政务办理、交通调度、便民咨询）
- E. 娱乐创作（如内容生成、创意辅助、互动娱乐）
- F. 司法领域（如案件分析、证据梳理、量刑参考）
- E. 日常生活（如购物推荐、行程规划、语言翻译）

13. 整体上，您对 AI 大模型的信任程度是【单选题】

- A. 信任

B.半信半疑

C.不信任

14. 在使用 AI 大模型的过程中，您遇到过哪些问题？【最多选 5 项】

A.生成的结果不令人满意

B.数据隐私和安全问题

C.偏见歧视和错误判断

D.技术复杂性高，难以学习和使用

E.错误和无用信息

F.系统的成本过高，难以承担

G.上手的门槛比较高，存在“智能鸿沟”

H.不使用 AI 独立完成任务的能力下降

I.个人创造性下降

J.对 AI 产生过度情感依赖

K.其他

15. 当 AI 被广泛使用，对于此，您最常产生的情绪是？【最多选 5 项】

A.兴奋感（体验新技术）

B.新鲜感（赶上现在潮流）

C.成就感（快速完成复杂任务）

D.安全感（解决紧急问题）

E.挑战性（赋能感，能做以前不会的事，如画画/写代码）

F.迷茫感（无所适从）

G.焦虑（担心依赖 AI 后自身能力退化）

H.挫败感（AI 无法理解需求）

I.恐慌（担心被 AI 取代工作）

J.不确定感（无法判断 AI 结果的可靠性）

K.无明显情绪

16. 您认为 AI 大模型对人类的风险在于？【最多选 5 项】

A.被用作武器

B.人类过度使用 AI 工具，能力退化

C.人类与交互式 AI 频繁聊天产生情感依赖，影响正常社交

- D.工作岗位被 AI 替代
- E.机器人主导世界
- F.不受控制地搜集和处理个人信息，隐私问题
- G.政府过度使用 AI 进行监管
- H.AI 的偏见与歧视
- I.使用 AI 拉大数字鸿沟
- J.AI 幻觉，生成错误信息，影响人类决策
- K.其他

智能体

17. 您对“智能体”（AI Agent）的了解和使用情况是？（智能体指能自主理解用户需求、规划任务步骤并执行复杂操作的 AI 程序，如自动订票并规划行程的旅行助手、独立完成报告撰写的办公助理等）【单选题】

- A. 非常了解
- B. 有点了解
- C. 完全不了解

18. 您对“具身智能”（Embodied AI）的了解和使用情况是？（具身智能指能与物理世界交互、具备感知和行动能力的 AI 系统，如自动驾驶汽车、家政机器人、工业机器人等）【单选题】

- A. 非常了解
- B. 有点了解
- C. 完全不了解

如果 14、15 两题都选择 C，那么跳转第 19 题；否则继续

19. 您对“智能体”或者“具身智能”最大的期待有哪些？【多选题】

- A.减少重复劳动、提高工作效率
- B.激发个人的创造力
- C.提高生活的便利度
- D.替代人类从事高风险活动
- E.提供感情陪伴

F.其他

20. 您对“智能体”或者“具身智能”最大的担心有哪些？【多选题】

- A. 技术脱离人类的掌控
- B. 直接威胁人身安全和财产安全
- C. 导致大规模的失业
- D. 放大单个智能体的固有缺陷、形成系统性风险
- E. 事故责任难以界定
- F. 其他

21. 您认为是否要大力发展“智能体”“具身智能”等人工智能新形态？【单选题】

- A. 需要（能极大提升生产力和生活质量）
- B. 不确定（影响尚不明确，保持观望）
- C. 不太需要（现有技术已足够，或担心相关风险）

伦理&治理

22. 您认为 AI 大模型在未来的发展中，最应该关注的问题是什么？【最多选 3 项】

- A. 技术的创新和突破
- B. 服务应用的拓展
- C. 法律法规的完善
- D. 伦理道德的讨论
- E. 社会接受度的提高
- F. 消费者数字素养的提升
- G. 其他

23. 为避免 AI 大模型应用带来的不利影响，您个人通常会？【最多选 3 项】

- A. 尽可能减少 AI 大模型产品的使用
- B. 主动学习基础知识，增强对 AI 大模型的认知和理解
- C. 注意保护个人隐私和数据安全
- D. 加强对自身能力的培养
- E. 其他

24. 您认为应该采取哪些措施以降低 AI 大模型可能带来的风险？【最多选 3 项】

- A. 鼓励技术创新以解决可能的风险
- B. 企业建立内部伦理审查及自我监督机制
- C. 加强 AI 监管方面的立法和司法实践
- D. 限制 AI 大模型运用的空间和领域
- E. 加大对企业的监管和处罚力度
- F. 构建全社会开放、合作、协商的多元治理机制
- G. 其他

25. 您更关注以下哪些人工智能领域的伦理原则？【最多选 5 项】

- A. 公平与非歧视（AI 公平地对待所有人，避免偏见和歧视）
- B. 透明性与可解释性（AI 决策过程透明且可被人类理解）
- C. 隐私保护与数据治理（尊重和保护个人隐私，数据处理安全合规）
- D. 强化责任担当（坚持人类是最终责任主体，建立人工智能问责机制）
- E. 安全与可靠（AI 安全、稳健，能抵御恶意攻击和意外故障）
- F. 增进人类福祉（以人为本，遵循人类共同价值观）
- G. 包容性与公共利益（AI 应惠及全人类，关注弱势群体）
- H. 人类监督与控制（人类始终拥有监督权和控制权）
- I. 环境福祉（AI 模型消耗算力和能源，应考虑其环境影响）
- H. 其他
- I. 都不关心

26. 总体而言，您认为当前对 AI 发展进行规范和治理的紧迫程度如何？【单选题】

- A. 非常紧迫
- B. 比较紧迫
- C. 不太紧迫

27. 您是否了解国家有关算法安全、AI 大模型安全的相关规定和公民自身所拥有的合法权利？（如《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》、《互联网信息服务算法推荐管理规定》、《生成式人工智能服务管理暂行办法》等）【单选题】

- A. 亲自阅读过

- B. 听说过部分内容
- C. 只听说过名字
- D. 不了解

28. 您是否了解，在使用 AI 产品时，您依法享有要求企业进行算法解释，并在特定情况下拒绝自动化决策的权利？【单选题】

- A. 非常了解
- B. 大致了解
- C. 听说过，但不清楚具体内容
- D. 完全不了解

29. 如果您认为某 AI 大模型提供的答案存在严重偏见或歧视，对您造成了不利影响，您首先会考虑采取哪种行动？【最多选 2 项】

- A. 向提供该服务的公司投诉和反馈
- B. 向网信办、消协等监管机构举报
- C. 在社交媒体上曝光，提醒他人
- D. 不再使用该产品，但不会主动维权
- E. 不知道能做什么

30. 如果未来所有由 AI 生成的内容（如文字、图片、视频、音频）都必须带有清晰的“AI 生成”标识（如水印、标签），当您看到这个标识时，您的第一反应和后续行动最可能是什么？【单选题】

- A. 理性谨慎，核实后使用（我会将其视为一个风险提示，会更仔细地辨别内容真伪，但核实后仍可能使用）
- B. 强烈排斥，避免使用（我将视其为“非人工、不可靠”的标签，会主动避免使用任何带此标识的内容）
- C. 产生戒备，减少使用（我会对带标识的内容天然地产生不信任感，并倾向于减少接触和使用）
- D. 基本无视，习惯照旧（这个标识不会影响我的判断，我依然像以前一样根据自己的经验来使用这些内容）

其他问题

31. 您的学历是

- A. 高中及以下
- B. 专科
- C. 本科
- D. 研究生及以上
- E. 不便透露

32. 您的职业是

- A. 学生
- B. 机关和事业单位工作人员
- C. 企业工作人员
- D. 自由职业者
- E. 其他

33. 您的年收入是

- A. 2 万以下
- B. 2-5 万
- C. 6-20 万
- D. 21-50 万
- E. 50 万以上
- F. 不便透露