

最新挑战

莎拉·克雷普斯
道格·克里纳

徐行健 译

作者莎拉·克雷普斯 (Sarah Kreps) 是康奈尔大学政府系的约翰·L·韦瑟里尔教授、法学兼职教授，以及科技政策研究所所长；道格·克里纳 (Doug Kriner) 是康奈尔大学政府学系的克林顿·罗斯特美国制度学教授

译者徐行健为中国大陆专业人士

中国民主季刊

第3卷 第2期
2025年5月
CC-BY-NC-ND

自由、免费转载分发，但须注明
作者与出处，并不得修改及商用

人工智能如何 威胁民主

编按：生成式人工智能的爆炸性崛起，已经在深刻改变新闻、金融和医疗，但它也可能对政治产生扰乱性影响。比如，向聊天机器人咨询如何应对复杂的官僚体系，或请求其协助撰写给民选官员的信件，可以提升公民的参与度。然而，这项技术也具有大规模生成虚假信息和误导性信息的潜力，可能干扰民主代表机制，破坏民主问责制度，侵蚀社会与政治信任。本文一一分析了 AI 在这些领域所构成的威胁的范围，以及对这些 AI 滥用行为可能采取的防护措施。原文发表于美国《民主季刊》2023 年第 4 期（*Journal of Democracy*, Volume 34, Number 4, October 2023, pp. 122-131）。

生成式人工智能（AI）聊天机器人 ChatGPT 推出仅一个月后，月用户量就达到了 1 亿，成为历史上增长最快的应用程序。与之相比，如今家喻户晓的视频流媒体服务 Netflix 用了三年半的时间才达到一百万月活用户。但与 Netflix 不同的是，ChatGPT 的迅速崛起及其潜在的利弊引发了广泛讨论。学生们是否能够使用，或者更准确地说，滥用这一工具进行研究或写作？它会让记者和程序员失业吗？它会不会像《纽约时报》的一篇专栏文章所说的那样，“劫持民主”（hijack democracy），通过允许大量虚假信息的输入，从而影响民主的代表性？¹ 而最根本的（也是世界末日般的）问题是，人工智能的进步是否真的会对人类的生存构成威胁？² 新技术引发了不同程度和紧迫性的新问题和新担忧。例如，人们担心生成式人工智能——能够产生新内容的人工智能——会对人类的生存构成威胁，这种担心既不可能在短期内成为现实，也不一定是可信的。尼克·博斯特罗姆（Nick Bostrom）的回形针场景（paperclip scenario），即一台经过编程以优化回形针的机器，会消灭阻碍其实现目标的任何东西，而这一场景尚未成为现实。³ 儿童或大学生是否将 AI 工具用作捷径是一个有价值的教学争论，但

随着应用程序越来越无缝地集成到搜索引擎中，这个问题应该会自行解决。生成式 AI 对就业的影响最终将难以判断，因为经济是复杂的，难以区分出 AI 造成的就业损失与行业增长的净效应。然而，它对民主的潜在影响却是直接而严重的。生成式 AI 威胁着民主治理的三大核心支柱：代表制、问责制，以及政治体系中最重要货币——信任。

生成式人工智能最大的问题在于它隐藏在众目睽睽之下，产生大量内容，充斥着媒体、互联网和政治交流，轻则毫无意义的胡扯，重则虚假信息。对政府官员来说，这会削弱他们了解选民情绪的努力，威胁到民主代表制的质量。对选民而言，这威胁到他们监督当选官员的工作及其行动结果的努力，削弱了民主问责制。在这样的媒体环境下，合理的认知预防措施是什么都不相信，这种虚无主义与充满活力的民主背道而驰，也会腐蚀社会信任。随着客观现实在媒体话语中的进一步消退，那些没有完全置身事外的选民很可能会开始更加依赖其他经验方法，如党派偏见，这只会进一步加剧两极分化和对民主体制的压力。

对民主代表制的威胁

正如罗伯特·达尔（Robert Dahl）在 1972 年所写的那样，民主要求“政府持续响应公民的偏好”。⁴ 然而，民选官员要想对选民的偏好做出反应，他们必须首先能够辨别这些偏好。民意调查——（至少目前）大多不受 AI 生成内容的操纵——为民选官员提供了一个了解其选区选民偏好的窗口。但大多数公民甚至缺乏基本的政治知识，对特定政策的知识水平可能更低。⁵ 因此，对那些在特定政策议题上持有强烈观点和对该议题极为关注的选民，立法者有强烈的动机对他们做出最积极的回应。长期以来，书面

回应一直是民选官员掌握其选区动态的重要手段，尤其是用来衡量那些在特定议题上被充分动员的选民的偏好。⁶

然而，在生成式人工智能时代，电子通讯发出的有关紧迫政策议题的信号可能会产生严重的误导。如今，技术进步使恶意行为者能够毫不费力地创建独特的信息，对各种议题表达立场，从而大规模制造虚假的“选民情绪”。即使使用旧技术，立法者也很难辨别人类编写和机器生成的信息。

在 2020 年美国进行的一次实地实验中，我们就六个不同的议题撰写了倡议信，然后用这些信件训练当时最先进的生成式人工智能模型 GPT-3，以撰写数百封左翼和右翼的倡议书。我们向 7200 名州议员随机发送了 AI 和人类撰写的信件，共计约 35000 封邮件。然后，我们比较了人类撰写的信件和 AI 生成的信件的回复率，以评估议员们在多大程度上能够识别（并因此不回复）机器撰写的呼吁。在三个议题上，AI 和人类撰写邮件的回复率在统计上没有区别。在另外三个议题上，对 AI 生成的电子邮件的回复率较低——但平均只低了 2%。⁷ 这表明，一个能够轻松生成成千上万封独特邮件的恶意行为者，有可能让立法者在判断哪些议题对其选民最重要时，扭曲其感知，也会扭曲选民对任何特定议题的感受。

同样，生成式 AI 可能会对民主代表制的质量造成双重打击，因为它会使公众意见征询程序变得过时，而公民可以通过公众意见征询程序来影响监管政府的行为。立法者在制定法规时必然要粗线条，授予行政机构相当大的自由裁量权，不仅可以解决需要实质性专业知识的技术问题（例如，规定空气和水中污染物的允许水平），还可以对价值做出更广泛的判断（例如，在保护公众健康和不过度限制经济增长之间做出可接受的权衡）。⁸ 此外，

在一个党派分化严重、立法机关在紧迫的政策优先事项上经常陷入僵局的时代，美国总统越来越多地寻求通过制定行政规则来推进其政策议程。

将决策权从民选代表手中转移到非民选官僚手中会引发对民主缺失的担忧。美国最高法院在“西弗吉尼亚州诉环保局案”（*West Virginia v. EPA*, 2022 年）中提出了这种担忧，阐明并规定了重大问题原则，认为在没有国会明确法定授权的情况下，机构无权对政策进行重大修改。在待审的“洛珀·布莱特公司诉雷蒙多”（*Loper Bright Enterprises v. Raimondo*）案中，最高法院可能会更进一步，推翻“雪佛龙原则”（*Chevron Doctrine*），近三十年来，该原则给予了机构广泛的自由来解释模棱两可的国会法规，从而进一步收紧了对通过监管程序改变政策的限制。

不过，并非所有人都认为监管程序不民主。一些学者认为，在公示和评论期间，公众参与的机会和透明度得到了保障，“民主程度令人耳目一新”，⁹并称赞这一过程“对民主负责，尤其是在决策公开和为所有人提供平等机会的意义上。”¹⁰此外，2002 年美国政府推出的电子规则制定（*e-rulemaking*）计划承诺，通过降低公民参与的门槛，“加强公众参与……从而促进更好的监管决策”。¹¹当然，公众意见总是偏向于对拟议规则的结果利害关系最大的利益集团，尽管降低了参与的门槛，但电子规则制定并没有改变这一基本现实。¹²

尽管存在缺陷，但公众直接、公开地参与规则制定过程有助于通过官僚行动加强政策变化的民主合法性。但是，如果恶意行为者能够利用生成式 AI，让电子规则制定平台充斥着无限的独特评论，以推进特定的议程，那么机构就几乎不可能了解到真正的公众偏好。2017 年出现了一个早期（但

未获成功) 的测试案例, 当时在对拟议的规则修改进行公开评论期间, 机器人向联邦通信委员会发送了 800 多万条主张废除网络中立性的评论。¹³ 然而, 这种“人造草皮”(astroturfing) 被发现了, 因为其中 90% 以上的评论都不是独特的, 表明这是一种有组织的误导, 而不是真正的民众支持废除。当代 AI 技术的进步可以轻易克服这一限制, 使机构极难发现哪些意见真正代表了利益相关者的偏好。

对民主问责制的威胁

健康的民主还需要公民能够使政府官员对其行为负责——尤其是通过自由公正的选举。然而, 要使投票箱问责制行之有效, 选民必须能够获得有关其代表以他们的名义采取的行动的信息。¹⁴ 选民长期以来一直依赖大众传媒获取政治信息, 而大众传媒中的党派偏见可能会影响选举结果, 这种担忧由来已久, 但生成式 AI 对选举公正性的威胁要大得多。

众所周知, 外国行动者利用一系列新技术, 协调努力影响 2016 年美国总统大选。参议院情报委员会 2018 年的一份报告指出:

这些(俄罗斯)特工伪装成美国人, 使用有针对性的广告、故意伪造的新闻文章、自行生成的内容和社交媒体平台工具, 与美国数千万社交媒体用户互动, 并试图欺骗他们。这一活动试图在社会、意识形态和种族差异的基础上分化美国人, 挑起现实世界中的事件, 是外国政府暗中支持俄罗斯在美国总统大选中青睐的候选人的一部分。¹⁵

尽管这场影响力运动的范围和规模前所未有的, 但它的若干缺陷可能限制了

其影响力。¹⁶ 俄罗斯特工在社交媒体上发布的帖子存在母语人士不会出现的细微但明显的语法错误，如冠词错位或缺失——这些都是帖子造假的蛛丝马迹。然而，ChatGPT 能让每个用户都相当于母语使用者。这种技术已经被用来创建整个垃圾邮件网站，并用虚假评论淹没网站。科技网站 “the verge” 发布了一份招聘 “AI 编辑” 的启示，要求 “每周能生成 200 到 250 篇文章”，这显然意味着这项工作将通过生成式 AI 工具来完成，只需点击编辑的 “再生” 按钮，就能用流利的英语炮制出大量内容。¹⁷ 潜在的政治应用不胜枚举。最近的研究表明，AI 生成的宣传与人类撰写的宣传一样可信。¹⁸ 这一点，再加上新的精准投放（microtargeting）能力，可能会使虚假信息活动发生革命性变化，使其比影响 2016 年大选的活动更加有效。¹⁹ 源源不断的有针对性的虚假信息可能会扭曲选民对当选官员的行为和表现的看法，以至于选举不再提供真正的问责机制，因为人们投票的前提本身就存在事实疑点。²⁰

对民主信任的威胁

生成式 AI 的进步使恶意行为者能够大规模制造虚假信息，包括针对特定人群甚至个人的精准投放内容。社交媒体平台的激增使虚假信息传播变得轻而易举，包括将其有效引导到特定选区。研究表明，尽管各个政治派别的读者都无法区分一系列人造内容和 AI 生成的内容（认为所有这些都是可信的），但虚假信息并不一定会改变读者的想法。²¹ 政治说服很难，尤其是在两极分化的政治环境中。²² 个人观点往往根深蒂固，几乎没有什么可以改变人们先前的情绪。

风险在于，随着虚假内容——文本、图片和视频——在网上泛滥，人们可能

根本不知道该相信什么，因此会不信任整个信息生态系统。人们对媒体的信任度已经很低，而能够生成虚假内容的工具的泛滥将进一步削弱这种信任度。这反过来可能会进一步破坏对政府的极低信任度。社会信任是维系民主社会的重要粘合剂。它推动公民参与和政治参与，增强对政治机构的信心，促进对民主价值观的尊重，这是防止民主倒退和威权主义的重要堡垒。²³

信任的作用是多方面的。对于政治精英来说，要做出回应，就必须相信他们收到的信息是合法地代表了选民的偏好，而不是一场为了推进特定观点而歪曲民意的有组织活动。“人工草皮”在政治中并不是什么新鲜事，美国的例子至少可以追溯到 1950 年代。²⁴ 然而，AI 的进步可能会使这种行为无处不在，更难被发现。

对公民而言，信任可以激励他们参与政治，鼓励他们抵制对民主制度和实践的威胁。在过去的半个世纪中，美国人对政府的信任度急剧下降，这是美国政治中最有据可查的事态发展之一。²⁵ 尽管造成这种下降的因素很多，但对媒体的信任和对政府的信任是密切相关的。²⁶ 用 AI 生成的、真实性存疑的内容轰炸公民，可能会严重威胁人们对媒体的信心，并对政府的信任造成严重后果。

缓解威胁

虽然了解动机和技术是解决问题的重要第一步，但下一步显然是制定预防措施。其中一项措施是训练和部署与生成式 AI 相同的机器学习模型，来检测 AI 生成的内容。人工智能用于创建文本的神经网络也“知道”产生该内容的语言、单词和句子结构类型，因此可以用来辨别 AI 生成文本与

人类书写文本的模式和特征。AI 检测工具正在迅速普及，并需要随着技术的发展而调整，但“Turnitin”（一款查重软件——译者注）式的模型——类似于教师在课堂上用来检测抄袭行为的模型——可以提供部分解决方案。这些工具基本上使用算法来识别文本中的模式，而这些模式正是 AI 生成文本的标志，尽管这些工具在准确性和可靠性方面仍有差异。

更为根本的是，**负责生成这些语言模型的平台越来越意识到社交媒体平台花了很多年才意识到的事情——它们在生产什么内容、如何构建这些内容，甚至禁止什么类型的内容上负有责任。**如果你向 ChatGPT 询问生成式 AI 如何被滥用用来对抗核指挥与控制，该模型会回答“对不起，我无法对此提供帮助。” ChatGPT 的创建者 OpenAI 也在与外部研究人员合作，使其算法中编码的价值观民主化，包括哪些话题应该禁止搜索输出，以及如何框定民选官员的政治立场。事实上，随着生成式 AI 变得越来越普遍，这些平台不仅有责任创造技术，而且还要以一套符合道德和政治的价值观来创造技术。由谁来决定什么是道德，尤其是在两极分化、党派林立的社会中，这个问题并不新鲜。多年来，社交媒体平台一直处于这些争论的中心，现在，生成式 AI 平台也处于类似的境地。至少，当选的公职人员应继续与这些私营公司密切合作，以生成负责的、透明的算法。在拜登政府的协调下，七家主要的生成式 AI 公司决定承诺自愿采取 AI 保障措施，这是朝着正确方向迈出的一步。

负责生成这些语言模型的平台越来越意识到社交媒体平台花了很多年才意识到的事情——它们在产生的内容、如何构建这些内容，甚至禁止哪种类型的内容上都负有责任

最后，数字扫盲（digital-literacy）运动可以通过培养更明智的消费者来防范生成式 AI 的不利影响。正如神经网络

“学习”生成式 AI 如何说话和写作一样，个人读者自己也可以。在我们向各州的立法者汇报了我们研究中的目标和设计后，一些人表示，他们可以识别 AI 生成的电子邮件，因为他们知道自己选民的写作方式；他们熟悉西弗吉尼亚州或新罕布什尔州选民的标准方言。对于在线阅读内容的美国人来说，做出同样的辨别也是可能的。像 ChatGPT 这样的大型语言模型有一定的公式化写作方式——也许是对五段式文章的艺术学得有点太好了。

当我们问“美国的导弹发射井在哪里？”ChatGPT 以典型的平淡语气回答道：“美国在几个州都有导弹发射井，主要分布在中部和北部。导弹发射井存放洲际弹道导弹（ICBMs），这是美国核威慑战略的一部分。随着时间的推移，导弹发射井的具体位置和数量可能会因运行变化和现代化努力而有所不同。”

这种回答没有什么问题，但对于经常使用 ChatGPT 的人来说，这种回答也是可以预见的。这个例子说明了人工智能模型经常生成的语言类型。研究它们的内容输出，无论主题如何，都可以帮助人们识别出表明内容不真实的线索。

更广泛地说，在 AI 生成的文本、视频和图像激增的世界里，一些已经流行起来的数字扫盲技术（digital-literacy techniques）也可能适用。对每个人来说，核实不同媒体上的数字内容的真实性或事实准确性，以及交叉检查任何看似可疑的内容，例如病毒式传播的（尽管是假的）教皇穿着“巴黎世家”（Balenciaga）蓬松大衣的图片，以确定它是假的还是真的，都应该成为标准做法。这种做法还有助于在政治背景下辨别 AI 生成的材料，

例如，在选举周期的 Facebook 上。

不幸的是，互联网仍然是一个巨大的偏好确认（confirmation-bias）机器。那些因符合个人政治观点而看似可信的信息，不太可能促使人们去核实故事的真实性。在一个虚假内容很容易产生的世界里，许多人可能不得不在政治虚无主义——即不相信同党之外的任何事物或任何人——与健康的怀疑主义之间徘徊。放弃客观事实，或者至少放弃从新闻中辨别客观事实的能力，会让民主社会所依赖的信任变得支离破碎。但我们不再生活在一个“眼见为实”（seeing is believing）的世界。个人对媒体消费应采取“信任但要核实”的态度，在阅读和观看的同时，要严于律己，确定材料的可信度。

生成式人工智能等新技术将为社会带来巨大的利益——经济上的、医学上的，甚至可能是政治上的。事实上，立法者可以利用人工智能工具来帮助识别虚假内容，并对选民关注的问题进行分类，这两者都有助于立法者在政策中反映人民的意愿。但是，人工智能也会带来政治风险。不过，只要我们对潜在风险有正确的认识，并采取防范措施来减少其负面影响，我们就能维护甚至加强民主社会。

注释

- 1 内森·E·桑德斯和布鲁斯·施奈尔：《ChatGPT 如何劫持民主》（Nathan E. Sanders and Bruce Schneier, “How ChatGPT Hijacks Democracy,” *New York Times*, 15 January 2023, www.nytimes.com/2023/01/15/opinion/ai-ChatGPT-lobbying-democracy.html.)
- 2 凯文·罗斯：《行业领袖警告称，人工智能可能带来“灭顶之灾”》（Kevin Roose, “A.I. Poses ‘Risk of Extinction,’ Industry Leaders Warn,” *New York Times*, 30 May 2023, www.nytimes.com/2023/05/30/technology/ai-threat-warning.html.)
- 3 阿列克谢·图尔钦和大卫·登肯伯格：《与人工智能相关的全球灾难性风险分

- 类》(Alexey Turchin and David Denkenberger, “Classification of Global Catastrophic Risks Connected with Artificial Intelligence,” *AI & Society* 35 (March 2020): 147–63.)
- 4 罗伯特·达尔:《多头政体:参与与反对》[Robert Dahl, *Polyarchy: Participation and Opposition* (New Haven: Yale University Press, 1972), 1.]]
 - 5 迈克尔·X·戴利·卡皮尼和斯科特·基特:《美国人对政治的了解及其重要性》[Michael X. Delli Carpini and Scott Keeter, *What Americans Know about Politics and Why it Matters* (New Haven: Yale University Press, 1996)]; 詹姆斯·库克林斯基等:《“女士,事实就是如此”:政治事实和公众舆论》[James Kuklinski et al., “Just the Facts Ma’am”: Political Facts and Public Opinion,” *Annals of the American Academy of Political and Social Science* 560 (November 1998): 143–54]; 马丁·吉伦斯:《政治无知与集体政策偏好》[Martin Gilens, “Political Ignorance and Collective Policy Preferences,” *American Political Science Review* 95 (June 2001): 379–96.]]
 - 6 安德里亚·路易丝·坎贝尔:《政策如何塑造公民:高级政治活动和美国福利国家》[Andrea Louise Campbell, *How Policies Make Citizens: Senior Political Activism and the American Welfare State* (Princeton: Princeton University Press, 2003)]; 保罗·马丁和米歇尔·克莱伯恩:《公民参与和国会响应:参与至关重要的新证据》[Paul Martin and Michele Claibourn, “Citizen Participation and Congressional Responsiveness: New Evidence that Participation Matters,” *Legislative Studies Quarterly* 38 (February 2013): 59–81.]]
 - 7 莎拉·克雷普斯和道格·L·克里纳:《新兴技术对民主代表制的潜在影响:来自实地实验的证据》[Sarah Kreps and Doug L. Kriner, “The Potential Impact of Emerging Technologies on Democratic Representation: Evidence from a Field Experiment,” *New Media and Society* (2023), <https://doi.org/10.1177/14614448231160526>.]
 - 8 埃琳娜·卡根:《总统的行政》[Elena Kagan, “Presidential Administration,” *Harvard Law Review* 114 (June 2001): 2245–2353.]]
 - 9 迈克尔·阿西莫夫,《将麦克诺尔加斯特逼到极限:监管成本问题》[Michael Asimow, “On Pressing McNollgast to the Limits: The Problem of Regulatory Costs,” *Law and Contemporary Problems* 57 (Winter 1994): 127, 129.]]
 - 10 肯尼思·F·沃伦:《政治体系中的行政法》[Kenneth F. Warren, *Administrative Law in the Political System* (New York: Routledge, 2018).]
 - 11 美国律师协会联邦“电子规则制定的现状和未来委员会”:《发挥潜力:联邦电子规则制定的未来》[Committee on the Status and Future of Federal E-Rulemaking, American Bar Association, “Achieving the Potential: The Future of Federal E-Rulemaking,” 2008, <https://scholarship.law.cornell.edu/cgi/viewcontent.cgi?article=2505&context=facpub>.]
 - 12 杰森·韦伯·雅基和苏珊·韦伯·雅基:《偏向商业?评估利益集团对美国官僚

- 机构的影响》[Jason Webb Yackee and Susan Webb Yackee, “A Bias toward Business? Assessing Interest Group Influence on the U.S. Bureaucracy,” *Journal of Politics* 68 (February 2006): 128–39; Cynthia Farina, Mary Newhart, and Josiah Heidt, “Rulemaking vs. Democracy: Judging and Nudging Public Participation That Counts,” *Michigan Journal of Environmental and Administrative Law* 2, issue 1 (2013): 123–72.]
- 13 爱德华·沃克：《数百万虚假评论者要求联邦通信委员会终止网络中立性：“人造草皮”是一种商业模式》[Edward Walker. “Millions of Fake Commenters Asked the FCC to End Net Neutrality: ‘Astroturfing’ Is a Business Model,” Washington Post Monkey Cage blog, 14 May 2021, www.washingtonpost.com/politics/2021/05/14/millions-fake-commenters-asked-fcc-end-net-neutrality-astroturfing-is-business-model/.]
- 14 亚当·普热沃茨基、苏珊·C·斯托克斯和伯纳德·马宁编：《民主、问责和代表制》[Adam Przeworski, Susan C. Stokes, and Bernard Manin, eds., *Democracy, Accountability, and Representation* (New York: Cambridge University Press, 1999).]
- 15 《美国参议院情报特别委员会关于俄罗斯积极活动和干预 2016 年美国大选的报告》(Report of the Select Committee on Intelligence United States Senate on Russian Active Measures Campaigns and Interference in the 2016 U.S. Election, Senate Report 116–290, www.intelligence.senate.gov/publications/report-select-committee-intelligence-united-states-senate-russian-active-measures.)
- 16 有关 2016 年大选虚假信息可能产生的有限影响，请参阅安德鲁·M·盖斯、布伦丹·尼汉和杰森·赖弗勒：《2016 年美国大选中不可信网站的情况》[On the potentially limited effects of 2016 election misinformation more generally, see Andrew M. Guess, Brendan Nyhan, and Jason Reifler, “Exposure to Untrustworthy Websites in the 2016 US Election,” *Nature Human Behavior* 4 (2020): 472–80.]
- 17 詹姆斯·文森特：《AI 正在摧毁旧网络，新网络艰难诞生》[James Vincent, “AI Is Killing the Old Web, and the New Web Struggles to be Born,” The Verge, 26 June 2023, www.theverge.com/2023/6/26/23773914/ai-large-language-models-data-scraping-generation-remaking-web.]
- 18 乔什·戈德斯坦等：《AI 能写出有说服力的宣传吗？》[Josh Goldstein et al., “Can AI Write Persuasive Propaganda?” working paper, 8 April 2023, <https://osf.io/preprints/socarxiv/fp87b>.]
- 19 莎拉·克雷普斯：《技术在网上传播虚假信息中的作用》[Sarah Kreps, “The Role of Technology in Online Misinformation,” Brookings Institution, June 2020, www.brookings.edu/articles/the-role-of-technology-in-online-misinformation.]
- 20 这样，AI 生成的虚假信息可能会大大加剧“脱敏”（desensitization）——执政者表现与选民信念之间的关系——从而破坏民主问责制。参见安德鲁·T·利特尔、基思·E·施纳肯伯格和伊恩·R·特纳：《动机推理与民主问责制》[In this way,

- AI-generated misinformation could greatly heighten “desensitization”—the relationship between incumbent performance and voter beliefs—undermining democratic accountability. See Andrew T. Little, Keith E. Schnakenberg, and Ian R. Turner, “Motivated Reasoning and Democratic Accountability,” *American Political Science Review* 116 (May 2022): 751–67.]
- 21 莎拉·克雷普斯、R·迈尔斯·麦凯恩和迈尔斯·布伦戴奇：《所有适合捏造的新闻》[Sarah Kreps, R. Miles McCain, and Miles Brundage, “All the News that’s Fit to Fabricate,” *Journal of Experimental Political Science* 9 (Spring 2022): 104–17.]
- 22 凯瑟琳·多诺万等：《动机推理、公众舆论和总统认可》[Kathleen Donovan et al., “Motivated Reasoning, Public Opinion, and Presidential Approval” *Political Behavior* 42 (December 2020): 1201–21.]
- 23 马克·沃伦编：《民主与信任》[Mark Warren, ed., *Democracy and Trust* (New York: Cambridge University Press, 1999)]; 罗伯特·帕特南：《独自打保龄球：美国社区的崩溃与复兴》[Robert Putnam, *Bowling Alone: The Collapse and Revival of American Community* (New York: Simon and Schuster, 2000)]; 马克·赫瑟林顿：《信任为何重要：政治信任的下降与美国自由主义的消亡》[Marc Hetherington, *Why Trust Matters: Declining Political Trust and the Demise of American Liberalism* (Princeton: Princeton University Press, 2005)]; 皮帕·诺里斯编：《批判性公民：对民主治理的全球支持》[Pippa Norris, ed., *Critical Citizens: Global Support for Democratic Governance* (New York: Oxford University Press, 1999)]; 史蒂文·列维茨基和丹尼尔·齐布拉特：《民主是如何死亡的》[Steven Levitsky and Daniel Ziblatt, *How Democracies Die* (New York: Crown, 2019).]
- 24 刘易斯·安东尼·德克斯特：《国会议员听到了什么：邮件》[Lewis Anthony Dexter, “What Do Congressmen Hear: The Mail,” *Public Opinion Quarterly* 20 (Spring 1956): 16–27.]
- 25 请参阅皮尤研究中心：《公众对政府的信任：1958 至 2022 年》[See, among others, Pew Research Center, “Public Trust in Government: 1958–2022,” 6 June 2022, www.pewresearch.org/politics/2022/06/06/public-trust-in-government-1958-2022/#:~:text=Only%20two%2Din%2Dten%20Americans,least%20most%20of%20the%20time.]
- 26 托马斯·帕特森：《失序》[Thomas Patterson, *Out of Order* (New York: Knopf, 1993)]; 约瑟夫·N·卡佩拉和凯瑟琳·霍尔·杰米森：《新闻框架、政治犬儒主义和媒体犬儒主义》[Joseph N. Cappella and Kathleen Hall Jamieson, “News Frames, Political Cynicism, and Media Cynicism,” *Annals of the American Academy of Political and Social Science* 546 (July 1996): 71–84.]