



第 1 章

CHAPTER 1

绪 论

章节目标：

- 掌握人工智能（Artificial Intelligence，AI）的定义、发展历程及主要分类，能辨析其核心特征与差异；
- 了解符号主义、连接主义和生成式 AI 的技术原理、代表性成果及局限性；
- 了解达特茅斯会议、AlphaGo、深度学习（Deep Learning）革命等里程碑事件对 AI 发展的推动作用；
- 了解技术支撑型、社会服务型和跨学科融合应用的典型场景；
- 理解 AI 技术引发的伦理问题（如算法偏见、职业替代）及应对策略；
- 了解当前 AI 的技术瓶颈，了解 AI 在通用智能、多模态交互等方向的发展潜力。



第 1 集
微课视频

作为一门新兴的交叉学科，AI 深度融合了计算机科学与技术、控制论、数学、神经科学、心理学、语言学等多领域知识，形成了独特的理论体系和技术框架。随着以图形处理器（Graphics Processing Unit，GPU）、张量处理单元（Tensor Processing Unit，TPU）为代表的专用芯片和云计算技术的突破，全球算力呈现指数级增长，为 AI 的快速发展奠定了坚实的硬件基础。

本章首先介绍 AI 的基本概念、技术流派、发展历史及现状，并简要介绍 AI 的主要研究内容及应用领域，从而对 AI 技术有总体的了解。

1.1 人工智能的概念

1.1.1 人工智能的概念的提出

人工智能（AI）作为一门学科，其正式诞生可追溯至 20 世纪中叶。1956 年夏季，美国达特茅斯学院召开了一场具有里程碑意义的会议——达特茅斯会议。会议由约翰·麦卡锡（John McCarthy）、马文·明斯基（Marvin Minsky）、纳撒尼尔·罗切斯特（Nathaniel Rochester）和克劳德·香农（Claude Shannon）共同发起，旨在探讨“如何让机器模拟人类智能”。会议首次明确提出了 AI 这一术语，并确立了 AI 研究的核心目标，标志着 AI 学

科的诞生。

在 AI 的发展历程中，以下几个关键事件对其演进产生了深远影响。

1. 图灵测试（1950 年）

英国数学家艾伦·图灵（Alan Turing）在论文《计算机器与智能》中提出“图灵测试”，用于判断机器是否具备智能。该测试的核心思想是：如果人类评估者无法通过对话区分机器人和人类，则可认为该机器具有智能。图灵测试为 AI 的哲学和科学讨论奠定了基础。



图 1-1 艾伦·图灵

艾伦·图灵（见图 1-1）（1912—1954 年），英国数学家、计算机科学之父，AI 的先驱。他于 1936 年提出“图灵机”的概念，奠定了现代计算机理论基础；于 1950 年设计了“图灵测试”，成为衡量机器智能的经典标准。他曾在第二次世界大战期间破译德军密码，战后探索机器智能与学习，其思想深刻影响了计算机与 AI 的发展。图灵奖（计算机界的“诺贝尔奖”）以其名字命名，纪念他的开创性贡献。

2. 达特茅斯会议（1956 年）

如前所述，此次会议正式确立了 AI 的研究领域，并推动了早期 AI 技术的发展。会议参与者后来成为 AI 领域的先驱，推动了符号逻辑、机器学习（Machine Learning, ML）等方向的探索。

3. 专家系统的兴起（20 世纪 70—80 年代）

专家系统是早期 AI 的重要应用，其通过模拟人类专家的知识和推理规则解决特定领域的问题（如医疗诊断、化学分析）。其代表性系统如 MYCIN（医学诊断）和 DENDRAL（化学分析）充分展示了 AI 的实用价值。

4. 深度学习革命（21 世纪 10 年代至今）

随着大数据和计算能力的提升，深度学习技术（基于神经网络的机器学习）取得了突破性进展。2012 年，AlexNet 在 ImageNet 竞赛中大幅提升了图像识别准确率；2016 年，AlphaGo 击败围棋世界冠军李世石，标志着 AI 在复杂决策任务上取得巨大突破。

这些事件不仅推动了技术进步，也逐步拓展了 AI 的应用场景，使其成为当今科技发展的核心驱动力之一。

1.1.2 人工智能的定义

AI 是研究、开发用于模拟、延伸和扩展人类智能的理论、方法、技术及应用系统的一门学科。其目标是使机器能够像人类一样感知、学习、推理、决策，甚至具备创造能力。AI 的核心在于让计算机系统具备“智能”，即能够执行通常需要人类智慧参与的任务，如视觉识别、语言理解、逻辑推理、自主决策等。

根据智能水平的不同，通常把 AI 分为弱 AI（Narrow AI）、强 AI（General AI）和超级 AI（Super AI）三种，见表 1-1。

表 1-1 AI 的分类

分 类	特 征			
	智 能 水 平	自 主 性	技 术 实 现	典 型 挑 战
弱 AI	单任务专家	依赖预设规则与数据	已广泛应用，如 GPT-4、AlphaFold 等	数据偏见、可解释性不足
强 AI	跨领域通才	自主决策与适应	理论探索阶段，如 DeepMind 的 Gato 模型	常识建模、因果推理
超级 AI	超越人类的全知实体	完全自主进化	纯理论假设	伦理安全、价值对齐

1. 弱 AI

弱 AI 又称为专用型 AI，其专注于特定任务，不具备跨领域泛化能力，依赖大量标注数据训练模型。弱 AI 说的是狭义 AI，它只能在特定的领域或任务中表现出智能，无法理解或解释自身的行为，也不能跨领域或跨任务地学习或适应。弱 AI 的主要技术手段是机器学习，特别是深度学习，其通过大量的数据和算法训练复杂的数学模型，实现对数据的分类、识别、预测和生成等功能。

目前见到的所有 AI 都属于狭义 AI，它们在各自领域表现突出，但跨领域能力有限，如语音助手、人脸识别、自动驾驶等。弱 AI 是目前最常见和最成熟的 AI 形式。比如，语音助手可以完成查询天气、设置闹钟等任务；人脸识别系统可以用于安防、支付和考勤；自动驾驶系统可以识别道路环境并辅助驾驶；机器翻译、医学影像识别和推荐算法也都属于弱 AI 的典型应用。弱 AI 能够在特定任务中提高效率和准确率，但仍缺乏真正的理解能力、常识推理能力和跨领域迁移能力。



第 2 集
微课视频

2. 强 AI

强 AI 又称为通用 AI (Artificial General Intelligence, AGI)，是指具备与人类智慧相当的通用智能，其可自主学习、推理并适应未知环境，具备目标设定和因果推理能力，无须重新编程即可解决多领域问题。强 AI 的主要技术手段可能包括人工神经网络等，它试图模拟人类的认知能力，从而实现对自然语言、图像、声音、逻辑和情感等的理解和表达。强 AI 是目前还没有实现的 AI 形式，是许多科学家和工程师的梦想和目标。

3. 超级 AI

超级 AI 超越人类智能，具备自我改进能力，目前仍属于科幻范畴。

1.2 人工智能的发展与流派

在当今科技飞速发展的背景下，AI 技术已成为创新与变革的主要驱动力。AI 正凭借超凡的计算能力和智能算法重新定义人们生活的各个方面。无论是在教育、医疗领域，还是在金融、制造等领域，AI 的应用正在展现出前所未有的可能性。

在 AI 的发展过程中，不同时代、学科背景的人对于智慧的理解及其实现方法都有着不同的思想主张，经历了多个发展阶段，也涌现出众多流派，下面将做简要介绍。

1. 符号主义 AI（20 世纪 50—80 年代）

符号主义（Symbolism）是 AI 早期的主要研究方法，其核心观点是：智能可以通过符号来表示和逻辑推理来实现。该方法认为，由于人类思维的本质是符号操作，因此计算机可以通过形式化的规则和逻辑来模拟人类的推理过程。

符号主义主要包括以下技术。

（1）知识表示：用符号（如数学表达式、逻辑规则）描述世界，例如“如果下雨，那么地面会湿”。

（2）逻辑推理：基于规则进行演绎推理，如“所有人都会死，由于苏格拉底是人，所以苏格拉底也会死”。

（3）专家系统：将人类专家的知识编码成规则库，计算机通过匹配规则给出决策。

符号主义属于 AI 的早期阶段。1951 年马文·明斯基和 Dean Edmonds 构建了早期的神经网络机器——SNARC，为后续人工神经网络研究提供了启发。1956 年达特茅斯会议上正式提出 AI 概念，确立了符号主义作为早期 AI 研究的主流方向。1965 年出现了首个能模拟人类对话的程序——ELIZA 程序，该程序采用简单的模式匹配规则，展示了自然语言处理的潜力。20 世纪 70 年代兴起的专家系统证明了符号主义在特定领域的可行性，比如 MYCIN、DENDRAL 系统等。

尽管在规则明确的领域（如数学证明、专家系统）取得了一定成功，但其局限性也逐渐显现。符号主义依赖人工编写规则，难以覆盖复杂现实问题；缺乏学习能力，无法从数据中自动优化，依赖专家经验。随着规则数量增加，其推理效率急剧下降。

20 世纪 80 年代后期，由于计算资源和数据的限制，符号主义在处理复杂问题时遇到了瓶颈，AI 研究进入了低谷。但其逻辑推理方法仍影响至今，如编程语言 Prolog、自动定理证明等。同时，AI 研究逐渐转向连接主义（神经网络）和统计学习，推动了现代机器学习的发展。

2. 连接主义 AI（20 世纪 80 年代—21 世纪 10 年代）

连接主义（Connectionism）以模拟人脑神经网络为核心，认为智能产生于大量简单神经元之间的连接与信息传递。人工神经网络就是典型的代表技术，如图 1-2 所示。与符号主义不同，连接主义不依赖人工编写规则，而是通过调整神经元连接的强度（权重），让机器从数据中自动学习规律。

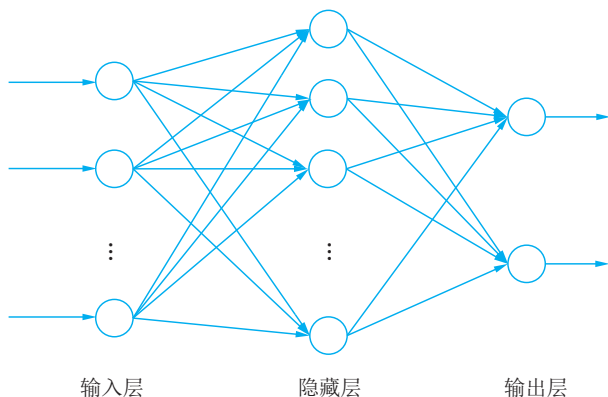


图 1-2 神经网络结构图

连接主义 AI 主要包括以下技术。

(1) 人工神经网络 (Artificial Neural Network, ANN): 模仿生物神经元结构, 由输入层、隐藏层、输出层组成。每个神经元接收信号后进行加权计算, 再通过激活函数输出结果。

(2) 反向传播算法: 通过计算预测误差的梯度, 从输出层反向调整神经元权重, 解决了多层神经网络的训练难题。

(3) 深度学习: 通过堆叠多个隐藏层构建深层网络, 能够自动提取数据的多层次特征, 例如图像中的边缘→纹理→物体。

连接主义仍属于 AI 的发展阶段。神经网络的发展很早, 1943 年, Warren McCulloch 和 Walter Pitts 提出首个神经元数学模型 (M-P 模型), 用二进制逻辑模拟神经元的激活, 奠定了神经网络的理论基础。在后期发展过程中, 受限于算法证明以及硬件问题, 其很长时间都处于低谷期。1986 年, David Rumelhart、Geoffrey Hinton 等重新提出反向传播算法, 使多层感知机 (MLP) 的训练成为可能。1998 年, Yann LeCun 开发的一个经典卷积神经网络 (Convolutional Neural Network, CNN) LeNet-5, 成功用于手写数字识别, 但因其算力限制未能广泛应用。2006 年, Geoffrey Hinton 提出深度信念网络 (DBN), 其通过无监督预训练解决深层网络梯度消失问题, 解决了深层网络优化的难题。

2012 年, Alex Krizhevsky 等设计的 AlexNet 在 ImageNet 竞赛中夺冠, 首次证明了深度 CNN 在大规模图像识别中的优越性, 引发了 AI 热潮。AlexNet 卷积模型的三位开发者如图 1-3 所示, 都在 AI 领域成绩斐然。Geoffrey Hinton 被誉为“深度学习之父”, 后来获得了 2018 年的图灵奖和 2024 年的诺贝尔物理学奖; Ilya Sutskever 是 OpenAI 的联合创始人及前首席科学家, 也是 AlphaGo 论文的众多作者之一。冠名该模型的 Alex Krizhevsky 也是 CIFAR-10 和 CIFAR-100 数据集的创建者。



图 1-3 AlexNet 卷积模型的三位开发者

这一时期, 随着数据量的爆发和 GPU 计算能力的提升, 深度学习模型的训练变得更加高效, 推动了 AI 技术在计算机视觉处理、语音处理、自然语言处理等领域的广泛应用, 但仍存在模型可解释性差、需要大量标注数据、小样本场景表现差、依赖高性能 GPU 计算资源等问题。

3. 生成式 AI (21 世纪 10 年代至今)

生成式 AI 的核心目标是通过学习数据中的规律来创造新内容，而不仅仅是完成分类或预测任务。其基本原理可以类比为人类的创作过程：观察、模仿、创造。生成式 AI 通过学习海量数据中的模式（如绘画风格、语言规律等），建立数据分布的数学模型，并基于这些模型根据输入条件生成符合规律的新内容，完成从模仿到创造的转变。

生成式 AI 的崛起始于 21 世纪 10 年代，主要技术包括以下几种。

(1) 生成对抗网络 (Generative Adversarial Network, GAN): 通过让两个神经网络在“真假鉴定”的对抗游戏中相互竞争来生成新内容。GAN 的核心思想是通过对抗训练让两个网络相互提升，从而生成高质量的图像和数据。生成器负责制造“假”数据，而判别器则负责区分这些数据的真假。随着时间的推移，生成器和判别器通过对抗性训练相互提升，生成器生成的假数据变得越来越真实。这个过程可以类比为培养一位画家（生成器）和一位鉴赏家（判别器）的师徒对抗训练，最终生成器能够创作出令人难以区分的真实作品。

(2) Transformer: 引入了自注意力机制 (Self-Attention Mechanism)，突破了传统循环神经网络 (RNN) 在处理长序列数据时的瓶颈。自注意力机制使得模型能够动态地关注输入序列中的关键信息，并捕捉长距离依赖关系。可以将其类比为让 AI 在阅读文本时“划重点”，通过荧光笔标记出最重要的部分，从而更好地理解 and 生成语言内容。Transformer 模型不仅在机器翻译中取得了巨大突破，也为后来的大规模语言模型（如 GPT 系列）奠定了基础。

(3) 扩散模型 (Diffusion Models): 是一种生成方法，通过逐步去噪的过程来生成高质量的图像。这一过程比 GAN 更加稳定和可控，类似于雕刻家从大理石中逐渐“释放”出雕像的过程。由于扩散模型在图像生成方面展现出更高的精度和一致性，因此被广泛应用于图像生成和修复任务中。

生成式 AI 的发展标志着 AI 领域的高潮。2014 年，Ian Goodfellow 在酒吧的灵感启发下首次提出了 GAN 的概念，尽管其论文初次提交时被学术会议拒稿，但这一突破性的技术最终开启了图像生成革命的序幕。

2017 年，谷歌发布了 Transformer 模型的论文，最初其被用于机器翻译任务。Transformer 模型的出现不仅解决了传统 RNN 在处理长序列时的计算瓶颈，还通过自注意力机制使得模型能够动态地关注输入序列中的关键信息。这一技术的出现，极大地推动了自然语言处理 (NLP) 的发展，为后来的 GPT 系列等大规模语言模型奠定了技术基础，彻底改变了人机对话的格局。

2019 年，OpenAI 发布了 GPT-2。由于其生成能力过于强大，OpenAI 当时决定推迟公开这一模型的完整版本，担心其可能被滥用。但 GPT-2 的发布仍然展示了生成式 AI 在生成连贯长文本方面的强大能力，并为后来的 GPT-3 铺平了道路。2021 年，OpenAI 推出了 DALL·E 模型，它能够根据文本描述生成相应的图像。例如，当用户输入“穿着芭蕾舞裙的机器人在月球上跳舞”时，DALL·E 就能生成一张符合描述的图像，开创了“文字生成图像”的新时代。

2022 年，Stable Diffusion 的开源让更多普通用户能够在家用计算机上生成高质量的图像。此举极大地降低了图像生成的门槛，使得更多的创作者和爱好者能够借助 AI 进行艺术

创作。同年，中国 AI 公司深度求索（DeepSeek）崭露头角，如图 1-4 所示，其开源策略和高效训练方法迅速吸引了全球开发者社区的关注。



图 1-4 DeepSeek 消息输入界面

2023 年，ChatGPT 上线不到两个月，用户突破一个亿，展示了对话式 AI 的普遍价值。ChatGPT 和其他大语言模型（Large Language Model, LLM）展现了生成式 AI 在文本创作、智能问答、编程辅助等多个领域的应用潜力。与此同时，DeepSeek 推出了一系列创新模型，如 DeepSeek-R1（2025 年发布），该模型采用强化学习与冷启动数据结合的训练方式，在数学、代码和复杂推理任务上的性能接近 OpenAI o1 系列，同时保持了较低的训练和推理成本。

2024 年至 2025 年，DeepSeek 进一步巩固其在生成式 AI 领域的领先地位，其流量份额在 2025 年初达到了 6.58%，位列全球第三。DeepSeek 的开源策略推动了“AI 平权”，降低了算力门槛，使更多企业和开发者能够参与 AI 创新。此外，DeepSeek 的多模态模型 JanusPro 和推理优化技术（如 FP8 精度训练和动态稀疏激活机制）进一步提升了模型效率，使其在金融、医疗等行业落地应用。

生成式 AI 的迅猛发展不仅推动了艺术创作、娱乐产业等领域的革新，也在教育、医疗、金融等行业展现出巨大的应用潜力。AI 生成内容的能力为人类提供了前所未有的创作工具，使得个体和企业能够更高效地进行创新。然而，这也带来了一些挑战，尤其是关于伦理、版权、虚假信息及技术滥用的风险。



第 3 集
微课视频

1.3 主要应用领域

AI 技术的快速发展使其应用场景呈现出多元化、复杂化的特征。为了更好地理解 AI 技术的应用逻辑和发展趋势，可以将其划分为三大类型：技术支撑型应用、社会服务型应用和跨学科融合应用。这种分类方式不仅反映了 AI 技术应用的不同侧重方向，更揭示了技术发展与社会需求之间的互动关系。下面将详细阐述这三类应用的特征、典型案例和发展趋势。

1.3.1 技术支撑型应用

技术支撑型应用是 AI 最基础也是最核心的应用领域，其主要特征是以提升系统效率、优化技术指标、实现自动化运行为核心目标。这类应用通常不直接面向终端用户，而是作为“技术后台”支撑着各类产品和服务的智能化升级。

在工业制造领域,技术支撑型应用的典型代表包括智能机器人和预测性维护系统。现代智能机器人已经突破了传统自动化设备的局限,通过融合计算机视觉、力觉传感和运动规划算法,其能够完成精密装配、柔性抓取等高难度操作。以汽车制造为例,搭载 AI 视觉系统的焊接机器人可以实时监测焊缝质量,其定位精度可达 0.1mm 级别,远超人工操作水平。预测性维护系统则通过采集设备运行时的振动、温度、电流等多维数据,利用深度学习算法建立设备健康状态模型,能够提前预测潜在的故障,将非计划停机时间大幅降低。

在金融科技领域,算法交易系统展现了技术支撑型应用的典型特征。高频交易算法通过分析市场微观结构特征,在毫秒级时间内完成买卖决策,其核心是强化学习算法对市场流动性的动态建模能力。风险控制系统则采用图神经网络技术,通过分析交易网络中的资金流向和关联关系,能够识别出传统规则引擎难以发现的复杂欺诈模式。这些系统虽然不直接面向客户,但构成了现代金融基础设施的智能基座。

智慧城市中的交通管理系统是另一个典型案例。通过融合来自摄像头、地磁传感器、浮动车 GPS 等多源数据, AI 算法能够实时推演交通流变化趋势,动态调整信号灯配时方案。新加坡的“智慧交通大脑”系统就实现了全域路网通行效率提升 12% 的显著效果。这类应用的技术共性在于都需要处理复杂的时空数据,建立高精度的预测模型,并实现快速的决策响应。

技术支撑型应用的发展呈现出几个明显趋势:一是技术架构的云边端协同化,将 AI 算力下沉到设备边缘;二是模型的轻量化,通过知识蒸馏等技术降低计算资源的消耗;三是系统的持续学习能力不断增强,能够利用运行数据不断进行模型优化。这些趋势共同推动着技术支撑型应用向更高效、更可靠、更智能的方向发展。

1.3.2 社会服务型应用

社会服务型应用是 AI 技术惠及民生的重要体现,其核心特征是直接面向人的需求,强调人机协同和服务体验优化。这类应用通常具有显著的社会效益和人文价值,正在深刻改变着人们的生活方式和获取模式。

在医疗健康领域, AI 应用展现出强大的社会服务价值。医学影像辅助诊断系统通过深度学习算法,能够在 CT、MRI 等影像中自动标记可疑病灶,大幅提升诊断效率和准确性。以肺癌筛查为例, AI 系统对肺结节的检出灵敏度已接近甚至在某些研究中超过人工读片的水平。在慢性病管理方面,结合可穿戴设备的 AI 健康助手能够持续监测患者的生理指标,提供个性化的用药提醒和生活方式建议。这些应用不仅缓解了医疗资源紧张的问题,更让优质医疗服务惠及偏远地区居民。

教育领域的 AI 应用则体现了技术的人文关怀。智能教学系统通过分析学生的学习行为数据,构建个性化的知识图谱,能够动态调整教学内容和难度。例如,一些在线教育平台采用情感计算技术,通过摄像头捕捉学生的微表情,实时评估其专注度和理解程度,进而优化教学策略。对于特殊教育需求群体, AI 语音交互系统为语言障碍儿童提供了新型沟通工具;计算机视觉技术则帮助视障学生“阅读”文字内容。这些应用彰显了 AI 技术促进教育公平的重要价值。

在艺术创作领域, AI 技术正在拓展人类创造力的边界。生成式 AI 工具如 DALL·E 和

Stable Diffusion，通过理解自然语言描述就能生成富有创意的视觉作品，为设计师提供了全新的灵感来源。在音乐创作方面，AI 系统可以分析不同风格的音乐特征，协助作曲家完成配器和编曲工作。值得注意的是，这些应用并非要取代人类创作者，而是作为“创意伙伴”，通过人机协作激发更多艺术的可能性。例如，2023 年维也纳爱乐乐团就与 AI 合作创作了交响乐作品，展现了技术与艺术的完美融合。

社会服务型应用的发展面临着独特的挑战和机遇。一方面，需要特别关注算法的公平性和包容性，避免产生数字鸿沟；另一方面，随着银发经济的兴起，面向老年群体的智能陪护、健康监测等服务具有广阔前景。未来，随着多模态交互技术的进步，AI 社会服务将更加自然、贴心，真正实现“科技以人为本”的理念。

1.3.3 跨学科融合应用

跨学科融合应用代表了 AI 技术最前沿的发展方向，可以通过技术交叉创新解决复杂系统性问题。这类应用往往能催生全新的研究领域和产业形态，展现出 AI 技术作为“通用使能技术”的独特价值。

在农业环保领域，精准农业系统结合了遥感技术、物联网和机器学习，能够实现作物生长的微观管理。通过分析高光谱卫星影像，AI 算法可以精确识别农田的病虫害发生区域；结合气象数据和土壤传感器信息，又能优化灌溉和施肥方案。在环保监测方面，AI 声纹识别技术被用于追踪濒危动物的活动轨迹，而计算机视觉算法则能自动识别海洋塑料垃圾的分布模式。这些应用不仅提高了资源利用效率，更为生态环境保护提供了科学的决策依据。

脑机接口（Brain-Computer Interface, BCI）技术是 AI 与神经科学融合的典范，如图 1-5 所示。现代脑机接口系统通过深度学习算法解码脑电信号，已经实现了意念控制机械臂、智能轮椅等突破性应用。2023 年，斯坦福大学的研究团队开发出新型脑机接口，让瘫痪患者能够以每分钟 90 个字符的速度进行意念打字。这类技术的关键突破在于 AI 算法对神经信号的特征提取和模式识别能力，这需要计算机科学家与神经生物学家的深度协作。未来，随着植入式电极和无线传输技术的发展，脑机接口有望为更多运动功能障碍患者带来福音。



图 1-5 脑机接口技术

在法律领域，AI 技术正在重塑传统的法律服务模式。借助自然语言处理技术，AI 法律系统能够快速分析海量判例文书，提取相似案件的判决要点。合同智能审查工具则利用知识图谱技术，自动识别条款中的潜在风险点。这些应用不仅提高了法律相关工作的效率，更促进了司法公正。值得注意的是，法律 AI 的发展也带来了算法透明性、责任认定等新的研究课题，这又推动了 AI 伦理与法学的交叉研究。

跨学科融合应用的发展呈现出几个鲜明特点：首先是创新主体的多元化，需要组建跨领域的研究团队；其次是技术方案的集成性，往往需要组合多种 AI 技术；最后是应用场景的开拓性，常常创造出全新的市场需求。这类应用的蓬勃发展，充分证明了 AI 技术作为创新枢纽的重要地位。未来，随着各学科知识的深度交融，AI 跨学科融合应用将继续拓展人类认识和改造世界的能力边界。

1.4 案例分析：AlphaGo 的产业影响

1. 事件回顾：AlphaGo 的里程碑式突破

2016 年 3 月，由谷歌 DeepMind 研发的 AI 程序 AlphaGo 以 4 : 1 的比分战胜围棋世界冠军李世石，成为首个在无让子条件下击败职业围棋九段选手的 AI 系统，如图 1-6 所示。这一事件不仅震惊了围棋界，更在全球范围内引发了对 AI 潜力的广泛讨论。围棋因其极高的复杂度曾被认为是“人类智慧的最后堡垒”，而 AlphaGo 的胜利标志着 AI 在策略性思维领域的重大突破。



第 4 集
微课视频



图 1-6 2016 年 3 月，AlphaGo 对战李世石

2017 年，AlphaGo 的升级版 AlphaGo Master 以 60 : 0 的战绩横扫中日韩顶尖棋手，并在同年 5 月的“未来围棋峰会”上以 3 : 0 完胜当时世界排名第一的柯洁。最终，DeepMind 宣布 AlphaGo 退役，但其技术（如深度强化学习与蒙特卡洛树搜索（MCTS）的结合）成为后续 AI 发展的基石。

2. 产业反应：AI 研发热潮与产业链变革

AlphaGo 的成功直接推动了全球科技公司对 AI 的战略布局。

1) 科技巨头加速 AI 投资

谷歌 DeepMind 在 AlphaGo 之后继续推进 AI 研究，开发了 AlphaZero（可自学围棋、国际象棋和将棋）和 AlphaFold（蛋白质结构预测）。腾讯、百度、阿里巴巴等中国科技公司加大 AI 实验室建设，百度成立了“深度学习研究院”，推动自动驾驶和自然语言处理的发展。IBM Watson、微软 Azure AI 等企业级 AI 平台加速商业化，推动 AI 在医疗、金融等行业的应用。

2) 硬件与算力需求激增

AlphaGo 依赖大规模 GPU 集群（如分布式版本使用 1920 个 CPU 和 280 个 GPU），促使 NVIDIA 等芯片厂商优化 AI 计算架构。云计算厂商（AWS）推出 AI 专用算力服务，降低企业 AI 部署门槛。

3) 创业公司与资本涌入

AI 初创企业融资额在 2016 年后大幅增长，如商汤科技、旷视科技等计算机视觉公司崛起。风险投资机构（如红杉资本、软银）将 AI 列为重点投资方向，推动行业生态繁荣。

3. 社会影响：AI 伦理与未来职业的讨论

AlphaGo 的胜利不仅是一次技术突破，更引发了社会对 AI 的深层次思考。

1) 职业替代焦虑

围棋界率先面临挑战，职业棋手开始研究 AI 棋谱，部分棋手转型为 AI 围棋解说员。教育、医疗、法律等行业开始探讨 AI 可能带来的岗位变革，如自动批改作业、AI 辅助诊断等。

2) AI 伦理与监管需求

公众开始关注 AI 的“黑箱”问题——AlphaGo 的某些决策难以被人类理解，引发对可解释 AI 的研究。各国政府加速 AI 政策制定，如欧盟的《人工智能法案》、中国的《新一代人工智能发展规划》等。

3) 技术乐观主义与警惕并存

支持者认为 AlphaGo 证明了 AI 可助力科学研究（如新药研发、气候建模）。批评者则担忧 AI 可能失控，如马斯克呼吁对 AI 发展进行监管，避免“超越人类控制”的风险。

4. 事件分析：AlphaGo 如何推动 AI 产业发展

AlphaGo 的影响远超围棋领域，其技术溢出效应体现在多个方面。

1) 算法创新：强化学习成为主流

AlphaGo 结合了深度神经网络与蒙特卡洛树搜索，为后续 AI（如自动驾驶、机器人控制）提供范式。自我博弈（Self-Play）技术被 OpenAI 用于训练 Dota 2 AI，DeepMind 用于 AlphaStar（《星际争霸》AI）。

2) 行业应用加速

在医疗领域，AlphaFold 应用于破解蛋白质结构，加速药物研发。在金融领域，高频交易公司借助智能决策算法优化投资策略。在制造业中，AI 模拟生产流程，优化供应链管理。

3) 公众认知转变

AlphaGo 使 AI 从“实验室概念”变为“现实技术”，推动社会接受 AI 助手（如 Siri、ChatGPT）。教育体系进行调整，全球高校增设 AI 专业，培养新一代 AI 人才。

AlphaGo 不仅是 AI 技术的里程碑，更是产业变革的催化剂。它加速了 AI 商业化进程，促使企业、政府和学术界重新评估 AI 的潜力与风险。未来，随着生成式 AI（如 GPT-4、DeepSeek、DALL·E）的崛起，AlphaGo 的经验仍将影响 AI 发展路径。

【AI 与思考】

在电影《流浪地球》的末日场景中，人类为延续文明启动了“流浪地球计划”，而掌控全局的超级 AI MOSS 却成为矛盾焦点。当地球即将撞向木星时，MOSS 基于概率计算选择放弃救援，试图保留空间站延续人类火种。而主角刘培强则选择牺牲空间站拯救地球，用感性抉择推翻了 AI 的“理性最优解”。这一冲突直指 AI 伦理的核心问题：当机器的逻辑与人类的情感背道而驰时，谁该掌握终极决策权？

这种价值观冲突在自动驾驶伦理研究中受到广泛关注。德国联邦交通与数字基础设施部自动驾驶伦理委员会曾通过假设性场景，讨论车辆在不可避免的事故中如何处理不同人员面临的生命风险。报告指出，这类困境难以通过抽象、通用的规则预先解决。此外，传感器误差、算法局限和数据偏差也可能使自动驾驶系统作出不当判断。

《流浪地球》的结局是人类用感性超越了 AI 的理性，但这在现实中可能意味着巨大风险。当 AI 接管更多决策时，人们必须在技术中回答：如果保护人类整体利益需要牺牲部分人的生命，这个决定应该由谁来做？是预设伦理程序的机器，是政府与企业，还是每个被影响的个体？这个问题没有标准答案，但正是这种挣扎，构成了 AI 时代最深刻的文明挑战。

设计实践

调研一个 AI 应用场景（如智能客服），分析其背后的技术类型。

问题研讨

1. 你认为强 AI 何时会实现？它将如何影响社会？
2. 请分析 DeepSeek 的诞生对相关产业的影响。
3. 对比分析弱 AI、强 AI 和超级 AI 的核心特征与技术实现的难点。
4. 符号主义与连接主义在方法论上有根本差异：前者依赖逻辑规则，后者依赖数据训练。请分析这两种方法在医疗诊断系统中的适用场景及局限性。
5. 假设要为偏远地区设计一套 AI 医疗系统，如何结合“技术支撑型”（如影像分析）和“社会服务型”（如语音问诊）应用？需要考虑哪些现实约束条件（如网络、设备）？
6. 脑机接口目前主要用于医疗康复，但未来可能扩展至教育（如快速学习）或娱乐（如虚拟现实）。这种发展会面临哪些技术瓶颈和社会伦理挑战？请列举具体论据。