

数据与知识双驱动式人工智能

▶ 本章导读

人工智能技术发展至今经历了数次革命,大数据、机器学习和深度学习的飞速发展催生了数据驱动的人工智能新范式,推动了包括图像识别、自然语言处理和语音识别等领域的显著进步。然而,数据驱动的人工智能依赖昂贵且耗时的标注数据,同时缺乏可解释性。相对地,知识驱动的人工智能试图通过对领域知识的表示和推理,实现计算机对复杂问题的理解,但面临适应新情境和泛化能力的挑战。数据和知识双驱动的人工智能融合了两者的优势,致力于实现智能化、可解释的系统,同时考虑人类安全和伦理道德。本章将深入介绍这三种人工智能范式的理论、技术和应用,提供全面视角,探索未来的潜力和挑战。

5.1 引言

自 21 世纪初以来,到如今以 ChatGPT 为代表的大模型掀起了新一轮的浪潮,人工智能领域经历了数次革命,这些变革背后的核心是大数据、机器学习和深度学习技术的飞速发展,催生了数据驱动的人工智能的新范式。这一范式的核心是利用海量数据和复杂算法让计算机自主学习、识别模式并作出预测,在图像识别、自然语言处理和语音识别等领域取得了令人瞩目的成果。然而,数据驱动的人工智能也暴露出了一些明显的局限性。例如,它依赖大量标注数据,而这些数据的获取成本高昂且耗时;此外,深度学习模型往往缺乏可解释性,这在一些关键领域,如医疗、金融和法律等,可能导致潜在的风险。相对地,知识驱动的人工智能则试图通过对领域知识的表示和推理,实现计算机对复杂问题的理解和解决。知识驱动的人工智能关注于将人类的知识形式化地表达在计算机系统中,从而使这些系统能够像人类一样进行推理和解决问题。然而,知识驱动的人工智能在面对不断变化的现实世界时,也存在难以适应新情境和缺乏泛化能力的挑战(张钹等, 2020)。与此同时,随着人工智能技术在智能手机、自动驾驶汽车、医疗诊断以及金融交易领域的普及,越来越多的关注点开始聚焦于保护人类安全、伦理道德等方面(Yang et al., 2021)。

数据和知识双驱动人工智能融合了大数据技术和知识表示方法,旨在实现更加智能化和可解释的人工智能系统(吴飞, 2022)。在这个范式中,核心是将数据和知识的优势相互结合,以提高人工智能系统的泛化能力、可解释性和适应性。更重要的是,数据和知识双驱动人工智能关注于在技术发展中充分考虑人类安全和伦理道德因素,能够在数据隐私、算法公平与透明、可解释性和可靠性、人机协

作、可持续发展,以及法律法规与道德规范等方面取得平衡(许为等,2021)。在本章的后续章节中,将详细介绍数据驱动的人工智能、知识驱动的人工智能,以及数据和知识双驱动人工智能的最新研究理论、技术及应用。希望能为读者提供一个全面的视角,了解数据和知识双驱动人工智能的潜力和挑战,同时激发读者对人工智能未来发展的思考和探讨。

5.2 数据驱动的人工智能

数据驱动的人工智能是一种依赖大量数据进行训练和优化的方法,主要是通过从数据中学习模式和规律,从而进行预测、分类、推荐等任务。数据驱动的人工智能与传统的基于规则的方法不同,侧重于利用数据本身来提高预测、分类、推荐等任务中机器学习算法的性能,从而使计算机系统能够自主学习、识别模式并做出决策。在数据驱动的人工智能早期阶段(1950—1980年),人们尝试通过连接主义的方法来模拟人类智能,即基于神经网络模型的计算范式模拟生物神经网络的结构和功能,从而实现智能行为,代表性的工作为 Rosenblatt(1958)提出的感知机(perceptron)模型,其通过设计模拟人类感知能力的神经网络来实现接近人类学习过程(迭代、试错)的学习算法。连接主义和感知机取得了一定程度的成功,为后续的神经网络模型奠定了基础。然而,由于感知机模型的局限性(如无法处理异或问题)以及受限计算资源,再加上缺乏有效的学习算法和足够的训练数据,在实际应用中的表现并不理想(Smolensky et al., 2022a; Smolensky et al., 2022b)。

随着计算能力的提高和统计学习理论的发展,机器学习已成为主流的 AI 方法。在 1980—2000 年,基于数据的学习方法,如支持向量机(Hearst et al., 1998)、决策树和集成学习等开始取得显著的成功,人们借此来自动发现有用的规律和模式。而随着时间进入 2000 年,基于神经网络的机器学习方法——深度学习——开始崭露头角。深度学习通过多层的神经网络结构有效地表示复杂数据模式,尤其在语音识别(Hinton et al., 2012)和图像分类(Krizhevsky et al., 2017)等任务上取得了卓越的成果,展现出了卓越的数据表征能力和泛化性能。随着计算资源的增长和大数据的普及,深度学习也逐渐在自然语言处理等多个领域取得了突破性的进展。近几年,大规模语言模型彻底改变了自然语言处理领域的面貌,它们通过从海量文本数据中的学习,能够生成连贯、自然且具有逻辑性的文本。

5.2.1 关键技术

数据驱动的人工智能核心是从大量数据中学习知识和规律的原理,涉及多种机器学习方法、神经网络结构以及最新的大规模语言模型,下面对重要概念和代表性的技术做简要介绍。

1. 传统机器学习

有监督学习是指从带有标签的训练数据中学习预测模型的方法。在这种情况下,训练数据包含输入特征和对应的目标输出(标签)。典型的有监督学习任务包括分类(输出是离散的)和回归(输出是连续的)。常见的有监督学习算法包括线性回归、支持向量机、决策树和神经网络等。

半监督学习是指利用大量未标记数据和少量标记数据来学习预测模型的方法,其利用未标记数据的结构信息来提高学习效果。半监督学习可以进一步细分为纯半监督学习和直推学习(transductive learning)。纯半监督学习的基本前提是训练数据中的未标注样本,并非预测目标数据,而直推学习则认为在学习过程中涉及的未标注样本正是需要预测的数据,其目标是在这些未标注样本上实现最佳的泛化性能。换言之,纯半监督学习遵循“开放世界”的假设,旨在训练出一个能适应训练过程中未曾观测到的数据的模型;而直推学习则遵循“封闭世界”的假设,仅专注于对学习过程中观

察到的未标注数据进行预测(周志华,2016)。半监督学习通常用于标记数据成本较高的场景。常见的半监督学习算法包括标签传播、生成式对抗网络(GAN)(Creswell et al., 2018)和自编码器等。

无监督学习旨在从未标记数据中学习预测模型的方法。这类学习算法直接从原始数据中学习,试图发现数据中的潜在规则而不依赖人工标签或反馈等指导信息。相较于监督学习旨在建立输入与输出之间的映射关系,无监督学习的目标是发现数据中潜藏的有价值信息,如有效特征、类别、结构以及概率分布等(邱锡鹏,2020)。典型的无监督学习任务包括聚类 and 降维。常见的无监督学习算法包括 K-means(Hartigan & Wong, 1979)、主成分分析(PCA)(Wold et al., 1987)和深度生成模型等。

强化学习是一种学习决策策略的方法,致力于研究智能体(agent)如何在与环境的交互过程中学习到最优策略,从而使得累积奖励最大化。在强化学习中,智能体通过与环境互动来学习最佳行为策略。在每个时间步,智能体选择一个动作,环境根据该动作给出新的状态和奖励。强化学习与有监督学习的主要区别在于,它不依赖于预先给定的标签,而是通过试错学习来寻找最优策略。强化学习的主要模型和算法有三类:值函数方法(Mnih et al., 2013)、策略搜索方法(Konda & Tsitsiklis, 1999)以及模型法(Silver et al., 2017)。

迁移学习是一种利用已有知识来提高新任务学习效果的方法(Pan & Yang, 2009)。在迁移学习中,模型首先在源任务上进行训练,然后将部分或全部知识应用于目标任务。迁移学习可以减少目标任务所需的标记数据量和训练时间。典型的迁移学习方法包括预训练和微调、知识蒸馏和多任务学习等(Zhuang et al., 2020)。

集成学习是一种将多个学习器组合以提高预测性能的方法(Sagi & Rokach, 2018)。这种方法基于多样性和投票原则,试图通过整合多个模型的预测结果来降低过拟合和提高泛化能力。常见的集成学习方法包括 Bagging、Boosting 和 Stacking 等。

2. 深度神经网络

深度神经网络是一类多层次的神经网络模型,可以对数据进行非线性表示和复杂特征提取。近年来,深度神经网络在各种机器学习任务中取得了显著的成功。本节将介绍深度神经网络的几个重要类别,包括多层感知器(MLP)、循环神经网络(RNN)、卷积神经网络(CNN)以及自注意力机制和基于 RLHF 的 LLM 模型。

多层感知器(multilayer perceptron, MLP)是一种基本的前馈神经网络,由输入层、隐藏层和输出层组成。MLP 的每一层都由多个神经元组成,相邻层之间的神经元通过权重连接,如图 5.1 所示。MLP 通过逐层计算和激活函数实现非线性表示,并使用反向传播 BP 算法(Rumelhart et al., 1985)来更新权重。尽管 MLP 在处理简单问题上表现良好,但其扩展性和深度受限,因此在处理复杂任务和大规模数据时,性能可能不尽人意。

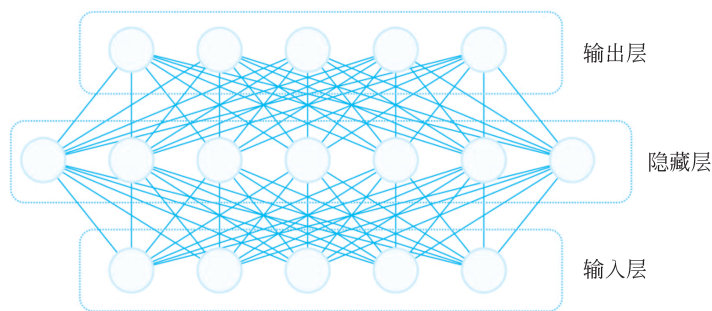


图 5.1 多层感知器示意

循环神经网络(recurrent neural network,RNN)是一种适用于处理序列数据的神经网络模型。循环神经网络具有内部循环连接,使得网络能够在处理当前输入时考虑之前的输入信息,如图 5.2 所示,简单循环网络在时刻 t 的更新公式为:

$$\begin{cases} z_t = Uh_{t-1} + Wx_t + b \\ h_t = f(z_t) \end{cases}$$

其中, x_t 表示 t 时刻的网络输入, h_t 表示 t 时刻的隐藏状态。

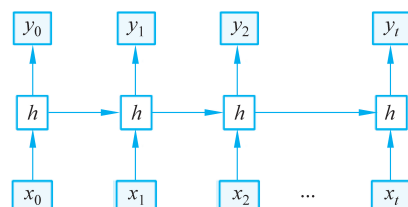


图 5.2 循环神经网络示意

尽管 RNN 具有处理时间序列和自然语言等任务的能力,但是在学习长期依赖时面临梯度消失或梯度爆炸问题。为此,长短时记忆网络(long short-term memory,LSTM)(Hochreiter & Schmidhuber, 1997)和门控循环单元(gated recurrent unit,GRU)(Cho et al., 2014)这两种改进版本的 RNN 相继提出。具体来说,LSTM 引入三个门(门控单元)——输入门、输出门和遗忘门,其可以让 LSTM 在不同时间步骤对输入信息进行选择性记忆或遗忘,从而提高了模型的表达能力。输入门控制着新的输入信息的输入,遗忘门控制着过去信息的遗忘,输出门控制着当前状态信息的输出。GRU 模型与 LSTM 模型相似,都引入了门控机制来控制当前状态信息和记忆信息的更新。GRU 有两个门控单元——重置门和更新门。重置门可以决定如何将前一时刻的隐藏状态和当前输入相结合,更新门可以决定如何将当前输入和前一时刻的隐藏状态相结合。GRU 中还引入了一个候选隐藏状态,用于更新当前的隐藏状态,从而实现了记忆功能。GRU 的参数数量少于 LSTM,因为它将输入门和遗忘门合并成一个更新门,同时减少了细胞状态的数量。相比于 LSTM,GRU 的计算速度更快,但在处理某些任务时可能会稍逊一筹。

卷积神经网络是一种具有局部感受野、权值共享和池化操作的神经网络,尤其适用于处理图像数据。卷积神经网络能够自动学习空间层次结构的特征表示,从而在计算机视觉任务中取得了显著的成功。卷积神经网络的基本组成包括卷积层、激活函数、池化层和全连接层,如图 5.3 所示。通过堆叠这些层,卷积神经网络能够捕捉图像中的局部特征和全局信息。

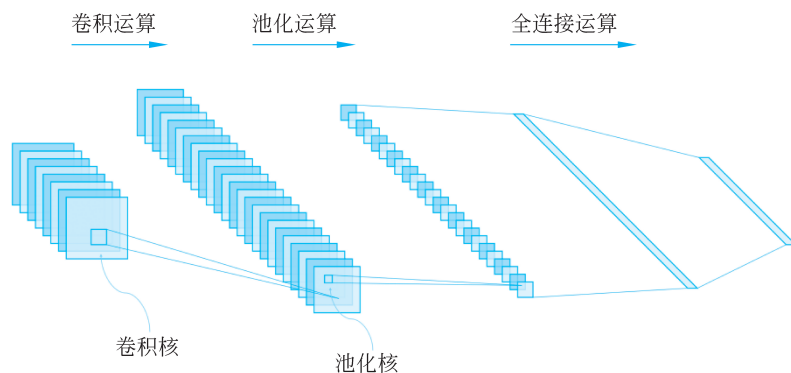


图 5.3 卷积神经网络示意

3. 大规模语言模型

Transformer 是一种基于自注意力机制(self-attention mechanism)的神经网络架构,摒弃了 RNN 和 CNN 的序列和局部结构,提供了一种全新的处理序列数据的方法。Transformer 通过多头自注意力(multi-head self-attention)和位置编码(positional encoding)实现了并行计算和长距离依赖的捕捉。Transformer 在自然语言处理(natural language processing, NLP)任务中取得了显著的成

功,成为现代 NLP 的基石(Vaswani et al., 2017)。Transformers 中的核心操作是自注意力(self-attention)机制,它基于从输入片段序列中获取的查询向量(query)、键向量(key)和值向量(value),如图 5.4 所示。

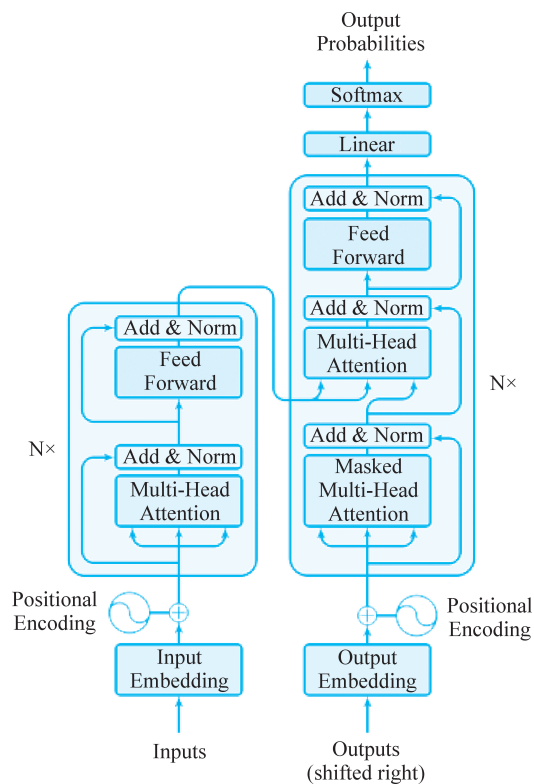


图 5.4 Transformer 结构示意(Vaswani et al., 2017)

BERT(bidirectional encoder representations from transformers)是一种基于 Transformer 编码器结构的预训练语言模型,通过大量无标签文本数据进行训练,学习上下文相关的词向量表示。BERT 使用了掩码语言模型(masked language model)和下一个句子预测(next sentence prediction)作为预训练任务。在微调阶段,BERT 可以轻松地适应各种 NLP 任务,如文本分类、命名实体识别、问答等(Devlin et al., 2018)。从 2019 年开始,Google 就在其搜索引擎中开始使用 BERT;从 2020 年开始,Google 所有的英文输入几乎都使用了 BERT 处理。BERT 的发布对自然语言处理具有非常重要的意义,BERT 消除了许多特定于任务的高度工程化的模型结构的需求,是第一个基于微调的表示模型,它在大量的句子级和标记级任务上实现了最先进的性能,优于许多特定于任务的结构模型。BERT 使得预训练大模型成为自然语言处理领域的主导技术。一般来说,当前最先进的预训练语言模型可以归类为掩码语言模型(编码器)、自回归语言模型(解码器)以及编码器-解码器语言模型,如图 5.5 所示。解码器模型广泛用于文本生成,而编码器模型主要用于分类任务。通过结合两种结构的优点,编码器-解码器模型可以利用上下文信息和自回归特性来提升各种任务的性能。在接下来的部分,我们将深入探讨解码器和编码器-解码器架构的最新进展。

GPT(generative pre-trained transformer)是一种基于 Transformer 的预训练语言模型,在 2018 年由 OpenAI 公司发布。与 BERT 不同,GPT 将 Transformer 和无监督的预训练结合在一起,采用单向的自回归语言建模(Brown et al., 2020; Radford et al., 2018; Radford et al., 2019)。GPT 及其后

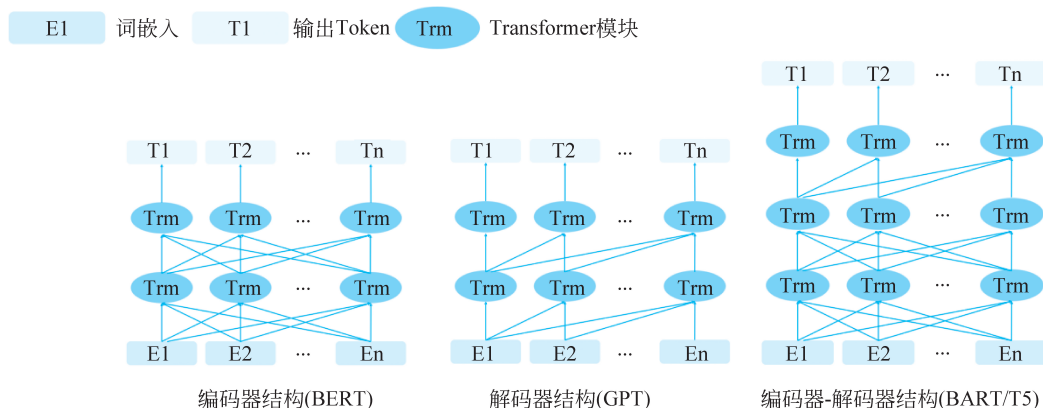


图 5.5 预训练大规模语言模型的分类

续版本(如 GPT-2、GPT-3、ChatGPT、GPT4 等)可以通过预训练-微调的方式应用于多种 NLP 任务,在生成任务方面的表现尤为出色,例如文本生成、摘要、对话等,带来了许多应用上的可能性,例如自动写作、编程帮助、学术研究、语言翻译等。从整个人工智能发展的角度看,GPT 的出现极大地推动了人工智能领域的发展,标志着人工智能领域正式步入大规模语言模型的时代。

具体地,GPT 使用多层 Transformer decoder 作为模型,使用标准语言模型(standard language model)作为预训练目标,给定未标注的语料 $U = \{u_1, \dots, u_p\}$,预训练目标是让如下的似然最大化:

$$L_i(U) = \sum_i \log P(u_i | u_{i-k}, \dots, u_{i-1}; \theta)$$

其中, k 代表上下文窗口的大小,条件概率密度 P 在参数为 θ 的神经网络上建模。GPT 在一些用于词向量表示的简单模型(如 GLoVe、word2vec)的基础上重新审视了在机器学习中流行的有监督的机器学习方法。有监督的机器学习方法需要大量人工标注的数据集,构建这些数据集需要消耗大量的人力和物力成本,高成本会限制数据集的大小,使得现有的有监督学习方法使用的数据集的大小都较为有限,这逐渐成为机器学习发展的瓶颈。GPT 采用的无监督的机器学习方法使用大量未经标注的数据进行训练。而 GPT 的成功也证明了大量未经标注的数据比少量经过标注的数据更加有效。

BERT 使用的双向语言模型在当时的数据量和参数量上比 GPT 的标准语言模型更具优势,2019 年 BERT 在同样的参数规模上性能超过了初代 GPT。OpenAI 随后在同一年发布了 GPT-2,GPT-2 继续沿用了初代 GPT 的标准语言模型,但是使用了更大更好的预训练数据和更大的模型参数(Radford et al., 2019)。值得注意的是,与初代 GPT 不同,GPT-2 展示出了一定的上下文学习(in-context learning, ICL)能力(Dong et al., 2022),和 GPT、BERT 需要微调来适应下游任务不同,GPT-2 可以在预训练之后直接适应下游任务,GPT-2 在无样本学习(zero-shot learning)的任务上取得了较好的成绩。

2020 年,OpenAI 发布了更大的模型 GPT-3,它可以像 GPT-2 一样完成无样本学习的任务,但是 GPT-3 将重点放在了上下文学习上(Brown et al., 2020)。通过在上下文中的几个例子,GPT-3 就可以比无样本学习更好地完成任务。上下文学习是一种范例,它允许语言模型以演示的形式仅给出几个示例来学习任务(Dong et al., 2022)。上下文学习有两个相关的概念:少样本学习(few-shot learning)和提示学习(prompt learning)。上下文学习和少样本学习的区别在于少样本学习一般使用微调模型的方法,而上下文学习冻结了模型参数。而上下文学习其实可以看作提示学习的一个子集,但是上下文学习一般仅指大语言模型的学习能力。

在 GPT-3 的基础上,通过微调可以进一步增强模型的能力。Codex 是在 GPT-3 的基础上微调得到的模型,它的训练数据包含数十亿行开源代码(Chen et al., 2021)。微调后的 GPT-3 可以辅助程序员编写 Python、JavaScript、Go 等数十种编程语言的代码,并且在编写代码的过程中可以显示一些逻辑思维的能力。Codex 之后的大语言模型都会将代码加入它们的训练数据,让模型在支持辅助编程的同时提高模型的逻辑思维能力。

除此之外,在预训练大模型的基础上进行的指令微调(instruction fine-tuning)可以改善预训练模型在特定任务上性能。指令微调是一种监督微调方法(supervised fine-tuning, SFT),预训练模型已经从大型数据集中学习了一般模式和关系,在指令微调中,进一步训练时用较小的、特定任务的人工标注数据集微调大语言模型,以使其适应该任务。指令微调使用的数据集来自人类实际使用过程中的问题或者指令,答案也由人类标注。指令微调给大语言模型带来了响应人类指令和泛化到新任务的能力。代码训练将代码引入文本数据来微调大语言模型,在赋予大语言模型补全代码的能力的同时,也让大语言模型拥有了更强的推理能力。

InstructGPT 是一种基于 GPT-3 的语言模型,它采用了指令微调强化学习(reinforcement learning, RL)和人类反馈(human feedback, HF)相结合的训练方法,旨在提高模型的任务执行能力和理解力(Ouyang et al., 2022)。在 InstructGPT 的训练过程中使用监督学习方法对模型进行预训练,在微调阶段收集一组人类的评分数据。评分通常基于帮助性、无害性等因素,在训练过程中生成一个奖励信号,指导模型优化输出结果。具体来说,首先需要训练一个奖励模型(reward modeling, RM)来预测人类的反馈。有的工作通过让标注人员在多个答案中选择最佳答案来构建 RM 数据集,通过根据偏好来对答案进行排序的方式构建的数据集可以减少模型训练过程中的过拟合。对 K 个结果进行评价时, RM 的损失函数表示为

$$\text{loss}(\theta) = -\frac{1}{\binom{x}{2}} E_{(x, y_w, y_l) \sim D} [\log(\sigma(r_\theta(x, y_w) - r_\theta(x, y_l)))]$$

其中, $r_\theta(x, y)$ 表示有参数 θ 的 RM 对指令 x 和回复 y 的标量输出, y_R 代表在一对结果 y_R 和 y_S 之间较好的一个, D 代表 RM 数据集。最后通过强化学习算法 PPO 对模型进行优化。强化学习的目标函数为

$\text{objective}(\phi) = E_{(x, y) \sim D_{\pi_\phi^{\text{RL}}}} [r_\theta(x, y) - \beta \log(\pi_\phi^{\text{RL}}(y | x) / \pi_\phi^{\text{SFT}}(y | x))] + \gamma E_{x \sim D_{\text{pretrain}}} [\log(\pi_\phi^{\text{RL}}(x))]$
其中, π_ϕ^{RL} 表示学习到的强化学习策略(RL policy), π_ϕ^{SFT} 表示经过指令微调后的模型, D_{pretrain} 表示预训练分布。KL 奖励系数 β 和预训练损失系数 γ 分别控制 KL 惩罚和预训练梯度。RLHF(reinforcement learning from human feedback)是 InstructGPT 训练过程中的关键环节,通过在模型训练中引入人类的评估和指导,提高大语言模型的准确性,并且降低大语言模型的有害信息(toxicity)和偏见,让大语言模型更贴合用户的需求,即对齐(alignment)。

ChatGPT 是在 InstructGPT 的基础上使用基于人类反馈的强化学习进一步微调得到的对话式模型,虽然 OpenAI 没有披露全部的技术细节,但是可以确定 ChatGPT 使用了不同的数据集。使用对话数据对 InstructGPT 使用基于人类反馈的强化学习进行微调,让原来的模型拥有了建模对话历史的能力(Fu et al., 2022)。GPT-4 在 2023 年 3 月由 OpenAI 发布,与之前的 GPT 模型不同, GPT-4 是一个多模态模型,可以处理文字和图像的输入,并生成文字作为输出。GPT-4 在训练的过程中使用了可预测缩放(predictable scaling)的方法,通过在比 GPT-4 小 1000 倍到 10000 倍的模型上进行实验可以预测训练方法在 GPT-4 上的效果。为了测试 GPT-4 在复杂场景下的自然语言理解和生成能力,

OpenAI 在多种为人类设计的考试上验证 GPT-4 的能力, GPT-4 在一些考试中可以取得比多数人类更好的成绩(OpenAI, 2023)。GPT-4 相对于 GPT-3.5 系列模型幻觉(hallucinations)明显降低, 在回答的准确性上有了明显提升。GPT-4 的基础模型与 GPT-3.5 相比差距不大, 但是在引入了基于人类反馈的强化学习之后, GPT-4 有了较大的提升。GPT-4 仍然有很多 GPT 系列模型的局限性, 比如 GPT-4 仍然不具有时效性知识, 有时会犯一些简单的推理错误, 仍然存在偏见等。

5.2.2 数据驱动的人工智能应用现状与挑战

随着深度学习和计算能力的飞速发展, 大语言模型, 如 OpenAI 的 ChatGPT, 已经开始在多个领域显示出其超越传统方法的能力。由于大语言模型的预训练和微调策略, 它们能够有效地从大量的非结构化数据中提取知识, 从而在无须标签的情况下实现零样本学习或者少样本学习。所以, 大语言模型对于那些难以获得大量标注数据的任务或领域提供了解决方案。例如, 在医疗领域, Nuance 公司基于 GPT-4 研发的 DAX 使临床医生能够在护理点上准确、有效地自动记录病人的情况(Overman, 2022)。在金融领域, Wu 等(2023)训练的 BloombergGPT 是一个有 500 亿个参数的 LLM, 使用英语金融文档数据集 FinPile 训练, 该数据集在 LLM 训练中常用的数据的基础上, 加入了包括新闻、文件、出版物、网页、社交媒体等金融领域的文档数据。在与其他通用 LLM 比较时, BloombergGPT 在 5 个金融领域任务中的 4 个上达到 SOTA。在法律领域, Yu 等(2022)提出的 Legal Prompting 方法使用 COLIEE 数据集构建指令微调 GPT-3 模型, COLIEE 数据集是一系列在给定背景文章上进行法律推理的数据集。同时, Legal Prompting 还使用了法律推理提示(legal reasoning prompt)来让 GPT-3 在思维链推理的过程中更接近专业律师。大语言模型已经被实际应用到各种垂直领域中, 但是同时面临一些挑战。

- **垂直领域的應用風險。** 尽管以 ChatGPT 为代表的大语言模型在文本生成方面取得了巨大成功, 但开放式对话问答领域的应用通常可以容忍错误。相反, 对于包括医疗、金融服务、自动驾驶车辆和科学发现在内的高风险应用来说, 数据驱动的人工智能模型仍然充满挑战。在这些领域, 任务的重要性决定了它们需要高度的准确性、可靠性、透明度, 而且错误容忍度接近或等于零。例如, 为自动组织科学知识而设计的大型语言模型 Galactica 可以执行知识密集型的科学任务, 并在几个基准任务上表现出前景。然而, 由于它生成的结果带有权威语气但又存在偏见和错误, 其公共演示版本在初始发布仅 3 天后就被下架。对于这些高风险应用的生成模型, 提供置信度得分、推理和源信息与生成结果同样重要。只有专业人员理解这些结果的来源, 他们才能信心满满地在任务中使用这些工具。
- **专业化与泛化。** 基于大语言模型的应用依赖于基础模型的选择, 这些模型是在不同的数据集上进行训练的。然而, Bommasani 等(2021)指出“在更多样化的数据集上进行训练并不总是比使用更专业化的基础模型得到更好的下游性能”。然而, 构建高度专业化的数据集可能既耗时又成本高昂。对跨领域的数据进行合理的表示以及探索测试数据分布的差异, 能够更好地指导人们设计既考虑专业化又考虑泛化的训练数据集。
- **持续学习与再训练。** 人类的知识库不断扩大, 新的任务不断出现。要生成包含最新信息的内容, 不仅需要模型“记住”学到的知识, 而且需要能够学习并推理新获取的信息。对于一些情景, 只需要在下游任务上进行持续学习, 同时保持预训练的基础模型不变, 但在必要时, 也可以在基础模型上进行持续学习。然而, 持续学习可能并不总是胜过再训练的模型, 这就需要理解什么时候应该选择持续学习策略, 什么时候应该选择再训练策略。此外, 从头开始训练

基础模型可能是不现实的,所以在下一代数据驱动的人工智能基础模型中,考虑模块化设计可能会阐明哪些部分的模型应该被再训练。

- **推理。**推理是人类智能的重要组成部分,使人们能够推导出结论,做出决定,解决复杂问题。然而,即使是使用大规模数据集训练的大语言模型,有时仍然会在常识推理任务中失败。最近,越来越多的研究者开始关注这个问题。链式思维(chain-of-thought, CoT)提示(Wei et al., 2022)是一种应对生成式 AI 模型中推理挑战的有前景的解决方案,它旨在提高大型语言模型在问答场景中学习逻辑推理的能力。通过向模型解释人类用来得出答案的逻辑推理过程,它们可以遵循人类在处理推理时走的道路。通过合并这种方法,大型语言模型可以在需要逻辑推理的任务中取得更高的准确度和更好的性能。CoT 也被应用到其他领域,如代码生成(Zheng et al., 2023)。然而,如何根据具体任务构建这些 CoT 提示仍然是一个挑战。
- **扩大规模。**扩大规模一直是大规模预训练的常见问题。模型训练总是受限于计算预算、可用数据集和模型的大小。随着预训练模型大小的增加,训练所需的时间和资源也显著增加,这对于寻求利用大规模预训练完成各种任务的研究者和组织提出了挑战。另一个问题是大规模数据集预训练的有效性,如果实验超参数,如模型大小和数据量,没有得到深思熟虑的设计,可能不会产生最优结果。因此,次优的超参数可能会导致资源的浪费,以及通过进一步的训练无法达到期望的结果。已经有一些工作被提出来解决这些问题,Hoffmann 等(2022)介绍了一个正式的扩展规律,以预测模型性能基于参数数量和数据集大小。这项工作为理解扩展过程中这些关键因素之间的关系提供了有用的框架。Aghajanyan 等(2023)进行了实证分析,验证了 Hoffmann 扩展定律,并提出了一个额外的公式,探讨了在多模态模型训练环境中不同训练任务之间的关系。这些发现为理解大规模模型训练的复杂性和优化不同训练领域性能的细微之处提供了宝贵的见解。
- **社会问题。**随着数据驱动的人工智能在各个领域的广泛应用,关于其使用的社会问题日益突出。这些问题涉及偏见、道德,以及 AI 生成内容对各个利益相关者的影响等问题。一个主要的关注点是 AI 生成内容中的潜在偏见,特别是在自然语言处理领域。AI 模型可能在无意中继续或放大现有的社会偏见,特别是如果用来开发模型的训练数据本身就存在偏见(Zhuo et al., 2023),这可能产生重大的负面后果,如在招聘、贷款审批和刑事司法等领域继续传播歧视和不平等。使用 AI 生成内容也会引起道德问题,特别是在技术被用来生成深度伪造或其他形式的操纵媒体的情况下。这种内容可以被用来传播虚假信息、煽动暴力,或者伤害个人或组织。此外,还有关于 AI 生成内容可能侵犯版权和知识产权的问题,以及关于隐私和数据安全的问题。总体来说,虽然 AI 生成的内容在各个领域都有巨大的潜力,但解决这些社会问题至关重要,以确保其使用是对社会负责的,对社会整体有益的。

5.3 知识驱动的人工智能

知识驱动的人工智能最早起源于符号主义,尤其是基于知识图谱的方法,其由麦卡锡(J. McCarthy)和明斯基(M. L. Minsky)等学者于 1955 年提出(张钹等, 2020)。他们认为符号 AI (artificial intelligence)的基本思路是“人类思维的很大一部分是按照推理和猜想规则对‘词’(words)进行操作所组成的”。根据这一思路,他们提出了基于知识与经验的推理模型,因此又把符号 AI 称为知识驱动方法。斯坦福大学的费根堡姆发表了特定领域的知识是人类智能的基础这一观点,提出知

识工程(knowledge engineering)与专家系统(expert systems)等一系列强 AI 方法,为符号主义开辟了新的道路,他们开发了专家系统 DENDRAL(有机化学结构分析系统)。

早期的专家系统规模较小,直到 1997 年,IBM 的深蓝国际象棋程序打败世界冠军卡斯帕罗夫标志着符号 AI 可以真正解决大规模复杂系统的开发问题。符号主义的实现基础是纽威尔和西蒙提出的物理符号系统假设,其假设人类认知和思维的基本单位是符号,而认知过程就是在符号表示上的一种运算。因此,计算机可以用来模拟人的智能行为,用计算机的符号操作来模拟人的认知过程。这种方法的实质是模拟人脑的抽象逻辑思维,通过研究人类认知系统的功能机理,用符号来描述人类的认知过程,然后将符号输入计算机进行计算以模拟人类的认知过程,从而实现人工智能。可以把符号主义的思想简单地归结为“认知即计算”(Zhang et al., 2021)。随着技术的进步,知识图谱已成为连接符号主义与现代 AI 技术的桥梁,将丰富的领域知识与数据驱动的方法相结合,可以提供更加强大和灵活的解决方案。

5.3.1 关键技术

从符号主义的观点来看,知识是信息的一种形式,是构成智能的基础,知识表示、知识推理、知识运用是知识驱动的人工智能的核心,下面对重要概念和代表性的技术做简要介绍。

1. 规则引擎

规则引擎是一种在计算领域中用于实现和执行业务逻辑的软件系统,其核心目的是将业务决策逻辑与应用程序代码分离,从而使非技术人员也能定义、修改和管理业务规则,同时确保这些规则的一致性和准确性。规则引擎通常包含一个规则库,其中存储了一系列的业务规则以及一个推理机,它负责对给定的数据或情境应用这些规则。

20 世纪 60 年代初,出现了运用逻辑学和模拟心理活动的一些通用问题求解程序,它们可以证明定理和进行逻辑推理。但是这些通用方法有一个显著的缺点,即它们无法把实际问题改造成适合于计算机解决的形式,并且由于硬件的限制,难于处理对于解题所需的巨大的搜索空间。1965 年,费根·鲍姆等在总结通用问题求解系统成功与失败的经验的基础上,结合化学领域的专门知识,研制了世界上第一个专家系统,可以推断化学分子结构。多年来,专家系统的理论和技术不断发展,在各个领域都产生了很有影响力的应用,包括化学、物理、生物医学等领域,某些专家系统甚至超过同领域中人类专家的水平,并在实际应用中产生了巨大的经济效益。以往的专家系统的局限性促使研究人员开发新类型的方法,进而可以更容易地纳入新的知识,使专家系统自身的更新变得更容易。这样的系统可以更好地从现有的知识中归纳和处理大量的复杂数据。有时,这些类型的专家系统被称为智能系统。

由推理引擎和知识库构成的专家系统通常由人机交互界面、知识库、推理引擎、解释器和知识获取设备这些部分组成,其中,核心部分是知识库和推理引擎。专家系统的体系结构随专家系统的类型、功能和规模的不同而有所差异。知识库代表关于世界的事实信息。在早期的专家系统(如 Mycin 和 Dendral)中,这些事实主要表示为关于变量的平面断言(flat assertions)。在后来为了商业应用而开发的专家系统中,知识库有了更多的结构,并使用了面向对象编程的概念。世界被表示为类、子类和实例,断言被对象实例的值所取代。规则通过查询和断言对象的值来工作。推理引擎是一个自动推理系统,它评估知识库的当前状态,应用相关的规则,然后将新的知识断言到知识库中。使用规则明确表示知识,使专家系统具有可解释性。一个重要的研究领域是用自然语言从知识库中生成解释,而不是简单地显示太直观的规则。人机交互界面是系统与用户进行交流时的界面。通过该界面,用户输入基本信息、回答系统提出的相关问题,并输出推理结果及相关的解释等。知识获取是专家系统