

人工智能基础算法与应用

人工智能的三大核心基石是数据、算力和算法。构建或者选用合适的算法是人工智能技术各个领域真正落地的重要工作。本章首先介绍人工智能的一些基础常用算法,然后以路径规划和海量数据挖掘为应用场景,分别阐述人工智能典型算法应用,主要包括:

任务一:人工智能基础算法简介。

任务二:汽车自动导航路径规划。

任务三:个性化智能推荐系统。

本章教学目的

通过对具体案例的剖析,了解人工智能的一些基本算法,熟悉人工智能技术应用的基本过程和应用模式。

3.1 人工智能基础算法简介

3.1.1 算法模型的两大类划分

人工智能相关的算法模型数量众多,大致上可以划分为两大类:基于统计的机器学习算法(machine learning)和深度学习算法(deep learning)。一般认为,深度学习是一种端到端(end-to-end)的学习,属于黑箱(black box)系统,它是机器学习领域中的一个新研究方向。机器学习算法涉及概率论、统计学、逼近论、矩阵论和算法复杂度理论等多门学科。Sklearn(全称 scikit-learn)是基于 Python 语言的开源机器学习算法库,为了降低算法实践与应用的门槛,Sklearn 提供了 6 个常用算法模块:分类、回归、聚类、降维、模型选择和预处理。Sklearn 官网页面如图 3.1 所示。

人工智能算法的一个重要特点就是具备学习的能力,简单来说就是统计数据中的规律,得到算法模型的一系列最佳参数。例如,人工智能训练师的工作相当于助教,先拿出一个红苹果给机器,并教会它说“苹果”;再拿出一个绿苹果给它辨别,因为颜色的差异机器无法认出来,但助教同样可以教会它说“苹果”,甚至也可以教会它识别被咬了一口的苹果。因此,机器的学习过程是离不开样本训练数据的,期间需要不断修正机器自身的各种模型参数。

针对数据样本的学习方式可以分成三种:有监督学习、无监督学习和半监督学习。有监督学习的数据样本都贴有标签(例如,无论红苹果还是绿苹果都有“苹果”标签),根据有标

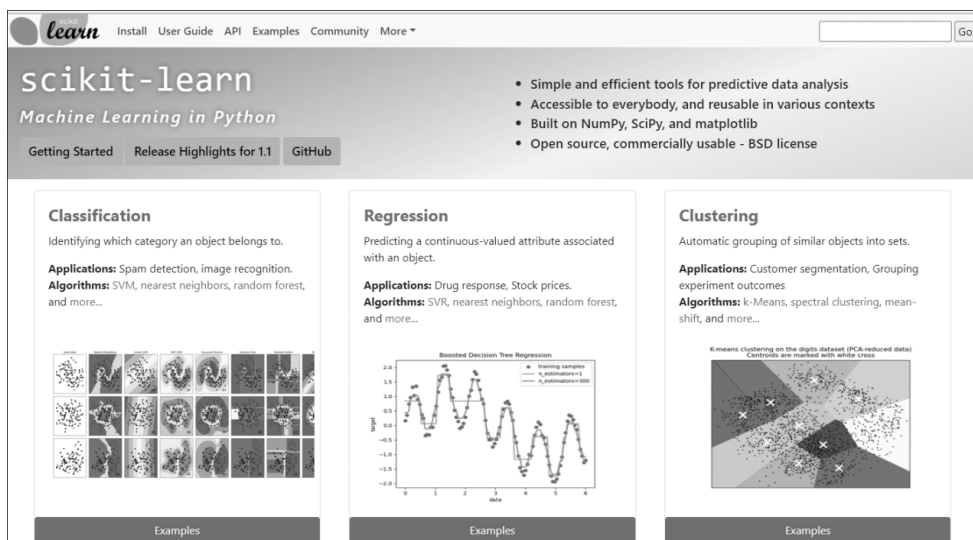


图 3.1 Sklearn 官网

签数据尽最大可能拟合输入和输出之间的关系,其目标是学习一个函数关系或网络结构。无监督学习的输入样本可以具备标签,但学习过程中不存在标注过的样本输出标签,因此其目标是推断一组数据样本的内部结构。例如,无监督学习中最常见的任务是聚类,就是学习数据的内在密度分布情况。半监督学习是有监督学习和无监督学习相结合的一种学习方式,半监督分类任务旨在利用无类标签样例的帮助来训练有类标签的样本,获得性能更优的分类器。可类比一名老师给学生讲一两道有答案的例题,然后再给学生布置没有答案的课后习题,供学生自学巩固。实际情况中,由于有些信息量太大且人工标注成本高昂,而无法获得所有信息的标记,这也是半监督学习的应用场景。

3.1.2 基于统计的机器学习算法

这六大算法模型包含了 8 个纯算法:回归算法、分类算法、聚类算法、降维算法、概率图模型算法、文本挖掘算法、优化算法和深度学习算法;还包含了两个建模方面的算法:模型优化和数据预处理。下面介绍如下几个代表算法模型。

1. 线性回归模型

线性回归(又称线性拟合)模型属于经典的统计学模型,该模型的应用场景是根据已知的变量(自变量)来预测某个连续的数值变量(因变量)。例如,餐厅会根据每天的营业数据 X (包括菜谱价格、就餐人数、预定人数、特价折扣等)来预测营业额 y ,就可以采用线性回归模型。它属于一种有监督学习,在模型建立过程中必须同时具备自变量 X 和因变量 y 。

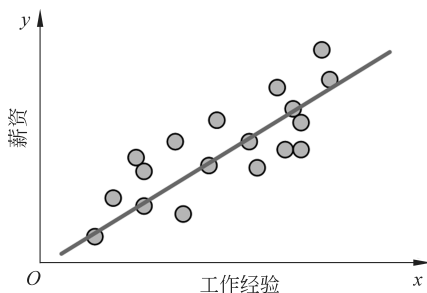


图 3.2 工作经验和薪资关系的散点图

“线性”是指两个变量之间是一次函数关系,函数图像是直线。线性回归算法就是要找到一条直线,让这条直线尽可能地拟合散点图中的数据点,图 3.2 所

示的工作经验和薪资就是一元线性关系,可以形成一元线性回归模型。一元线性回归模型中只含有一个自变量和一个因变量,模型的数学公式可以表示成: $y = a + bx + \epsilon$, 其中, a 为模型的截距项; b 为模型的斜率项; ϵ 为模型的误差,模型的求解量是 a 和 b ,将它们统称为回归系数。求解该模型回归系数的常用方法是最小二乘法(least of squares),其思想是通过最小化平方误差来拟合模型。

2. KNN 分类算法

KNN(k-nearest neighbor)分类算法又称 K 最邻近法,最初由 Cover 和 Hart 于 1968 年提出,属于有监督学习中的分类算法。KNN 算法的原理就是当预测一个新的值 x ,依据距离 x 最近的 K 个邻居点的类别来判断 x 属于哪个类别,充分体现了“物以类聚,人以群分”。

如图 3.3 所示,当 $K=3$ 时, x 周围邻居以三角形居多,所以 x 属于三角形类。由此可见,除了数据点自身的分布情况, K 的大小也十分关键。因为若 K 值逐步增大,数目占多的邻居类型可能会发生变化。进一步优化 KNN 分类模型,可考虑两种解决方案。第一种:设置 K 个邻居样本的投票权重,如果有的邻居样本距离 x 比较远,则该邻居的投票权重就可设置低一些,否则权重就高一些,通常将权重设置为距离的倒数。第二种:采用多重交叉验证法,该方法是目前比较流行的方案,其核心就是将 K 取不同的值,然后在每种值下执行 m 重的交叉验证,最后选出平均误差最小的 K 值。

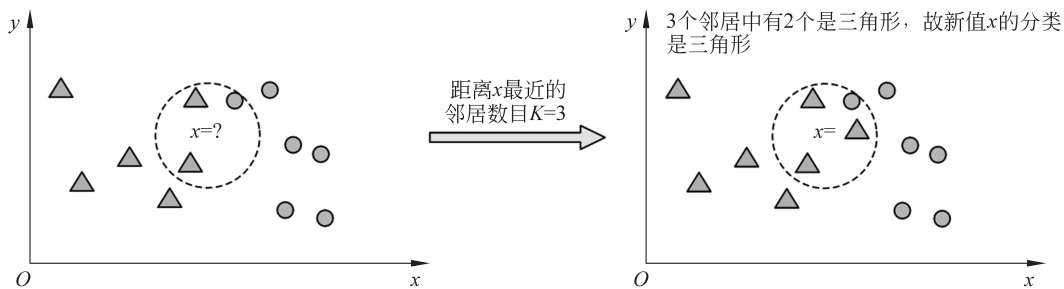


图 3.3 KNN 分类算法示意图

3. K-means 聚类算法

K-means 聚类(即 K 均值聚类)算法是一种迭代求解的聚类分析算法,它是无监督聚类算法中的典型代表。该算法将相似的样本数据自动归聚到一个类别中,下面通过 K-means 聚类图例(如图 3.4 所示)进一步说明其算法原理。

聚类步骤:①随机选取 K 个数据对象作为初始的聚类质心(表示数据集经过聚类将得到 K 个集合);②计算其他每个数据与这 K 个质心之间的距离,把每个样本数据分配给距离它最近的聚类质心;③重新计算每个集合的质心,若新的质心和原质心之间的距离小于某个阈值(表示新质心的位置相比原来的变化不大,可终止算法);④如果新质心和原质心距离变化很大,需要迭代步骤②和③,直至收敛。

4. 人工神经网络中的单层感知器

单层感知器是人工神经网络的一种典型结构(如图 3.5 所示),包含输入层和输出层。输入层负责接收外部信息,每个输入节点接收一个输入信号。输出层也称为处理层,具有信息处理和输出结果的能力。图 3.6 所示为感知器对二维样本的分类示例。

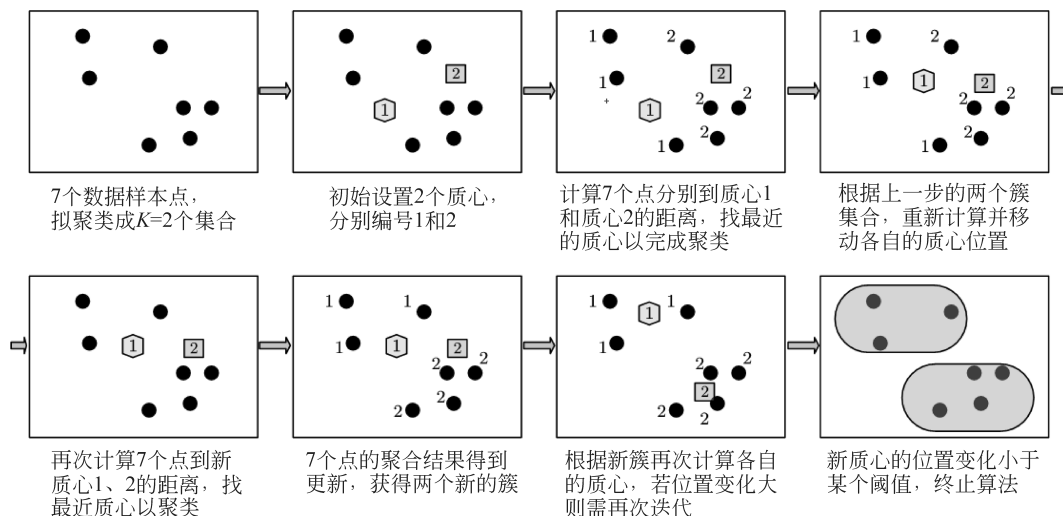


图 3.4 K-means 算法的聚类过程

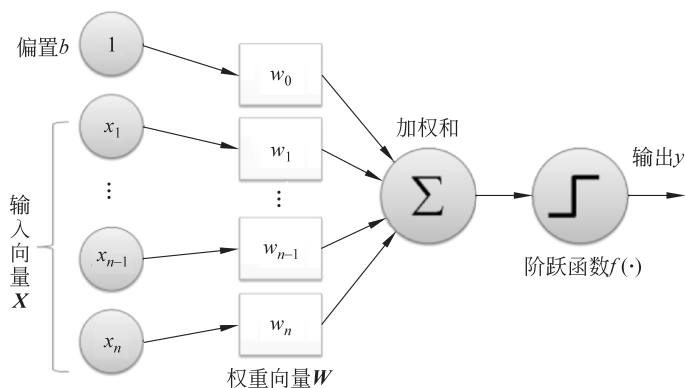


图 3.5 单层感知器的网络结构

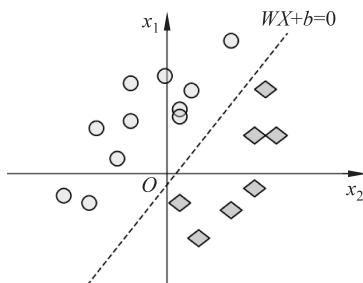


图 3.6 感知器对二维样本的分类

感知器的输出信息： $y=f(WX+b)$ 。其中， W 和 b 为感知器模型参数， $W=[w_1, w_2, \dots, w_n]$ 表示权值向量， b 表示偏置，它们都是需要通过数据样本训练来求解的模型未知量。在训练过程中，感知器的输入信号 $X=[x_1, x_2, \dots, x_n]$ 是每一个样本的特征值向量，输出层的最终输出是该样本的类别。如果需要分类的样本数据是线性可分的，而模型输出类别 y 与

真实类别 y^* 不同时,则可以通过调整突触权值 $\mathbf{W}$ 和偏置值 $\mathbf{b}$,直到每个样本的输出类别与期望类别相同。感知器又可称为线性二元分类器,因为通过改变感知器的权值和偏置值的大小,可改变分界线($\mathbf{WX} + \mathbf{b} = 0$)或分界面的位置,通过阶跃函数(或激活函数) $f(n)$,如式 3.1 所示,最终将所有输入样本分为两类。

$$f(n) = \begin{cases} +1, & \text{当 } n > 0 \text{ 时} \\ -1, & \text{其他} \end{cases} \quad (3.1)$$

3.1.3 深度学习算法

相对传统机器学习的“divide and conquer(分而治之)”,深度学习属于“end-to-end(端到端)”的特征学习。输入端是原始数据,然后输出端的信息直接就是最终目标,也就是说,两端之间的具体转化过程不可知。例如,基于深度学习和摄像头视频信号的自动驾驶系统,输入端是图像像素数据,而输出端直接就是针对方向盘的操作指令。这种近似“黑箱”操作的实现依赖于深度神经网络结构可以自行提取数据特征,中间不再需要人工的特征提取介入。虽然利用深度学习算法不用在特征提取上花费力气,但是人们依然要根据经验和海量数据样本,不断地调整和优化大量的网络参数,这同样是不小的挑战。但由于大数据的加持,深度学习网络模型往往能够收获更加优越的性能。图像识别应用中的代表网络——深度卷积神经网络如图 3.7 所示。传统浅层网络结构如图 3.8 所示。

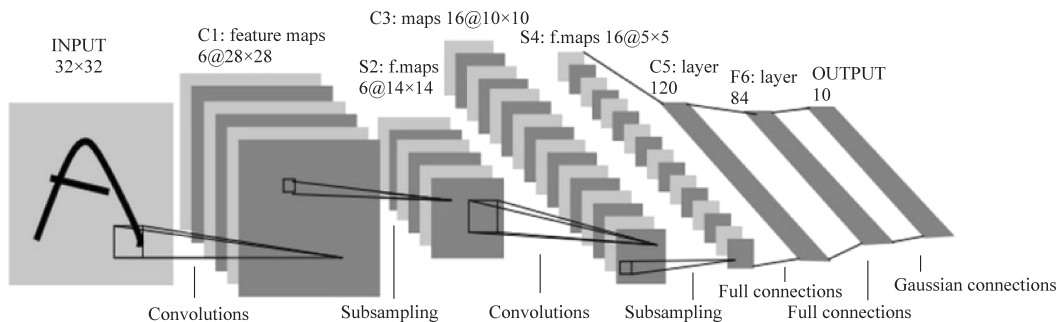


图 3.7 深度卷积神经网络

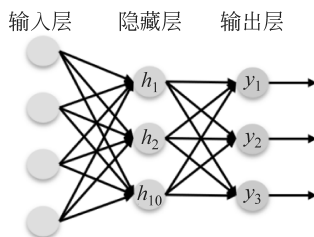


图 3.8 传统浅层网络

可是为什么一定要“深度”网络,“浅度”可否?

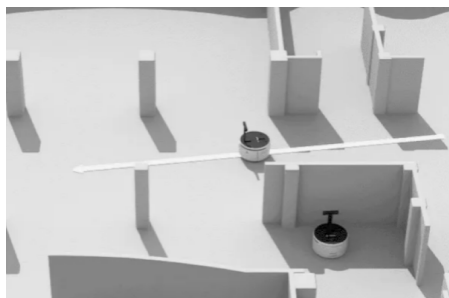
有理论指出,人脑中的神经元组成了不同的层次,多个层次之间相互连接,就形成了一个过滤体系。各层神经元在其所处的环境中获取一部分信息,经过处理后再向更深的层级传递,这有助网络进行更少量的参数调节以适应快速变化的外部环境。以深度卷积神经网络

络(Convolutional Neural Networks,CNN)为例,它由一个或多个卷积层和末端的全连接层组成,同时还包含池化层(pooling layer)。这些结构使得卷积神经网络能够充分提取输入数据的二维信息,因此它在图像处理方面(图像属于二维数据)能够表现出更优的性能。然而,传统的浅层网络针对复杂分类问题,其泛化能力受到一定制约;另外,其训练参数众多而不利于学习和更新。

3.2 汽车自动驾驶路径规划

3.2.1 路径规划介绍

人工智能的应用场景涉及生产、生活的方方面面。基于人工智能技术的路径规划应用十分广泛,包括机器人的自主移动、无人机的避障飞行、GPS 导航、物流管理中的车辆问题(VRP)、通信技术领域的路由问题等。凡是可拓扑为点线网络的规划问题,基本上都可以采用路径规划的方法解决。路径规划的目标是使路径与障碍物的距离尽量远,同时路径的长度尽量短(图 3.9 所示为路径规划的典型应用场景)。机器人的导航规划一般分为构建地图、自定位、路径规划和轨迹规划四大部分。自动驾驶汽车的路径规划算法最早源于机器人的路径规划研究,但是实际操作过程中却比机器人的路径规划复杂得多,因为需要考虑车速、道路的复杂情况、车辆最小转弯半径、外界环境变化等因素。



(a) 机器人底盘路径规划



(b) 汽车自动驾驶路径规划

图 3.9 路径规划应用场景

路径规划可以分为全局路径规划和局部路径规划两类问题。全局路径规划是道路级别的导航(例如高德地图和百度地图导航),它根据全局的地图数据库信息,规划出起点至终点的一条“可通过”的路径。因此,全局路径规划所生成的路径是一条粗略的道路路径。但是,汽车自动驾驶系统在辅助车辆行驶过程(如换道行驶、行人避让等)中,更多关注的是路径的宽度、方向、曲率、路障以及道路交叉口等局部环境。由此可见,结合车辆自身状态信息和诸多细节信息,要规划出一段无碰撞的、平滑的、理想局部路径非常具有挑战性。

汽车自动驾驶任务可以分为三层,包括上层路径规划、中层行驶行为规划和下层轨迹规划。每层完成不同的任务需求,所采用的算法也不相同。

3.2.2 上层路径规划

上层路径规划在获取宏观交通信息、路网和数字地图信息等先验数据信息后,根据某个

优化目标得到两点之间的最优路径。完成该全局路径规划的环境传感信息主要来自于GPS或北斗定位信息以及数字地图信息。主要方法包括栅格(网格)法、概率路线图法、视图法和拓扑法等,下面对栅格法和概率路线图法进行介绍。

1. 栅格法

栅格法首先对地图建模,就是将汽车行驶的路况环境,或者移动机器人的工作环境进行单元分格,将障碍物模拟成一个个的小方格集合。此时场景中的所有物体将进行二值化处理,即障碍物为1,非障碍物为0。栅格法的优点是原理十分简单,编程容易实现。但是在寻找最优路径过程中,栅格大小的选取是影响规划算法性能的一个很重要的因素。若栅格较小,则环境信息将会非常清晰,但会增大存储开销,规划速度就会大大降低,实时性不佳,同时环境干扰信号也会随之增加;反之,若栅格太大,虽然规划速度随之提升了,但环境信息变得较为模糊,不利于有效最优路径的搜寻。因此,通常栅格法仅作为环境建模技术来使用。若作为路径规划算法,它很难解决复杂环境信息变化的问题,一般需要与其他智能算法相结合,例如基于栅格法的蚁群路径规划、基于栅格法的遗传算法路径规划等。如图3.10所示,利用栅格法计算了从一条左上角到右下角的避障路径。

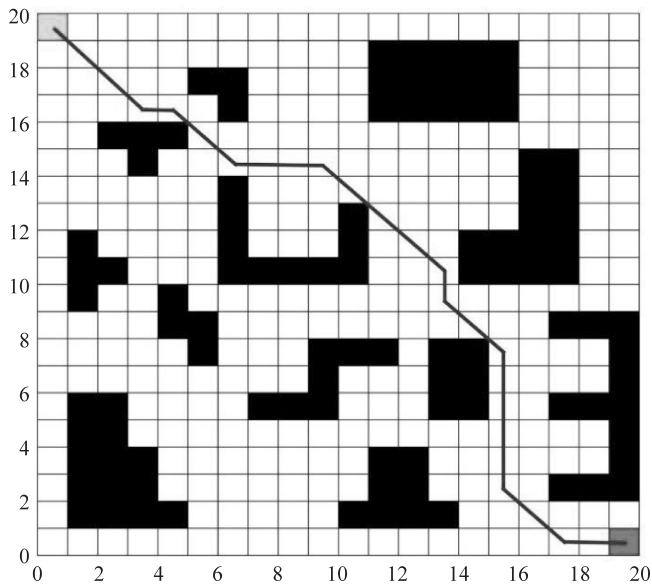


图 3.10 栅格法路径规划

2. 概率路线图法

概率路线图(probabilistic roadmaps)法是一种基于采样和图搜索的方法,它将连续空间转换成离散空间,再利用A*算法或Dijkstra(迪杰斯特拉)算法等搜索算法在路线图上寻找路径,以提高搜索效率。概率路线图法将规划分为“学习”和“查询”两个阶段。在学习阶段,建立一个路线图。在查询阶段,利用搜索算法在路线图上寻找最优路径。其中用于查询阶段的Dijkstra算法是计算最短路径的经典算法之一。该算法适于求解道路权值为非负的最短路径问题,最终可以给出图中某一结点到其他所有结点的最短路径,优点是算法思路清晰,搜索准确。但是由于其输入为大型稀疏矩阵,容易耗时较长,占用空间较大。因此许多

基于 Dijkstra 的改进算法相继提出。例如, A * 算法(又称 A 星算法)是一种基于启发式搜索的算法,该算法结合 BFS(广度优先搜索)和 Dijkstra 算法的优点,在进行启发式有导向性的搜索同时,优化了底层的搜索空间,大大提高了路径的寻优速度。

3.2.3 中层行驶行为规划

中层行驶过程中的行为规划的内容包括导航系统如何生成安全的、可行驶的轨迹,以到达目的地。行为规划是指根据驾驶员感兴趣区域道路、交通车等周围环境信息,决策出当前时刻满足道路环境约束、遵守交通法规的最优行驶行为。一系列的动态规划行驶行为将组成宏观的行驶路径。中层阶段的行为规划所涉及的传感数据,主要来自车载传感器(如雷达、摄像机等),其中将包含用以识别路障、车道线、道路标识和交通信号灯的所有行为指导信息。因此周围环境信息和障碍物的实时检测,将大大影响驾驶行为决策(如停车、换道、超车和避让等)。

针对环境多变性、交通复杂性、交规约束性等诸多车辆行驶不利因素,如何降低其产生的不利影响是行为决策算法模型的研究重点。目前应用较广的是基于有限状态机的行为决策模型和基于深度强化学习的行为决策模型。

有限状态机模型是经典的智能车辆驾驶行为决策方法,模型结构简单、控制逻辑清晰,大多应用于一些较封闭的场景中(如工业园区、港口等)。这些场景中的道路环境变化小、障碍物较固定,可预先设计行驶规则。因此,有利于状态机模型通过构建有限的有向连通图,并且描述不同的驾驶状态转移关系,进而根据状态迁移响应式地生成驾驶操作决策。但是,当车辆行驶环境比较复杂时,场景划分比较困难,各种状态集将大量增加,致使模型结构变得复杂而难以胜任行为决策任务。随着深度学习在图像处理、视频分析分类等方面的巨大成功,基于深度强化学习的行为决策模型发展迅速。深度强化学习是一种端到端的系统。强化学习是指智能体通过与环境的交互获得反馈,在试错中不断进步并强化自身。端到端是指整个系统直接从感知到控制,不需要人工编码,智能体(agent)完全依靠自身学习与环境交互信号。如图 3.11 所示的英伟达(NVIDIA)公司基于 CNN(深度卷积神经网络)的自动驾驶算法训练架构,和以往需要划分感知、检测、决策控制等过程的无人驾驶不同,全程仅通过摄像头采集周围环境图像来完成行为规划。

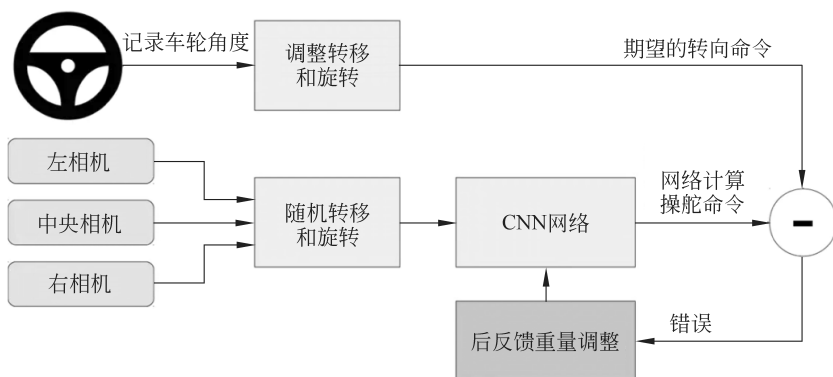


图 3.11 英伟达(NVIDIA)自动驾驶算法训练架构

3.2.4 下层轨迹规划

下层轨迹规划是指在当前时刻的汽车微观动态轨迹规划,就是针对当前已经确定的行驶行为,同时结合周围交通环境,加上车辆的动力学约束,实时做出最优运动轨迹的决策。因此,下层轨迹规划除了要考虑外部环境信息,还需要对车辆状态信息(如速度、车轮转角方向等)进行测量或估计。最关键的是选定一条路线后,用什么样的速度行驶,即速度规划。因此,轨迹规划算法将结合刚体车运动学模型,以及二次规划、变量的边界约束和引导线平滑算法等进行综合建模分析。

3.3 个性化推荐系统

3.3.1 推荐系统介绍

个性化推荐系统是一种高级商务智能平台,为平台顾客提供个性化的信息服务与决策支持,是互联网和电子商务发展的产物。系统综合用户的偏好兴趣、商品属性以及用户之间的社交关系等,发掘用户需求,并且主动推荐商品。推荐系统目前已经成为 AI 成功落地的标志性产品之一,它是许多互联网产品的核心智能组件。例如,电商(淘宝、京东)、资讯(今日头条、微博)、音乐(网易云音乐、QQ 音乐)、短视频(抖音、快手)等热门应用中都配备了个性化推荐系统。推荐系统产生的背景是人们已经从信息匮乏时代走入了信息过载的时代,面对爆炸式的海量数据信息,无论是用户还是生产商都“无所适从”和“不堪重负”,如图 3.12 所示。中文互联网数据研究资讯中心的一份统计显示,每天数以亿计的网络信息被产生、被分享、被接收,其中只有 20%的搜索结果可靠而有用,94%的人感觉“信息过载”。



图 3.12 推荐系统产生背景：信息超载

搜索引擎是一个比推荐系统更早出现的信息过滤系统,爬虫和索引是搜索引擎的基础模块。搜索引擎与推荐系统都是帮助用户快速发现有用信息的工具,但二者的不同之处却很多,如表 3.1 所示。

表 3.1 搜索引擎与推荐系统的区别

比 较 项	搜 索 引 擎	推 荐 系 统
用户行为方式	主动	被动
用户意图	明确	模糊
个性化	弱	强
流量分布	马太效应	长尾效应
目标	快速满足	持续服务
评估指标	简明	复杂

搜索引擎需要用户主动提供准确的关键词来搜索和筛选信息,推荐系统则不需要用户提供明确的信息,用户甚至不知道已经被平台推荐信息。因此,搜索引擎可以满足主动地查找需求;推荐系统能够在用户不明确自身需求的时候,帮助他们发现可能感兴趣的内容。推荐系统是一个综合性很强的工程系统,既需要配置大容量动态随机存取存储器的推理服务器,还需要强大的推荐算法支撑,例如文本分析、用户意图识别、行为分析等。

3.3.2 流行的推荐算法

推荐系统发展之初,传统经典的基于协同过滤、矩阵分解和聚类的推荐算法,在电商推荐系统中扮演着非常重要的角色。1994年,美国明尼苏达大学 GroupLens 研究组首次提出了基于协同过滤(Collaborative Filtering, CF)来完成推荐任务的思想。

基于协同过滤的推荐算法思想是“物以类聚,人以群分”,同时采用两个重要假设。

(1) 基于用户的协同过滤(User-based CF)推荐,例如,和你爱好相似的人喜欢的东西,你也可能会喜欢;

(2) 基于物品的协同过滤(Item-based CF)推荐,例如,和你喜欢的东西相似的东西你也可能会喜欢。此外,矩阵分解(Matrix Factorization, MF)也是一种经典且应用广泛的推荐算法,在基于用户行为的推荐算法里,矩阵分解推荐算法表现效果较为优异,它相较于协同过滤,泛化能力有所加强。基于聚类(如 K-means 聚类)的推荐算法通常与协同过滤相结合,可以有效降低数据稀疏度和提高推荐准确率。

近年来,基于深度学习的推荐系统(如图 3.13 所示)的评价表现尤为突出,与传统的机器学习模型相比,深度学习模型表达能力更强,能够挖掘出数据中更多潜在的模式。

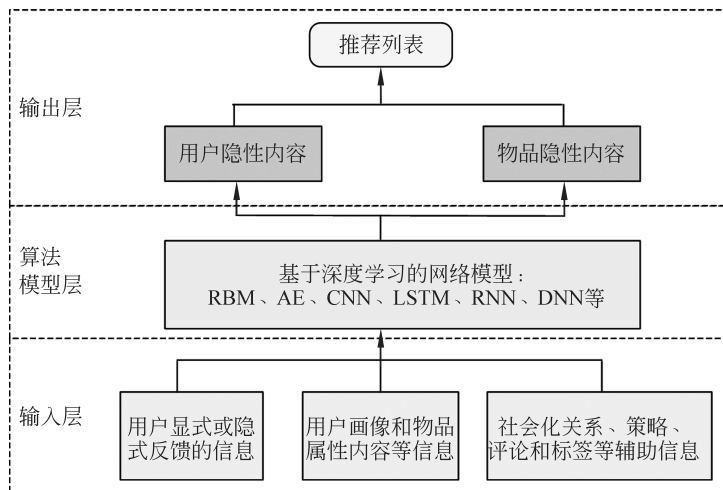


图 3.13 基于深度学习的推荐系统架构

其中,循环神经网络(Recurrent Neural Network, RNN)模型(如图 3.14 所示)和长短时记忆模型(Long Short-Term Memory, LSTM)具有“记忆”能力(如图 3.15 所示),可以“模拟”时序数据间的依赖关系,因而在推荐算法框架设计中被广泛采用。

循环神经网络 RNN 最大特点在于神经网络中的各个隐藏层之间结点是有连接的,故能对过去的信息进行一定时间的存储和记忆。长短时记忆模型 LSTM 具有与循环神经网络

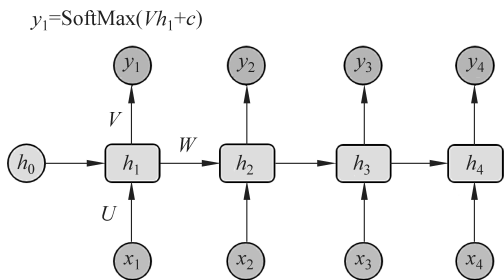


图 3.14 循环神经网络 RNN

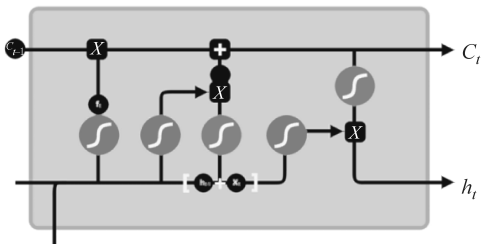


图 3.15 长短时记忆模型 RNN

络相似的控制流,二者的区别在于单元内的处理过程不同,LSTM 有三个门:忘记门、输入门、输出门。在训练过程中,通过门控制可以自主学习到哪些信息是需要保存或遗忘的。这两种算法模型多用于预测评分、图像推荐、文本推荐和基于社交网络的兴趣点推荐等。

习题 3

1. 人工智能对数据样本的学习方式主要有哪几种? 简述它们各自的特点。
2. 汽车自动驾驶可以分哪几层路径规划任务? 每层的任务要求是什么?
3. 简述搜索引擎和推荐系统的共同点和主要区别。