

第1章 绪论

人工智能（Artificial Intelligence, AI）是当前科学技术迅速发展及新思想、新理论、新技术不断涌现的形势下产生的一个学科，也是一门涉及数学、计算机科学、哲学、认知心理学、信息论、控制论等学科的交叉和边缘学科。人工智能主要研究用人工的方法和技术，模仿、延伸和扩展人的智能，实现机器智能。人工智能的长期目标是实现人类水平的机器智能。人工智能自从诞生以来，取得了许多令人瞩目的成果，并在很多领域得到了广泛的应用。本章主要介绍人工智能的概念、发展历史、研究方法和主要应用领域。

1.1 人工智能的概念

1.1.1 智能的定义

什么是智能？智能的本质是什么？这是古今中外许多哲学家、脑科学家一直在努力探索和研究的问题，但至今仍然没有完全解决，以至被列为自然界四大奥秘（物质的本质、宇宙的起源、生命的本质、智能的发生）之一。

目前，人们大多是把对人脑的已有认识与智能的外在表现结合起来，从不同的角度、不同的侧面、用不同的方法对“智能”进行研究，提出的观点亦不相同，其中影响较大的主要有思维理论、知识阈值理论、进化理论。

1. 思维理论（Thinking Theory）

思维理论来自认知科学，认知科学又被称为思维科学。思维理论是研究人们认识客观世界的规律和方法的一门科学，目的在于揭开大脑思维功能的奥秘。思维理论认为，智能的核心是思维，人的一切智慧或智能都来自大脑的思维活动，人类的一切知识都是人们思维的产物，因而通过对思维规律与方法的研究有望揭示智能的本质。

2. 知识阈值理论（Knowledge Threshold Theory）

知识阈值理论强调知识对于智能的重要意义和作用，认为智能行为取决于知识的数量及其一般化的程度，一个系统之所以有智能是因为它具有可运用的知识。在此认识的基础上，知识阈值理论把“智能”定义为：智能就是在巨大的搜索空间中迅速找到一个满意解的能力。知识阈值理论在人工智能的发展史中有着重要的影响，知识工程、专家系统等都是在这个理论的影响下发展起来的。

3. 进化理论（Evolutionary Theory）

进化理论是由美国麻省理工学院（Massachusetts Institute of Technology, MIT）的布鲁克（R. A. Brook）教授提出的。进化理论认为，人的本质能力是在动态环境中的行走能力、对外界事物的感知能力、维持生命和繁衍生息的能力，正是这些能力对智能的发展提供了

基础，因此智能是某种复杂系统所浮现的性质。进化理论的核心是用控制取代表示，从而取消概念、模型及显式表示的知识（Intelligence without Representation, Intelligence without Reasoning），否定抽象对于智能及智能模拟的必要性，强调分层结构对于智能进化的可能性和必要性。

综合上述各种观点，“智能”可以被认为是：智能是知识和智力的总和。其中，知识是一切智能行为的基础，智力是获取知识并运用知识求解问题的能力，即在任意给定的环境和目标的条件下，正确制订决策和实现目标的能力，它来自人脑的思维活动。

智能具有下列特征。

（1）感知能力（Perceiving Ability）。

感知能力是指人们通过感知器官感知外部世界的能力，是人类最基本的生理、心理现象，也是人类获取外界信息的基本途径，一般认为：

感知能力=视觉（80%）+听觉（10%）+触觉+嗅觉+…

也就是说，80%以上的信息通过视觉得到，10%的信息通过听觉得到。

（2）记忆和思维能力（Memorizing and Thinking Ability）。

记忆和思维是人脑最重要的功能，也是人类智能最主要的表现形式。

记忆是对感知到的外界信息或由思维产生的内部知识的存储过程。

思维是对所存储的信息或知识的本质属性、内部规律等的认识过程。人类基本的思维方式有形象思维、抽象思维和灵感思维。

形象思维也称为直感思维，是一种基于形象概念，根据感性形象认识材料，对客观现象进行处理的一种思维形式，如视觉信息加工、图像或景物识别等。神经生理学认为，形象思维是由右半脑实现的。

形象思维一般具有下列特征：

- ① 依据直觉。
- ② 思维过程是并行协同的。
- ③ 形式化困难。
- ④ 在信息变形或缺少的情况下仍有可能得到比较满意的结果。

抽象思维也称为逻辑思维，是一种基于抽象概念，根据逻辑规则对信息或知识进行处理的理性思维形式，如推理、证明、思考等活动。神经生理学认为，抽象思维是由左半脑实现的，如推理、证明、思考等活动。

抽象思维一般具有下列特征：

- ① 依靠逻辑进行思维。
- ② 思维过程是串行的。
- ③ 容易形式化。
- ④ 思维过程具有严密性、可靠性。

灵感思维也称为顿悟思维，是一种显意识与潜意识相互作用的思维方式。平常，人们在考虑问题时往往会因获得灵感而顿时开窍。这说明人脑在思维时，除了那种能够感觉到的显意识在起作用，还有一种感觉不到的潜意识在起作用，只不过人们意识不到而已。

灵感思维一般具有下列特征：

- ① 不定期的突发性。
- ② 非线性的独创性和模糊性。

③ 穿插于形象思维与逻辑思维之中。

记忆和思维能力可表示为：

记忆和思维能力=分析+计算+对比+判断+推理+关联+决策+...

(3) 学习和自适应能力 (Learning and Self-adapting Ability)。

学习是一个具有特定目的知识获取过程。学习和自适应是人类的一种本能，一个人只有通过学习，才能增加知识、提高能力、适应环境。尽管不同人在学习方法、学习效果等方面有较大差异，但学习是每个人都具有的一种基本能力。

(4) 行为能力 (Acting Ability)。

行为能力是指人们对感知到外界信息做出动作反应的能力。引起动作反应的信息可以由感知直接获得的外部信息，也可以是经思维加工后的内部信息。完成动作反应的过程一般通过脊髓来控制，并由语言、表情、体姿等来实现。

1.1.2 人工智能的定义

1956 年，4 位年轻学者——约翰·麦卡锡 (J. McCarthy)、马文·明斯基 (M. Minsky)、纳撒尼尔·罗彻斯特 (N. Rochester)、克劳德·香农 (C. Shannon)，在美国新罕布什尔州的达特茅斯 (Dartmouth) 大学共同发起和组织了用机器模拟人类智能的夏季专题研讨会。会议邀请了包括数学、神经生理学、精神病学、心理学、信息论和计算机科学领域的 10 名学者参加，为期两个月。

在研讨会上，科学家们运用数理逻辑和计算机的成果，提供关于形式化计算和处理的理论，模拟人类某些智能行为的基本方法和技术，构造具有一定智能的人工系统，让计算机完成需要人的智力才能胜任的工作。其中，明斯基的神经网络模拟器、麦卡锡的搜索法、赫伯特·西蒙 (H. Simon) 和艾伦·纽厄尔 (A. Newell) 的逻辑理论成为研讨会的 3 个亮点。

在达特茅斯夏季讨论会上，麦卡锡提议用 AI (Artificial Intelligence) 作为这一交叉学科的名称，标志着人工智能学科的诞生，具有十分重要的意义。麦卡锡也被称为人工智能之父。从那以后，研究者们发展了众多理论和原理，人工智能的概念也随之扩展。人工智能是当前科学技术迅速发展及新思想、新理论、新技术不断涌现的形势下产生的一个学科，也是一门涉及数学、计算机科学、哲学、认知心理学、信息论、控制论等学科的交叉和边缘学科。人工智能的发展虽然已走过了几十年的历程，但是对“人工智能”至今尚无统一的定义，人们试图用下列 4 种方法给出其定义。

1. 类人行为系统 (Systems That Act Like Human)

定义 1.1 人工智能是制造能够完成需要人的智能才能完成的的任务的机器的技术 (The art of creating machines that perform functions that requires intelligence when performed by people), R. Kurzweil, 1990。

定义 1.2 人工智能是研究如何让计算机做现阶段人类才能做得更好的事情 (The study of how to make computers do things at which, at the moment, people are better), Rick and Knight, 1991。

这种观点与图灵测试的观点吻合，是一种类人行为定义的方法。1950 年，阿兰·图灵 (Alan Turing) 提出了图灵测试，为智能提供一个满足可操作要求的定义。图灵测试用人类的表现来衡量假设的智能机器的表现，这无疑是评价智能行为的最好且唯一的标准。

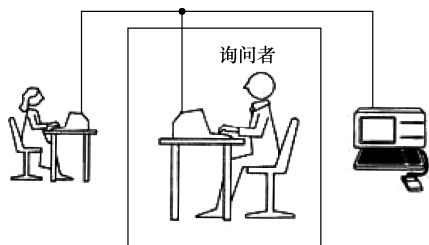


图 1-1 图灵测试

称为“模仿游戏”的图灵测试是这样进行的：将一个人与一台机器置于一间房间中，而与另一个人分开，并把后一个人称为询问者，如图 1-1 所示。

询问者不能直接看到屋中任一方，也不能与他们说话，因此他不知道到底哪一个实体是机器，只可以通过一个类似终端的文本设备与他们联系。然后，让询问者仅根据通过这个终端提问收到的答案辨别出哪个是机

器，哪个是人。如果询问者不能区别出机器和人，那么根据图灵的理论，就可以认为这个机器是智能的。

一台机器要通过图灵测试，它需要具有如下能力。

- ① 自然语言处理：实现用自然语言与计算机进行交流。
- ② 知识表示：存储它知道的或听到的、看到的。
- ③ 自动推理：能根据存储的信息回答问题，并提出新的结论。
- ④ 机器学习：能适应新的环境，并能检测和推断新的模式。

当然，要通过完全图灵测试，机器还需要具有如下能力。

- ⑤ 计算机视觉：可以感知物体。
- ⑥ 机器人技术：可以操纵和移动物体。

这 6 个领域构成了人工智能的大部分内容。

图灵测试的重要特征如下。

① 它给出了一个客观的智能概念，就是根据对一系列特定问题的反应来决定是否是智能体的行为。这为判断智能提供了一个标准，从而避免了有关人工智能“真正”特征的必然争论。

② 这项实验使我们免于受到诸如以下目前无法回答的问题的牵制：机器使用的内部处理方法是否恰当，或者机器是否真地意识到其动作。

③ 通过使询问者只关注回答问题的内容，消除了有利于生物体的偏置。

因为图灵测试具有这些重要特征，所以它已成为许多现代人工智能程序评价方案的基础。如果一个程序已经有可能在某个专业领域实现了智能，那么可以通过把它对一系列给定问题的反应与人类专家的反应相比较，来对其进行评估。这种评估技术只是图灵测试的一个变体：请一些人对程序和人类专家对特定问题的反应结果做出封闭式的比较。

尽管图灵测试具有直观上的吸引力，但是这种方法仍然受到了很多批评。其中最重要的质疑是它偏向于纯粹的符号问题求解任务，并不测试感知技能或要实现手工灵活性所需的能力，而这些都是人类智能的重要组成部分。有人提出，图灵测试没有必要把机器智能强行套入人类智能的模具之中。或许机器智能就是不同于人类智能，试图按照人类的方法来评估它，可能根本上就是一个错误。

2. 类人思维系统 (Systems That Think Like Humans)

定义 1.3 人工智能是一种使计算机能够思维、使机器具有智力的激动人心的新尝试 (The exciting new effort to make computers think...machines with minds, in the full and literal sense), Haugland, 1985。

定义 1.4 人工智能是那些与人的思维、决策、问题求解和学习等有关活动的自动化 (The automation of activities that we associate with human thinking, activities such as decision

making, problem solving, learning ...), Bellman, 1978。

这个定义主要采用认知模型的方法——关于人类思维工作原理的可检测的理论。认知科学是研究人类感知和思维信息处理过程的一门学科，把来自人工智能的计算机模型和来自心理学的实验技术结合在一起，目的是对人类大脑的工作原理给出准确和可测试的模型。

3. 理性思维系统 (Systems That Think Rationally)

定义 1.5 人工智能是用计算模型对智力行为进行的研究 (The study of mental faculties through the use of computational models), Charniak and McDermoth, 1985。

定义 1.6 人工智能是研究那些使理解、推理和行为成为可能的计算 (The study of the computations that make it possible to perceive, reason, and act), Winston, 1992。

一个系统如果能够在它所知范围内正确行事，它就是理性的。古希腊哲学家亚里士多德 (Aristotle) 是首先试图严格定义“正确思维”的人之一，他将其定义为不能辩驳的推理过程。他的三段论方法给出了一种推理模式，当已知前提正确时总能产生正确的结论。例如，专家系统是推理系统，所有的推理系统都是智能系统，所以专家系统是智能系统。

4. 理性行为系统 (Systems That Act Rationally)

定义 1.7 人工智能是一门通过计算过程力图解释和模仿智能行为的学科 (A field of study that seeks to explain and emulate intelligent behavior in terms of computational processes), Schalkoff, 1990。

定义 1.8 人工智能是计算机科学中与智能行为自动化有关的一个分支 (The branch of computer science that is concerned with the automation of intelligent behavior), Luger and Stubblefield, 1993。

行为上的理性指的是已知某些信念，执行某些动作以达到某些目标。主体 (Agent) 是可以进行感知和执行动作的某个系统，在这种方法中，人工智能可以被认为是研究和建造理性主体。

简言之，人工智能主要研究用人工的方法和技术，模仿和扩展人的智能，实现机器智能。人工智能的长期目标是实现人类水平的机器智能。

1.2 人工智能的产生和发展

“人工智能”自从在 1956 年的达特茅斯夏季研讨会上被提出后，研究者们发展了众多理论和原理，人工智能的概念也随之扩展。人工智能是一门极富挑战性的科学，虽然它的发展比预想的要慢，但一直在前进，从出现到现在，已经出现了许多人工智能程序，并且它们影响到了其他技术的发展。人工智能的产生和发展过程可大致分为孕育期、形成期和发展期。

1.2.1 孕育期 (20 世纪 50 年代中期以前)

人工智能的孕育期大致可以认为是 1956 年以前的时期。这个时期的主要成就是数理逻辑、自动机理论、控制论、信息论、神经计算、电子计算机等学科的建立和发展，为人工智能的诞生准备了理论和物质的基础。

这个时期最早可以追溯到公元前的古希腊时代，古希腊著名哲学家亚里士多德写了《工具论》一书，为后来“人工智能”的形式逻辑奠定了基础。在书中，亚里士多德总结了以

三段论为核心的演绎法，这应该是一切推理活动最早和最基本的出发点。三段论的推理由大前提、小前提和结论 3 个判断构成，大前提是一个一般性的原则，小前提是一个附属于前面大前提的特殊化陈述，如果同时满足了大前提和小前提，那么一定会满足结论。后来，英国哲学家和自然科学家培根（F. Bacon）系统地提出了归纳法并强调了知识的作用，成为与亚里士多德的演绎法相辅相成的思维法则。至此，形式逻辑系统已经比较严密了，而将它进一步符号化，从而能对人的思维进行运算和推理，最终形成数理逻辑的是德国数学家和哲学家莱布尼茨（G. W. Leibniz），数理逻辑也成为日后人工智能符号主义学派的重要理论基础。而后，英国数学家、逻辑学家布尔（G. Boole）进一步将莱布尼茨的数理逻辑思维符号化和数学化的思想发扬光大，提出了一种崭新的代数系统——布尔代数。布尔代数已经成为现代计算机软件、硬件中逻辑运算的基础。这个时期还有一个重要人物是美籍奥地利数理逻辑学家哥德尔（K. Gödel），他提出了著名的哥德尔不完备定律，研究了数理逻辑中的根本性问题：形式系统的完备性和可判断性。这些理论基础对人工智能的创立发挥了重要作用。

虽然电子计算机为人工智能提供了必要的技术基础，但直到 20 世纪 50 年代早期人们才注意到人类智能与机器之间的联系。维纳（N. Wiener）是最早研究反馈理论的美国人之一。最熟悉的反馈控制的例子是自动调温器。它将收集到的房间温度与希望的温度比较，并做出反应，将加热器开大或关小，从而控制环境温度。这项对反馈回路的研究的重要性在于：维纳从理论上指出，所有的智能活动都是反馈机制的结果，而反馈机制是有可能用机器模拟的。这项发现对早期人工智能的发展影响很大。

图灵证明了使用一种简单的计算机机制从理论上能够处理所有问题，从而奠定了计算机的理论基础。不仅如此，在 1950 年的杂志上，他预言简单的计算机能够回答人的提问，并且能够下棋。MIT 的香农（Claude E. Shannon）于 1949 年提出了国际象棋的计算机程序的基本结构。CMU（Carnegie Mellon University，卡内基·梅隆大学）的纽厄尔和西蒙从心理学的角度研究人是怎样解决问题的，提出了问题求解的模型，并用计算机加以实现。他们发展了香农的设想，编写了下国际象棋的程序。

1955 年年末，纽厄尔和西蒙编写了一个名为“逻辑专家”（Logic Theorist, LT）的程序，被许多人认为是第一个人工智能程序。它将每个问题表示成一个树形模型，然后选择最可能得到正确结论的那一支来求解问题。LT 对公众和人工智能研究领域产生了巨大影响，从而成为人工智能发展中一个重要的里程碑。

1.2.2 形成及第一个兴旺期（20 世纪 50 年代中期至 60 年代中期）

人工智能诞生于一次历史性聚会。为了使计算机变得更“聪明”，或者说使计算机具有智能，1956 年夏季，当时在达特茅斯大学的年轻数学家、计算机专家麦卡锡（后为 MIT 教授）和他的 3 位朋友：哈佛大学数学家和神经学家明斯基（后为 MIT 教授）、IBM 公司信息中心负责人罗彻斯特、贝尔实验室信息部数学研究员香农共同发起，并邀请 IBM 公司的莫尔（T. More）和塞缪尔（A. L. Samuel）、MIT 的塞尔弗里奇（O. Selfridge）和索罗蒙夫（R. Solomonff）、CMU 的纽厄尔和西蒙共 10 人，在美国达特茅斯大学举行了一个为期两个月的夏季学术研讨会（Workshop）。

自这次会议之后的 10 多年间，人工智能的研究取得了许多引人注目的成就。虽然这个领域还没明确定义，但会议中的一些思想已被重新考虑和使用了。特别是这次会议后，在

美国很快形成了 3 个从事人工智能研究的中心：以纽厄尔和西蒙为首的 CMU 研究组，以塞缪尔为首的 IBM 研究组，以麦肯锡和明斯基为首的 MIT 研究组。

CUM 和 MIT 组建了人工智能研究中心，研究面临新的挑战：下一步需要建立能够更有效解决问题的系统，如在“逻辑专家”中减少搜索、建立可以自我学习的系统。

1957 年开发了一个新程序，即对“通用解题机”（General Purpose Solver, GPS）的第一个版本进行了测试。这个程序是由制作“逻辑专家”的同一个小组开发的。GPS 扩展了维纳的反馈原理，可以解决很多常识问题。两年以后，IBM 成立了一个 AI 研究组。同时，赫伯特（G. Herbert）花了 3 年时间，制作了一个解几何定理的程序。

麦卡锡在理论研究的基础上，于 1960 年设计了 LISP 程序设计语言，适合字符串处理。字符串处理的重要性是在纽厄尔等人编写问题求解程序时得到认识的。他们那时使用的语言后来成为 LISP 的前身。LISP 成为后来人工智能研究所用语言的基础。

在 MIT，专家使用 LISP 开发了几个问答系统。博布罗（D. Bobrow）开发了解决用英文书写的代数应用问题的 STUDENT 系统。问题本身是高中程度的，采用自然语言描述。拉斐尔（B. Raphael）开发了能够存储知识并回答问题的语义信息检索系统（Semantic Information Retrieval, SIR）。如果告诉它“人有两个胳膊”“一个胳膊连着一只手”和“一只手上有五个手指”，它能够正确地回答“一个人有几个手指”等问题。虽然输入句型受到严格的限制，但它能够通过推理来回答问题。

在逻辑学方面，鲁滨逊（J. A. Robinson）发表了使用逻辑表达式表示的公理，机械地证明给定的逻辑表达式的方法，被称为归结原理，对后来的自动定理证明和问题求解的研究产生了很大的影响。现在有名的程序设计语言 PROLOG 也是以归结原理为基础的。

当人工智能各领域的基础建立起来时，美国各主要研究所开始研究综合了各种技术的智能机器人。以明斯基为指导者的 MIT、麦卡锡所在的斯坦福大学、从 MIT 转来的拉斐尔率领着 SRI（当时的斯坦福研究所，现为国际 SRI）成为研究的中心。在各研究所，智能机器人的研究目标有所不同。MIT 和斯坦福大学着重观察、识别积木和制作简单的结构件等，而 SRI 研究的机器人 Shakey 能够观察房间、躲开障碍物移动与推运物体等。给机器人下达简单的命令，如“把物体 B 拿到房间 A 去”，机器人自己就能制订详细的作业计划。研究智能机器人的目的不在于创造能够代替人工作的机器人，而在于证实人工智能的能力。同时，问题求解的理论研究也在发展，与机器人没有直接关系的复杂作业过程的研究也在发展。此外，利用积木的边线确定三维积木的理论也建立了。

这时期，最大的人工智能研究成果是涉及语义处理的自然语言处理（英语）的研究。MIT 的研究生威诺格拉德（T. Winograd）开发了能够在机器人世界进行会话的自然语言系统 SHRDLU。它不仅能分析语法，还能够分析语义解释不明确的句子，对提问通过推理进行回答。恰好在第一届人工智能国际会议召开之际，人工智能作为一个学术领域得到了承认。

斯坦福大学成立了人工智能实验室，SRI 也成立了推进人工智能课题的组织。CMU 稍微晚一些，1970 年左右开始在计算机系内研究人工智能。MIT、斯坦福大学和 CMU 被称为人工智能和计算机科学的三大中心。

1.2.3 萧条波折期（20 世纪 60 年代中期至 70 年代中期）

科学的发展往往不是一帆风顺的，人工智能也不例外。许多人工智能理论和方法未能得

到通用化，在推广和应用方面也存在重重困难。

上一个 10 年的成果丰富，初战告捷的欢乐只是暂时的，当人们进行了比较深入的工作以后，发现这里的困难比原来想象的要严重得多。就定理证明来说，1965 年发明的消解法曾给人们带来了希望，可是很快就发现了消解法的能力有限。用消解法证明两个连续函数之和还是连续函数，推导 10 万步也还没有推出来。塞缪尔的跳棋程序打败了州冠军后并没有进一步打败全国冠军。GPS 的研究也只持续了 10 年，从神经生理学角度研究人工智能的人们也发现他们遇到了几乎是不可能逾越的困难。

最糟糕的恐怕还是机器翻译。原先，人们曾以为只要用一部双向字典和某些语法知识即可很快地解决自然语言之间的互译问题。结果发现机器翻译的文字阴差阳错、颠三倒四。最著名的例子是：The spirit is willing but the flesh is weak（心有余而力不足）翻译后，会变成 The vodka is good but the meat is spoiled（酒是好的，肉变质了）。再如，英语句子“Out of sight, out of mind”（眼不见心不烦），译成俄文却成了“又瞎又疯”。

因此有人挖苦说，美国花 2000 万美元为机器翻译立了一块墓碑。1971 年，剑桥大学的应用数学家詹姆斯（James）在应政府要求起草的一份报告中指责人工智能的研究即使不是骗局，至少也是庸人自扰。这种情况使英国、美国政府撤销了所有对于学术翻译项目的资助。甚至在人工智能研究方面颇有影响的 IBM 公司也取消了该公司所有的人工智能研究项目。人工智能在世界范围内陷入困境，处于低潮。

但是，这个时期还是取得了一些重要成果。

① 1968 年，在美国斯坦福大学费根鲍姆（E. A. Feigenbaum）的主持下，第一个成功的专家系统 DENDRAL 投入使用。

② 在世界范围内成立国际人工智能联合会（International Joint Conference on Artificial Intelligence, IJCAI），从 1969 年开始，每两年召开一次国际会议。这是人工智能发展史上的一个重要里程碑，标志着人工智能这门新兴学科已经得到了世界的肯定。

③ 1970 年，国际性的人工智能杂志 *Artificial Intelligence*（AI）创刊，对推动人工智能的发展、促进研究者的交流起到了重要作用。

④ 1972 年，法国马赛大学的 A. Coherner 和他领导的研究小组研制成功了第一个 PROLOG 系统，成为继 LISP 后的另一种重要的人工智能程序语言；斯坦福大学的肖特利夫（E. Shortliff）研制了用于诊断和治疗感染性疾病的专家系统 MYCIN。

⑤ 1974 年，明斯基提出了框架理论；肖特利夫于 1975 年提出并在 IMYCIN 中应用了不精确推理；杜达（Duda）于 1976 年提出并在 PROSPECTOR 中应用了贝叶斯（Bayes）方法等。

1.2.4 第二个兴旺期（20 世纪 70 年代中期至 80 年代中期）

尽管社会的压力很大，却没有动摇人工智能研究先驱者的信念。经过认真地反思、总结前一时期的经验和教训，费根鲍姆重新举起了培根的旗帜：“知识就是力量！”他的关于以知识为中心开展人工智能研究的观点被大多数人接受。从此，人工智能的研究又迎来了蓬勃发展的新时期，即以知识为中心的时期。

从 1970 年初到 1979 年左右，人工智能得到了广泛的研究。在计算机视觉的研究中，人工智能不仅为机器人研究积木和室内景物的识别方法，还处理机械零件、室外景物、医学用相片等对象所使用的视觉信息。这种视觉信息不但包括颜色深度，而且包括不同的颜色和

距离。在机器人控制方面，人工智能通过触觉信息和受力信息来控制机械手的速度和力度。

受威诺格拉德研究的影响，自然语言研究也多了起来。与 SHRDLU 那样局限于机器人世界的系统相比，后来的研究把重点放在处理较大范围的自然语言。人在使用语言交流的时候，是以对方具有某种程度的知识为前提的。因此，会话中省略了对方能够正确推断的内容。而计算机为了理解人的语言，需要具备很多知识。因此，要研究如何在计算机中有效地存储知识，并且根据需要使用它。

在自然语言理解和计算机视觉领域，明斯基考查了知识表示和使用方法的各种实现方法，于 1975 年提出名为“框架”的知识表示方法，作为各种方法共同的基础。框架理论为许多研究者所接受，出现了适合使用框架的程序设计语言（Frame Representation Language, FRL）。转入斯坦福大学的威诺格拉德和 Xerox 研究所的博布罗共同开发了基于框架的知识表示语言（Knowledge Representation Language, KRL）。作为其应用，他们开发了用自然语言回答问题、制订旅行计划的系统。

以知识利用为中心的另一个研究领域是知识工程，如通过把熟练技术人员或医生的知识存储在计算机中，进行故障诊断或者医疗诊断。1973 年，费根鲍姆在斯坦福大学开始研究启发式程序设计计划（Heuristic Programming Plan, HPP），研究在医学方面的应用，并成功研制出了几个系统。其中最有名的是肖特利夫开发的 MYCIN 系统。MYCIN 系统是一种帮助医生对住院的血液感染患者进行诊断和用抗菌素类药物进行治疗的专家系统，用 LISP 编写。同时，与费根鲍姆等人协作，肖特利夫用 3 年时间完成了 MYCIN 系统的研究。MYCIN 系统采用与自然语言相近的语言进行对话，具有解释推理过程的机能，为后来的研究提供了一个样本。

在这样的背景下，在 1977 年第五届人工智能国际会议上，费根鲍姆提议使用“知识工程（Knowledge Engineering）”这个名词。他说，“人工智能研究的知识表示和知识利用的理论不能直接地用于解决复杂的实际问题。知识工程师必须把专家的知识变换成易于计算机处理的形式并存储。计算机系统利用知识进行推理来解决实际问题。”从此之后，处理专家知识的知识工程和利用知识工程的应用系统（专家系统）大量涌现。专家系统可以预测在一定条件下某种解的概率。当时计算机已有巨大容量，专家系统有可能从数据中得出规律，因此专家系统的市场应用很广。很快，专家系统被用于股市预测，帮助医生诊断疾病，以及指示矿工确定矿藏位置等，这一切都因为其自身存储规律和信息的能力而成为可能。

进入 20 世纪 80 年代，人工智能的各种成果已经作为实用产品出现。在实用上，出现最早的是工厂自动化中的计算机视觉、产品检验、集成电路芯片的引线焊接等方面的应用，70 年代后期开始普及。但这些都是各公司为了在公司内部使用，作为一种生产技术所开发的，而作为一种产品进入市场还是 80 年代以后的事情。例如，进入 80 年代以后，SRI 开发的计算机视觉系统被风险投资企业机器智能公司商品化。

作为典型的人工智能产品最早要数 LISP 机，其作用是用高速专用工作站把以往在大型计算机上运行的人工智能语言 LISP 加以实现。MIT 从 1975 年左右开始试制 LISP 机，作为一个副产品，一部分研究者成立了公司，最先把 LISP 机商品化。美国主要的人工智能研究所最先购入 LISP 机，用户的范围逐渐扩大。同时，各种程序设计语言和作为人机接口的自然语言软件（英语）、CAI（Computer Aided Instruction）、具有视觉的机器人等都被商品化了。在各公司内部使用的产品中，GE 公司的机车故障诊断系统和 DEC 公司的计算机辅助系统是很有名的。

此外，随着专家系统应用的不断深入，专家系统自身存在的知识获取难、知识领域窄、推理能力弱、智能水平低、没有分布式功能、实用性差等问题逐步暴露。日本、美国和欧洲制订的那些针对人工智能的大型计划多数执行到 20 世纪 80 年代中期就开始面临重重困难，已经看出达不到预想的目标。1992 年，FGCS 正式宣告失败。进一步分析便发现，这些困难不只是个别项目的制订有问题，而是涉及人工智能研究的根本性问题。

这时期，人工智能的发展涉及两个问题：一是交互（Interaction）问题，即传统方法只能模拟人类深思熟虑的行为，而不包括人与环境的交互行为；二是扩展（Scaling Up）问题，即大规模问题，传统人工智能方法只适合建造领域狭窄的专家系统，不能把这种方法简单地推广到规模更大、领域更广的复杂系统。这些计划的失败对人工智能的发展是一个挫折。于是到了 20 世纪 80 年代中期，人工智能特别是专家系统热大大降温，进而导致了一部分人对人工智能前景持悲观态度，甚至有人提出人工智能的冬天已经来临。

1.2.5 稳步增长期（20 世纪 80 年代中期至今）

尽管 20 世纪 80 年代中期人工智能研究的淘金热跌到谷底，但大部分人工智能研究者还保持着清醒的头脑。一些老资格的学者呼吁不要过于渲染人工智能的威力，应多做些脚踏实地的工作，甚至在这个淘金热到来时就已预言其很快就会降温。也正是在这批人的领导下，大量扎实的研究工作不断进行，从而使人工智能技术和方法论的发展始终保持了较快的速度。

20 世纪 80 年代中期的降温并不意味着人工智能研究停滞不前或遭受重大挫折，因为过高的期望未达到是预料中的事，不能认为是挫折。自那以来，人工智能研究便呈稳健的线性增长，而人工智能技术的实用化进程也逐步成熟。

20 世纪 90 年代以来，随着计算机网络、通信技术的发展，关于智能体（Agent）的研究成为人工智能的热点。1993 年，肖哈姆（Y. Shoham）提出了面向智能体的程序设计。1995 年，罗素（S. Russell）和诺维格（P. Norvig）出版了《人工智能》一书，提出“将人工智能定义为对从环境中接收感知信息并执行行动的智能体的研究”。所以，智能体应该是人工智能的核心问题。斯坦福大学计算机系的海斯·罗斯（R. B. Hayes）在 IJCAI1995 的特约报告中谈到，“智能体既是人工智能最初的目标，也是人工智能最终的目标。”

在人工智能研究中，智能体概念的回归并不只是因为人们认识到了应该把人工智能各领域的研究成果集成为一个具有智能行为概念的“人”，更重要的原因是人们认识到了人类智能的本质是一种“社会性的智能”。要对社会性的智能进行研究，构成社会的基本构件“人”的对应物“智能体”理所当然地成为人工智能研究的基本对象，而社会的对应物“多智能体系统”也成为人工智能研究的基本对象。

在这个时期，数据量激增推动了人工智能进一步发展，从推理、搜索升华到知识获取阶段后，又一次进化到了机器学习阶段。1996 年，人们已经系统定义了机器学习，它是人工智能的一个研究领域，其主要研究对象是人工智能，特别是在经验学习中如何改进具体算法的性能。到了 1997 年，随着互联网的发展，机器学习被进一步定义为“一种能够通过经验自动改进计算机算法的研究”。数据是载体，智能是目标，而机器学习是从数据通往智能的技术途径。Boosting、支持向量机（Support Vector Machine, SVM）、集成学习和稀疏学习是机器学习界也是统计界在近 10 年或者 20 年来最活跃的方向，这些成果是统计界和计算机界共同努力成就的。例如，数学家瓦普尼克（Vapnik）等人在 20 世纪 60 年代就提出了

支持向量机的理论，但直到 20 世纪 90 年代末才发明了非常有效的求解算法，并随着后续大量优秀实现代码的开源，支持向量机现在成为分类算法的一个基准模型。再如，核主成分分析（Kernel Principal Component Analysis, KPCA）是由计算机学家提出的一个非线性降维方法，其实它等价于经典的多维尺度分析（Multi-Dimensional Scaling, MDS）。

2006 年，多伦多大学杰弗里·辛顿（G. Hinton）教授在前向神经网络的基础上，提出了深度学习。深度学习在 AlphaGo、无人驾驶汽车、人工智能助理、语音识别、图像识别、自然语言理解等方面取得了很好进展，对工业界产生了巨大影响。

随着深度学习的兴起，人工智能迎来了它的第三波发展热潮。近年来，谷歌、微软、百度、Facebook 等拥有大数据的科技公司争相投入资源，占领深度学习的技术制高点。在大数据时代，更加复杂且更加强大的深度学习模型能深刻揭示海量数据所承载的复杂而丰富的信息，并对未来或未知事件做出更精准的预测。今天，人工智能领域的研究者几乎无人不谈深度学习，很多人甚至高喊“深度学习=人工智能”的口号。当然，深度学习绝对不是人工智能领域的唯一解决方案，二者之间无法画上等号。但说深度学习是当今乃至未来很长一段时间内引领人工智能发展的核心技术，则一点不为过。

1.2.6 中国的人工智能发展

中国的人工智能研究起步较晚，纳入国家计划的研究（“智能模拟”）开始于 1978 年；1984 年，召开了智能计算机及其系统的全国学术讨论会；1986 年起，把智能计算机系统、智能机器人和智能信息处理（含模式识别）等重大项目列入国家高技术研究计划；1993 年起，把智能控制和智能自动化等项目列入国家科技攀登计划。进入 21 世纪后，已有更多的人工智能与智能系统研究获得各种基金计划支持。1981 年起，相继成立了中国人工智能学会（Chinese Association for Artificial Intelligence, CAAI）、全国高校人工智能研究会、中国计算机学会人工智能与模式识别专业委员会、中国自动化学会模式识别与机器智能专业委员会、中国软件行业协会人工智能协会、中国智能机器人专业委员会、中国计算机视觉与智能控制专业委员会，以及中国智能自动化专业委员会等学术团体。1989 年，首次召开了中国人工智能联合会议（China Joint Conference on Artificial Intelligence, CJCAI）。1987 年，《模式识别与人工智能》杂志创刊。中国科学家在人工智能领域取得了一些在国际上有影响的创造性成果，如吴文俊院士关于几何定理证明的“吴氏方法”。

2006 年是符号逻辑（功能模拟）人工智能诞生 50 周年，中国人工智能学会和美国人工智能学会、欧洲人工智能协调委员会合作，在北京召开了“2006 人工智能国际会议”，系统总结了 50 年来人工智能发展的成就和问题，探讨了未来研究的方向。会议期间，中国人工智能学会提出了以“高等智能”为标志的研究理念和纲领，得到了与会者的普遍认同，表明中国人工智能研究已在国际上崭露头角。

高等智能的理念如下。

① 结构、功能、行为是研究智能的重要侧面，但更具本质意义的研究途径是“智能生成的共性核心机制”（探索智能生成机制的研究方法称为人工智能的“机制主义”方法）。

② 机制主义方法的技术实现是信息-知识-智能转换。

③ 通过机制主义方法，应当把人工智能的结构主义、功能主义和行为主义方法有机、和谐统一起来。

④ 通过机制主义方法可以发现和沟通意识、情感、智能的内在联系，从而打破人工智

能与自然智能之间的壁垒。

中国人工智能大会（China Conference on Artificial Intelligence, CCAI）由中国人工智能学会创办于 2015 年，每年举办一届。该会议是我国最早发起举办的人工智能大会，目前已经成为我国人工智能领域规格最高、规模最大、影响力最强的会议之一。

人工智能自 2016 年起进入国家战略地位，国家相关支持政策进入爆发期：2016 年 3 月，国务院发布《国家经济和社会发展的第十三个五年规划纲要（草案）》，人工智能概念进入“十三五”重大工程；5 月，国家发展改革委、科技部、工业和信息化部、中央网信办发布《“互联网+”人工智能三年行动实施方案》，明确提出到 2018 年国内要形成千亿元级的人工智能市场应用规模，规划确定了在 6 个具体方面支持人工智能的发展，包括资金、系统标准化、知识产权保护、人力资源发展、国际合作和实施安排；2017 年 3 月，在十二届全国人大五次会议的政府工作报告中，“人工智能”首次被写入政府工作报告；2017 年 7 月，国务院发布《新一代人工智能发展规划》；2018 年 1 月 18 日，“2018 人工智能标准化论坛”发布了《人工智能标准化白皮书（2018 版）》。

2018 年 9 月 17 日，世界人工智能大会在上海开幕，习近平致信祝贺：“新一代人工智能正在全球范围内蓬勃兴起，为经济社会发展注入了新动能，正在深刻改变人们的生产生活方式。希望与会嘉宾围绕‘人工智能赋能新时代’这一主题，深入交流、凝聚共识，共同推动人工智能造福人类。”

人工智能作为一项基本技术，在国家相关政策的支持下正在被全面推进和高质量发展，为大力提高经济社会发展智能化水平、有效增强公共服务和城市管理能力做出努力。

1.3 人工智能的主要学派

由于人们对“智能”本质的不同理解和认识，形成了人工智能研究的不同途径。不同的研究途径拥有不同的研究方法、不同的学术观点，逐步形成了符号主义、连接主义和行为主义三大学派。目前，这三大学派正在由早期的激烈争论和分立研究逐步走向取长补短和综合研究。

1.3.1 符号主义学派

符号主义（Symbolicism），又称为逻辑主义（Logicism）、心理学派（Psychologism）或计算机学派（Computerism），是基于物理符号系统假设和有限合理性原理的人工智能学派。符号主义认为，人工智能起源于数理逻辑，人类认知（智能）的基本元素是符号（Symbol），认知过程是符号表示上的一种运算。符号主义还认为，知识是信息的一种形式，是构成智能的基础。人工智能的核心问题是知识表示、知识推理和知识运用。知识可用符号表示，也可用符号进行推理，因而有可能建立起基于知识的人类智能和机器智能的统一理论体系。基于以上认识，符号主义学派的研究方法以符号处理为核心，通过符号处理来模拟人类求解问题的心理过程。符号主义主张用逻辑的方法来建立人工智能的统一理论体系，但是有“常识”问题、不确定事物的表示和处理问题等，因此受到了其他学派的批评。

尽管不是所有的人都支持物理符号系统假设，“经典的人工智能”却大多是在此指导下产生的，如启发式算法、专家系统、知识工程理论和技术等。这类方法的突出特点是将逻辑操作应用于说明性知识库中，即用说明性的语句来表达问题领域的“知识”，这些语句基

于或实际上等同于一阶逻辑中的语句，并且可以采用逻辑推导对这种知识进行推理。当遇到实际领域中的问题时，该方法需要具有该问题领域的足够的知识，以对其问题进行处理，通常被称为基于知识的方法。

符号主义的代表性成果是1957年纽厄尔和西蒙等人研制的称为逻辑理论机的数学定理证明程序（Logic Theorist, LT），成功说明了可以用计算机来研究人的思维过程，模拟人的智能活动。符号主义诞生的标志是1956年夏季的那次历史性聚会，符号主义者最先正式采用了“人工智能”这个术语。几十年来，符号主义走过了一条“启发式算法→专家系统→知识工程”的发展道路，并一直在人工智能领域处于主导地位，即使在其他学派出现后，也仍然是人工智能的主流学派。符号主义学派的主要代表人物有纽厄尔、西蒙、尼尔森（N. J. Nilsson）等。

符号主义的主要特征如下。

① 立足于逻辑运算和符号操作，适合模拟人的逻辑思维过程，可以解决需要逻辑推理的复杂问题。

② 知识可用显式的符号表示，在已知基本规则的情况下，不需要输入大量的细节知识。

③ 便于模块化，当个别事实发生变化时，易于修改。

④ 能与传统的符号数据库进行连接。

⑤ 可对推理结论进行解释，便于对各种可能性进行选择。

符号主义的主要特点是可以解决逻辑思维问题，但对于形象思维难以模拟，信息在表示成符号后，在处理或转换过程中，信息有丢失的情况发生。

1.3.2 连接主义学派

连接主义（Connectionism），又称为仿生学派（Bionicsism）或生理学派（Physiologism），是基于神经网络和网络间的连接机制与学习算法的人工智能学派。以网络连接为基础的连接主义是近年来研究比较多的一种方法，属于非符号处理范畴。这种研究能够进行非程序的、可适应环境变化的、类似人类大脑风格的信息处理方法的本质和能力。持这种观点的人认为，人的思维基元是神经元，而不是符号处理过程；对物理符号系统假设持反对意见，认为人脑不同于计算机，并提出了连接主义的大脑工作模式，用于取代符号操作的计算机工作模式。连接主义主张，人工智能应注重结构模拟，即模拟人的生理神经网络结构，并且功能、结构与智能行为是密切相关的，不同的结构表现出不同的功能和行为。目前，他们已经提出了多种人工神经网络结构和众多的学习算法。

连接主义的代表性成果是1943年麦卡洛克（W. S. McCulloch）和皮茨（W. Pitts）提出的一种神经元的数学模型，即M-P模型，并由此组成一种前馈网络。可以说，M-P是人工神经网络最初的模型，开创了神经计算的时代，为人工智能创造了一条用电子装置模拟人脑结构和功能的新途径。从1982年约翰·霍普菲尔德（J. Hopfield）提出用硬件模拟神经网络和1986年鲁梅尔哈特（D. Rumelhart）等人提出多层网络中的反向传播（Back Propagation, BP）算法开始，神经网络理论和技术研究的不断发展，并在图像处理、模式识别等领域的重要突破，为实现连接主义的智能模拟创造了条件。

连接主义的主要特征如下。

① 通过神经元之间的并行协作实现信息处理，处理过程具有并行性、动态性、全局性。

② 可以实现联想功能，便于对有噪声的信息进行处理。

③ 可以通过对神经元之间连接强度的调整实现学习和分类等。

④ 适合模拟人类的形象思维过程。

⑤ 在求解问题时，可以较快地得到一个近似解。

但是连接主义不适合解决逻辑思维问题，而且结构固定和组成方案单一的系统不适合多种知识的开发。

1.3.3 行为主义学派

行为主义(Actionism)，又称为进化主义(Evolutionism)或控制论学派(Cyberneticsism)，是基于控制论和“动作-感知”控制系统的人工智能学派。行为主义认为：人的本质能力是在动态环境中的行走能力、对外界事物的感知能力、维持生命和繁衍生息的能力，正是这些能力对智能的发展提供了基础。因此，智能行为只能在与环境的交互下表现出来。他们认为，机器由蛋白质还是各种半导体器件构成无关紧要，智能行为是由所谓的“亚符号处理”即“信号处理”而不是“符号处理”产生的。如识别熟悉的人的面孔对人来说易如反掌，但是对机器来说很困难。最好的解释是，人类把图像或图像的各部分作为多维信号而不是符号来处理，因而不需要知识表示和知识推理。

行为主义的代表性成果是布鲁克斯(Brooks)研制的机器虫。在1991年和1992年，布鲁克斯提出了不需要知识表示的智能和不需要推理的智能。他认为，智能体现在与环境的交互中，不应采用集中模式，而需要具有不同的行为模块与环境交互，以此产生复杂的行为。他认为，任何一种表达方式都不能完善地代表客观世界中的真实概念，因而用符号串表示智能过程是不妥的。这在许多方面是行为心理学在人工智能中的反映。以这些观点为基础，布鲁克斯研制出了一种机器虫，用一些相对独立的功能单元，分别实现避让、前进、平衡等基本功能，组成分层异步分布式网络，取得了一定的成功，特别是为机器人的研究开创了一种新的方法。但行为主义的研究方法同样受到其他学派的怀疑和批判，认为行为主义最多只能创造出智能昆虫行为，而无法创造出人的智能行为。

行为主义的主要特征如下。

① 智能取决于感知和行动，应直接利用机器对环境作用，以环境对作用的响应为原型。

② 智能行为只能在现实世界中，通过与周围环境交互而表现出来。

③ 人工智能可以像人类智能一样逐步进化，分阶段发展并增强。

目前，符号处理系统和神经网络模型的结合是一个重要的研究方向。如模糊神经网络系统是将模糊逻辑、神经网络等结合在一起，在理论、方法和应用上发挥各自的优势，设计出具有一定学习能力、动态获取知识能力的系统。

总之，以上3种人工智能学派将长期共存和合作，取长补短，并走向融合和集成，共同为人工智能的发展做出贡献。

1.4 人工智能的主要研究内容

人工智能是一门新兴的边缘学科，是自然科学和社会科学的交叉学科，吸取了自然科学和社会科学的最新成就，以智能为核心，形成了具有自身研究特点的新体系。人工智能的研究涉及广泛的领域，如各种知识表示模式、不同的智能搜索技术、求解数据和知识不确定性问题的各种方法、机器学习的不同模式等。人工智能的应用领域包括专家系统、博弈、定理证明、自然语言理解、模式识别、计算智能、机器人等。人工智能也是一门综合性学

科，是在控制论、信息论和系统论的基础上诞生的，涉及哲学、心理学、认知科学、计算机科学、数学和各种工程学，这些学科为人工智能的研究提供了丰富的知识和研究方法。图 1-2 给出了人工智能的研究和应用领域及其相关学科。

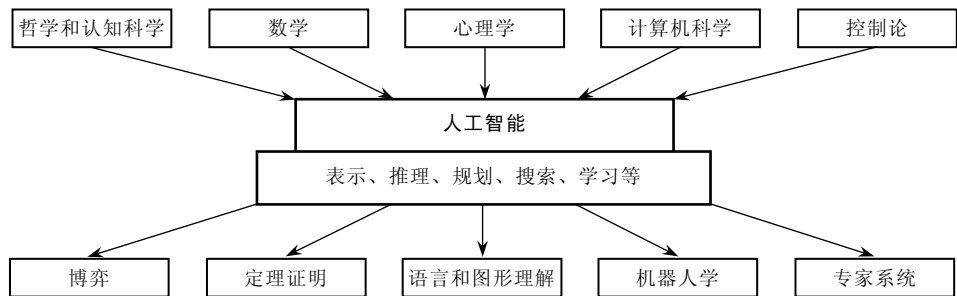


图 1-2 人工智能的研究和应用领域及其相关学科

1. 知识表示

人类的智能活动过程主要是一个获得并运用知识的过程，知识是智能的基础。人们通过实践，认识到客观世界的规律性，经过加工、整理、解释、挑选和改造而形成知识。为了使机器具有智能，使它能模拟人类的智能行为，就必须使它具有适当形式表示的知识。关于知识的表示问题是人工智能中十分重要的研究领域。

“知识表示”实际上是对知识的一种描述，或者是一组约定，是机器可以接受的用于描述知识的数据结构。知识表示是研究机器表示知识的可行的、有效的、通用的原则和方法。知识表示问题一直是人工智能研究中最活跃的部分之一。目前，常用的知识表示方法有逻辑模式、产生式系统、框架、语义网络、状态空间、面向对象、连接主义等。

2. 推理

推理是人工智能中的最基本问题之一。所谓推理，是指按照某种策略，从已知事实出发，利用知识推出所需结论的过程。根据所用知识的确定性，机器推理可以分为确定性推理和不确定性推理。确定性推理是指推理所使用的知识和推出的结论都是可以精确表示的，其真值要么为真，要么为假。不确定性推理是指推理所使用的知识和推出的结论可以是不确定的。不确定性是对非精确性、模糊性和非完备性的统称。

推理的理论基础是数理逻辑。逻辑是一门研究人们思维规律的学科，数理逻辑则用数学的方法去研究逻辑问题。确定性推理主要基于一阶经典逻辑，包括一阶命题逻辑和一阶谓词逻辑。确定性推理的主要方法包括：直接运用一阶经典逻辑中的推理规则进行推理的自然演绎推理，基于鲁宾逊归结原理的归结演绎推理，基于规则的演绎推理等。由于现实世界中的大多数问题是不能精确描述的，因此确定性推理能解决的问题很有限，更多的问题应该采用不确定性推理方法。

不确定性推理的理论基础是非经典逻辑和概率等。非经典逻辑是泛指除一阶经典逻辑以外的其他各种逻辑，如多值逻辑、模糊逻辑、模态逻辑、概率逻辑、默认逻辑等。最常用的不确定性推理方法包括：基于可信度的确定性理论，基于贝叶斯（Bayes）公式的主观贝叶斯方法，基于概率的证据理论，基于模糊逻辑的可能性理论等。

3. 搜索与规划

搜索也是人工智能中的基本问题。搜索是指为了达到某个目标，不断寻找推理路线，以

引导和控制推理，使问题得以解决的过程。根据问题的表示方式，搜索可以分为状态空间搜索、与/或树搜索。其中，状态空间搜索是一种用状态空间法求解问题的搜索方法；与/或树搜索是一种用问题规约法求解问题的搜索方法。

对搜索问题，人工智能最关心的是如何利用搜索过程尽快得到对目标有用的信息来引导搜索过程，即启发式搜索方法，包括状态空间的启发式搜索方法、与/或树的启发式搜索方法等。

规划是一种重要的问题求解技术，是从某个特定问题状态出发，寻找并建立一个操作序列，直到求得目标状态为止的一个行动过程的描述。与一般问题求解技术相比，规划更侧重问题求解过程，并且要解决的问题一般是真实世界的实际问题，不是抽象的数学模型问题。

比较完整的规划系统是斯坦福研究所的问题求解系统（Stanford Research Institute Problem Solver, STRIPS），它是一种基于状态空间和 F 规则的规划系统。所谓 F 规则，是指以正向推理使用的规则。整个 STRIPS 系统由以下 3 部分组成。

① 实际模型：用一阶谓词公式表示，包括问题的初始状态和目标状态。

② 操作符（F 规则）：包括先决条件、删除表和添加表。其中，先决条件是 F 规则能够执行的前提条件；删除表和添加表是执行一条 F 规则后对问题状态的改变，删除表包含的是要从问题状态中删除的谓词，添加表包含的是要在问题状态中添加的谓词。

③ 操作方法：采用状态空间表示和中间-结局分析的方法。其中，状态空间包括初始状态、中间状态和目标状态；中间-结局分析是一个迭代过程，每次都选择能够缩小当前状态与目标状态之间的差距的先决条件可以满足的 F 规则执行，直至达到目标状态为止。

4. 机器学习

机器学习（Machine Learning, ML）是机器获取知识的根本途径，也是机器具有智能的重要标志。有人认为，一个计算机系统如果不具备学习功能，就不能称为智能系统。机器学习是人工智能研究的核心问题之一，是当前人工智能理论研究和实际应用的非常活跃的研究领域。

机器学习的研究尚需大力加强，只有机器学习的研究取得进展，人工智能和知识工程才会取得重大突破。目前，机器学习领域的研究工作主要围绕以下 3 方面进行。

① 面向任务的研究：研究和改进一组预定任务的执行性能的学习系统。

② 认知模型：研究人类学习过程并进行计算机模拟。

③ 理论分析：从理论上探索各种可能的学习方法和独立于应用领域的算法。

机器学习有不同的分类方法，如果按照对人类学习的模拟方式，机器学习可以分为符号学习和神经学习等。

（1）符号学习。

符号学习是指从功能上模拟人类学习能力的机器学习方法，是一种基于符号主义学派的机器学习观点。按照这种观点，知识可以用符号来表示，机器学习过程实际上是一种符号运算过程。根据学习策略及学习中所使用推理的方法，符号学习可以分为记忆学习、演绎学习和归纳学习。

记忆学习也叫死记硬背学习，是一种基本的学习方法，原因是任何学习系统都必须记住它们所获取的知识，以便将来使用。归纳学习是指以归纳推理为基础的学习，是机器学习中研究得较多的一种学习类型，其任务是从关于某个概念的一系列已知的正例和反例中归纳出一般的概念描述。示例学习和决策树学习是两种典型的归纳学习方法。演绎学习是指

以演绎推理为基础的学习。解释学习是一种演绎学习方法，是在领域知识的指导下，通过对单个问题求解例子的分析，构造出求解过程的因果解释结构，并对该解释结构进行概括化处理，得到可用来求解类似问题的一般性知识。

（2）神经学习。

神经学习也称为连接学习，是一种基于人工神经网络的学习方法。现有研究表明，人脑的学习和记忆过程都是通过神经系统来完成的。在神经系统中，神经元既是学习的基本单位，也是记忆的基本单位。神经学习有多种分类方法，比较典型的有感知器学习、BP 网络学习和 Hopfield 网络学习等。

感知器学习实际上是一种基于纠错学习规则，采用迭代思想对连接权值和阈值进行不断调整，直到满足结束条件为止的学习算法。BP 网络学习是一种误差反向传播网络学习算法，学习过程由输出模式的正向传播过程和误差的反向传播过程所组成。其中，误差的反向传播过程用于修改各层神经元的连接权值，以逐步减少误差信号，直至得到所期望的输出模式为止。Hopfield 网络学习实际上是寻求系统的稳定状态，即从网络的初始状态开始，逐渐向其稳定状态过渡，直至达到稳定状态。网络的稳定性则通过一个能量函数来描述。

（3）知识发现和数据挖掘。

知识发现（Knowledge Discover in Database, KDD）和数据挖掘（Data Mining, DM）是在数据库的基础上实现的一种知识发现系统，通过综合运用统计学、粗糙集、模糊数学、机器学习和专家系统等多种学习手段和方法，从数据库中提炼和抽取知识，从而可以揭示出蕴含在这些数据背后的客观世界的内在联系和本质原理，实现知识的自动获取。

传统的数据库技术仅限于对数据库的查询和检索，不能从数据库中提取知识，使得数据库中蕴含的丰富知识被白白浪费。知识发现和数据挖掘以数据库作为知识源去抽取知识，不仅可以提高数据库中数据的利用价值，还为各种智能系统的知识获取开辟了一条新的途径。目前，随着大规模数据库和互联网的迅速发展，知识发现和数据挖掘已从面向数据库的结构化信息的数据挖掘，发展到面向数据仓库和互联网的海量、半结构化或非结构化信息的数据挖掘。

1.5 人工智能的主要应用领域

尽管目前人工智能的理论体系还没有完全形成，不同研究学派在理论基础、研究方法等方面存在一定差异，但是这些并没有影响人工智能的发展，反而使人工智能的研究更加客观、全面和深入。目前，人工智能的研究是与具体领域相结合进行的。

1. 专家系统

专家系统（Expert System, ES）是依靠人类专家已有的知识建立起来的知识系统，目前是人工智能研究中开展较早、最活跃、成效最多的领域。专家系统是在特定的领域内具有相应的知识和经验的程序系统，应用人工智能技术，模拟人类专家解决问题时的思维过程，来求解领域内的各种问题，达到或接近专家的水平。专家系统的研究起源于前述 DENDRAL 系统，与后来研制的 MYCIN 系统一起推动了专家系统技术的大发展。进入 20 世纪 80 年代后期，专家系统加快了实用化步伐。例如，据 1988 年美国的统计资料显示：1987 年的实用

专家系统为 50 个，而 1988 年达 1400 个，不包括正在研制开发中的专家系统。

目前，专家系统已广泛用于工业、农业、医疗、地质、气象、交通、军事、教育、空间技术、信息管理等方面，大大提高了工作效率和工作质量，创造了可观的经济效益和积极的社会效益。

2. 模式识别

模式识别（Pattern Recognition, PR）是人工智能最早的研究领域之一。“模式”一词的原意是指供模仿用的完美无缺的一些标本。在日常生活中，那些客观存在的事物形式可以被称为模式，如一幅画、一个景物、一段音乐、一幢建筑等。在模式识别理论中，通常把对某一事物所做的定量或结构性描述的集合称为模式。

所谓模式识别，就是让计算机能够对给定的事务进行鉴别，并把它归入与其相同或相似的模式中。其中，被鉴别的事物可以是物理的、化学的、生理的，也可以是文字、图像、音频等。为了能使计算机进行模式识别，通常需要给它配上各种感知器官，使其能够直接感知外界信息。模式识别的一般过程是先采集待识别事物的模式信息，然后对其进行各种变换和预处理，从中抽出有意义的特征或基元，得到待识别事物的模式，再与机器中原有的各种标准模式进行比较，完成对待识别事物的分类识别，最后输出识别结果。

根据标准模式的不同，模式识别技术有多种识别方法，经常采用的方法有模板匹配法、统计模式法、模糊模式法、神经网络法等。

模板匹配法是把机器中原有的待识别事物的标准模式看成一个典型模板，并把待识别事物的模式与典型模板进行比较，从而完成识别工作。

统计模式法是根据待识别事物的有关统计特征构造出一些彼此存在一定差别的样本，并把这些样本作为待识别事物的标准模式，然后利用这些标准模式及相应的决策函数对待识别事物进行分类识别。统计模式法适合那些不易给出典型模板的待识别事物。例如，对手写体数字的识别，其识别方法是先请很多人来书写同一个字，再按照它们的统计特征给出识别的标准模式和决策函数。

模糊模式法是模式识别的一种新方法，建立在模糊集理论上，用来实现对客观世界中那些带有模糊特征的事物的识别和分类。

神经网络法是把神经网络与模式识别相结合所产生的一种新方法。这种方法在进行识别之前，首先需要用一组训练样例对网络进行训练，确定连接权值，然后才能对待识别事物进行识别。

3. 自然语言处理

自然语言处理（Natural Language Processing, NLP）一直是人工智能的一个重要领域，主要研究如何使计算机能够理解和生成自然语言。自然语言是人类进行信息交流的主要媒介，但由于它的多义性和不确定性，使得人类与机器之间的交流主要依靠那种受到严格限制的非自然语言。要真正实现人机的自然语言交流，还有待于自然语言处理研究的突破性进展。

自然语言处理可分为声音语言理解和书面语言理解两大类。其中，声音语言的理解过程包括语音分析、词法分析、句法分析、语义分析和语用分析 5 个阶段；书面语言的理解过程除了不需要语音分析，其他 4 个阶段与声音语言理解的相同。自然语言处理的主要困难在语用分析阶段，原因是它涉及上下文知识，需要考虑语境对语言的影响。

与自然语言处理密切相关的另一个领域是机器翻译,即用机器把一种语言翻译成另一种语言。尽管自然语言处理和机器翻译都已取得了很大进展,但离计算机完全理解人类自然语言的目标还相距甚远。自然语言处理的研究不仅对智能人机交互有着重要的实际意义,也对不确定人工智能的研究具有重大的理论价值。

4. 智能决策支持系统

智能决策支持系统(Intelligent Decision Support System, IDSS)是指在传统决策支持系统(Decision Support System, DSS)中增加了相应智能部件的决策支持系统。IDSS把人工智能技术与决策支持系统相结合,综合运用决策支持系统在定量模型求解与分析方面的优势,以及人工智能在定性分析和不确定推理方面的优势。20世纪80年代以来,专家系统在许多方面取得了成功,将人工智能的智能和知识处理技术应用于决策支持系统,扩大了决策支持系统的应用范围,提高了系统解决问题的能力,利用人类在问题求解中的知识,通过人机对话的方式,为解决半结构化和非结构化问题提供决策支持。

智能决策支持系统通常由数据库、模型库、知识库、方法库和人机接口等主要部件组成。目前,实现系统部件的综合集成和基于知识的智能决策是IDSS发展的必然趋势,结合数据仓库和OLAP技术构造企业级决策支持系统是IDSS走向实际应用的一个重要方向。

5. 神经网络

神经网络(Neural Network, NN),也称为神经计算(Neural Computing, NC),是通过大量人工神经元的广泛并行互连所形成的一种人工网络系统,用于模拟生物神经系统的结构和功能。神经计算是一种对人类智能的结构模拟方法,其主要研究内容包括:人工神经元的结构和模型,人工神经网络的互连结构和系统模型,基于神经网络的连接学习机制等。

人工神经元是指用人工方法构造的单个神经元,有抑制和兴奋两种工作状态,可以接受外界刺激,也可以向外界输出自身的状态,用于模拟生物神经元的结构和功能,是人工神经网络的基本处理单元。

人工神经网络的互连结构(或称拓扑结构)是指单个神经元之间的连接模式,是构造神经网络的基础。从互连结构的观点,神经网络可分为前馈网络和反馈网络两种。网络模型是对网络结构、连接权值和学习能力的总括。在现有的网络模型中,最常用的有感知器模型、具有误差反向传播功能的BP网络模型、采用反馈连接方式的Hopfield网络模型等。

神经网络具有自学习、自组织、自适应、联想、模糊推理等能力,在模仿生物神经计算方面有一定优势。目前,神经计算的研究和应用已渗透到许多领域,如机器学习、专家系统、智能控制、模式识别、计算机视觉、信息处理、非线性系统辨识、非线性系统组合优化等。

6. 自动定理证明

自动定理证明(Automatic Theorem Proving, ATP)是让计算机模拟人类证明定理的方法,自动实现像人类证明定理那样的非数值符号演算过程,既是人工智能的一个重要研究领域,又是人工智能的一种实用方法。实际上,除了数学定理,很多非数学领域的任务,如医疗诊断、信息检索、难题求解等,都可以转化成一个定理证明问题。自动定理证明的主要方法包括自然演绎法、判定法、定理证明器和人机交互定理证明。

自然演绎法的基本思想是依据推理规则,从前提和公理中推出一些定理。这种方法的突出代表是纽厄尔等人研制的数学定理证明程序逻辑理论机LT等。

判定法的基本思想是对某一类问题找出一个统一的、可在计算机上实现的算法，突出代表是我国数学家吴文俊院士提出的证明初等几何定理的算法。其基本思想是把几何问题代数化，即先通过引入坐标，把几何定理中的假设和求证部分用一组代数方程表达出来，再利用代数几何中的代数簇理论求解代数方程，以证明定理的正确性。

定理证明器是一种研究一切可判定问题的证明方法，典型代表是 1965 年鲁滨逊提出的归结原理。

人机交互定理证明是一种通过人机交互方式来证明定理的方法，把计算机作为数学家的辅助工具，用计算机完成手工证明中难以完成的计算、推理、穷举等。其典型代表是 1976 年 7 月，美国的阿佩尔（K. Appel）等人合作使用该方法解决了长达 124 年之久未能证明的四色定理。这次证明使用了 3 台大型计算机，花费了 1200 小时的 CPU 时间，并对中间结果反复进行了 500 多处的人为修改。

7. 博弈

博弈（Game Playing, GP）是一个有关对策和斗智问题的研究领域。例如，下棋、打牌、战争等竞争性智能活动都属于博弈问题。博弈是人类社会和自然界中普遍存在的一种现象，博弈的双方可以是个人或群体，也可以是生物群或智能机器，各方都力图用自己的智力击败对方。

到目前为止，人们对博弈的研究主要以下棋为对象，其代表性成果是 IBM 公司研制的超级计算机“深蓝”和“小深”。“深蓝”被称为世界上第一台超级国际象棋计算机，有 32 个独立运算器，其中每个运算器的运算速度都在每秒 200 万次以上，安装了一个包含 200 万个棋局的国际象棋程序。“深蓝”于 1997 年 5 月 3 日至 11 日在美国纽约曼哈顿与当时的国际象棋世界冠军卡斯帕罗夫（Garry Kasparov）对弈 6 局，结果以 2 胜 3 平 1 负的成绩战胜卡斯帕罗夫。“小深”是一台比“深蓝”功能更强大的超级计算机，于 2003 年 1 月 26 日至 2 月 7 日，接替“深蓝”与国际象棋世界冠军卡斯帕罗夫对弈。比赛结果为平局，即在 6 局比赛中“小深”1 胜 1 负 4 平。

国内有关学者正在积极研究中国象棋的机器博弈，并于 2006 年 8 月在北京举行了首届中国象棋人机大赛，其比赛结果是机器以 3 胜 5 平 2 负的微弱优势战胜了人类象棋大师。

其实，机器博弈的目的并不完全是让计算机与人下棋，主要是为了给人工智能研究提供一个试验场地，同时证明计算机具有智能。试想，连国际象棋世界冠军都能被计算机战败或者战成平局，可见计算机已具备了何等的智能水平。

8. 分布式人工智能与 Agent

人工智能的研究和应用出现了许多新的领域，它们是传统人工智能的延伸和扩展。进入 21 世纪后，这些新研究已引起人们更为密切的关注。其中，较为突出的一个新领域是分布式人工智能与 Agent。

分布式人工智能（Distributed Artificial Intelligence, DAI）是分布式计算与人工智能结合的结果。分布式人工智能系统以鲁棒性作为控制系统质量的标准，并具有互操作性，即不同的异构系统在快速变化的环境中具有交换信息和协调工作的能力。

分布式人工智能的研究目标是要创建一种能够描述自然系统和社会系统的精确概念模型。分布式人工智能中的智能并非独立存在的概念，只能在团体协作中实现，因而其主要问题是各 Agent 之间的合作和对话，包括分布式问题求解（Distributed Problem Solving, DPS）

和多 Agent 系统（Multi-Agent System, MAS）两个领域。其中，分布式问题求解把一个具体的求解问题划分为多个相互合作和知识共享的模块或结点。多 Agent 系统则研究各 Agent 之间智能行为的协调，包括规划、知识、技术和动作的协调。这两个研究领域都研究知识、资源和控制的划分问题，但分布式问题求解往往含有一个全局的概念模型、问题和成功标准，而 MAS 含有多个局部的概念模型、问题和成功标准。

MAS 更能体现人类的社会智能，具有更大的灵活性和适应性，更适合开放和动态的环境，因而备受重视，已成为人工智能乃至计算机科学和控制科学与工程的研究热点。当前，Agent 和 MAS 的研究包括 Agent 和 MAS 理论、体系结构、语言、合作与协调、通信和交互技术、MAS 学习和应用等。

9. 智能检索

智能检索（Intelligent Retrieval, IR）是指利用人工智能的方法从大量信息中尽快找到所需要的信息或知识。随着科学技术和信息手段的迅速发展，在各种数据库中，尤其是因特网上存放着大量的甚至海量的信息或知识。面对这种信息海洋，如果还用传统的人工方式进行检索，已经很不现实，因此迫切需要相应的智能检索技术和智能检索系统来帮助人们快速、准确、有效地完成检索工作。

智能信息检索系统的设计需要解决的主要问题包括：

① 具有一定的自然语言理解能力，能理解用自然语言提出的各种询问；

② 具有一定的推理能力，能够根据已知的信息或知识，演绎出所需要的答案。

③ 拥有一定的常识性知识，以补充学科范围的专业知识，系统根据这些常识，将演绎出更一般询问的一些答案。

需要特别指出的是，互联网的海量信息检索既是智能信息检索的一个重要研究方向，也对智能检索系统的发展起到了积极的推动作用。

10. 机器人学

机器人（Robot）是一种具有人类的某些智能行为的机器，是在电子学、人工智能、控制论、系统工程、精密机械、信息传感、仿生学、心理学等学科或技术发展的基础上形成的一种综合性学科。从某种意义上说，在社会公众的认识中，机器人是一个比人工智能更容易被接受的概念。

机器人学研究的主要目的有两个。一是从应用方面考虑，可以让机器人帮助或代替人们去完成一些人类不宜从事的特殊环境的危难工作，以及一些生产、管理、服务、娱乐等工作。二是从科学研究方面考虑，机器人可以为人工智能理论、方法、技术研究提供一个综合试验场地，对人工智能各领域的研究进行全面检查，以推动人工智能学科自身的发展。可见，机器人既是人工智能的一个研究对象，又是人工智能的一个很好的试验场，几乎所有的人工智能技术都可以在机器人中得到应用。

从 20 世纪 60 年代世界上第一台工业机器人诞生以来，机器人得到了快速发展和广泛应用。从数量上，到 2005 年年底，全球机器人达千万台。从技术上，机器人的研究和发展已经历了遥控机器人、程序机器人、自适应机器人和智能机器人 4 个阶段。

遥控机器人和程序机器人是两种最简单的机器人，它们只能靠遥控装置或事先安装好的程序来控制其活动，一般用来从事一些简单或重复性的工作。

自适应机器人是一种自身具有感知能力，并能根据外界环境改变自己行动的机器人。这

种机器人配备的感知装置通常有视觉传感器、触觉传感器、听觉传感器等，机器人可通过这些感知装置获取工作环境和操作对象的简单信息，然后由计算机对这些信息进行分析和处理，并根据处理结果控制机器人的动作。自适应机器人主要用于焊接、装配等工作。

智能机器人是具有感知能力、思维能力和行为能力的新一代机器人，能够主动适应外界环境变化，通过学习丰富自己的知识，提高自己的工作能力。目前已研制出了肢体和行为功能灵活，能根据思维机构的命令完成许多复杂操作，能回答各种复杂问题的机器人。

当然，目前所研制的智能机器人只具有部分智能，真正具有像人那样的智能还需要一个相当长的时期，尤其是在自学习能力、分布协同能力、感知和动作能力、视觉和自然语言交互能力、情感化和人性化等方面，离人类的自然智能还有相当的距离。

11. 机器视觉

机器视觉 (Machine Vision, MV) 是一门用计算机模拟或实现人类视觉功能的新兴学科，其主要研究目标是使计算机具有通过二维图像认知三维环境信息的能力。这种能力不仅包括对三维环境中物体形状、位置、姿态、运动等几何信息的感知，还包括对这些信息的描述、存储、识别和理解。

视觉是人类各种感知能力中最重要的一部分，在人类感知到的外界信息中，有 80% 以上是通过视觉得到的，正如“百闻不如一见”。人类对视觉信息获取、处理与理解的大致过程是：人们视野中的物体在可见光的照射下，先在眼睛的视网膜上形成图像，再由感光细胞转换成神经脉冲信号，经神经纤维传入大脑皮层，最后由大脑皮层对其进行处理和理解。可见，视觉不仅指对光信号的感受，还包括对视觉信息的获取、传输、处理、存储与理解的全过程。

目前，机器视觉已在人类社会的许多领域得到了成功应用，如在图像、图形识别方面有指纹识别、染色体识别、字符识别等，在航天与军事方面有卫星图像处理、飞行器跟踪、成像精确制导、景物识别、目标检测等，在医学方面有 CT 图像的脏器重建、医学图像分析等，在工业方面有各种监测系统和生产过程监控系统等。

12. 进化计算

进化计算 (Evolutionary Computation, EC) 是一种模拟自然界生物进化过程与机制，进行问题求解的自组织、自适应的随机搜索技术。它以达尔文进化论的“物竞天择，适者生存”作为算法的进化规则，并结合孟德尔 (G. J. Mendel) 的遗传变异理论，将生物进化过程中的繁殖 (Reproduction)、变异 (Mutation)、竞争 (Competition) 和选择 (Selection) 引入算法，是一种对人类智能的演化模拟方法。

进化计算主要包括遗传算法 (Genetic Algorithm, GA)、进化策略 (Evolutionary Strategy, ES)、进化规划 (Evolutionary Programming, EP) 和遗传规划 (Genetic Programming, GP) 4 个分支。其中，遗传算法是进化计算中最初形成的具有普遍影响的模拟进化优化算法。

遗传算法的基本思想是使用模拟生物和人类进化的方法来求解复杂问题，从初始种群出发，采用优胜劣汰、适者生存的自然法则选择个体，并通过杂交、变异来产生新一代种群，如此逐代进化，直到满足目标为止。

13. 模糊计算

模糊计算 (Fuzzy Computing, FC)，也称为模糊系统 (Fuzzy System, FS)，通过对人类处理模糊现象的认知能力的认识，用模糊集合和模糊逻辑去模拟人类的智能行为。模糊

集合和模糊逻辑是美国加州大学扎德（L. A. Zadeh）教授提出的一种处理因模糊而引起的不确定性的有效方法。

通常，人们把那种因没有严格边界划分而无法精确刻画的现象称为模糊现象，并把反映模糊现象的各种概念称为模糊概念。例如，人们常说的“大”“小”“多”“少”等都属于模糊概念。

在模糊系统中，模糊概念通常是用模糊集合来表示的，而模糊集合是用隶属函数来刻画的。一个隶属函数描述一个模糊概念，其函数值为 $[0, 1]$ 的实数，用来描述函数自变量所代表的模糊事件隶属于该模糊概念的程度。目前，模糊计算已经在推理、控制、决策等方面得到了非常广泛的应用。

14. 人工心理、人工情感和人工生命

在人类神经系统中，智能并不是一个孤立现象，它往往与心理和情感联系在一起。心理学的研究表明，心理和情感会影响到人的认知，即影响到人的思维，因此在研究人类智能的同时，应该开展对人工心理和人工情感的研究。

人工心理（Artificial Psychology, AP）是利用信息科学的手段，对人的心理活动（重点是人的情感、意志、性格、创造）更全面地再一次机器（计算机、模型算法）模拟，目的是从心理学广义层次上研究情感、情绪与认知、动机与情绪的人工机器实现问题。

人工情感（Artificial Emotion, AE）是利用信息科学的手段对人类情感过程进行模拟、识别和理解，使机器能够产生类人情感，并与人类自然和谐地进行人机交互的研究领域。目前，对人工情感的研究的两个主要领域是情感计算（Affective Computing, AC）和感性工学（Kansei Engineering, KE）。

AP 和 AE 有着广阔的应用前景。例如，支持开发有情感、意识和智能的机器人，实现真正意义上的拟人机械研究，使控制理论更接近于人脑的控制模式，人性化的商品设计和市场开发，以及人性化的电子化教育等。

人工生命（Artificial Life, AL）是美国洛斯·阿拉莫斯（Los Alamos）非线性研究中心克里斯·兰顿（C. Langton），在研究“混沌边沿”的细胞自动机时，于 1987 年提出的一个概念。他认为，人工生命是研究能够展示人类生命特征的人工系统，即研究以非碳水化合物为基础的、具有人类生命特征的人造生命系统。

人工生命研究并不关心已经知道的以碳水化合物为基础的生命的特殊形式，即“生命之所知（Life as we know it）”，其最关心的是生命的存在形式，即“生命之所能（Life as it could be）”。应该说，生命之所知是生物学研究的主题，生命之所能才是人工生命研究关心的主要问题。按照这种观点，如果能从具体的生命中抽象出控制生命的“存在形式”，并且这种存在形式可以在另一种物质中实现，就可以创造出基于不同物质的另一种生命——人工生命。

人工生命研究主要采用自底向上的综合方法，即只有从“生命之所能”的广泛内容中去考察“生命之所知”，才能真正理解生命的本质。人工生命的研究目标是创造出具有人类生命特征的人工生命。

人工生命的研究内容主要包括计算机病毒、计算机进程、细胞自动机、人工脑和进化机器人等。其中，进化机器人不同于传统意义上的机器人，它是一种利用计算机和非有机物质构造出来的具有人类生命特征的人工生命实体。

小 结

本章首先讨论了什么是人工智能的问题。人工智能是研究可以理性地进行思考和执行动作的计算模型的学科，是人类智能在机器上的模拟。人工智能作为一门学科，经历了孕育、形成和发展几个阶段，还在不断发展。尽管人工智能也创造出了一些实用系统，但我们不得不承认这些远未达到人类的智能水平。

目前，人工智能的主要研究的学派有符号主义、连接主义和行为主义。人工智能的研究是与具体领域相结合进行的，主要包括机器学习、问题求解、专家系统、模式识别、自然语言处理、智能决策支持系统、人工神经网络、自动定理证明和机器人学等方面。

习 题 1

- 1.1 什么是智能？什么是人工智能？
- 1.2 什么是图灵测试？它有什么重要特征？
- 1.3 一台机器要通过图灵测试，它必须具备哪些能力？
- 1.4 人工智能的发展经历了哪几个阶段？
- 1.5 人工智能研究有哪几个主要学派？其特点是什么？
- 1.6 人工智能的主要研究内容和应用领域是什么？