



关于我们 (About SOP)

SOP是一个致力于心理学的专业组织。通过提供一个多维平台 (SOP BBS, SOP World, SOP Blog, SOP wiki, SOP QQ etc.), 在此之上, 广泛讨论所有与心理学有关的话题。我们希望: 无论您是一个心理学研究者, 或是一个心理咨询师, 或是一个心理学爱好者; 无论您是一个入门者, 还是一个资深的从业人士, 或者仅仅是一个试图认识自身、谋求心灵成长的朋友, 都能在SOP找到自己的地方来寻求帮助, 互相交流、互助共赢、增进了解, 从而最终推动中国心理学事业的发展。我们相信这个目标不仅仅是我们SOP GROUP, 也是所有喜欢心理学的、每一个人的心愿。。。。

SOP 搜普:

搜索最新、最专业的心理学资讯

普及以科学为导向的心理学知识

We are SOPers. We are from SOP Group.

*Our aim is to search for the newest and most professional information of psychology, to promote the development of **Scientifically Oriented Psychology** in China.*

SOP World: <http://www.sciopsy.com>

SOP BBS: <http://bbs.sciopsy.com>

SOP BLOG: <http://blog.sciopsy.com/>

SOP wiki: <http://wiki.sciopsy.com>



第一章 绪 论

现代认知心理学把人看作为一个积极的知识探求者和信息加工者，同时，在人脑与计算机之间进行类比，把人脑看作为一个类似于计算机的信息加工系统。它认为，人的认知过程，就是人对外部的或内部的信息的加工过程。因此，把人脑看作为一个信息加工系统，从信息加工的观点来探讨人类是怎样获取信息的，是怎样存储、操作和使用信息的，简言之，人类是怎样对信息进行加工的，这是现代认知心理学的基本出发点。

本章将简要地介绍一下关于信息加工系统和信息加工模型的一般知识，作为进一步阅读本书的一个引子。

一、信息加工系统的一般概念

一般而言，能够接收、存储、处理和传送信息的系统，叫做信息加工系统。关于信息加工系统究竟有哪些部分构成这个问题，人工智能和现代认知心理学的奠基者纽厄尔(A. Newell)和西蒙(H. A. Simon)在“人的问题解决”一书中，提供了一个简要的描述。他们认为，大致说来，一个信息加工系统是有以下几个部分组成的：记忆装置、加工器、接收器和效应器。图 1-1 表示这几个部分之间的关系。

记忆装置，或者叫做长时记忆，是信息加工系统的一个重要组成部分。在记忆里，存放着许许多多的符号结构。这些符号结构，是由各种各样的符号，按照一定的关系联结起来

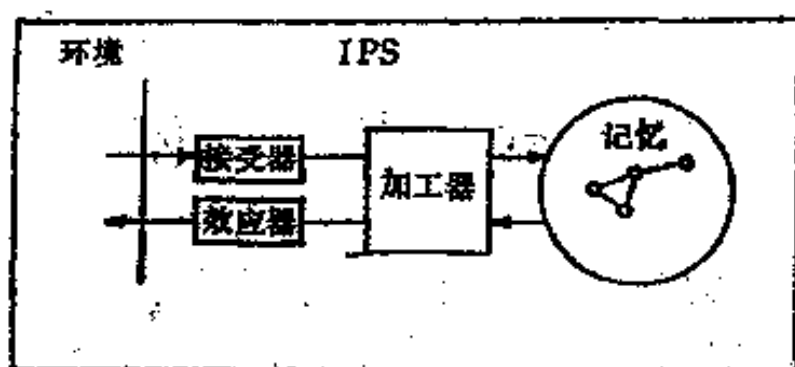


图 1-1 信息加工系统的一般结构

注：IPS为information processing system的缩写，即“信息加工系统”

组成的。它们就是信息。信息加工过程也就是对这些符号结构的操作和处理过程。所以，他们有时也把信息加工系统叫做“符号操作系统”。在这里，信息有两个不同的方面，一是作为数据，就是要加以处理的对象；二是作为程序。所谓程序，就是完成某个特定动作的一系列指令。这些指令是用来指导加工器的活动的。实际上，数据和程序没有什么本质的差别。信息加工系统的一个重要特点就是：存储在记忆中的程序也可以看作为数据，程序也是可以检查、评价和修改的。同时，程序也可以执行，这时，程序就作为指导整个系统进行工作的一系列指令。

信息加工系统的另一组成成分是加工器，它包括：（1）一套基本的信息加工过程；（2）短时记忆，它容纳进行加工的输入和输出的符号结构；（3）解释程序，它确定信息加工系统中要执行的基本的信息加工过程的序列，也就是确定先做什么，后做什么。

加工器的实际作用是执行系统的操作。它可以检查记忆中的信息，也可以执行某类操作。例如，比较两个记忆项目，看看它们是否有差别；或者决定下一步应该执行哪个操作；在记忆中进行搜索，或者进行一系列推理。一般来说，加工

器控制和执行着系统的动作。但是，加工器在执行这个任务时，是根据存储在记忆（长时记忆）中的指令来进行的，或者说，每当加工器工作时，它总是遵照特定的由指令构成的程序进行的。

信息加工系统与外部环境的相互作用，是通过“输入”和“输出”这两个过程实现的。输入是指把信息输进系统，也就是在记忆系统中建立那种代表外部事件的内部的符号结构。输出是一个相反的操作过程，它是在外部环境中建立一个反应。这种与外部环境的相互作用，是通过信息加工系统中的接收器和效应器这两个成分来实现的。信息加工是一个过程，可以分成一个个连续的阶段，其中包括信息的传入和输出。

信息加工系统这个概念，是随着计算机的发展而出现的。信息加工理论不仅把计算机看作为一个信息加工系统，而且也把人脑看作为一个信息加工系统。人脑与计算机在物质结构上是完全不同的，但它们在功能上大致是类似的，它们的工作原则，即信息加工的原则，是一致的。

二 人的信息加工系统

上面已经指出，人脑也是一个信息加工系统。人只要醒着，就不断地通过眼睛或耳朵等感觉器官接收外界的信息，进行一系列复杂的加工活动。在谈到人的信息加工系统时，也可以说，人的记忆系统就是人的信息加工系统。这是因为所有进入人脑的各种信息，都是在人的记忆系统的各个不同的部位上进行存储和加工的。图 1-2 是描述人的信息加工系统的一个概图。它简要地指明了，当一个来自客观世界的

刺激（信息）进入系统时，这个系统是怎样工作的。当然，它只是用来说明系统如何工作的一个理论模型。

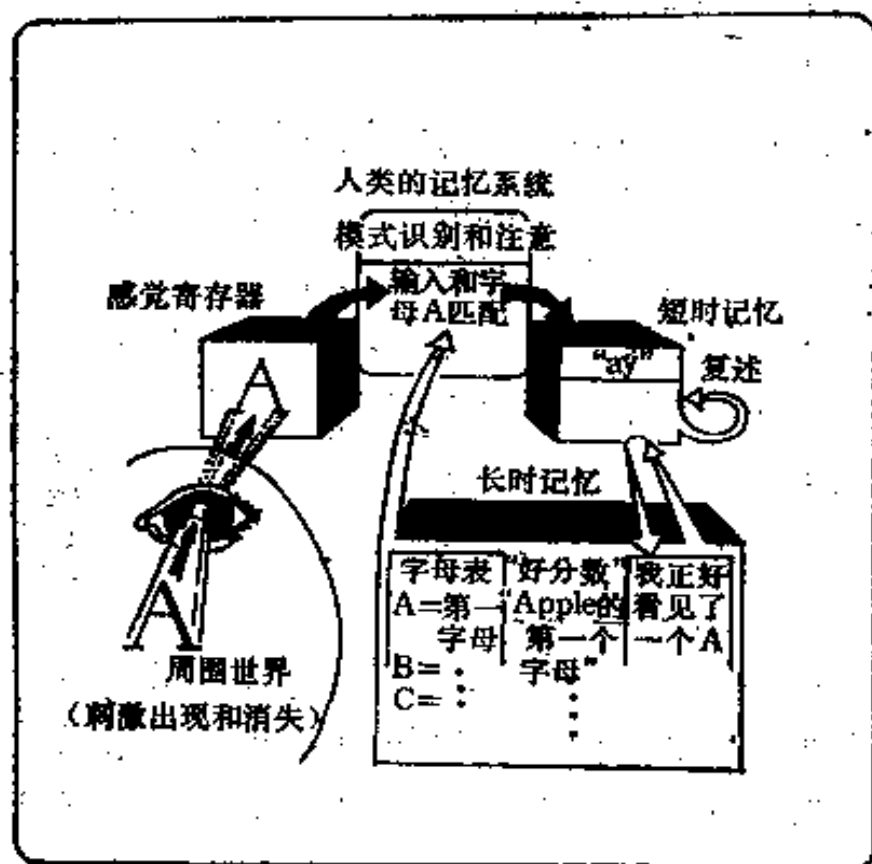


图 1-2 人的信息加工系统的一个模型

下面简要地谈一下系统包括的各个成分及其特点。

(一) 感觉寄存器

外间世界的刺激呈现后，通过人的感官而进入系统工作的第一个阶段。在这个阶段中，关于刺激的一定信息量，就被记录在系统的感觉寄存器即感觉记忆中。在图 1-2 中，刺激（字母）A 通过人的视觉通道，进入并被记录在系统的视觉寄存器中。除了视觉寄存器外，还有听觉、触觉、嗅觉和味觉等寄存器。

一般来说，一个感觉寄存器的作用是：以真实的形式（即与原初呈现的刺激几乎相同的形式），把刺激的特征短暂

地保存在系统中。接着，刺激就将转化为新的形式（通过模式识别过程），并传递到系统的另一个成分。刺激的信息停留在寄存器中，会自动地逐渐变弱，直至完全消失。这种逐渐减弱的过程，我们把它叫做“衰变”。在感觉寄存器中，这种衰变过程是非常迅速的。所以，无论在什么情况下，刺激信息在感觉寄存器中保留的时间是很短暂的，大约1秒钟左右。另一方面，由于新的刺激信息进入感觉寄存器，原有的刺激信息也会从寄存器中退出（被掩蔽或抹掉）。感觉寄存器的这种特性是很重要的。这是因为，如果感觉寄存器中刺激的“映象”没有上述这种特性，那么，我们就会经常看到一集重叠的关于客观世界的图象，而不是一些分离的景象，其结果将是，我们会对什么东西也看不清楚。

（二）模式识别

模式识别是介于感觉寄存器和短时记忆之间的一个过程。它是把进入系统的感觉信息与先前掌握的、存储在长时记忆中的信息进行匹配的过程。模式识别的目的，是把粗糙的、对系统来说相对地无效用的感觉信息，转换成某种对系统来说是有意义的东西。模式识别的一种比较特殊的含义是“命名”。它的意思是说，当我们给某个特定的视觉刺激一个“名称”时，我们也就接受了它的视觉信息，并且把它与已知的概念（例如“字母A”）等同起来，也就意味着识别了这个刺激。当然，模式识别并不始终意味着“命名”，因为在某些情况下，我们虽然不能给刺激一个适当的名称，但我们也能识别某些模式。所以，在谈论模式识别时，我们最好是在“给刺激指定意义”这个比较一般的意义上来理解它。

(三) 短时记忆

通过模式识别，刺激信息进行编码以后，就被传送到系统的另一个阶段，即短时记忆。短时记忆与感觉记忆不同。首先，在这里存储的刺激信息，已经不是关于刺激的一种粗糙的、感觉的形式了。刺激“A”已经不单是某种纯粹的视觉形象，而是作为字母A被保存下来（它可能是“A等分数”的A或英文“苹果”第一个字母的A）。其次，短时记忆与感觉记忆之间还有另一个差别，即信息保持的时间长度不同。在视觉寄存器中，一个项目衰变得很快，比如说，在1秒钟之内。但是，短时记忆中的一个项目，借助于那种叫做“复述”的过程，可以长时间地保存下去。这种“复述”过程使记忆项目一次又一次地穿过短时记忆（见图1-2），反复循环，从而使信息的强度得到更新而不产生衰变。“复述”过程本身实际上也是一种有趣的短时记忆现象。它好象是每次都在重复地把某个记忆项目重新存入短时记忆。如果没有这种“复述”过程，短时记忆中的信息也会相当快地衰变和丢失。短时记忆能保持信息的时间，大约不到1分钟。

“复述”的第二个功能是有助于信息向长时记忆传送。有人指出，记忆项目在短时记忆中逗留的时间愈长，它就愈可能被记住。也就是说，“复述”过程能够加强记忆项目向长时记忆的转移和存储，加强记忆项目在长时记忆中的强度，使该记忆项目在日后能经得起回忆（提取）的检验。

实际上，这种存储器有几个不同的名称。除了短时记忆这个名称外，人们也把它叫做“直接记忆”或“工作记忆”。其所以把它叫做“直接记忆”和“工作记忆”，是因为我们正在直接地加以思考或加工的那些东西，就处在短时记忆之

中。因此，也有人把短时记忆与人的意识同等地看待。

大多数信息加工心理学家认为，人的信息加工系统是根据产生式规则来进行加工活动的。例如：“如果A则B”是一条产生式规则，它的意思是说，如果条件A得到满足，那么，就采取B的行动。当然，这里的条件不是外部环境，而是系统内部保持的信息。这些作为人的信息加工基础的产生式活动，都是与短时记忆的活动有关的。这些产生式活动包括：（1）从短时记忆引起外部动作；（2）从感觉输入经过模式识别再送到短时记忆中保存；（3）信息在短时记忆中经过加工后，仍存入短时记忆（从第一步组合到第二步组合）；（4）在短时记忆中经过加工得出的结果，传送到长时记忆中保存；（5）从长时记忆中取出东西来，在短时记忆中进行加工。（见图1-3）由此我们可以看出，短时记忆在人的信息加工的活动中，起着非常重要的作用。

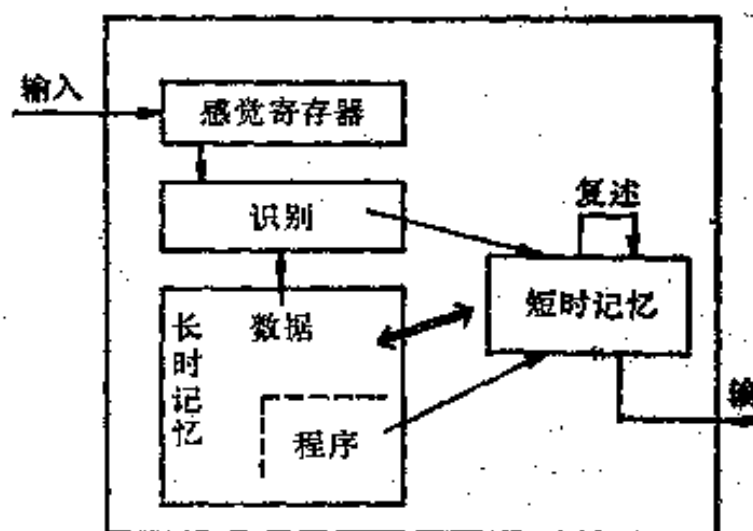


图1-3 以短时记忆为中心的系统各成分的协调活动

（四）长时记忆

可以说，长时记忆实际上是一个有关世界知识的永久性

的仓库。长时记忆以两种方式进入加工活动。第一，进入系统的刺激被识别而传递到短时记忆以后，再转移到长时记忆中。被转存到长时记忆中的信息，就可以长期地保留在人的头脑中，并成为我们的关于世界的永久性知识；第二，系统在进行加工活动时，要从长时记忆中提取（或检索）有关知识（包括数据和程序），以供加工活动使用。例如，进行模式识别时，就要从长时记忆中提取有关数据，与刺激进行匹配，以便把刺激识别为某个已知的东西。存储在长时记忆中的信息，对识别一个客体来说，起着决定性的作用。

长时记忆是一个复杂而庞大的信息库，里面存储的东西是我们所知道的关于客观世界一切事物的知识。“狗必须吃东西才能活着”，“鞋穿在脚上”等，都是存储在长时记忆中的信息。

在一个人的长时记忆中存储的信息量是非常巨大的。当人在解决问题的时候，在进行逻辑推理时，或者在回答问题和回想事实时，都要检索和使用长时记忆中存储的信息。一般说来，长时记忆中的信息是有组织的、有序的，是以一种使检索容易进行的方式来安排的。例如，对“笑”这个词的刺激，我们能在长时记忆中轻而易举地找到有关该词意义的信息，而且还能很快地找到其他有关的信息（如“露出牙齿的笑”、“哧哧地笑”）。从长时记忆中检索信息如此之快，这既表明检索不是一种偶然的过程，也表明长时记忆必定是系统中一个高度有组织的成分。

人们认为，长时记忆与感觉记忆或短时记忆不同，它是一种“永久性的”存储器。那么，为什么我们不能回忆起我们曾经知道的每一件事物（发生“遗忘”过程）呢？持这种看法的人认为，遗忘并不意味着长时记忆中的信息消失，而

只是信息的提取或检索发生了障碍。信息还是在长时记忆中，我们只是不能提取它。人们经常碰到这样的情况，有些事物一下子回忆不起来，但过了一段时间以后，又重新想起来了。这表明，信息并没有从长时记忆中消失。

另外，有人认为，在长时记忆中，存储着有关某个项目的各种不同类型的信息。例如，“火车”这个词，在我们的长时记忆中，有该词的字形和语音的信息；有关于该词意义的信息；有关于火车外形和鸣响方式的信息等等。所以，有些人认为，长时记忆应该分成各种不同的、保存这些不同类型信息的存储器：词典或词的存储器、存储概念定义的语义存储器、声音和视觉表象的存储器等等。

（五）注 意

一般来说，人们把注意看作是系统的过滤器或瓶颈。例如，当我们读一本书时，除了接受从纸张上来的视觉刺激以外，也在接受来自触觉的、听觉的和嗅觉的信息。其中，有些信息是重要的，有些则是不重要的。“注意”起着这样的作用，即把不重要的输入过滤掉，选取重要的输入作进一步的加工。由于这个原因，注意被看作为心理活动“有选择”地指向某种对象。

注意可以看作为是在系统的任何位置上都可能发生的容量的限度。这种见解是一个重要的假设，它可以把需要注意的加工与那种叫做“自动的”加工区别开来。“自动的”加工能处理所呈现的全体信息。例如，我们的耳朵自动地记录声波，不管是呈现一个、二个或更多个声音。但是，系统不能对其所有的输入进行加工，因为它的容量是有限的。这就是注意发生的所在点。当然，在不同的情境下，这个点也是

可以变化的。

语言和理解、学习和概念形成、判断和推理，以及问题解决和决策等，都是在人的信息加工系统中进行的复杂的信息加工活动，是人类的高级的认知活动。

上面，对人的信息加工系统作了简要的描述。在描述过程中，提到了三个术语：即编码、存储和提取（检索）。所谓编码，是指为了把信息放入人的信息加工系统，把信息加以修改、转换，使它适合于人的信息加工系统的过程。这就好象是，为了把一大堆数据输入计算机，就得在卡片上穿孔，使之变成计算机能识别的代码，然后再输入计算机的记忆一样。存储是指把编了码的信息放到记忆中去保持。提取指的是把存储的有关信息，从记忆中取出来，以供使用。这些术语以后将会经常用到的。

三、信息加工模型的含义

人对信息的加工是一个过程，它可以分成一系列的阶段。换句话说，正在思维着的人的行为，是由一套过程和机制产生的。这是现代认知心理学即信息加工心理学的一个重要的假定。人从接受外部世界的刺激，到作出可以观察到的反应，这一段时间可以分成一个个连续的阶段。每一个阶段相应于发生在刺激和反应之间的整个事件的一个子集。我们把这种假定的事件的一系列阶段，叫做内部加工过程的模型。

为了说明这个问题，我们可以举一个简单的例子。比如说，所有的汽车司机都记得，碰上红灯就要停住汽车。也许有人会说，这不是很简单的事吗，有什么信息加工过程可言呢？其实不然，我们可以把这个事件的整个过程分成以下几个阶段：（1）灯光记录在司机的视觉系统中；（2）司机

识别出感觉到的是什么东西——它是一盏红色交通灯。为了达到这一点，司机必须检索和利用存储在记忆中的信息，即关于红色交通灯是什么样子的知识；（3）司机要应用在他的记忆中已有的一条规则——当看见红色交通灯时，就要停车。当然，这三个阶段只是对整个过程的粗略的区分。

在上面描述的各个阶段的进程中，它的原始信息（即视觉刺激）经历了一定的转换。原始信息从视觉刺激变成一个被识别的范畴（即“红色交通灯”），然后又成为一条应用规则（当……时，就停车）的条件部分。所以说，加工过程的一个阶段，通常相应于刺激信息的某种特定的表示形式。当信息从一个阶段传送到另一个阶段时，信息的表示形式也将相应地发生变化。信息加工模型经常是用语言来表述的。我们上面所举的这个例子就是如此。但是，有时为了对行为作更精细的量的描述，人们就用数学术语来详细说明一个模型；有时也用图来表示一个模型。计算机是一个信息加工系统，人们常常用它作为手段来构造人类行为的模型，即给计算机编写程序，使它象一个模型所指出的那样去工作。这就是我们通常所说的计算机模拟。怎样编写程序，使计算机象人那样动作，或者至少类似于人的行为，这主要属于人工智能的领域。但是，人工智能与现代认知心理学是紧密联系在一起的。用计算机模拟来构造人的心理行为的模型，或者模拟人的心理行为，使计算机表现出人的智能，均已成为它们共同的目标。

现代认知心理学的主要目的，不在于去叙述关于人的一般特征，而是要用信息加工的模型，比较深入地去表述正在从事特定任务的人的行为。这种模型不是生理学的，也不是化学的或其他的，而是一种心理学的模型。当然，相对人的实际的信息加工过程来说，它仍然是一种近似的、假设性的东西。

第二章 现代认知心理学及其 主要研究方法

我们通常提到的认知心理学是一个广义的概念，它包括了许多以不同的方式进行研究的领域，例如：有些人把皮亚杰学派（Piagetian）也包括在内。我们这里所说的认知心理则是一个狭义的概念，指的是认知心理学中持信息加工观点的那一部分人。有人也把它称为现代认知心理学，或者更明确一点，称为信息加工心理学。

一、什么是现代认知心理学

给现代认知心理学下一个精确的定义，在目前似乎还无法做到，这主要是因为作为一个新的研究方向它还处在不断发展之中。这里，我们想通过对它的一般理论观点、研究的内容和通常采用的研究手段的概括性描述来勾画出它的大致轮廓。

现代认知心理学起源于五十年代的美国，有些人认为它是心理学研究中的一个新学派，也有些人认为它是一个新方向，而近年来，越来越多的心理学者更倾向于把它看成是一个新范式（Paradigm）。

范式的概念出自美国研究科学哲学的专家托马斯·库恩（Thomas Kuhn）的著作《科学革命的结构》一书。他认为，范式是一种科学研究的理论原则，是一种程式。它包括

一系列指导科学活动的规范，并决定了科学研究的对象、方法和推导方式。任何科学的发展进程在一段时期只遵循一种或主要遵循一种范式，科学的进步表现为新范式的出现、发展和最终确立从而替代旧的范式，这也就是所谓科学上的革命。正是从这样的观点出发，一些心理学家认为，现代认知心理学的出现是心理学发展中的一次革命。现代认知心理学从信息加工的观点作为新的范式代替了行为主义观点这一旧的范式。这样一次革命导致了心理学家们对心理学的研究对象和研究方法等基本问题的看法发生了重大的改变。

心理学研究中的行为主义观点曾经作为一个主要的范式，统治了心理学研究约五十年之久。它的研究重点是人的外在的、可观察到的行为，而不考虑人的内部所发生的心理过程。但是，这种内部过程确实是存在的，并对人的外在行为表现起着某种调节和控制的作用；现代认知心理学则把研究的重点放在人的内部发生的这种心理过程上，放在人对外界信息的内部加工过程上。在研究方法上，行为主义强调严格的实验室控制方法，排斥一切来自被试的主观报告；而现代认知心理学则把实验室的研究和被试的主观报告结合起来进行实验研究。行为主义的实验注重外部条件的改变所引起的人的外部行为的变化，并以此作为研究的最终目的；而现代认知心理学则只把这种外部条件的改变作为揭示人的内部心理结构的辅助手段。现代认知心理学强调人的知识和经验的重要性，认为人的行为取决于他的知识和经验。在这一点上，现代认知心理学和早期的一些心理学派有相似之处。受五十年代初期信息论、控制论、系统论和计算机科学的影响，现代认知心理学把人的知识解释为信息，把人看成是一个主动的信息加工者；把人的信息加工过程和计算机的数据处理过程

相比拟。

上面简要地描述了什么是现代认知心理学。为了更容易理解起见，下面举一个日常生活中的例子，来具体地说明从现代认知心理学家的观点出发，是如何来分析这个问题的。大多数人都会骑自行车，当我们骑着自行车在路上行进的时候，我们已经历了许多认知加工的过程，尽管这些过程中的绝大多数我们都没有意识到，但是它们确实在影响着我们的行为。我们下班骑车回家，对这条路线是这样的熟悉，我们毫无困难地回到了家里。在路上，我们似乎根本没有做任何辨认和选择，或者说我们没有意识到做了这些。根据现代认知心理学的信息加工观点，我们肯定是已经把路途中的一些标记存储在我们的记忆中了，当我们再次经过时，我们肯定是已经将一路上的这些标记同我们所存储的标记的表征 (representation) 进行了比较，这种过程叫做识别 (recognition)。

正因为我们进行了识别，所以我们才能选择正确的路线。可以认为，“表征”的问题乃是现代认知心理学的中心研究课题之一。现代认知心理学家们认为，不同的信息在人们的大脑中的表征形式是不一样的，也就是说，不同的信息有着不同的编码形式。这意味着，外界事物是根据它们的不同特征被转换成了某种内部形式而存储起来的。此外，在我们骑车的路上，我们是处在一个具有大量的不同信息的环境中，我们遇到了许多不同的汽车、其他的骑车人和行人，各种建筑物、树木和各种各样的信号等。它们中的一大部分我们几乎一过去就忘记了，如一般骑车人的正常的行为。但是，我们常常能够回忆起路上所遇到的某个骑车人的奇怪的、与众不同的行为。现代认知心理学家认为，我们对自己所处的环境中的几乎所有的信息都进行了某种程度的加工，尽管这

种加工并没有导致我们对它们的记忆，但它确实是进行了，否则，我们就不可能从众多的信息中区分出那奇怪的、与众不同的信息。这个问题表明，存在着加工信息的不同方式，它们中的某些将导致记忆，而另一些则不产生记忆。现代认知心理学家对于存在哪些不同的信息加工方式和人的认知系统是怎样驱动这些不同的方式进行工作的问题，产生了极大的兴趣。这里还要提到的是，在我们的信息加工过程中还涉及大量的决策（decision）。我们骑车时，是加速还是减速，是超过前面的骑车者还是跟在他后面，是绕道避开前面交通堵塞的路还是跟在大多数人后面慢慢推车过去，这些都需要我们做出决策，而决策又和各种各样的因素有关。例如，当我们骑车到十字路口时，交通灯由绿灯变为黄灯，我们是停下来还是一直冲过去呢？这个决策的做出可能会涉及这样的一些因素：我们可能对黄灯将持续多久有一个大致的估计，另一方向上的车辆和行人是否已经向交叉路口接近，自己的骑行速度如何，离交叉口的距离怎样；另外还可能涉及我们是否有急事或者在路口是否有一个交警在值班等等。现代认知心理学家们认为，大量的认知活动依赖于内部的决策过程，某些决策过程是有意识地进行着的，有一些是无意识的，另有一些甚至是超出人的控制以外的，但是，它们是人的心理过程的一个重要的组成部分。

根据现代认知心理学的信息加工观点，任何认知加工都是占据了时间的。许多信息加工的心理实验，都是涉及各种不同的活动占用了多少时间，得到的结论也涉及某一认知操作（operation）到底花费了多长时间。这种时间的测量被用来估计一个特定的任务到底有多复杂或者被用来估计可能涉及的有关信息加工的子操作（suboperation），如：

决策、重新编码、记忆搜索等等。

总之，现代认知心理学的信息加工观点作为心理学研究中的一个新的范式，已经对心理学的发展产生了重大的影响。尽管我们还不能说它是心理学研究的唯一正确的方法，而且心理学仍然有其它的范式存在（心理学的多重范式状态已经存在了相当长的时间）。但是，从目前的研究来看，它似乎是更便于解释许多现存的心理现象，从而使我们更容易理解许多问题和探索新的问题。现代认知心理学试图用人是一个信息加工系统这一统一的观点来解释人的知觉、注意、思维、学习甚至情绪、动机、个性等等人的心理的各个方面，进而趋向于把这些领域引向一个单一的结构中去。

上面描述了什么是现代认知心理学，在下面的几节里我们将简单地介绍一下现代认知心理学的形成和其它学科对它的影响，进一步加深我们对它的理解。

二、现代认知心理学和早期的心理学

现代认知心理学既不同于行为主义心理学，又继承了某些行为主义心理学的东西，并且还受到了早期心理学的各种流派的影响。

五十年代初期，行为主义心理学达到了它的鼎盛时期，成了心理学研究的一个主要的范式。它完全承袭了西方科学的经验主义传统，强调严格的实验室控制研究和一般性的规律的探索。行为主义心理学考虑的是人类一般行为的模式而不是某一个体的特定的行为。现代认知心理学家接受了这样的观点，他们也尽可能地把各种现象的观测建立于实验室的控制条件下，把对外部世界的分析依据于实验室中得来的结论；

他们以一种规划了的、系统的方式来探求人的行为过程的原因，认为这样得到的结论是更精确、更可信的；他们这样进行工作的目的，是要寻求人类行为产生的一般规律，而不是某个个体行为产生的含糊的解释。

但是，现代认知心理学家却反对行为主义的另一些东西。首先，他们反对行为主义的逻辑实证主义的哲学观点。本世纪三十年代，爱因斯坦的相对论导致了物理学研究上的革命，牛顿物理学的许多概念和结论失去了存在的意义。为了扫除这些错误的、过时的观念，科学哲学中的逻辑实证主义观点兴盛起来。这种观点认为，任何科学的事物只在两种形式下存在：一种是逻辑的形式，另一种是数学的形式。从这样的观点出发，科学的理论方面和经验方面就被割裂开来了，一切非形式化的科学内容便被排斥于科学大门之外。行为主义因其自身发展的需要而完全接受了这样的观点。在行为主义之前，心理学在讨论到人的身心关系时往往采用哲学上的论述，对于人的感知过程和思维过程的研究主要采用的是所谓内省的方法。不幸的是，这样的方法常常导致无结果的争辩和自相矛盾的结论。行为主义心理学家们为了扭转这种局面，采纳了已经较完善地建立起来的科学——如物理学、生物学——的工作方式。他们否定了过去心理学研究中的许多概念和结论，认为那些都是不科学的、无意义的。他们把过去给心理学家带来了巨大困难的一些纠缠不清的概念——心理、思想、意识——排除在他们的研究之外，仅仅考虑那些可观察到的、在外界刺激条件下人们所表现的外在行为。正是在逻辑实证主义的原则中，行为主义心理学家们找到了对他们这种反心理主义观点的强有力的依据，同时他们试图通过发展他们的观点使心理学从哲学的轨道中摆脱出来。

但是，随着时间的推移，逻辑实证主义的观点在科学的哲学上表现出了难以克服的逻辑困难。现代认知心理学家抛弃了它的绝大多数原则，反对行为主义的反心理主义的观点；在新的基础上，再一次走上了探索和分析隐藏在人类行为之下的思维、理解等等内部心理过程的艰难历程。现代认知心理学的信息加工观点所导致的心理学研究的革命使那些具有人类心理特征的东西再一次成为心理学研究的合理对象。

学习问题一直是许多行为主义心理学家所关注的中心课题。根据行为主义的观点，学习过程就是某种条件反射的形成过程。条件反射是最简单的、最基本的学习形式，一切复杂的行为都是由较简单的、中间的、条件反射链构成。这种原则适用于任何种类、任何形式的学习。因此，由动物实验所得来的数据便被推广到人类学习的研究中来。现代认知心理学家并不这样看待学习问题，不同意把人的学习和动物的学习混为一谈的作法。他们认为，人作为自然界的一个特殊的生物是有其自身特点的。尽管动物和人类可能具有相同的某种低级的认知能力，但是，认知心理学家更倾向于直接观察和研究人类被试者。

早期心理学研究中，语言学习问题曾经吸引了一大批实验心理学家。从事这个领域研究的心理学家们带有早期机能主义心理学的痕迹，同时，在行为主义盛行之时，他们又大量地接受了许多行为主义的观点。但是，由于他们的研究重点集中于记忆方面，因此他们更容易转向现代认知心理学的信息加工观点。他们的许多研究结果现在确实已经成为现代认知心理学的一个组成部分。

语言学习研究的鼻祖艾宾浩斯（H·Ebbinghaus）最初是采用无意义音节作为记忆的材料，而后来的语言学习心

理学家逐渐使用了有意义的词，由此，这些研究者在他们的实验中发现了学习上的组织策略，这是在动物实验中所没有的。这种发现使语言学习心理学家们认识到了人类与动物之间在学习和记忆上的不同之处，记忆的研究涉及到遗忘问题，而在对遗忘的机制进行研究时，语言学习心理学家发现了一个新的记忆过程，即现在称之为短时记忆的过程。这一发现促进了现代认知心理学的信息分阶段进行加工的思想萌生和发展。现代的语言学习的研究已经同长时记忆和短时记忆的研究溶为一体了。

另外，格式塔心理学也同样对现代认知心理学的发展产生了影响。格式塔心理学的研究侧重于人的知觉和高级心理过程，他们采用综合的方法，强调模式、组织、结构的作用，这正是现代认知心理学的基本观点。格式塔心理学反对行为主义把人看成是被动的刺激反应器，强调人对感觉输入的组织和解释的主动性，这也是现代认知心理学家的看法。

在方法论上，格式塔心理学主张研究直接的生活经验，主张把直接的生活经验的材料同实验室的数据结合起来。例如，他们重视被观察者直接描述自己的知觉内容，并把这种方法称为现象学方法。这既不同于冯特（W. Wundt）等人的严格训练被试的内省，也不同于只重视实验室研究的行为主义，这个原则正是现代认知心理学的基础。

“反应时”是传统实验心理学常用的一个概念和方法，在行为主义统治心理学的时期用得较少了。随着现代认知心理学的兴起，对反应时的研究又重新活跃起来。利用“反应时”作为参变量，现代认知心理学家在关于人的知觉、记忆等方面的研究中取得了出人意料的进展。

“反应时”作为心理学研究中的主要参变量之一，早在

上个世纪中叶就引起了心理学家的重视。1850年赫姆霍兹(H. Helmholtz)就成功地测定了蛙的运动神经的传导速度,后来又测定了人的神经传导速度。1868年,荷兰生理学家顿德斯(F. C. Donders)在对人的辨别和选择等心理过程的生理时间(生理学家有时称反应时间为生理时间)进行研究时,创立了选择反应时的实验方法。他把反应时研究划分为三种,一种是简单反应时,再一种就是所谓选择反应时,最后一种是在多种刺激条件下只作一种反应的反应时。例如:有三种条件,条件1出现,被试做一种反应,条件2出现,被试不做反应,条件3出现被试做与条件1相同的反应。我们可以把这三种反应时概括如下:

表2-1 反应时研究的分类

课题	刺激数	反应数	心理过程
1	一个	一个	简单反应时
2	多个	多个	简单反应时、刺激分类、反应选择
3	多个	一个	简单反应时、刺激分类

由表2-1中可以看出,第一种反应时是第二种反应时中的一部分,因此,第二种反应时比第一种反应时的时间要长,因为多了两个认知过程。同样可以看出第三种反应时也是第二种反应时的一部分,同时又包含了第一种反应时,由此把这三种反应时进行比较,采用相减的方法就能得出各个认知过程的反应时间。例如:刺激分类时=第三种反应时减第一种反应时,反应选择时=第二种反应时减第三种反应时。这就是顿德斯的减法法。这种方法后来被现代认知心理学家有效地用于研究人的认知加工过程,并有了新的发展。

冯特是现代心理学的创始人，他关于心理学的研究对象和方法的观点与现代认知心理学家的看法有其相似之处。他认为心理学的研究对象是经验，是人的意识的内容；研究的方法是实验控制下的内省。因此，有些心理学家在评论现代认知心理学时说，现在的心理学又返回到冯特时代了，不同之处只是实验方法更加可靠，更加精巧了。事实上，冯特的一些观点和某些实验方法确实为当今的现代认知心理学研究者们所采用了。但是，这些研究是建立在不同的理论基础之上的，冯特的构造主义思想已经被分阶段的信息加工思想所代替。

三、现代认知心理学与语言学

语言学和心理学是两个不同的研究领域，在历史上也有其不同的发展道路。语言学家大多注重语言的抽象的一面，即语言的形式和规则；心理学家注意的是语言的应用方面，即语言的功能。然而，在五十年代后期，语言学家乔姆斯基(N. Chomsky)根据他的研究提出了一个语言学的新理论，并对心理学中行为主义关于人类语言的产生和形成的理论提出了直接而尖锐的批判。由此，乔姆斯基的语言学理论开始对现代认知心理学的发展产生重大影响，并导致了当今已被认为是现代认知心理学的一个重要组成部分的新研究领域——心理语言学的产生。

乔姆斯基反对行为主义关于人类语言问题的两个基本假设。假设之一是，行为主义心理学家认为人类语言是一个词的反应链，在一句话中，每一个词都是其后继词的刺激物。这样前一个词的出现导致后一个词的出现，最终便产生出一

个完整的句子。这种基本的条件反射过程被认为是人类语言产生的基础。乔姆斯基认为，行为主义的语言理论过分强调了语言表述上的顺序性，忽视了语言在语法结构上所表现出的层次性。如果遵从行为主义的观点，那么支配词与词之间组成句子的某种句法规则就成了顺序联想词链的偶然产物。乔姆斯基指出，许多句子是受到某种规则支配的，这些句法规则并不要求词和词之间的顺序相连，例如：“如果……就……”，这两个词是相关的，我们由前一个词的出现可以预测到后一个词的出现，但是这两个词常常是被许多其它的词隔开的，而且，许多合乎语法的句子也没有高概率的顺序。

在语言获得问题上，行为主义认为，语言的获得是由于受到自己父母或他人的强化，儿童在学习说话的过程中，说对了时受到表扬和肯定，说错了时受到批评和纠正，儿童的语言就是通过这一过程的反复进行，按照强化、消退惩罚、泛化这一系列条件反射形成的总原则而获得的。乔姆斯基反对这种观点，他认为，儿童在进行语言学习的过程中，学到的是一系列复杂的语法规则而不仅是词链，正是这些语法规则才能很好的说明儿童在语言运用上的创新性，而行为主义的理论是难以解释这种创新现象的。

行为主义语言理论的第二个假设是，任何语言行为都能够在不需要建立抽象的、不可观察的系统的条件下而得到解释。乔姆斯基则恰恰相反，他认为，语言问题的解释必须用一整套的规则和某种抽象水平的，非直观的语言体系来概括。由此出发，乔姆斯基的语言理论首先区分了语言能力和语言作业。所谓语言能力是指关于语言本身的知识，而语言作业是指生活中所使用的话语。他认为，语言学理论应当解释语言能力而不是语言作业。语言学家的任务就是写出一部完善

的生成语法，它由一整套规则组成，使用这样一套规则可以产生出所有合乎语法的句子，而不只是实际生活中运用的语句。人们在语言运用上的好坏依赖于语言能力的高低。这些语言规则可以有效地运用于从未学过的新的语言形式上，这就体现出人们在语言运用上的创新性。

乔姆斯基的理论给语言学家们提出了一项艰巨的任务，这就是写一部完整的、抽象的规则系统来概括人类所使用的语言，这个系统应适用于所有的人类的信息交往方式。

在语言学理论上，乔姆斯基的另一个重要观点是区分句子的深层结构和表层结构。前者是指语言的意义，后者指的是表达语言意义的形式。例如，在英语中的一句话 “They are eating apples”，这是一个含糊的句子，我们既可以理解为“他们正在吃苹果”，也可以理解为“他们是就要吃的苹果”，为什么会出现这种现象呢？这是因为这个句子有一个表层结构——词序，但有两个不同的深层结构——含义。当然在不同的情景下我们是可以区分出这一句话到底要说的是什麼。例如，中国话中“我们完了”这样一个句子，在一群工人完成了他们的任务的情况下，我们理解这句话是在说“我们的工作结束了”；当在一伙敌人面临着失败的情况下，我们理解这句话是在说，“我们完蛋了”。另一方面，通过区分表层结构和深层结构，我们还可以了解到具有不同的表层结构的句子是如何表达同一种意义的。例如，“我打碎了瓶子”和“瓶子被我打碎了”这样两句话具有相同的意义，即具有相同的深层结构，但在词序上是不同的，即具有不同的表层结构。对这种情况，乔姆斯基提出了所谓转换规则，也就是把句子由主动式变为被动式。根据不同的转换规则，我们可以得到同一种意义的不同形式的句子。

乔姆斯基的语言理论通常也被人们称之为生成转换语法。所以叫生成的，是因为它是一套可以“生成”任何可能的句子的规则。所以叫转换的，是因为通过这样一套规则体系可以把一个句子转换成另外一种形式而不改变这个句子的意义，或者说深层结构。这个理论所阐述的一些原则是否具有心理意义上的现实性，还有待于心理学家们去逐步加以验证，但由此产生的心理语言学对现代认知心理学的影响和推动作用确实是巨大的。语言问题是现代认知心理学研究的核心问题之一。语言不仅作为人类的交际工具是重要的，而且，语言的研究与人的思维、知觉、记忆等方面的研究也是密切相关的。乔姆斯基的理论使心理学家们认识到行为主义在人类语言研究上的困境，由此推论行为主义也将难以解释人类的认知活动。虽然乔姆斯基的理论给了心理学家很多的启示，但是一个理论被推翻最重要的是要有一个新的理论去代替它，而乔姆斯基并没有给实验心理学家提供这样的替换物。他的理论主要是对语言学研究的而不是心理学的。导致现代认知心理学产生、发展并最终代替了行为主义而成为心理学研究的主要范式的决定性推动力，是二次世界大战之后一系列重要的技术性科学的出现，特别是计算机科学的产生和发展，它们促成了现代认知心理学的诞生。

四、现代认知心理学与技术科学

翻开心理学的历史，我们就会发现，心理学作为整个科学体系中的一部分始终受到其它科学发展的影响。1879年在德国的莱比锡，冯特建立了第一个心理学实验室，作为心理学发展史上的第一个里程碑，它的出现正是当时物理学、生

物理学发展在心理学研究中的反映。前面我们已经提到，作为心理学研究中的一个重要派别——行为主义的出现也同当时物理学研究中的革命有着密切的关系。科学发展到今天，各个学科之间的互相渗透，互相影响，互相推动，以各种不同的方式，在各个不同的方面更加明显的表现出来。四十年代信息论、控制论、系统论以及计算机科学的出现和发展对心理学产生了重大的影响。现代认知心理学的信息加工观点正是在这个基础上发展起来的。

美国科学家申农 (C. E. Shannon) 在关于通讯效率研究的基础上创立了信息论。与此同时，美国另一位从事火炮追踪目标研究的科学家维纳 (N. Wiener) 创立了控制论，控制论后来又成为系统论的基础。

信息论是一种给信息和它的不确定性进行定量分析的理论。它能够说明输入和输出之间的某种关系，并可以用来确定一定情景下信息量的大小，它涉及到通信信道、信道容量和传递函数等一系列问题。

可否把人看成为一个通信信道呢？早期的认知心理学家正是采用了这样一种比拟。人可以接受各种各样的环境信息，这相当于通信信道接受了输入；人根据他所接到的信息而做出各种各样的反应，这相当于通信信道的输出；人在一定环境条件下做出特定的反应，这似乎也可以用某种传递函数来表达；人把所接受的外界信息转换成了某种内部表征的形式来进行传递，这好象类似于通信中的编码和译码过程。心理学家们从信息论的研究中借用了许多有用的概念和方法，如信息、信息量以及测量信息量的数学方法，利用代码、编码、译码等概念来研究人的感知觉及其有关的性质。第一个把信息论的概念引入心理学研究中的是哈佛大学的米勒 (G.

A. Miller)。之后，信息论在心理学研究中曾风行一时，后来其影响有所下降。主要原因是，这个理论虽然能较好地说明一个被动的信息传输系统，但对于人这样一种特殊的信息传输系统的能动方面却不能做出很好的解释，忽视了人的内部心理过程的积极的、主动的方面。正是基于这样的认识——人是一个信息传输系统，但又不是一个被动的信息传输系统，人的内部表征和内部心理过程对外界接受的信息进行了积极、主动的加工，心理学家们才从单纯的信息论观点过渡到信息加工的现代认知心理学观点。因此，信息论的观点对于现代认知心理学的产生和发展确实起到了抛砖引玉的重要作用。米勒等一批在心理学研究中曾大力提倡过信息论思想的心理学家大都成为早期认知心理学的代表人物。时至今日，信息论的许多概念仍然是现代认知心理学中的基本的、重要的概念。

控制论主要研究系统的自我调节和控制过程，认为系统的控制是通过信息反馈而实现的。反馈的概念是控制论的一个基本概念，这个观点大大推进了自动化系统、计算机技术和人类生理心理过程的研究。控制论的创立者维纳本人曾直接论述过人类神经系统与计算和控制机器之间的类比，还谈到了神经生理学研究对控制论思想形成的重要影响。控制论的模型被引到心理学的研究中，改变了神经生理学和心理学研究中传统的反射弧概念，而用反射环的概念予以替代，从而能更好地解释人类行为的自我调节过程。

在美国，首先把控制论思想引入心理学的是认知心理学家米勒、加兰特(E. Galanter)和普利布莱姆(K. H. Pribram)。他们在1960年出版了《计划和行为结构》一书。书中用控制论的思想解释了人的问题求解过程，提出了问题解决

的测试操作(Test-Operate-Test-Exit,简称为SOTE模型。它已成为现代认知心理学关于问题解决研究的经典模型。

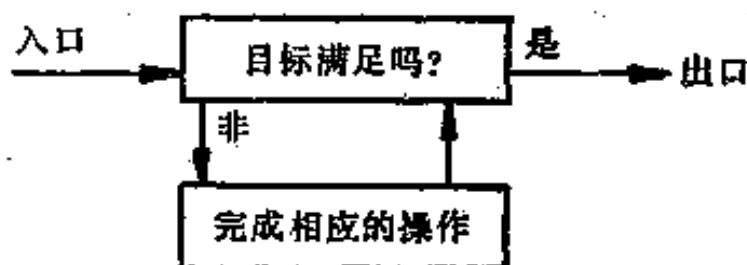


图 2-1 TOTE模型示意图

以控制论为基础的系统论在近代心理学研究中也产生了一定的影响。早在五十年代，美国一批心理学家就倡导把一般系统论的思想应用于心理学研究中。近年来苏联心理学家也开始大力提倡吸收系统论的思想，总的观点是，应当把人作为一个完整系统包括他周围的环境一起来进行研究，从生理的、神经的、心理的和行为的等各个方面进行综合考察。

信息论、控制论和系统论三者是密切相关的，它们的思想对心理学的影响是多方面的、深远的。它们从客观上推动心理学家走上了信息加工的道路，也就是现代认知心理学的道路。

计算机科学的思想对于现代认知心理学的产生发挥了直接的作用。利用机器来代替人，从事数值计算工作，这样的思想由来已久。但是，只是到了本世纪三十年代，数学家图灵(A·M·Turing)的自动机理论才为计算机器的问世打下了理论基础，而这些又都是建立在形式逻辑和数理逻辑的基础上的。

形式逻辑和自动机理论的符号和符号运算思想，给心理学家们以很好的启示。他们想到，人类的认知系统也可以看

成是一个符号运算系统。人类的思想和概念可以用符号来表示，而且这些符号可以通过确定的符号运算过程有意义的加以变换。这种把人类的心理过程抽象化，进而看成是符号运算系统的思想，正是现代认知心理学的理论核心。

在图灵的理论基础上，美国数学和计算机科学家冯·诺依曼（Von Neumann）设计了第一台通用电子计算机，物理符号的运算过程由此而具体化了。长期以来，许多心理学家面对人脑这个“黑箱”而束手无策，到了今天，他们突然发现人类的心理过程或者说认知系统可以看成是符号运算系统，在这一点上人类的大脑竟然可以和计算机相通起来，心理学家们不禁为之振奋起来。第一本认知心理学专著的作者奈瑟（U·Neisser）指出，计算机出现后，内部心理过程和状态的分析，突然不再是某种可疑的或矛盾的事了；认知心理学恢复了被行为主义割断了将近半个世纪之久的、与早期实验心理学的心理主义方向的联系，保持了新行为主义严格的假设演绎法，增加了机器模拟法。这样，就在认知过程的分析方面扩大了研究课题。

在心理学界首先倡导这种思想并从事研究的是卡内基—梅隆大学的纽厄尔（A·Newell）和西蒙（H·A·Simon）。五十年代后期，纽厄尔和西蒙等人召开了有关信息加工和问题解决的讲习会，宣传了信息加工的观点，与会者大都是当时在心理学研究中处于领先地位的心理学家，而这些人后来又都成了现代认知心理学的中坚人物。1956年纽厄尔和西蒙等人编制了第一个模拟人类解决问题过程的计算机程序“L T”，尔后他们又和肖（J·C·Shaw）一起发表了论文“一个人类问题解决理论的基本原理”，他们的工作证明了，利用机器来模拟人的心理过程并由此来揭示人的心理活动的规

律的可行性，从而使心理学研究在方法论上作出了新的突破。

最后需要指出的是，把计算机和人的认知过程进行类比，只是指在某一种水平上的类比，即在计算机程序水平上描述人的内部心理过程，它主要涉及的是人和计算机的逻辑功能方面，而不是计算机硬件和人脑的类比。不同的认知心理学家在不同的程度上使用计算机的类比。有些人明确认定，计算机和人的某些方面是同一种系统的两个事例。另一些人则只是把人的心理过程和计算机内部的信息流程作粗略的比拟，而大多数研究者采取的是中间立场。

五、现代认知心理学的主要研究方法

在心理学研究的历史上，各种学派的出现都曾以它特有的研究方法和研究角度表现出它的特征。现代认知心理学也不例外，它的研究以所谓抽象分析法（也有人称之为“会聚性证明法” convergent validation）为核心，采用计算机对人的心理过程进行模拟，并辅之以其它的实验分析，通过综合，以推理判断的方式得出某些推断性的结论。

我们可以用一个近似的比喻来说明这种研究方法。假如我们买来一架新机器，想了解这部机器的内部构造，但又不想把它拆开而损坏它，这时我们可以给这架机器一些输入来观察它的输出情况，并根据这些数据来推断这部机器的内部构造，再根据这种推断，制造相应的模型，然后在相同输入的情况下来检验这个模型是否也有同样的输出状态。如果情况不符，那么就说明我们制造的模型与它的原型是有差别的，然后再加以改进，重新进行上述的检验。这种过程反复多次

之后，就能够使我们制造的模型越来越接近于它的原型。抽象分析的研究方法正是基于这样的思想，只不过它的研究对象是人，是人的心理过程，它比机器更复杂。

了解人的心理机制和规律是心理学家们长期追求的目标。但是，由于其研究对象的特殊性和复杂性，使得这种研究困难重重。人们的内部心理过程是无法直接观察到的，心理学家们只能通过可观察到的对其研究对象的刺激输入和它的输出反应来推测内部发生的过程，而这种推测又往往采用描述性的语言来表达，这种表达的含糊性和不确定性常常导致重复验证时的困难。

现代认知心理学采用计算机模拟这一重要的工具，他们通过人的实验来积累数据，并经过各种分析得出某些确定性的推断，然后把所得到的这些结论用计算机程序的方式加以描述并输入计算机中。而后，在相同的输入刺激条件下来观察计算机的输出状态。如果计算机的输出状态同人类被试的输出反应一致，那么就说明用以描述人的心理过程的这种计算机程序在机能上与人的内部过程有其相似之处；如果情况不是这样，那么也就表明，据以构成这个程序的关于人的心理过程的心理学理论是不完善的，是需要修正的。通过这种方式来修改理论，往往使我们容易抓住问题的关键所在，从而集中注意力来研究这个不完善的地方，得出新的推断，进一步再重复我们的检验过程，最终得到比较精确的，可信度较高的理论。

下面我们举一个比较著名的计算机模拟的例子来说明采用这种方法所得到的令人惊讶的结果。

1964年，费肯鲍姆（F.A. Feigenbaum）等人发表了关于EPAM程序的报告，这是一个叫作初级的知觉者和记忆者

(Elementary Perceiver and Memorizer) 的程序。它是用来模拟包括人类的联想学习和记忆结构等方面内容的程序。在对偶联想学习的模拟中，程序是这样工作的：它的一部分首先模拟实验工作，类似一个记忆鼓。先呈现第一个对偶的刺激部分，如：TOZ；然后把刺激部分和反应部分一起呈现，如：TOZ-3；接着再呈现第二个对偶的刺激部分以及相应的刺激和反应部分；这样继续下去，一直到把字表上各个无意义音节和数字的对偶全部呈现完为止。模拟程序的另一部分的任务就是当刺激呈现时，预期（实际上是打出）正确的反应。

EPAM是通过建立起一个分类树或辨别网络来进行学习的。分类树是一个连续加工系统，所谓连续加工，就是说一次只考虑一个编码的信息。例如，一个程序要区分字母A、H、V和Y，它是根据三个特点的有无来进行的。这三个特点是：上部呈凹状，一条横线，一条竖线。这个连续的加工过程如图2-2所示。

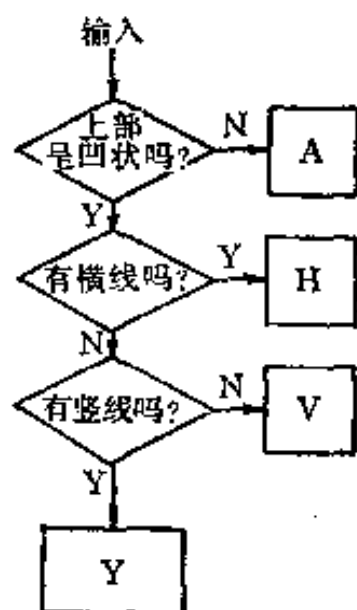


图 2-2 分辨A、H、V、Y四个字母的连续加工程序

它首先检查：“上部是不是呈凹状？”如果不是，那就是字母A；否则就继续检查，“有没有一条横线存在？”如果存在，那么就是字母H；如果没有，就再接着检查，“有没有一条竖线存在？”如果有，就是字母Y，若没有，就是字母V。这就是EPAM的分类原理。

现在让我们来说明一下EPAM的对偶联系学习的分类树是怎样建立的。假设对偶联系的字表是由TUZ-3，MEF-1，TOK-2，TUQ-4和PUZ-5组成。又设TUZ-3是第一个呈现的字母数字对偶，并且被试者是首先根据刺激部分的第一个字母来建立分类树的。图2-3中所示的在TUZ-3呈现以后所建立起来的分类树。也就是说，如果刺激部分的无意义音节的第一个字母是T，则以“3”反应就是对的。如果下一对呈现的是MEF-1，这时被试者（实际上就是计算机程序）就要把图2-3中的分类树加上一个负的分枝，这就是说，如果刺激部分的第一个字母是T，就以“3”反应，不是T就以“1”反应，如图2-4所示。当TOK-2呈现时，被试者又要进一步考虑：如果刺激部分的第一个字母是T，那还要看第二个字母是不是O，如果不是字母O，就以“3”反应，否则就以“2”反应，如图2-5所示。当TUQ-4被呈现时，如果根据已建立的分类树（第一个字母是T，并且第二个字母是U就以“3”反应），就应以“3”反应，但此时这是不对的，这时还必须检查刺激部分的第三个字母，如果是Q，就以“4”反应，如果不是Q，就以“3”反应，如图2-6所示。最后，当PUZ-5呈现时，如果认为第一个字母不是T（正如我们先前建立的那样）以“1”反应就不正确了。现在应对原来建立的分类树进行修改，如果刺激部分的第一个字母不是T并且也不是P时才能

以“1”反应，如果是字母P则应以“5”来反应，如图2-7所示。这就是EPAM的学习过程。

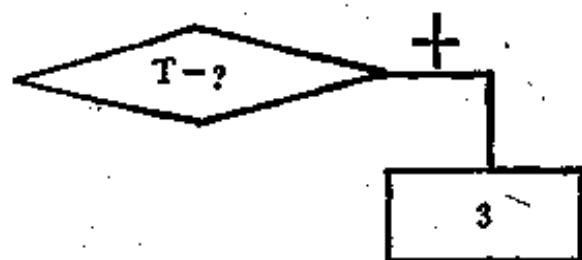


图 2-3 TUZ-3 呈现时建立的分类树

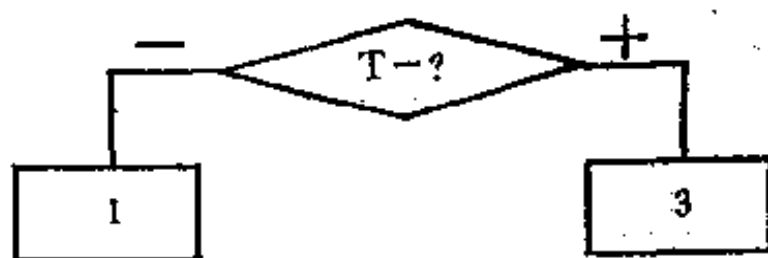


图 2-4 MEF-1 呈现后建立起来的分类树

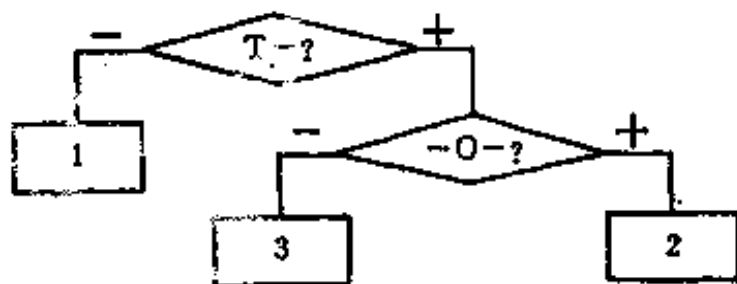


图 2-5 TOK-2 呈现后建立起来的分类树

EPAM程序通过建立这种分类树来进行学习，达到了人类被试那样的学习效率，但是在结果分析中发现，EPAM却不依赖于许多心理学实验中所使用的相似性、熟悉性和意义性等实验变量。它证明了，这些变量对于这个课题来说并不是至关重要的，最重要的是刺激单元必须是记忆中存储的综合模块。

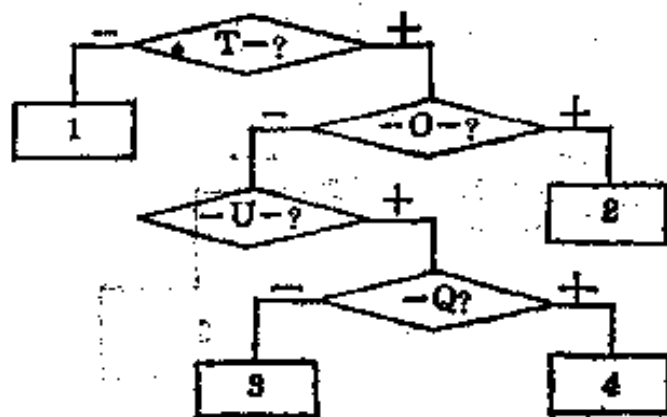


图 2-6 TUQ-4 呈现后建立起来的分类树

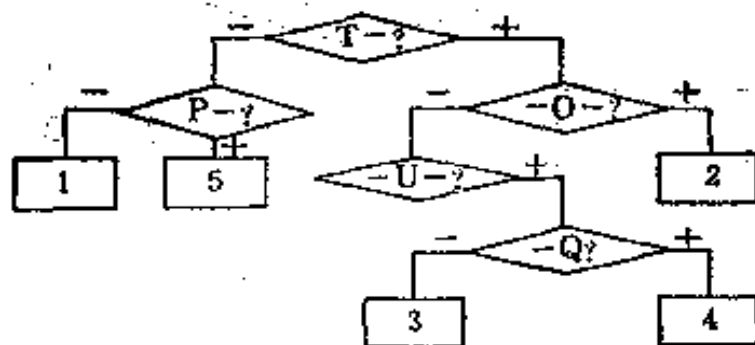


图 2-7 PUZ-5 呈现后建立起来的分类树

用模拟的方法来研究心理学问题，改变了过去实验心理学以分析为主的状况，如果对心理现象只进行分析，心理学家的的工作只不过做了一半，一个常常被忽略而非常重要的工作是用分析出来的单元把行为综合或者重建起来。计算机就是一个帮助我们进行这种重建或综合的工具。利用计算机模拟人的行为，在模型中不可能有含糊的假设，因为如果在程序上不把要求写清楚，计算机就不可能按要求运行。通过计算机的模拟可以看到，高级心理活动包括思维在内，并不是那样神秘和复杂，以致于不能确切的模拟。计算机和信息加工为心理学的解释工作提供了如此生动的比喻，编了程序的

机器可以对外界刺激进行觉察、确认、比较和辨别，进而能够思考、推理和解决问题。一系列的信息加工概念已经成了现代心理学的标准术语，完全替代了过去在实验心理学中占统治地位的“刺激—反应”语言。

在现代认知心理学中流程图的概念也被许多研究者所采用。流程图在计算机科学中也常常称之为粗框图，它被用来描述要编制的计算机程序的主要功能和特点，但并不具有计算机程序运行的细节。流程图可以通过细化而进一步变为计算机程序。在现代认知心理学的研究中，特别是一些从事心理过程的理论研究的学者们，常常用流程图来表示信息在人脑中的流动和加工过程，用以说明他们的理论，而且，许多认知心理学家也并不要走到把他们的理论编制成具体的计算机程序这一步。

下面我们用一个例子来说明，在现代认知心理学的研究中，心理学家是如何使用流程图说明他们的理论模型的。图 2-8 表示一个有代表性的记忆模型流程图。

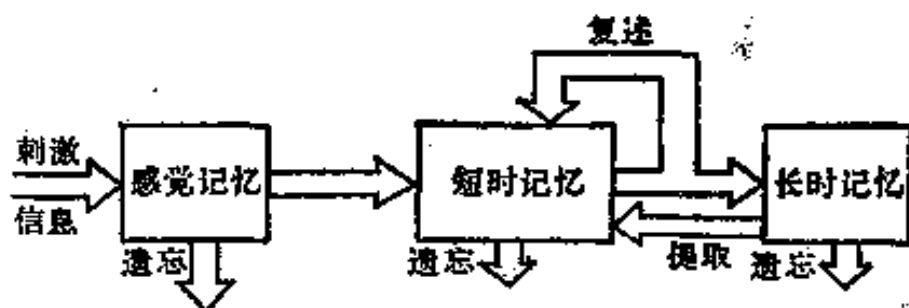


图 2-8 信息加工心理学中一个有代表性的记忆模型

从图中我们可以看到，这里说明了记忆有三个部分，这是不同于过去认为的只有一种记忆过程的观点。外界信息首先进入感觉记忆系统，而后转入短时记忆系统，最后进入长时记忆系统。这个流程图说明了现代认知心理学的关于记忆

研究的一种理论，即：有三种不同的记忆系统，而且，对于记忆的这三个不同的部分来说，它们各自具有不同的加工机制和信息转化规律。

图 2-9 是一个关于人类学习的信息加工模型，这个流程图表达了这样一种理论，它说明：人类的所有学习过程都

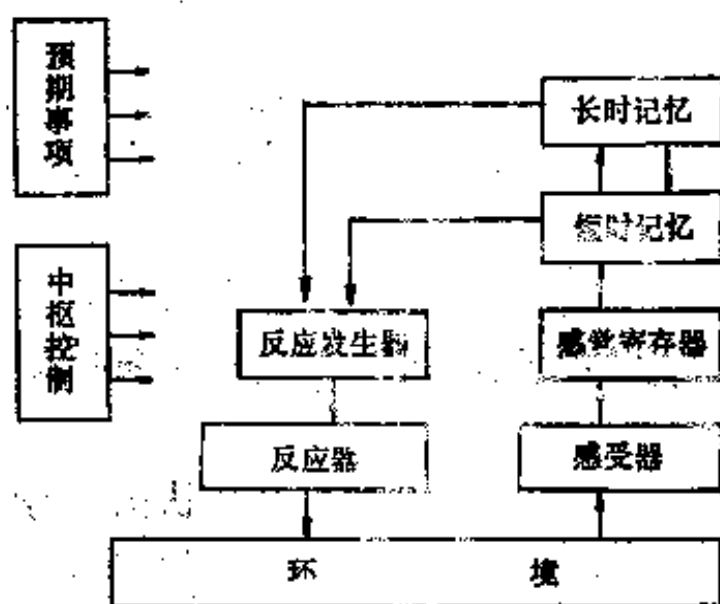


图 2-9 人类学习的一个信息加工模型

是通过内在的一系列心理操作，对外来信息或已存在于人类记忆中的信息进行不断的加工而完成的。这个过程主要包括三个环节：一是信息的接收，二是信息的加工处理，三是信息的输出和输出信息的反馈。图中的预期事项和中枢控制就是我们平时所说的学习得以发生和进行的主体因素，其中包括学习的需要和目的以及学习者已具有的身心发展状况。图中的环境是指学习者所处的学习环境，从环境作用开始，经过感受器→感觉寄存器→短时记忆→长时记忆→反应发生器→反应器，这一系列的过程就是上面提到的输入与输出环节之间的信息加工过程。由反应器→环境→感受器这一系列过程，

构成了学习的反馈环节。

上述的流程图同样说明了一种关于人类学习的信息加工理论。

现代认知心理学者往往把人类信息加工过程分解为若干阶段，引导他们去研究信息在人体内流动的时间过程，所以计时研究法是认知心理学家们最喜爱的方法之一。他们采用巧妙的方法去测量一个过程所需要的时间，并指出这个过程所具有的性质。我们举一个关于知觉过程研究的例子来说明这个方法的使用。

假定让被试看屏幕上投射的一个字母E。如果投射时间很短，比如1毫秒，那么被试者就不会看到什么。这说明知觉不是瞬时的。如果投射时间再长一点，比如5毫秒，那么被试者就会看到某种东西，但还不能确定看到的是什么。这说明知觉产生了，但是差别还未产生。如果投射时间再长一点，被试者也许能确定这个字母不是O和Q，但还不能看出它是E还是F或K，这说明被试者已经产生了部分辨别。这样一来，心理学家就可以分别确定完全辨别、部分辨别和刚刚看出有东西存在所需要的时间。这一切还说明，知觉过程是累积的，包括一系列特定的阶段。

计时研究法运用于记忆的研究上，使心理学家发现了人类短时记忆的存在。1959年在美国实验心理学杂志上，皮特森和皮特森（L·R·Peterson and M·J·Peterson）发表了一篇关于人类记忆研究的报告。他们给被试者念一个由三个辅音字母所组成的无意义音节，例如QCI，18秒钟之后让被试者再现。任务似乎很简单，但被试者却不能把这三个字母全都回忆起来。关键的一点是在这18秒钟的间隔期内，他们让被试者倒着数一个三位的数字，在倒数过程中要求逐

个减3，并且一秒钟数一个数。例如给被试者念309这个数字，让被试者重复309以后，就按已定的节奏向下数，306，303，300，297……一直持续到18秒钟红灯亮时为止。实验之所以这样安排是为了防止被试者在这期间进行复习。

如果在上述实验条件下，不仅让被试者在刺激消失后18秒钟再现三个字母的音节，而且在不同的时间间隔下再现类似的材料，其结果如图2-10中所示的曲线。

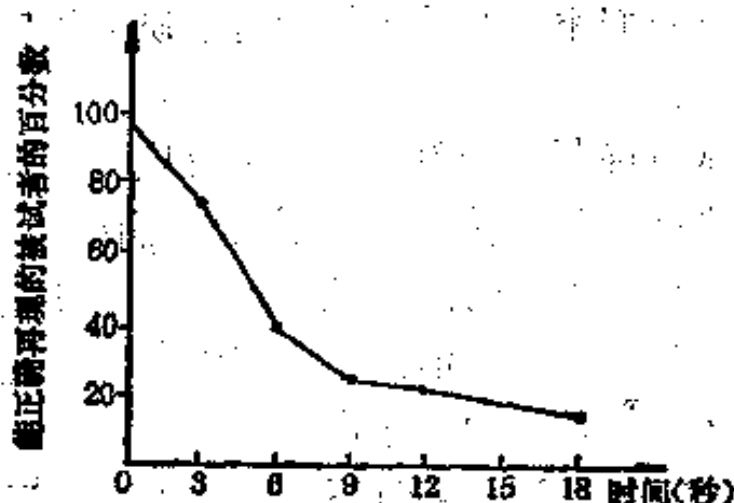


图2-10 短时记忆中信息的保存与时间的关系

由图中可以看出，时间间隔6秒钟时，只剩下40%的被试者能回忆起三个字母来，说明在这种记忆条件下，信息保持的时间还是很短的，一般在十几秒钟，最多不过30秒。显然这不是感觉记忆，因为在那里信息保持的时间更短，一般不超过1秒；它也不是长时记忆，因为长时记忆保持的时间往往是以小时、天、甚至月、年来计算的。由此心理学家们发现了这个介于感觉记忆和长时记忆之间的新的记忆过程，我们称之为短时记忆。由此可以看出，计时研究法被信息加工心理学家应用得是何等巧妙，它对于当代心理学的研究又是何等的重要。

最后谈一谈对于当代认知心理学研究有重要意义的反应

时的概念和方法。前面提到反应时的研究在早期实验心理学研究中已经受到了人们的重视，1868年荷兰生理学家顿德斯创立了选择反应时的实验方法，并利用这种方法对心理过程的时间进行了测量。在信息加工心理学发展起来以后，许多信息加工心理学家进一步扩展了选择反应时研究法的应用领域，把对人的心理过程的研究推进了一个新的阶段。

在过去的二十多年里，现代认知心理学家们作了许多选择反应时的实验，这里不能一一加以介绍，仅举一个比较有代表性的例子来说明当代心理学家是怎样运用选择反应时的实验方法来研究人的高级心理过程的。

美国心理学家斯顿伯格 (S. Sternberg) 曾经作过这样一个实验，他给被试者一组数字，例如：“1、3、4、7”，要求被试把它们记住，这一组数字被称之为正的数字集合，不同的数字集合包含有不同个数的数字，例如，另一组正的数字集合可以是“2、7”。然后，他给被试者呈现一个数字，这个数字被称为是检测刺激。如果这个检测刺激是正的数字集合中的一个成员，则被试就应回答“是”，否则就应回答“不是”，在这个过程中，要求被试应尽可能快而且准确地做出他的判断。在进行这个实验时，被试的反应通常是通过两个按键来实现的，给出肯定的回答时按一个键，给出否定的回答时按另一个键。

这个实验听起来似乎很简单，你可能觉得无论作出什么判断都将是很迅速的，特别是经过一段练习之后，这个过程就几乎是自动的过程了。但是，斯顿伯格的实验结果却证明这个过程是要花费时间的，而且反应时间的长短与正的数字集合的大小有关系，下面的表2-2列出了一个包含有四个正的数字集合的实验结果。

表2-2 一个包含了四个正的数串集合的实验结果

	第一次	第二次	第三次	第四次
正的数字集合	2, 7	9	3, 7, 6	8, 5, 2, 9, 1, 3,
检测刺激	2	8	7	4
正确的回答	是	不是	是	不是
平均的选择反应时	480毫秒	440毫秒	520毫秒	640毫秒

表中的结果表明正的数字集合的大小与选择反应时的长短是有关的。

为了解释上述的实验结果，斯顿伯格提出了一个新的记忆搜索理论。他认为：检测刺激与记忆在人的头脑中的正的数字集合中的每一个数字逐个地、顺序地进行了比较，他称这种比较过程是“穷尽式”的，也就是说，被试者在作出反应之前已经将作为检测刺激的数字同正的数字集合中的每一个数字进行了比较，正因为进行了这种“穷尽式”的比较，被试者在作业的反应时上就表现为正的数字集合越大他的反应时间就越长。图2-11是斯顿伯格理论的模型。

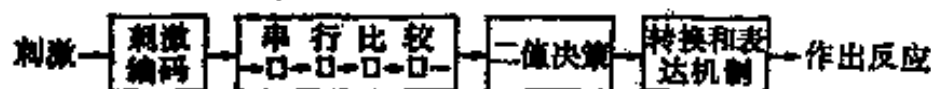


图2-11 斯顿伯格的记忆搜索串行比较模型

尽管斯顿伯格的理论并非完善，但是他所提出的这种研究思想和方法都是十分有意义的。它使心理学家们相信，人的高级心理过程是可以反应时来测量的，并且也可以用反应时实验来探测不能直接观测的人的内部心理状态，从而帮助人们掌握线索，提出理论，并进行可能的检验。

最后应当指出的是，人的信息加工过程是非常复杂的，一个控制参量的观测和量度往往并不能做出任何十分肯定的结论。因此，在心理学的研究上，我们应当更注重综合的观察和研究，应当从各个方面（心理的和生理的），各个角度（个体的和群体的），各种手段（实验的，数学的和计算机模拟的）等等综合起来进行推理和论证。

第三章 模式识别

人们对客观事物的认识，总是从对这些事物的感知和识别开始的。人们主要是通过眼睛和耳朵来获取外界的信息，并把这些信息传送给大脑；大脑则对这些信息进行分析，比较和解释。信息加工心理学把人类对外界客体的识别称之为“模式识别”。

一、模式识别的含义及其简单模型

（一）模式识别的概念

“模式”是由若干个元素集合在一起组成的一种结构。因此，一个物体、图象、语音、字符或人脸等等，都可以认为是一个模式，因为它们都是由若干个元素按照一定的关系集合在一起组成的。所以，通俗地讲，模式也就是我们周围世界中具有某种结构的各种客体，包括人在内。当然，一个较为复杂的模式，往往可以分成若干个部分，而这些部分本身也是由若干个元素按一定的关系组成的一种结构。所以，这些部分就叫做“子模式”。“模式识别”是人们把输入刺激（模式）的信息与长时记忆中的有关信息进行匹配，并辨认出该刺激属于什么范畴的过程。因此，对物体、图象、语音或字符等等的识别，都可称之为模式识别。

在一般的认知模型中，我们可以看到，要识别一个模式，感知到的信息必须与长时记忆中的信息进行匹配。前一

类信息来自于当时出现的某个刺激（模式）；而后一类信息却不同，它是先前获得的、有关这类刺激的知识。当模式的感觉特征与先前获得的有关知识相匹配的时候，我们就说这个模式被识别了。在这个时候，人们往往会给刺激一个名称。所以，一般地讲，模式识别是以刺激“命名”，即给刺激一个名称。例如，在一定的位置上，给出一个由三条线（一、/、\）组成的刺激，人们就可能说这是一个“大”字。在这种情况下，人识别一个刺激的同时，也对该刺激指定了一个名称。当然，这不是说，模式识别总是意味着要以刺激命名。有时候，我们可以识别一个模式，但不必要（或不能）给它命名。比如说，我们识别出一张脸是熟悉的，或者，某种气味可能使我们想起在什么地方闻到过，但没有把它们与某个确定的概念联系起来。

（二）模式识别的简单模型

图 3--1 是简单情境下的模式识别的一个模型。它指明了，模式识别的任何一种模型或理论，应该包括哪些基本的

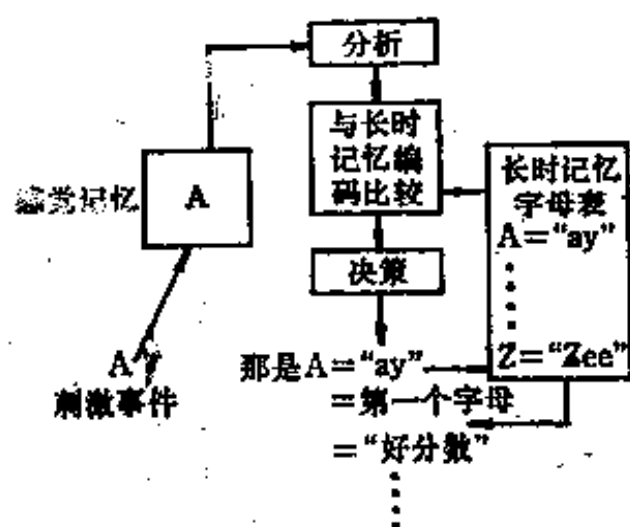


图 3-1 模式识别的基本组成部分

内容。这里规定的情境是简单的，只是一个孤立的刺激瞬时出现，然后又很快消失。但是，在实际生活中，大多数自然而然地发生的刺激，都出现在一种有意义的前后和左右的关系中。这种有意义的关系，对人识别某个刺激来说，是非常重要的。

从图 3-1 中可以看出，对一个模式的识别，需要利用两种记忆。一是感觉记忆，它主要存储输入的感觉信息；二是长时记忆，它保存着与上述输入信息进行匹配的、过去获得的关于各种事物的记忆编码。同时，对一个模式的识别，要使用三种基本过程。这些过程分别叫做“分析”、“比较”和“决策”。

1. 分析

识别过程的第一个阶段是“分析”，即把信息从感觉记忆中抽取出来。这是一种比较简单过程。在这个过程中，抽取出来的信息，实际上跟感觉记忆中的内容没有什么变化。或者，分析过程是把感觉记忆中关于刺激的信息分解成它的各个组成部分。例如，把字母 A 分解成 2 条斜线（左斜线和右斜线）及 1 条水平线三个组成成分。

2. 比较

在第二个阶段即比较过程中，分析阶段抽取出来的信息，与长时记忆中已有的各种信息进行比较。这意味着，长时记忆中的各种信息，必须处于一种当前刺激的信息能够跟它进行比较的状态。也就是说，长时记忆中存储的关于各种事物的记忆编码，必须与从实际刺激中抽取出来的信息是类似的。或者说，它必须描述从实际刺激中抽取出来的信息。

3. 决策

在刺激信息与长时记忆中的各种记忆编码进行比较以

后，究竟哪一个记忆编码提供了最好的匹配，需要作出判定。这就是决策过程所要做的事情。决策过程确定了识别系统的输出。模式被识别为相应于获得最佳匹配的记忆编码。

由此可见，即使在图 3-1 这种简化的模型中，模式识别也要经历几个阶段，只有当刺激的信息经历了所有这几个阶段以后，才能够说这个模式被识别了。如果关于一个刺激的信息经历了分析和比较阶段的加工，但过程却停顿下来了，识别系统对它不再作进一步的加工，或不能作出决策，在这种情况下，对模式的识别是不完全的，因为关于刺激的信息没有成功地经历模式识别的三个基本阶段。对模式的完全识别，这是第三个阶段即决策过程的产物。

某个模式一旦被识别以后，就会从长时记忆中诱发出与这个模式有关的更多的信息。比如说，一旦我们识别了刺激“A”，我们会由此而想到，它是英文字母表的第一个字母；或者，它是大家想要得到的一个好的分数等级，如此等等。

（三）模式识别的意义

经过模式识别以后，关于刺激的粗糙的、原始状态的信息（如视觉形状），就转换成某种对人来说有意义的东西。模式识别在人类适应环境和改造环境中是非常重要的。假如说，一个人不能把输入的刺激信息正确地识别为一只“虎”，而是错误地把它归结为一只“猫”，那么，他就会对刺激作出绝然不同的反应。人的识别系统一旦出现缺陷，就会危及到人的生存。

科学家们力图研制出一种能象人类那样进行模式识别的机器，以代替人的繁重的劳动。这方面的工作已取得了一定

的成绩。例如，科学家们已研制出了光学文字识别机，用它代替人的一部分识别任务。但是，它与人的识别能力相比，相差甚远。人的模式识别的机制是很复杂的。邮局里一位分信的工作人员，他要与千万人写的各种大小、方位和形状不同的文字打交道。他怎么识别这些形状各异、大小不一的手写体文字，从而把信分到各个地区，这是一个还没有弄清楚的问题。将来，科学家们把人的识别文字的机制研究清楚了，并且具备了模拟这种机制的技术手段，就可造出一部功能很强的文字识别机。这种机器不但能识别各种印刷体，而且也能识别各种大小、方位、形状不同的手写体汉字。这将大大促进生产力的发展。那时候，我们不必采用邮政编码，也可以用机器代替人去完成信件分类的任务。

二、模式识别的理论

为了对人的模式识别作进一步的描述，人们以模式识别时使用的记忆编码及这些编码的使用等问题，作了一些探讨，提出了三种假设。这三种假设是：模板假设、原型假设和特征假设。

（一）模板假设

1. 什么是“模板假设” 前面曾指出，存储在长时记忆中的、在模式识别中要使用的记忆编码，必须类似于（或描述）相应的刺激，否则就无法用来进行比较，识别过程就会产生障碍。有一种假设认为，记忆编码与相应的刺激之间是非常一致的，长时记忆中的编码，是该刺激的某种缩小了的拷贝（或模板、样板）。它们之间存在着一对一的对应关

系。也就是说，对我们识别的每一个刺激来说，都有一个用来与它进行匹配的**内部拷贝**。这就是“**模板 (template) 假设**”的内容。根据这个假设，刺激的感觉形象经过第一阶段的分析后，基本上不变地被传送到第二阶段去加工。在这个阶段上，刺激的感觉形象与长时记忆中的许多模板进行比较。通过比较，某一个模板可能与刺激的感觉形象匹配的**最好**。于是，决策过程把这一获得最佳匹配的模板选出来。这时，模式也就被识别了。比如说，我们识别一个“上”字。根据模板假设的规定，在长时记忆中有它的一个拷贝。每当“上”字作为一个刺激呈现时，它的视觉形象就与长时记忆中的许多模板进行比较。比较结果，记忆中关于“上”字的拷贝获得最好的匹配。因此，人就把这个模式识别为“上”。

这种模板假设看起来过于朴素和简单化，所以，用它来作为模式识别的理论基础有些不太合适。因为按照这个假设，为了识别每一个在形状、大小或方位上稍有变化的字，在长时记忆中就得具备一个相应的模板。这就是说，改变一下字的大小或形状，就需要一个不同的模板；如果把“上”稍微旋转一下，产生一个仍然可以识别的“上”字，则又需另一个模板，如此等等。如果对这些大小、形状和方位不同的刺激不具备相应的模板，那么，模式识别就会陷入困境。图 3-2 (a) 表明，当字母 A (深黑色) 的大小、方位和位置有变化时，模板 (浅色) 就不能成功地得到匹配；图 3-2 (b) 则表示发生错误匹配的情况。当字母的形状发生变化时，A 模板可能与字母 R 获得较好的匹配，而与字母 A 的匹配程度反而较差。同样，以一个倾斜的字母 A 来说，R 模板可能会获得比 A 模板更好的匹配。根据模板假设，当倾斜的 A



图 3-2 (a) 模板 (浅色) 不能成功地与字母 (深黑色) 相匹配; (b) 发生错误的匹配

出现时, 我们会错误地把它识别为 R, 而不是字母 A。如果要避免上述这种情况, 我们就需要无限数量的模板。毫无疑问, 这将会超出一个人的长时记忆的能力。

2. 模板假设的修正 对于人来说, 不管字母 A 是手写体或是印刷体, 也不管它在大小和方位上的某些变化, 都可以正确地识别出来。因此, 要使模板假设能比较满意地解释人的模式识别过程, 必须要做适当的修正补充。

有一种修正意见, 主张在模型中加上一个预处理过程。这个过程发生在比较阶段之前, 起一种“净化”(clean up)的作用。它能使刺激的感觉形象进入一种标准的大小和方位。由于这一过程缩减了刺激形象的不规则性, 使它进入一种比较正常的形式, 所以, 人们把它称之为对刺激的感觉形象的“标准化”(normalizing)。例如, 若有这样一个刺激 A, 那么, 在它与模板进行比较之前, “标准化”过程就要把它的位置重新摆正, 把尺寸缩小, 以及把它右边弯曲的一笔弄直。这样一个发生在比较阶段之前的预处理过程, 就会大大地减少识别字母 A 所需要的模板数量。图 3-3 表

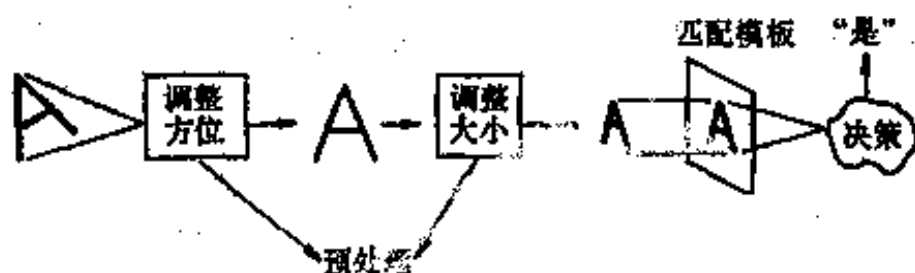


图 3-3 模板匹配中的预处理
示字母A的方位和大小的预处理过程。

但是，加上这样一个预处理过程，可能会出现一个逻辑上的毛病。要纠正刺激形象的方位和大小，就必须要知道合适的方位和大小是什么样的。这就等于说，“已经”知道刺激相当于什么模式了。例如，一个看起来象“目”这样的刺激，调正它的方位时，若把它横下来则象“四”，把它竖起来则象“目”。为了知道哪个方位是正确的，就必须首先确定它是哪个字。但是，这乃是模式识别系统要做的事情，而不是比较之前某个加工过程要完成的事情。它应该是识别的结果，而不是识别的前提。

当然，要处理这个逻辑上的问题，也不是不可能的。一则，如果明显偏离标准方位（如“上”偏离成“丄”），其结果或许形成一个确实不能识别的刺激。这就是说，当识别系统事实上不能处理这个刺激时，也就没有理由要求比较过程之前的“标准化”过程必须能够处理它。其次，正如前面已经指出的，被识别的刺激通常总是出现在某种有意义的关系中。这种关系能够指明一个刺激的合适的形状和方位应该是什么样的，从而会有助于这个标准化过程。

（二）原型假设

1. 关于原型的概念 “预处理”过程可以帮助填补模板假设

的不足,但它并不能完全解释人的识别问题。事实上,尽管刺激并不是出现在某种特定的前后关系中,但当它们的大小或形状有变化时,我们仍然能够识别它们。为了解释这种现象,人们就提出了“原型”(prototype)这个概念。原型与模板不同,它不是某个刺激的内部拷贝。原型是表示某类客体之基本成分的一种抽象的形式。例如,飞机的原型是有两个翅膀附在上面的一个长筒;狗的原型大致是一个有二尺高、四条腿、二只眼睛、尖的牙齿、较长的鼻子和身上长毛的动物形象。因此,原型是从某类客体的各个具体经验中抽取出来的。某一范畴的客体原型,代表了这个范畴中所有客体的集中趋势。

模式识别的原型假设认为,存储在长时记忆中的东西,不是模板,而是原型。原型与输入刺激之间的关系,不是一对一的关系。因此,人在识别模式时,尽管同一范畴的各个成员之间有形状或大小等方面的变化,该范畴的原型也能与它们相匹配。在人们的长时记忆中,存储着各种范畴的原型,象飞机、猫、狗或字母A的原型等等。人们依靠这些原型,可以有效地识别属于这些范畴的各个成员。每当一个新的刺激呈现时,它就与原型进行比较,并且只需要得到一种近似的匹配,而不是一种确切的、模板式的匹配。当刺激的感觉形象与某个原型获得最近似的匹配时,人们就确认这个刺激属于该原型代表的范畴。原型假设是对模板假设的改进,因为它表明了,即使长时记忆的编码不是某种刺激的确切的复写,人也能识别这个刺激。

2. 支持原型假设的实验 原型究竟是否存在?有一些实验表明,它们是存在的。波斯纳(M. I. Posner)和基尔(S. W. Keele)曾做过有关的实验。他们利用九个点

构成各种原型的模式：由九个点组成一个三角形；或者组成一个字母；或者排成一个随机的模式。同时，他们把原型中的某些点往某一方向移动一点儿，从而构成一个新的模式；再把某些点向另一方向移动一点儿，又构成一个新的模式。这些新的模式属于原型的变形。然后，他们就让被试做分类实验。

例如，他们先让被试者看三种随机点原型的变形。每一原型都有相应的4个变形，因此共有十二个变形。图3-4上面一排为三个随机点的原型；下面一排是右边一个原型的四种变形。被试者应该把这十二个变形分到三个不同的类别中去。但是，他们并不给被试者看图上部的三个原型，只是要求被试者把某个原型的变形，都归到同一个范畴中去。或

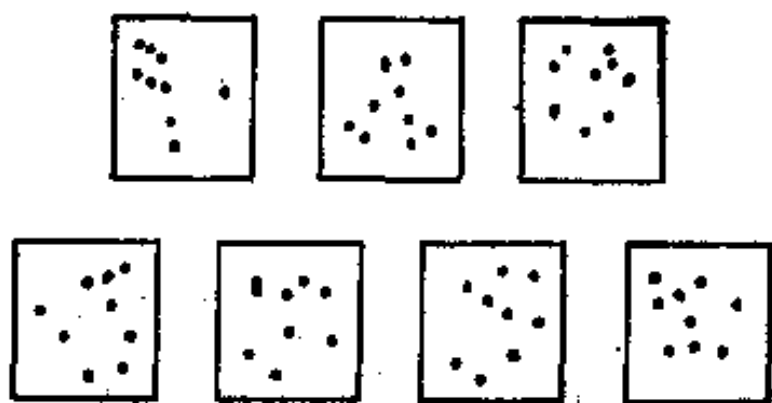


图3-4 随机点模式的例子

者说，要求被试者把某一原型的四个变形归在一组，而把另一原型的变形归入另一组。在被试者学习了对这些模式的分类以后，就开始第二步实验。他们给被试者看一系列的模
式，要求被试者把它们分别归入在第一步实验中学习过的三个范畴中去。这些模式中，有的是在第一步实验中没有呈现给被试者看的那几个原型本身；有的是属于这些原型的新的

变形。

结果发现,被试者对这些原型的分类非常精确,好象在第一步实验中已经学习过的那样。这一事实表明,在第一步实验中,尽管原型没有呈现给被试者看过,但是,被试者已经从那些看见过的作为变形的模式中,抽取出了原型。至于对那些新变形的分类,其精确性取决于它们与原型的类似程度。它们与原型越接近,则被试者就能越正确地把它们归入一个相应的范畴。这一事实说明,被试者在对这些新变形分类时,是把它们与被试者已经抽取出来的原型进行比较的。

另外,弗兰克斯 (J. J. Franks) 和布兰斯福德 (J. D. Bransford) 也做过有关实验,结果支持了原型假设。他们利用一些几何图形,如三角形、圆形、和星形等,组合成某种结构图。这种结构图作为原型模式,如图 3-5 左边第一个结构图。然后,对原型作某种(或几种)变换,从而构成该原型的各种变形,如图 3-5 右边的四个结构图。这种变换可能

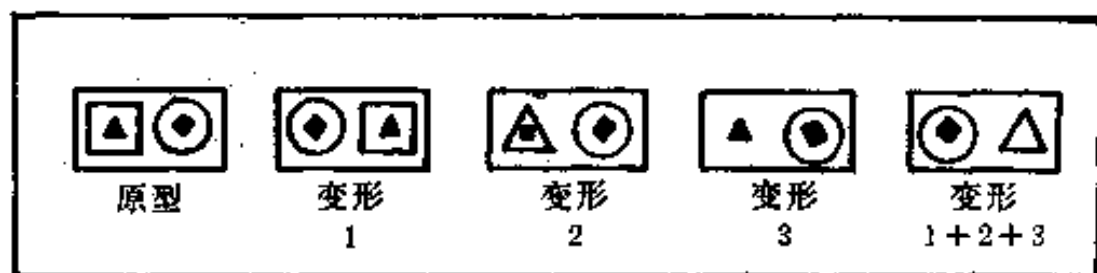


图 3-6 弗兰克斯和布兰斯福德在实验中使用的原型及其变形的例子

是: 左右位置变换(第一变形); 内外位置交换(第二个变形, 方块移至三角形里边); 删掉一个几何图形(第三个变形, 删掉了一个正方形); 而在第四个变形中, 包括了前面三种变换。他们先让被试者看一些属于变形的模式, 然后再

给被试者一个测验。在测验中，给被试者呈现一系列模式。其中一些是被试者曾见过的变形，另一些是没有见过的新变形，以及原型本身。针对每一个呈现给被试者看的图形，要求他确认以前是否见过这个图形。结果表明，被试者对原型图形的反应，往往比较肯定地认为是以前见过的。事实上，在实验的第一阶段中，并没有给被试者看过这些原型图形。而对于那些变形，不管是以前见过的或是没有见过的，被试者同样地认为见过的可能性较小。对于那些由几次变换而产生的变形，由于其偏离原型（与原型图形的差别）较大，被试者认为见过的可能性就更小。

这一实验同样表明，被试者从一组有关模式的经验中，抽取出了这组模式的原型。同时，也可以看到，被试者在识别那些变形时，利用了这个原型，是把它们与原型作比较的。这种实验结果支持了原型假设。

（三）特征假设

1. 什么是“特征假设” 前面说过，每一个模式（或客体）都是由若干个元素按照一定的关系结合在一起构成的。这就意味着，任何模式都可以分解成一些基本的成分或特征。例如，可以把英文字母看成是由垂直线、水平线等这样一些基本特征组成的。林德赛和诺曼（P. H. Lindsay和D. A. Norman）把构成英文字母的特征分成七种，它们是：垂直线、水平线、斜线、直角、锐角、连续曲线和不连续曲线。表3-1列出了每个字母具有上述哪些特征的详细情况。从表中可以看出，字母A包含有一条水平线、二条斜线和三个锐角；字母T包含一条水平线、一条垂直线和二个直角；而字母S则包含二条不连续曲线。事实上，不连续

表3-1 英语字母的特征分析

字母	特征	垂直线	水平线	斜线	直角	锐角	连续曲线	不连续曲线
A			1	2		3		
B		1	3		4			2
C								1
D		1	2		2			1
E		1	3		4			
F		1	2		3			
G		1	1		1			1
H		2	1		4			
I		1	2		4			
J		1						1
K		1		2	1	2		
L		1	1		1			
M		2		2		3		
N		2		1		2		
O							1	
P		1	2		3			1
Q				1		2	1	
R		1	2	1	3			1
S								2
T		1	1		2			
U		2						1
V				2		1		
W				4		3		
X				2		2		
Y		1		2		1		
Z			2	1		2		

曲线还可以分成两种，一是C曲线（即“C”形曲线），另一是D曲线（即“O”形曲线）。特征假设认为，各种刺激在长时记忆中的编码，是一张表，表上登记着该种刺激所具有的那些特征的名称。这种表叫做“特征表”。人的识别系统在接收一个输入刺激的信息以后，首先对它的特征进行分析，找出输入刺激具有哪些特征，然后把它们与长时记忆中的各种特征表进行比较。一旦长时记忆中的某张特征表与刺激的特征获得最佳的匹配，这个刺激也就被识别

了。

人的口头言语也是由各种基本特征组成的。口头言语的最小单位叫做“音素”。音素也可以说是一个声音。构成词的音素不同，词的意义也就会不同。例如，在汉语中，声母和韵母的各种组合，构成了不同的词。在bà（爸）、pà（怕）、mā（骂）这几个词中，相应于b、p、m的声音，就是不同的音素。当然，每个人的发音总是稍有不同。所以，单个音素实际上又相应于一定范围的一组声音。但是，由不同人说出的同一些音素构成的词，我们仍然能正确地识别。这与手写体文字的情况恰好是一样的。

科学家们也试图找出一组特征，能用它们来描述各种音素。开始，他们想按照人们在产生和发出声音时使用发音装置的方式来描述每一个音。这种发音装置包括舌头、鼻子、嘴唇和声带等各个部位。当然，人们在发各种音时，发音装置各部位的活动确是不同的。例如，Z音是从喉头发出的，声带参予了发音，而S音只是在口腔中产生的，声带并不振动。但是，用发音装置活动的特征来说明人对言语的识别，这是行不通的。听者看不见说话者的发音装置的活动，因此，他不可能利用发音装置活动的特点，去识别说话者发出的音。

另一条比较有效的途径是分析言语的声学特征。当我们说话的时候，就产生声波。声波有两个重要的特性，即频率和振幅。频率决定它的音高，而振幅决定它的响度。为了研究言语的声学特征，人们使用语谱仪（Speech spectrograph）把言语的声音模式转变成一种可见的记录，如“语谱”。图3-6是科学院声学所测量的普通话元音ü的语谱。图中F₁、F₂和F₃三个能量集中的区域，叫做共振峰（formant）不同元音的共振峰具有不同的特征。所以，共振峰

携带着语音识别所需要的大量信息。这方面的工作，对言语识别的研究具有重要的意义。

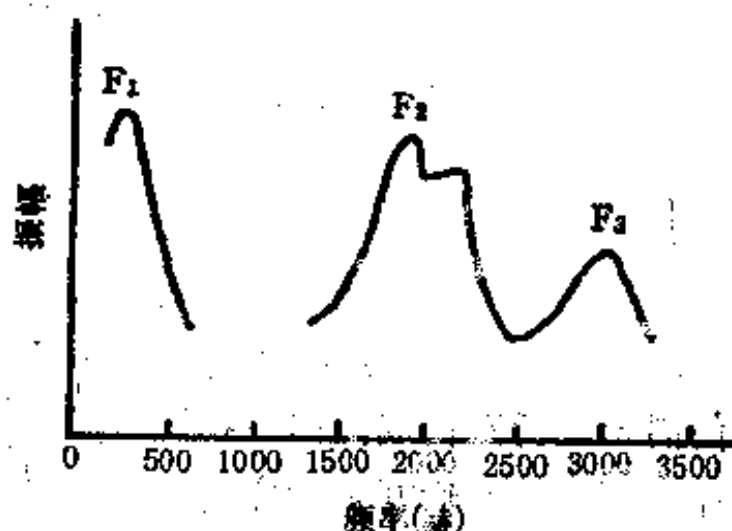


图 3-6 普通话元音 u 的语谱

特征假设确有它的优点。如果人们能对一组相当大数量的模式，抽取出少量的一些特征，而且用这少量特征足以描述这大量的模式（例如，用一些基本特征来描述人类的文字），那么，识别系统需要加以存储和处理的东西，就将大大减少。

但是，把特征假设用来解释模式识别时，也会碰到一些问题。有时候，两个模式的差别只是有或没有某一个特征，如E和F两个字母。字母E下面的一横，在F中是没有的。但长时记忆中的两张特征表（E特征表和F特征表），都可以与字母F相匹配。因为F的每个特征，在E特征表和F特征表中，都是存在的。在这种情况下，如何进行决策，以达到正确的识别，这似乎是一个问题。

2. “鬼域”模型 特征假设的一种形式叫做“鬼域”（pandemonium）。这是由塞尔弗里奇（O. G. Selfridge）提出来的。图 3-7 描述了鬼域模型的概貌。这个模型表明，

识别一个模式要通过几个阶段或层次。在每一层次上，都有一些小鬼（demon），它们肩负着不同的任务，起着不同的作用。“映象鬼”在第一个层次上，它要完成的工作，是把刺激作为感觉事件记录下来。然后，“特征鬼”在第二个层次上对刺激的感觉映象进行分析，把它分解成各种特征成分。每一个“特征鬼”代表某一个特征，如一条直线或一个直角。并且，它必须就输入刺激中是否存在它所代表的特征作

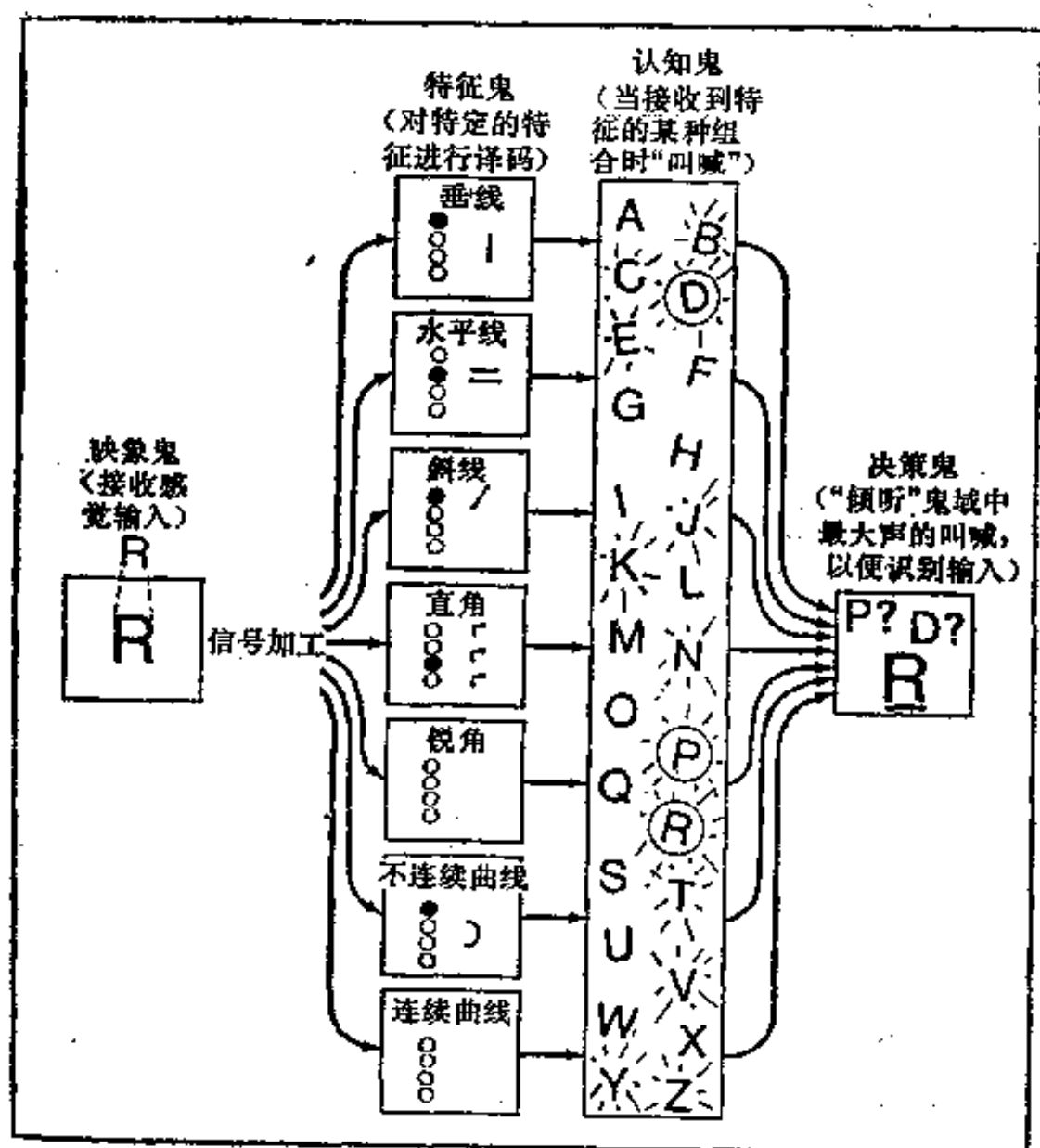


图 3-7 鬼城模型——小鬼工作图示

出明确的反应。“认知鬼”则相当于前面所说的特征表。每一个“认知鬼”代表某一种模式。“认知鬼”的任务是对“特征鬼”进行观察，每当由特征分析的结果表明它可能是一个候选者时，就尽可能大声地叫喊或发出信号。也就是说，“认知鬼”注视着发出反应的“特征鬼”及其数量。如果它发现许多作出反应的“特征鬼”是属于自己的特征，就大声地叫喊；否则，叫喊的声音就小，或者不作声。“决策鬼”在更深的—个层次上。它的工作是要确定在那些正叫喊着的“认知鬼”中，哪一个叫喊的声音最大(决策过程)，从而认知这个模式。

总之，在特征假设中，输入刺激被译码成一组特征，然后与记忆中的内部编码特征表进行比较，从而达到对输入刺激的识别。因此，它可以允许在刺激的大小、方位及其他因素变化的情况下达到识别的目的。

3. 模式识别中特征作用的证据 有不少事实证明，特征匹配是用于模式识别过程的。一些生理学的研究表明，视觉系统包含有专门化的细胞。它们的功能就是识别某种特定的特征。生理学家们(如D. H. Hubel及J. Y. Lettvin等人)发现，在青蛙、猫等动物的视觉系统中，有些神经细胞只是在视觉刺激的某种特征(如水平线、垂直线和移动着的点等)出现时才激活，即发生反应。生理学家们把这些功能专门化的神经细胞称之为特征觉察器。例如，在青蛙的视觉系统中，有一种特征觉察器仅对黑的、略呈圆形的、动的刺激发生反应。当这样的目标进入青蛙的视野时，相应的特征觉察器就开始活动。目标离青蛙越近，觉察器的活动也越强烈，最后引起视觉和运动联合行动，伸出舌头，捉住目标。这种目标往往是一些昆虫，所以，这种觉察器也叫做昆虫

觉察器。

人们推测，人也可能有这样的细胞。而且，这些细胞的激活，似乎不依赖于特殊的刺激特征，如长度。因此，它们好象可以觉察比较抽象的特性，如一个特定角度上的任意长的线。所以，尽管模式在大小方面有各种变化，但人们也能够识别它们。

特性在模式识别中的作用，还有其他的一些根据。例如，小孩常常用同一名称来称呼印刷字母b和d。这可能反映了小孩不能精确区分视觉特性c曲线和o曲线，也就是说，不能区别只是在方位上有差别的视觉特性。如果视觉刺激呈现得很快，那么，在成年人身上也会出现这种情况，这就是所谓的“混同”(confusion)现象。例如，在视觉上把一個字母(如Q)看成另一个字母(如D)，这种现象之所以发生，是因为这两个刺激共有一些视觉特征，在听觉上也有类似的现象。两个音节在发音特性方面愈接近，它们就愈容易混同。例如，被试很容易把B误听为是D。这是因为B和D在听觉上有类似的特征。所有这些情况表明，模式的特征在识别中确实是起作用的。

1966年，金内(G. C. Kinney)等人做过一个字母和数字识别的实验。他们在屏幕上给被试者瞬时呈现一个字母(或数字)，要求被试者即刻报告自己看见的是什么字母或数字。结果发现，那些具有共同特征愈多的字母或数字，彼此就愈容易混同，没有共同特征的，彼此就不发生混同。例如，对字母C的识别共犯了29次错误，其中误认为是“G”的就有21次，误认为“O”的有6次，而误认为“B”的只有1次。

这个实验说明，人们是依据特征分析去识别事物的，因

此，当二个刺激具有共同特征时，就易于犯混同的错误。

普里查德 (R. M. pritchard) 作了一个稳定映象的实验。平时，人的眼球是在不停地飘移的。眼球的这种运动对于人的视觉保持非常重要，，它使兴奋和抑制二种过程彼此不断地转换，以维持人眼的正常功能。普里查德设计了一种实验方法，使刺激映象在眼球运动的情况下仍投射在网膜同一位置不变。结果出现了非常有趣的现象。视觉映象一部分一部分地逐渐消退。这种消退可能是由于视觉细胞的疲劳引



图 3-8 普里查德实验用的刺激及其反应

起的。图 3-8 中每一行最左边的项目是给被试者在屏幕上呈现的刺激。在刺激呈现后，让被试者报告他的视觉中保留了什么，直到视觉映象全部消失为止。图中每一行右边的 4 个项目，是被试者的报告。例如，对于第一个刺激，被试者的报告说，在视觉中还保留了“H”，有的报告说保留了“B”、“3”或者“4”。

从这一实验结果中可以看出，网象的消退不是无规则的。它倾向于一部分一部分地消失。这种部分相当于刺激的特征。留下来的部分又构成一个整体。人对外界事物的识别就是先抽取特征，然后再加以组合的过程。当然，这个过程人们往往是意识不到的。

三、模式识别的比较过程

(一) 串 行 加 工

一个输入模式的信息必须与大量的记忆编码进行比较，然后才能选出那个与模式匹配得最好的编码（模板、原型或特征表）。模式识别的这一比较过程是以什么方式进行的？一种可能的方式是串行加工，即输入的模式与长时记忆中的编码一个接着一个地进行比较，然后再确定与哪一个编码匹配得最好。就拿模板假设而论，在能够作出决策之前，输入模式或许要与长时记忆中的全部模板进行比较。在长时记忆中存储的模板数量是很大的，所以，这种比较过程的任务将十分繁重。因此，如果人的识别系统在把输入刺激与长时记忆编码进行比较时，采用一个接着一个的串行加工方式，那末，为了识别某些模式，肯定会花很长的时间。可是，我们大家知道，人的模式识别往往是进行得很快的。所以，识别系统在把刺激与记忆编码进行比较时，在某种条件下并不是采用串行的加工方式。

(二) 平行加工

与串行加工相对应的另一个过程，叫做平行加工。“平行”来源于几何学的意义，是指二条或二条以上的线路并排地一起进行。平行加工恰好就是这个意思，在这里，就是指许多独立的比较同时进行。这就是说，在模式识别中，刺激可以同时地与许多内部记忆码进行比较，而整个过程所需的时间，差不多就等于刺激与单个记忆码进行比较的时间。在

刺激和记忆码进行比较时，识别系统如果采用平行加工的方式，那么，比较过程将会节省很多时间。

我们知道，在物理世界中，确有这样的平行加工。例如，我们拿一把频率未知的音叉，敲打它，使它嗡嗡地响，并把它放在一组频率已知的音叉旁边。这时，该组音叉中的某一把音叉也开始嗡嗡地作响。这把音叉的频率与频率未知的音叉的频率，是相匹配的。用这种办法，人们可以确定某个音叉的频率。在这里，被测音叉是同时地对着所有频率已知的音叉进行比较的，所以，这是一个平行的加工过程。

在心理活动领域中，也有平行加工的实验证据。奈瑟（U. Neisser）的视觉搜索实验就是一个例子。实验是这样做的：被试面前放着一张字母表，上面有50行，每行有几个字母，如UFCJ；要求被试从顶上一行开始，尽可能快地找出某些由主试规定的目标字母（目标字母放在表中任意确定的位置上）；当被试找到它时，就按一下钮；主试记录下整个搜索的时间。给训练好的被试10个目标字母，要求他们在发觉其中一个目标字母出现时就作出反应。结果发现，在这种情况下，被试的反应就象他们只要考察一个目标字母时同样的快。这个实验的结果成了平行加工的证据。因为，被试似乎是同时针对这10个目标字母来搜索字母表的。他所采取的加工方式是“平行比较”。此外，据那些对一定的目标字母接受过训练的被试说，当他们看字母表时，那些目标字母似乎是从纸上“跳出来”的。这很象测试音叉时不需要串行搜索，而从一组已知频率的音叉中“跳出来”一样。

用平行加工的方式，识别就可以加快，因为许多操作可以同时进行，从而比串行加工节省时间。如果比较过程是平行的（并且假定系统有足够的容量来同时执行所有这些操

作），那么，输入模式可以与每一个现存的长时记忆码进行比较，这在理论上是可能的。

但是，有证据表明，人们不必把输入模式与长时记忆中的每个编码都进行比较，而可以把比较限制在某些编码的范围内。事实上，任何刺激都是在一定的关系中发生的。刺激发生时所处的关系，对我们识别它是十分重要的。例如，假定我们的任务是确定某一情景中是否存在某一客体（如一只碗），而我们已经认出该情景是描述一个厨房的，这样，我们就有了刺激出现的有意义的关系，并且可以根据这种关系来限制客体可能是什么的范围。实验证明，在情景是一幅有条理的图象时，人们识别一个客体就容易些；当情景是一幅杂乱无组织的图象时，识别客体就差些；同时，在一个有条理的情景中，可能有的（与情景的关系一致的）客体，要比不可能出现的客体识别得好些。这是因为，有条理的情景中存在着有意义的关系。

四、识别模型的扩展及机器识别

（一）数据驱动和概念驱动

我们知道，人对外界事物的识别，是从外界刺激到达人的感官开始的。人们接受外界刺激并对它进行一系列阶段性的加工。这个加工过程受外界刺激信息的驱动，并且，每一加工阶段的输出信息都作为下一加工阶段的输入信息，驱动着下一阶段加工的进行。这是我们通常所说的数据驱动的加工过程。数据驱动的加工是一个自下而上的过程。

当然，人对外界刺激的识别，并不是单靠数据驱动来进行的。例如，若给你看图 3-9，但什么也不告诉你，只要



图3-9 雪地里的一条狗

求你说出该图画的是什么。你大概是识别不出来的。但是，如果告诉你该图画的是雪地里有一条狗，它的头向前探着，正嗅着地面上的东西，那么，你会很快就识别出来了。因此，我们具有的知识和经验对外界刺激的识别起着重要的作用。它们能推动对外界事物的识别过程。这就是概念驱动的加工过程。概念驱动的加工是一个自上而下的过程，我们对图3-9的识别，属于概念驱动的加工过程。

(二) 交互的识别模型

关于一个待识别的模式所处的情景的知识，在比较和决策过程中是很重要的。这种知识可以用来指导人的认识过程。仿佛是首先形成一个模式可能是什么的假设，然后再去

验证这个假设。从形成假设到验证假设的加工方式，在模式识别中起着重要的作用。

图 3-10 是一个扩展的识别模型。它把识别描述为感觉信息和假设相结合的过程：是来自现实世界的信息和预期的假设之间相互作用的过程。所以，人们把它叫做“交互的识别模型”（interactive model of recognition）。交互模型的核心是：识别既是从感觉输入开始的，又是受预期的假设控制的。从图中可以看到，预期的假设不仅指导着比较和决策过程，在分析阶段，它们也可以指导特征抽取。甚至在一个待识别的模式的信息进入感觉存储之前，预期的假设就可能发生影响，控制信息的最初记录。

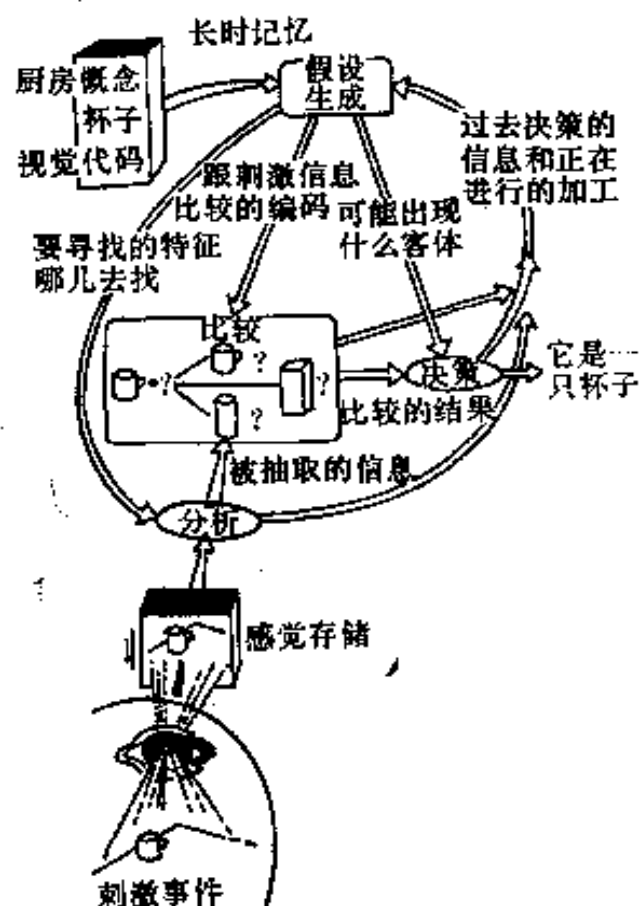


图 3-10 扩展的模式识别模型

这些预期的假设，是人们根据存储在长时记忆中的各种知识推导出来的。长时记忆中的各种知识，都可以用来形成假设。刺激的部分识别（如对刺激的某种典型特征的抽取），可以激活长时记忆中的有关知识，从而导致刺激可能是什么的假设；关于输入刺激物理特征的知识，可以用来形成在感觉存储中发现原始数据的假设；关于模式所处的背景的知识，也可以形成模式可能属于什么范围的假设；关于正在加工的概念的知识，可以导致形成有关正在呈现的概念的意义的假设，等等。识别系统几乎利用每一种可能的信息源来形成各种不同的假设。例如，在阅读或谈话的情况下，人们可以利用词的拼法规则、词的发音以及词的语义内容的知识；在识别句子时，人们能够运用词怎么结合起来的句法规则，以及有关的语义规则，形成各种有用的假设。

总之，人对外界事物的识别过程，是数据驱动加工和概念驱动加工相结合的过程，是自上而下和自下而上相结合的加工过程。

（三）机器识别

机器识别是指给计算机配上一定的软、硬件，使计算机能识别外界的事物。它包括机器视觉、听觉和触觉这几方面的识别。机器识别使计算机能直接感知并处理外界的原始信息，从而使计算机走向智能化。

1. 识别字母的一般方法 识别字母的一般方法，可用图 3-11 来说明。这种方法可以识别手写的大写字母，也就是说，当字母的粗细、大小和形状在一定范围内有变化时，也可以成功地识别。

在分析待识别的字母时，它把输入字母看成是由一组有

限的图象元素构成的。这些图象元素各有自己的特征。对每一个字母来说，它们都有一组规定好的图象元素，以及这些元素必须满足的一些关系。因此，在识别任何一个字母时，要从输入图象中发现两类信息：（1）出现了哪些图象元素；（2）这些图象元素之间的关系。

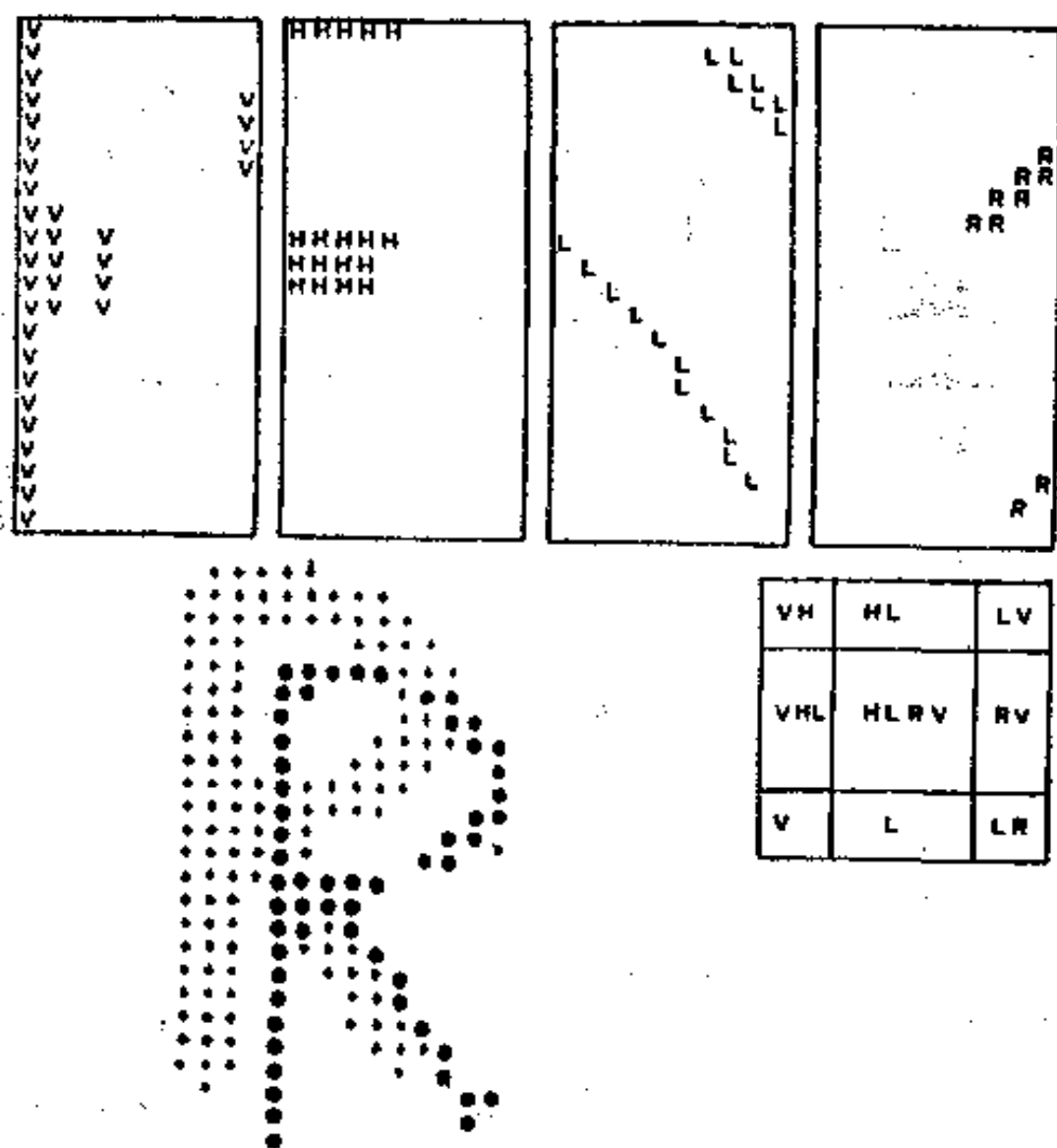


图 3-11 识别字母的一般方法

图 3-11 左下，是输入字母 R 和它的轮廓。这个轮廓由一个矩形框把它框起来，以划定字母 R 构成一个图的“图域”。

然后，把这个矩形图域分成九个区，如图3-11右下部所示，并对这九个区进行测试。例如，在输入字母是R的情况下，要测试的项目是：（1）图域的左边出现一条垂直线；（2）有一条“D”曲线，从左边顶端区开始，到左边中部区结束；（3）在图域的中底区和右底区出现左斜线。这里所说的“垂直线”等就是图象元素；“左边”就是元素应该满足的关系。

在识别时，沿着四个方向把轮廓分成一些线条，即垂直线（V），水平线（H）、右斜线（R）和左斜线（L）。在图3-11右上部，字母R用这些线段分开来表示。当然，分析的结果是产生一个合成的输出，它要指明图域的九个区的每个区中出现什么线段（见图3-11右下部）。再进一步的加工，就是利用一些专门的算法，确定图域中哪些线段的结合，意味着出现某个满足一定关系的图象元素。然后，再根据这些图象元素及其关系，计算出当前的刺激是一个什么字母。

2. 语音识别，查报电话号码 清华大学研制的“声控自动查报电话号码系统”，是汉语语音识别的一个实际应用。这个系统用微处理机（苹果II）代替人脑，记忆许多单位的电话号码，根据话务员的口呼命令，查报出相应单位的电话号码。它的工作原理见图3-12。

语音经过话筒后，由分析仪器进行语音的频谱分析，构成待识别单词的语音特征语谱。在使用前，计算机内已存储了许多单词的语谱的参照（标准）样板。这是由话务员给它“训练”而形成的。也就是说，话务员首先按顺序口呼所有单位的名称，系统就相应地产生有关各单位名称的语音语谱的样板，并存入系统的样板库中，在识别时作为比较标准使

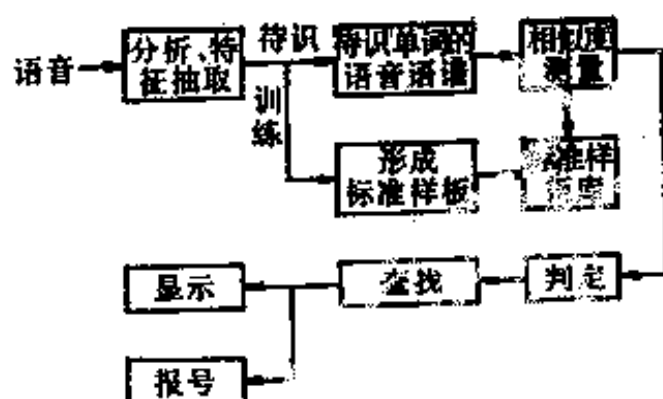


图3-10 系统工作原理框图

用。在进行比较时，待识别单词的语音语谱与某个参照样板最接近，就被判定为这是输入单词的最佳匹配对象。系统在识别输入模式时，不必与记忆中所有的样板进行比较，因为其识别过程是分二级进行的。第一级是大单位，第二级是小单位。话务员首先口呼大单位名称，该名称作为输入模式，只需与记忆中关于大单位的参照样板进行比较。大单位被识别并显示后，再呼下属小单位的名称。这时，输入模式仅与该大单位下属的小单位的语谱参照样板进行比较，而不必与所有大单位及其下属小单位的语谱样板去进行比较。这既省去了许多比较的操作，亦有利于识别的正确性。

从上述可以看出，机器的识别过程，与人的模式识别是类似的。事实上可以说，机器识别只有建立在对人的模式识别深入认识的基础上，才有广阔的前景。

第四章 注意和人类信息加工

意识问题曾经是早期心理学研究的主要课题，最初的研究主要是依赖于内省法，随着心理学研究的发展，内省法本身的不可靠性和不科学性越来越明显地暴露了出来，心理学家们对于用这种方法进行研究所得结论的可信程度产生了很大怀疑。取而代之的行为主义研究方式，抛弃了对意识的研究，使得关于意识问题的讨论长期处于不被重视的状态。认知心理学的兴起，使得有关意识的研究再次成为心理学家们感兴趣的问题。但是，今天的认知心理学家很少再宣称他们研究的是意识，而是说他们在研究注意。事实上，他们的研究内容许多方面都是在意识范畴之内的。注意就是知觉到某种事物；注意还意味着选择，当你关注某一事物时，你就选择它并或多或少地忽略了其它的事物；注意还涉及到了觉醒，一个人只有在一定的觉醒状态下才能有所注意并维持注意，因此，觉醒状态通过注意也将影响到信息加工。

一、人的注意

如果我们去参加一个晚会，就会处在一个充满了各式各样的讨论、交谈的环境中，我们将不得不选择和谁进行交谈或者倾听谁的谈论。我们将尽量使自己参加这个谈话并尽可能地排除其它谈话和声音的干扰。如果你想同时参加两个或更多的交谈，你会感到相当困难，往往会顾此失彼，这表明

了人们同时参与几件事的能力是有限的，这就是“注意”的某种限制。一方面，它允许我们在同时发生的众多事件中选取一组我们感兴趣的事件，去参与，去了解；另一方面，它又限制我们对同时出现的所有事物都进行干预的能力。尽管注意不允许我们同时去了解几件事物，但是，可以对同时发生的几件事情都保持某种程度的警觉。这样，当某一事物对你来说感兴趣的程度增加并且超过了你当前正在参与的事物时候，你的注意中心就会很快地转移到该事物上面来。这种能力的存在保证了我们能够在自然界里进行正常的交往和生活。

二、关于注意的“过滤器” 理论(Filter theory)

(一) 早期选择模型

1985年英国心理学家布鲁德本特(D. E. Broadbent)出版了他的《知觉和通讯》一书，书中详细地论述了有关注意的问题和他的研究成果。他把人看成是一个通讯信道，它的容量受到人对外界事物的注意容量的限制。注意被类比成某种类似于“过滤器”的电子设备，而这种设备的原理和内部结构是了解得比较清楚的。

在他的一个实验中，他让被试者通过耳机听一些数字，这些数字被分为三个一组，例如：一组为2，6，1；另一组为7，9，5。两组数字分别从左、右耳输入，例如：左耳输入2，6，1；右耳输入7，9，5。输入是左、右耳同时进行的，例如：左耳输入2，同时右耳输入7。这种刺激每隔半秒钟输入一对，一对一对连续进行，直到全部输入完，然后立即让被试者再现。通常被试者可以轻易地将这六

个数字全部正确的再现出来，但是要求被试者完全按照输入时的顺序来再现，即，2，7，6，9，1，5或者7，2，9，6，5，1，被试者就比较困难，准确度也比较低。被试者更喜欢首先再现来自一个耳朵的所有三个数字，然后是另一个耳朵的三个数字。

另一个实验是由美国心理学家E. C. 切里做的，这个实验模拟了前面介绍过的晚会上的情景。被试者通过耳机接受同时输入的两个不同的信息，通常是一些小故事，要求被试者只注意一只耳朵传来的信息，排斥另一只耳朵来的信息。为了保证被试者收听的效果，切里要求被试者听后立即重复从指定耳得来的信息。在这个实验中，当要求被试者重复的指定耳的信息快速传来时，被试者便反应说，他们几乎听不到来自另一个耳朵的信息了。在实验后的提问中，被试者对非指定耳的信息也基本上没有什么记忆。无论指定被试者用哪一只耳朵来收听，其结果都是一样的。此外，被试者也同样不能觉察，在非指定耳上信息发生的一些变化，例如，被试者注意不到非指定耳的信息在语言上的变化（从英语变为德语）。他们仅能注意到非指定耳的信息在声调上的变化（从男声变为女声）等简单的、物理特征的变换。

根据上述的实验结果，布鲁德本特提出了一个关于注意的早期选择的模型。他把人的两耳看成是分隔开的两个信息通道，每一次只能有一个进行工作，因此人是一个串行的加工器。但是，人具有一种补偿的能力，例如存在一个能够将原始数据存储几秒钟的缓冲记忆，这种缓冲记忆允许人们在一个通讯信道接收信息的时候将另一个信道接收到的信息保持住一会儿。随后，当一个信道的信息加工完毕后，所保持的另一个信道的信息能够被重新提取出来进行加工，因而人们

能够用这种缓冲记忆，把容量有限的过程用时间有限的过程来替换。布鲁德本特的模型如图 4-1 所示。

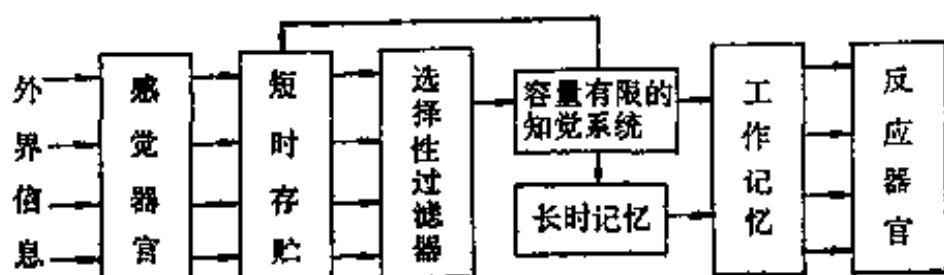


图 4-1 布鲁德本特的早期选择模型

从这个模型我们可以看出，外界信息经感觉器官到达短时存储器中进行暂存，然后经过选择性过滤（也就是注意）将无用的信息“滤掉”，仅让那些要进行认知分析的信息进入知觉。这就是布鲁德本特的注意的过滤器模型。这个模型是根据感觉特征来选择信息的，是按照“全或无”的原则工作的。由于他的这个模型是在知觉水平上的选择，所以被称为“早期选择模型”。

这个模型较好地解释了他本人的实验结果。对于来自一个耳朵的感觉信息首先进行加工，来自于另一个耳朵的信息暂存短时存储器中（相当于感觉记忆），一旦容量有限的信道中的内容加工结束，暂存的信息便通过选择性过滤器被调入并进行加工。从模型中我们可以发现，所谓完全并行的加工在这个模型上是不可能的，分时串行的原则被假定了。同样，这个模型也可以被用来解释切里的实验。此时过滤器已被调整成仅能通过要求重复的那些信息，但仍然保持着对所有输入信息的主要特征的监视，由于那些没有被要求重复的信息并没有进入意识，因此，这些信息在语义方面的变化是察觉不到的。当这些信息的主要物理特征发生变化时，如：

男声变为女声，因为受到监视，所以仍能够察觉。此外，布鲁德本特的模型允许对某些外界信息，进行某种程度的非意识状态下的知觉分析，只是这种分析是很简单的而且不涉及到它们的意义。

（二）中期过滤器模型

继布鲁德本特的理论提出之后，关于注意的大量研究相继出现了。在一系列的研究中，“早期选择模型”遇到了新的实验数据的挑战。有人在重复切里的实验时发现，对于进入非指定耳的信息来说，如果是被试非常熟悉的，比如，他的名字，则被试同样能产生知觉，并在事后报告出来。这种结果在用视觉刺激进行实验时也同样存在。这种情况与该模型对非注意信息不涉及其意义的预计相矛盾。六十年代初，特雷斯曼（A. M. Treisman）对“早期选择模型”进行了修正。她曾经作过这样一个实验，仍然是让被试者从左、右耳收听不同的信息，指定被试者报告从一个耳朵听到的内容，忽略掉从另一个耳朵听到的内容，收听的内容是一些小故事，在耳机上连接了一个开关，可以把一个耳朵听到的内容转换到另一个耳朵中去。在实验中，把被试者要报告的传入到指定耳中去的内容突然转换到非指定耳中去，同时，将另一新的内容立即送入指定耳，这种转换和连续过程是非常迅速的。按照布鲁德本特的理论，此时被试者将接着报告从指定耳听到的新内容，原来在指定耳现在被转换到非指定耳的内容应当被立即忽略掉了。而事实都不是这样，被试者往往是接着报告原来在指定耳听的那个小故事的内容，尽管现在它是从非指定耳输入的，要过一会儿，被试者才能开始报告指定耳所听到的新内容。这个实验表明，注意的选择性不仅

依赖于感觉特征，而且依赖于语义特征。而对外界刺激信息的语义分析是在信息加工的较后的阶段进行的，所以，还应该存在着另外一种过滤器，负责在语义方面进行选择。这就是说，在信息加工的不同阶段存在着两种过滤器，一个是负责感觉特征方面的选择，一个是负责语义特征方面的选择。前者很可能是对信息的感觉得特征进行分级性选择，而不是布鲁德本特的“全或无”方式，后者是对语义特征进行选择。所谓分级性选择是指，不是完全阻断感觉特征信息，而是程度不同的加强或减弱它们。这种加强和减弱过程受语义特征过滤器的调节。我们称特雷斯曼的模型为“中期过滤器模型”，因为她提出了在信息加工的中间阶段语义特征过滤器的作用。如图4-2所示。

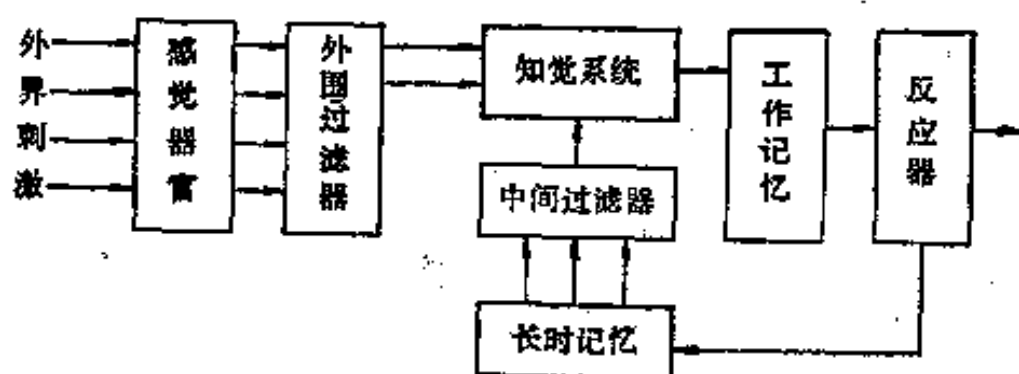


图4-2 特雷斯曼的中期过滤器模型

(三) 反应选择模型

与特雷斯曼提出“中期选择模型”几乎同一时间，格雷和韦德伯恩(J. A. Gray and A. A. Wedderburn)报告了他们的另一个实验发现。他们对布鲁德本特的实验稍稍作了改进，呈现给被试者的听觉刺激不是一对对数字，而是词

和数字构成的组合对。例如：耗子和 8，2 和吃，奶酪和 5。呈现的方式如下：

	左耳	右耳
第一对：	耗子	8
第二对：	2	吃
第三对：	奶酪	5

根据布鲁德本特的观点，被试者的再现应当是：耗子、2，奶酪；8，吃，5。而事实上，被试者是这样再现的：耗子吃奶酪；825。

格雷和韦德伯恩是这样解释他们的实验结果的：来自外界的所有信息都被登记并且保持下来，对这些信息全部进行了知觉分析，这种分析是在非意识状态下完成的，它并没有受到注意的限制，知觉分析的结果都被存入记忆之中。当知觉分析的结果被重新从记忆中提取出来以便得到一个反应的时候，某种选择的过程发生了。在实验中，被试者已经接收了全部六个刺激并把对于这些刺激的知觉分析结果都存储在记忆之中了，其中也包括对刺激信息的语义分析。当需要作出一个反应的时候，被试者便重新把这些刺激以及对它们的分析结果从记忆中提取出来，由于他们已经完全识别了这些刺激信息并且了解了它们的意义，因此，在他们构造反应时是根

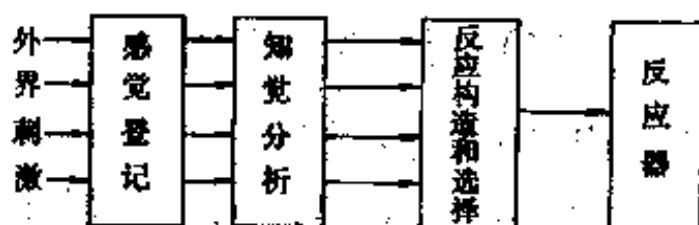


图 4-3 格雷和韦德伯恩的反应选择模型

据刺激信息的意义而不是它们到达耳朵时的顺序。从这样的角度来分析，注意并不是在对外界刺激信息进行选择而是对外界刺激的反应选择，因此，有人称之为反应选择理论。其模型可表达如图 4-3。

(四) 晚期过滤器模型

格雷和韦德伯恩已经把与注意有关的选择过程设立在整个信息加工过程的相当晚的阶段上。同样持有晚期选择观点的其他心理学家并不赞成这样一种解释。诺曼(N. A. Norman)等人提出了另外一种晚期选择的模型，如图 4-4 所示：

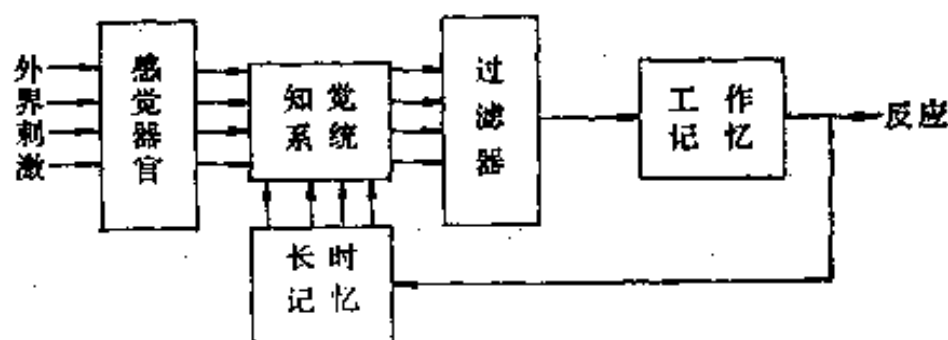


图 4-4 诺曼的晚期过滤器模型

从图 4-4 我们可以发现，此时代表注意选择的过滤器位于知觉系统和工作记忆之间。诺曼在他的理论中特别强调了在非意识条件下的知觉分析过程，也就是所谓“自动化”的过程，这个概念已经在有关注意和信息加工的研究中被广泛地接受和采用了。诺曼认为外界刺激可以完全无意识地激活人类信息加工系统中的相应知觉模式，而选择的过程仅仅对所有被激活的模式发生作用。诺曼曾经再一次用他的“鬼

域”模型来形象地解释他的理论，这里，诺曼在原有的映象鬼、特征鬼、认知鬼和决策鬼的基础上又加上了一个叫“监控鬼”的机制，用它来代表注意，起那种类似于过滤器的作用。当外界信息来到时，各种功能的小鬼便开始工作。两个外界刺激同时来到，这些小鬼便一分为二，一部分负责一种外界刺激的分析，另一部分负责另一种外界刺激的分析，而监控鬼则注视着所有小鬼的工作情况，并根据背景，外界刺激的意义性等各种因素来确定应进行主要分析的目标信息，指导绝大多数具有特定功能的小鬼去为这个主要分析服务。用这个理论可以解释特雷斯曼的实验结果，当要求报告的信息从指定耳被转换到非指定耳时，由于原先的主要分析过程已经形成，非指定耳得来的信息仍然适合这种主要分析过程的结构，内部期望也和它相适应，因此，那些具有特定功能的小鬼们仍在保持对它们的分析，对主要分析过程的结构进行更改是由监控鬼来控制的，这种更改需要时间。此外，这个理论同样能够解释格雷和韦德伯恩的实验，由于知觉分析已经完全自动化完成了，因此激活的刺激模式的意义影响到反应的构造。

上面介绍了关于注意的几种模型，这几种理论观点的共同之处在于，它们都认为在人类信息加工系统中存在某种类似于过滤器的机制，它们以某种方式对外界刺激信息进行选择，从而体现了人类注意的选择特点。所不同的是，在这种过滤器的位置上他们各自有不同的理解。近年来，上述不同的观点和解释仍然处于无法统一的状态下，并且研究已经转移到以视觉刺激为主的更广泛的领域中，一个更为流行的“瓶颈口”的概念已经代替了“过滤器”的概念。对于人类信息加工系统中“瓶颈口”的位置问题仍然悬而未决。

三、注意的资源分配理论

(Resource-allocating Theory)

从现在开始，我们讨论注意的分配问题，从另一个角度来讨论注意，把它和意识分开，把它看成是人类信息加工系统的某种可分配的资源。

我们先来观察一些日常生活中的现象。当你和一位朋友一起散步时，如果你突然问他17乘9等于多少？在他认真进行心算时，有时会不自觉地停下来。这表明，尽管行走的动作已经很“自动化”了，似乎它仍然需要人的某些注意或“管理”。我们在日常生活中可以发现许多类似的现象。当一个人同时完成两件事情的时候，常常不能立即完成，象仅做一件事那样顺利，这说明注意有一个分配的问题，也有一个限度的问题。我们需要研究的正是：注意是如何分配的和它的限度是什么？

首先我们来看一个具体的实验，这个实验要求被试者同时作两件事情，一是字母匹配，另一是对一种声调的反应。用右手来对字母的匹配作反应，即按一个键，反应精度被记录下来。在第一个字母呈现之前一秒钟有一个信号出现来提醒被试者，第一个字母呈现后半秒钟，第二个字母将被呈现，然后让被试者判断前，后两个字母是否一样，在作这种字母匹配的同时一个声音信号将会随机地被呈现，一旦被试者听到这个声音信号就要立即用左手去按一个键。图4-5是被试者对这个声音反应的时间曲线。

由图中可以看出，从提示信号出现直到这个信号出现后

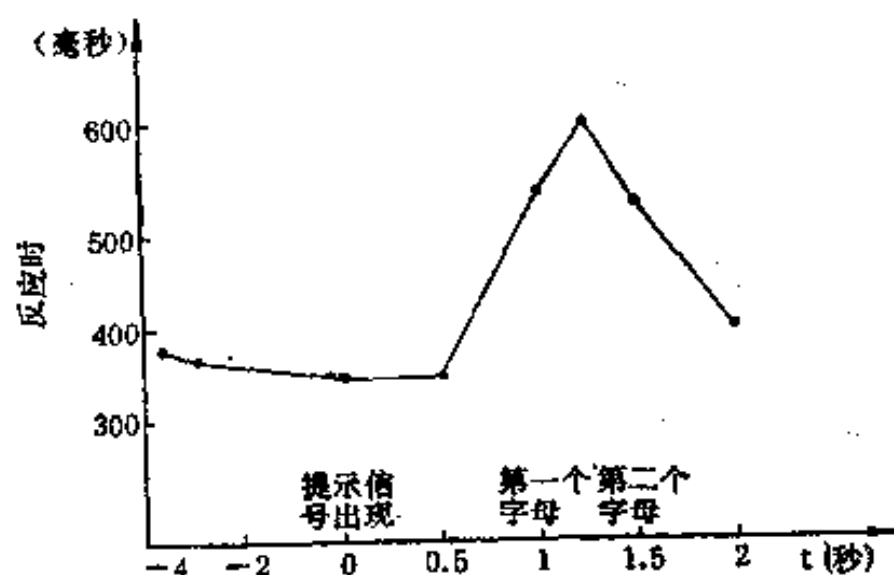


图 4-5 在进行字母匹配时，被试者对随机出现的
声音信号的反应时间曲线

不久，被试者对这个声音信号的侦察反应时基本是平稳的，但是，在第一个字母出现后到第二个字母出现前这个区间内，对这个声音信号的侦察反应时急剧地上升。这个结果表明，在这个时间范围内，被试者已将一部分注意分配给了字母匹配任务，所以，导致了对声音信号的侦察反应时增加。

1973年，凯恩曼 (D. Kahneman) 在他的《注意与努力》一书中提出了第一个关于注意分配的理论模型，如图 4-6 所示：

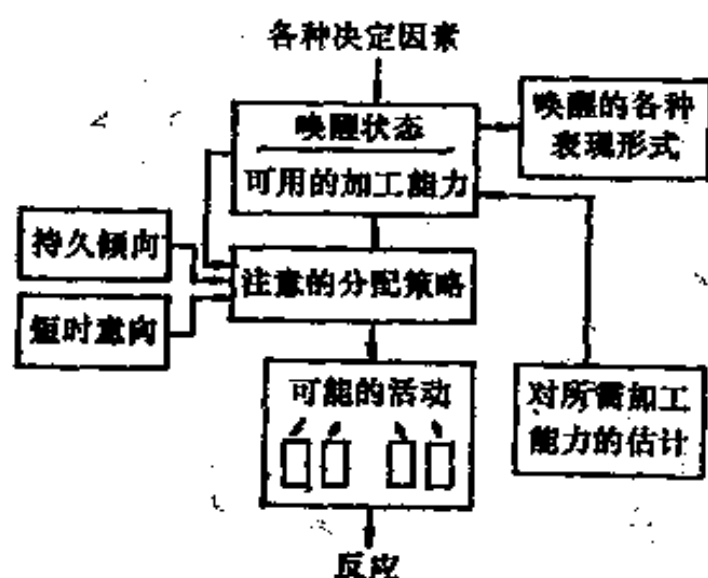


图 4-6 凯恩曼的注意分配理论模型

这个模型表明，人在各种因素影响下处于唤醒状态，外界刺激使人们将注意集中于某一件事物上或者在几件事物之间进行某种分配，分配的策略受到了几种因素的影响。持久倾向是指人们很熟悉的或者人们生而具有的认知事物的规则，例如：当有人叫你的名字时你常常能够立即知觉到。短时意向是指人们当时的主观决策，例如：人们要认真阅读一本书，尽量去排除各种外界噪音的干扰。分配策略还受到所注意的外界事物的难度的影响，模型中专门有一个机制来进行这种难度评价从而确定所需要的加工能力的大小，一般来说，越是困难的事物需要分配给它的加工能力越大。

凯恩曼的理论模型对注意的分配进行了一种定性的描述，但是，由此出发我们还不可能得到一种定量分析的途径，因此，诺曼和博布鲁(D. Norman and D. Bobrow)提出另一个理论模型，称之为“数据有限和资源有限的理论。”

他们认为，人类信息加工系统和一般信息加工系统一样，在它的内部存在着各种各样用于加工外界信息的资源，这些资源在数量上是有限的，这种资源上的有限导致了人们注意分配上的有限，当人们同时作两件事时，造成了系统对有限资源的竞争。图4-7是他们提出的理论模型。

图中显示，随着资源分配越来越多（即注意分配越来越多），被试者的操作水平越来越好，并最后到达数据有限区。在那里，虽然更多的资源被分配了，但被试者的操作水平将不能再进一步的改进。此外，诺曼和博布鲁还特别强调，在两个任务同时被操作的情况下，一个任务在难度上的变化并不一定必然导致另一个任务所用资源的减少，因此，另一个任务的操作水平可能不受同时进行的任务难度变化的影响。因为，它们之间在加工上可能并没有共享同一资源。

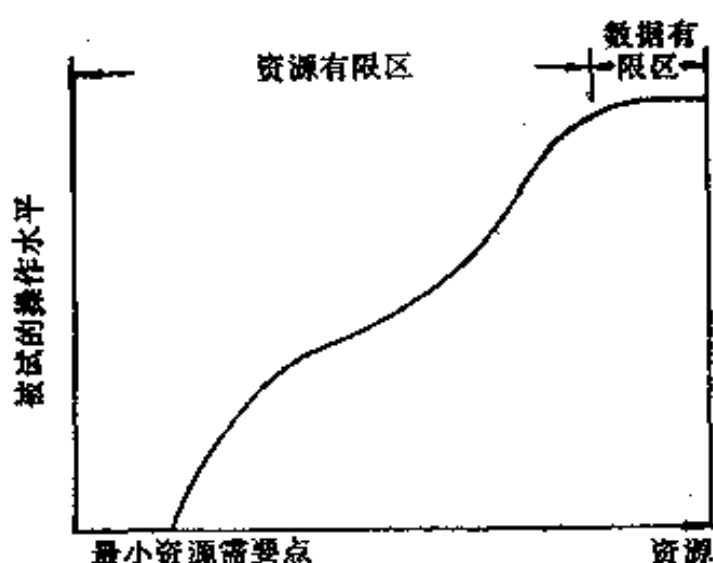


图4-7 诺曼和博布鲁的注意理论模型

诺曼和博布鲁的理论在有关注意的研究领域产生了巨大的影响，他们在理论上所提出的“POC”曲线（作业操作曲线，即Performance Operation Curve的缩写）已经被许多心理学家用来在实验中进行定量分析，从而确定某一特定任务在加工中所需的资源量。现今的许多心理学家越来越喜欢用资源的概念来解释他们的实验结果，而这些解释和说明从逻辑上来看与事实的吻合也是相当出色的。

尼翁和高佛 (D. Navon and D. Gopher) 进一步扩展了诺曼和博布鲁的理论。他们用一种微观经济学的方法，一方面定义了被试的主观倾向性，从而定量地描述了在两个任务实验中被试给予一个任务更高的优先权的问题；另一方面他们又阐述了资源使用的效率问题，进一步澄清了资源分配和操作水平之间的关系。

注意的资源理论的一般观点是，人类信息加工系统根据不同的任务目标来选择一定的输入进行加工，其它的刺激输入则被放弃掉。所放弃的信息已经进入了信息加工系统，只

是由于资源有限，它们不能被完全加工。

前面两节已经较系统地介绍了有关注意问题的两个大的理论，这两种理论尽管是从不同的角度，使用不同的实验方法来研究人的注意，但它们有一点是一致的，它们都认为，人类信息加工系统在注意上的限制或者是由于系统结构的限制（即，有一个“瓶颈口”存在），或者是由于系统资源的限制（即，可用于加工的资源有限或发生了对同一资源的竞争），也就是说，注意上的限制是由系统本身的固有属性所决定的。这样一种理论观点，当前已经受到一些实验证据的严重挑战，下面一节将简要地介绍奈瑟（U. Neisser）的注意理论和他的实验。

四、奈瑟的注意理论

奈瑟最初的实验是一种对布鲁德本特的实验的模拟，他采用视觉刺激，让被试者在一系列字母矩阵中搜索一个特定的目标。实验之前他首先对被试者进行训练，让他们熟悉实验方式和他们的任务。在搜索过程中，被试者看到的刺激似乎只是一个模糊的斑点，只有当目标字母进入视野后，才开始清晰起来。实验中的目标字母相当于双耳分听实验中要求报告的信息（也就是要求注意的信息），除目标字母以外的其它字母相当于双耳分听实验中的非注意信息。和前人得到的结果一样，奈瑟也发现对于非目标信息来说被试者基本上没有记忆，在整个搜索过程中这些非目标信息也没有被识别。但是对那些未进行识别的字母肯定已经进行了某种程度的加工，因为，如果不是这样，被试就不可能找到目标字母。那些非目标字母显然是在对它们进行了感觉特征的粗略分析后

被排除掉了。对注意中感觉特性分析的作用，奈瑟的观点在一定程度上与布鲁德本特的观点是一致的，但是在感觉特性分析的速度上，他们的见解不同。布鲁德本特认为这种分析在速度上是无差异的，而奈瑟认为它是有差别的，有时可能快一些，有时可能慢一点。例如，字母Z在由X、M、W等字母构成的矩阵中和在由O、C、Q等字母构成的矩阵中查寻的速度就不一样，前者要比后者慢一些。

奈瑟同样也发现了语义的影响。他把一段文章一行用红色印刷，一行用黑色印刷，这样红黑相间，要求被试者每隔一行来扫描这段文章，就是全看红色印刷的行或者全看黑色印刷的行。要求被试者扫描的行是目标信息，其它的就是非目标信息，但是被试者同样也能在非目标信息中发现他的名字。

奈瑟提出了“前注意加工”和“集中注意加工”两个重要的概念，他认为，前注意加工是一个快速的、总体的、很可能是并行的加工过程，这个过程负责那些非目标信息中的重要信息的加工，在使用过去存储的语义信息时是在无意识状态下进行的。“集中注意的加工”则是缓慢的、精细的、串行的加工过程，是从数据驱动和概念驱动两个方面同时进行的。

根据奈瑟的观点，目标信息和非目标信息之间的区别仅仅是在加工的程度，非目标信息的输入并不是完全被阻断了，而是它们没有被完全地加工。集中注意的过程导致意识的产生，它是将外界感觉信息和人们头脑中已有的经验相结合而出现的，是一个构造性的加工产物。在奈瑟看来，企图确定过滤器在人类信息加工系统中的位置是毫无意义的，因为知觉是主动的、灵活的。

奈瑟认为，注意中感觉特征的分析在速度上是不同的，在集中注意条件下，这种分析是受数据驱动和概念驱动两方面影响的，因此，复杂的心理活动有可能成为“自动化”的过程。奈瑟曾经训练两个被试者同时作两件对一般人来说是很困难的事情，即边进行快速阅读边听写单词，最初，被试者的成绩很差，大大低于他们仅完成一件事情时的精度和效率。但是，经过一个长达数周的训练过程，被试者的成绩逐渐地好起来，并且最终能够达到与仅做一件事的精度和效率完全一样的水平。在奈瑟看来，这是因为经过练习，复杂的信息加工过程进入了自动化的结果。下面的实验是另一个用以说明这种自动化作用的证据。

在这个实验中先用同类刺激的卡片20张（即刺激全由字母构成），让被试者搜索字母“J”，另一类刺激卡片是所谓异类目标刺激，即在字母中搜索一个数字，例如，数字“8”。实验结果表明：对异类目标刺激的扫描要比同类目标刺激快得多，搜索同类目标时，要想达到95%的正确率，每张卡片要呈现400毫秒，而搜索异类目标时达到同样的正确率，每张卡片只需呈现80毫秒。

对这个实验可以作如下的解释：人对于在字母中出现数字的情况已经在生活中形成了自动化，人看东西时先看整体，一看都是字母只有一个是数字，便一下子就挑了出来。但如果全是字母就不能用这种方式，需要逐一挑选。异类目标刺激下的扫描通常是处于前注意状态下的，而同类刺激下的扫描就不得不处于集中注意的状态下。但是，异类刺激概念的形成是人们的生活经验所导致的，那么，没有生活经验的事物经过训练也应能达到那种自动化的过程，下面的实验室训练说明了这种情况。

有两行字母：BCDFGHIJKL

QRSTUVWXYZ

要找的目标总是在第一行，干扰总是在第二行。这样训练二千一百次之后，被试渐渐将这些字母看成两类事物。最后也能达到前面所说的80毫秒呈现时间和95%的正确率的水平。这说明了，这种对事物的整体的反应确实是后天学习的结果。

以上简要地介绍了奈瑟的观点。奈瑟特别强调了练习对注意的影响，从而对那种认为注意是被系统本身的固有属性所限制的观点提出了挑战，尽管奈瑟的观点还不能被广泛的接受，但是，这些实验的事实至少应当引起我们的高度重视，并把练习的影响作为不可忽视的因素之一来加以考虑。

近几年来，有关注意的研究越来越成为广大心理学家关注的中心问题，尽管许多理论已经产生，一些模型也提出来了。但是，几个非常重要的问题依然不能做出定论，例如，

“瓶颈口”到底存在于何处？或者它根本就不存在或根本就不固定存在于某处。因为我们已经看到各种不同的“过滤器”模型。另一个问题是，什么是资源？尽管许多人都用资源的概念来解释他们的实验结果，甚至来测定加工资源的大小，但是，很少有人能够直接证实加工所使用的资源是什么。如果这个问题不能很好的解决，就会大大阻碍人们对信息加工系统本身的结构和特性的研究。

一般的认识是，信息加工系统内部有几个选择点，在加工的早期阶段，主要是根据刺激的物理特征来选择信息，从感觉输入中抽取的信息和从中枢处理器来的指令共同激活长时记忆中的相应成分。在各种供选择的知觉物中，一个或几个被选择出来进行进一步的加工，此时便产生了意识。有意识

的加工相当于与中枢处理器相联系并受其控制的活动，即相当于工作记忆中的加工活动。另一方面，注意的分配是受多种因素影响的，例如：信息的难度，人们对所接受的刺激信息的熟悉程度等等。此外，练习过程，特别是知觉学习都会或多或少地对注意分配产生影响。总之，注意的研究越来越受到重视，成为心理学研究的中心课题。我们仍然不能清楚地了解注意到底受什么机制控制，进一步探索是非常必要的。

第五章 短时记忆中的信息加工

前面已经讲到，外部刺激输入以后，进入人的感觉寄存器（感觉记忆），其中一部分经历了模式识别过程。这些被识别的模式或代码，将被传送到“短时记忆”（Short-term memory，简称STM）中去。除了短时记忆这个名称以外，也有人把它叫做“工作记忆”和“直接记忆”。这是因为，它除了有存储信息的功能以外，还是当前进行认知活动或对信息进行加工的地方。例如，词的意义的提炼，心算时涉及的符号操作以及推理活动等等。

一、短时记忆类似一个工作台

为了形象化地说明问题，我们可以把短时记忆比喻为木匠在那里干活的木工房里的一个工作台。所有的木工材料都端正地、有条不紊地放在房间墙壁四周的架子上。他从某个架子上取下马上就要使用的那些材料——工具、木板等，并把它们放到工作台上。同时，他为了能够在工作台上干活，台上还必须要留下一定的空间。当工作台太杂乱的时候，他可以把材料有次序地堆积起来，以便能在台上进行工作。如果工作台上材料的数量增加太多，有些东西就可能掉下来，或者，木匠可以把某些东西重新放回架子上去。

这种类比是适合于我们这里所说的短时记忆这个概念的。我们可以把木工房里的架子看作为长时间记忆（简称L

TM)，是木匠工作要用的大批材料的库房。而工作台（上面有木匠干活所需要的空间和容量有限的存放材料的区域）则就是短时记忆，木匠在工作台上进行的操作，就好象是短时记忆中进行的活动。当木匠把材料有次序地堆积起来，以创造更多的空间时，就好象是短时记忆中进行的组块化活动。组块化活动把几个记忆项目合并成一个，使它们只占用短时记忆中一个单位的空间。从工作台上掉下来的东西，相当于短时记忆中遗忘的项目。把材料从架子上取下来或把材料放回到架子上去，这就好象是把信息从长时记忆中提取出来，放到短时记忆中进行加工，或者，信息从短时记忆向长时记忆转移一样。当然，还应该假定，在长时记忆这个架子上，对任何（架子上有的、可供应的）材料都有复制品。材料从架子上拿到工作台上加工时，只是复制品的转移，材料本身并没有从长时记忆中移走。因此，长时记忆是永久性的。

当然，这种类比是有局限性的。短时记忆是人的信息加工系统中一个比较复杂的，也是很重要的部位，这不是用工作台的比喻能对它的性质加以概括无遗的。但是，从这种具体的类比中，我们可以知道，短时记忆是一个相当可变的实体。在短时记忆中，可以存储各种记忆项目，也可以对它们进行加工。同时，工作空间和存储空间之间，亦可以有一种协调的变动。工作空间占得多一些，存储空间就得少一些，反过来也是一样。

二、短时记忆的容量

（一）有限的容量

短时记忆究竟能容纳多少记忆项目？这是一个令许多心

理学家很感兴趣的问题。一些实验的结果表明，短时记忆的存储容量是很有限的，也就是说，它只能容纳少数几个记忆项目。

可以说，艾宾浩斯（H. Ebbinghaus）是首先对短时记忆的容量进行研究的心理学家。他利用一些“无意义音节”

（由“辅音——元音——辅音”组成的音节，它们没有什么意义，如Ruk、DAF等）来探讨学习和记忆问题时发现，如果给被试者呈现的记忆项目表上只有7个或少于7个无意义音节时，那么，被试者在阅读一次以后，就能把它们记住，并准确地回忆出来。但是，如果表上的项目增加到8个或8个以上时，那么，为了把它们记住，学习的次数就要增加。在7个项目这个数的周围，好象有一个突然的中断。刺激表上的项目低于这个数，被试者可以阅读一次而直接记住它们；若记忆项目高于这个数，被试者需要对它们阅读数次才能记住。为了达到记住的目的，学习的次数将随着记忆项目数量的增加而增加。这种阅读一次就能记住的项目的数量，叫做“短时记忆的广度”。人们把短时记忆的广度看作为短时记忆容量的指标。

从上述实验结果看来，7个无意义音节这个数是短时记忆广度的值，也就是短时记忆存储容量的限度。这好象是短时记忆中有一定数量的槽，大约有7个，每个无意义音节放进一个槽里。所有的槽填满以后，余下来的项目就没有地方存放了，因而就没有被人记住，在回忆时就出现“漏掉”的现象。

（二）组块

根据上面所述，短时记忆的容量似乎是7个无意义音节。但事实上，它也可以是7个字母（如果这些字母并不构

成音节或词的话），或者是7个词。这就是说，短时记忆的容量并不是以任何特定的单位——字母、音节或词——来加以定义的，而似乎是大约7个被呈现的无论什么样的单位。因此，被试者可以记住7个没有构成任何特定模式的字母，如：X、P、A、F、M、K和I。但是，如果字母构成了词，那么，被试者就能记住多得多的字母。这是因为他能够把多个字母的系列，再编码成一个单位，只要这种字母系列构成了有意义的词。以这种方式对呈现的刺激进行再编码的能力，即把一些单个刺激（如字母）结合成较大单位（如词）的能力，就是前面提到过的“组块化”。通过“组块化”而产生的新的记忆单位，称之为“组块”。这个科学术语是由米勒（G. A. Miller）在1956年首先提出来的。当时，他发表了一篇文章，题目叫做“奇妙的数字，7加或减2”。他在文章中证明，短时记忆的容量应该以“组块”为单位来表示，它的容量是“7加或减2”个单位，即5至9个组块。这是一个非常有名的论断。

米勒认为，短时记忆的容量可以用某种单位来测量，而且，这种单位的内部结构是可变化的。短时记忆中的一个单位，相当于一个组块。组块是一种可变的客体，它所包含的信息可多可少。那么，具体说来，组块究竟是什么呢？假定说，我们给被试者有次序地呈现几个“三字母词”（即由三个字母构成一个词）的字母，如：C、A、T、D、O、G、F、A、R等，那么，我们就会发现，在直接回忆即在刺激呈现以后立刻进行回忆时，被试者大约能够记住21个字母，即7个词的字母数量，而不是7个字母。这是因为，被试者在记忆这些字母时，不是孤立地进行的，而是把它们组块化成7个有意义的单位来记忆的。在这种情况下，我们可以

说，一个词就相当于一个组块。同样，如果给被试者有次序地呈现一系列“四字母词”的字母，那么我们可以预料，他能记住的字母数量将不只是21个，而是更多一些。这是因为，被试者把每4个字母组块化成一个有意义的单位来记忆的。

总之，组块化过程是一种把许多个别的信息单位结合成一个较大单位的操作。我们可以把组块化的操作分成彼此多少有关的两类。第一，每一组块是把那些在时间或空间上接近发生的单个项目结合成一个组，并且，结合成组的项目不一定需要形成一个有意义的单位。比如说，当某些项目有节奏地呈现时，对它们的直接回忆的成绩就较好，而当各项目以固定的速率呈现时，直接回忆的成绩就较差。显然，当某些项目作为一个时间(或空间)组出现时，就容易把它们组合成一个单位。而有节奏的呈现，正是起了这种功能。正是由于这个原因，人们常常把这类组块化叫做“分组”。由“分组”而形成的组块，其中的各单个项目不一定彼此构成有意义的联系。第二，在意义上把若干项目联结起来，构成一个已知的、单独的组块。这种组块化过程，要利用长时记忆中的信息。例如，单个字母C，A和T；可以结合成“cat”（猫）这个组块（词）。在进行这类组块化时，就要求被组块化的各项目之间，存在某种固有的、使它们能形成一个整体的联系。尤其是，如果一组刺激的结构刚好与长时记忆中的某个代码相匹配，那么，就可以说，这组刺激会形成一个与该代码相应的组块。

在组块化过程中，先前获得的知识具有重要的作用。比如说，当人们学会了仅仅为某种目的而设计的特定规则时，他们就可以利用这类规则以促进对客体的组块化。米勒曾报道过下面这样的事实。他要求被试者先学会把三位数字的模

式转换成单个数字,如:000=0;001=1;010=2;011=3;100=4;101=5;110=6;111=7。然后,他又给被试者呈现一系列数字,比如:001000110001110。被试者把这一长串数字转换成三位数字的模式(即001,000,110,001,110),并利用上述规则把这些三位数字转换成单个数字。结果就产生10616这个数字序列。这样,被试者就学会了把一长串0和1再编码成一串较短的数字。用这种方式经过良好训练的被试者,能够记住21个左右的0和1的数字系列。

对人来说,组块化的过程是非常重要的,因为它提供了一个超越短时记忆存储空间限度的手段。成功地对客体实现组块化,就意味着在短时记忆的每个“槽”中能装入较多的东西。这样,短时记忆能容纳的总的信息量也就增加了。

(三) 短时记忆容量的个别差异

对于不同的人来说,短时记忆容量的大小可能是不同的。你的短时记忆容量可能是9,而其他一些人或许只有5。短时记忆的容量不仅随着年龄而变化,而且在同一年龄的个体之间也是不同的。所以,早期有人把短时记忆容量在个体之间的这种差异,看作为心理能力的一种尺度。在测量心理能力时,给被试者呈现的记忆项目表上是一些单个的数字,要求他按顺序回忆这些数字。起初,记忆项目表较短,然后逐渐增加表的长度,直到回忆发生错误为止。在第一次出现错误之前的那个项目表中的项目的数量,可以看作是短时记忆的广度或容量。

人们确实有不同的记忆广度,这个事实使心理学家们感到奇怪,为什么会产生这种差异呢?现在有几种不同的解释,这些解释恐怕都不是完美的。第一种解释认为,儿童和

智商 (IQ) 低的个体, 他们的短时记忆中的“槽”较少, 所以, 他们具有较小的记忆容量。但是, 这种解释有点难于考证。一则, 当用不同类型的材料来测量记忆广度时, 即使是同一个个体, 其广度也是有变化的。既然说短时记忆中有固定数量的“槽”, 那为什么会出现这种变化? 这有点不好解释。二则, 要把短时记忆中“槽”的数量所产生的影响, 与个体对信息实现组块化的能力区分开来, 这是有困难的。也就是说, 一个儿童记住的字母少, 这究竟是儿童短时记忆中用来存储字母的“槽”比成人少, 还是因为成人在对字母加以组块化方面的能力比儿童强, 所以, 在每个“槽”中装了一比较多的字母呢? 某些成年人用特殊的诀窍, 来练习对信息进行组块化, 结果, 短时记忆能记住的项目数量大大增加。但是, 儿童却不象成人那样地利用组块化的策略, 所以, 可以肯定地说, 他们的短时记忆的容量就会相对地小。

因此, 第二种解释则是: 组块化能力的个别差异是记忆广度或短时记忆容量的个别差异的基础。在儿童的长时记忆中, 可以用来作为一个组块之基础的信息比较少, 例如, 他们掌握的词、知识比较少。所以, 他们的组块化的能力较差, 因而造成短时记忆容量较小。但是, 根据这个思想人们可以推测, 如果在对刺激项目进行“分组”时, 给被试者适当的帮助 (例如, 指导被试者如何对项目进行分组, 或者从实验者方面给予有节奏的分组), 那么, 个体之间的这种差异就会减小。因为, 组块化能力差的人, 会由于这种帮助而提高他们的能力。而组块化能力强的人, 没有这种帮助也能做到这一点。但是, 奇怪的是这种预测没有得到实验者的支持。告诉成年被试者如何对项目进行组块化, 这对低广度一组的被试者就象对记忆广度高的一组被试者一样, 都没有什么帮助。

同时，对年轻的儿童就象对成人一样，也没有使他们从中得到帮助，从而提高短时记忆的容量。这一事实说明，第二种解释也不是十分有力的。

第三种解释则认为，短时记忆容量较大的人，是由于他们对记忆项目进行复述的能力较好。这种解释也缺乏实验的支持。如果我们以很快的速度呈现刺激项目，其速度快到足以排除掉被试者对它们进行任何复述的可能性，结果仍然没有消除个体之间的差异。图 5-1 是二个被试组在二种不同速

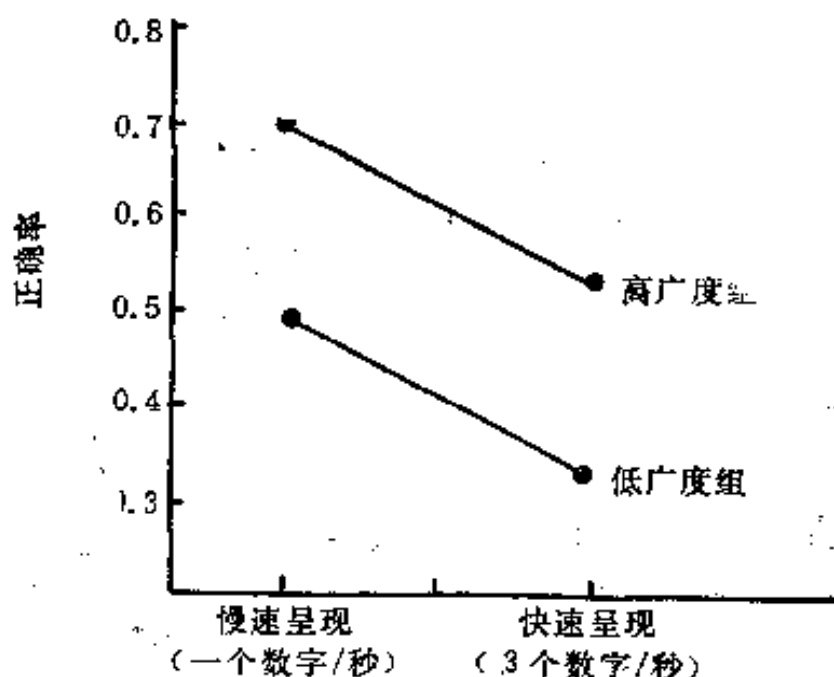


图 5-1 两类被试者对两种呈现速度的回忆成绩

度呈现刺激项目时的直接回忆的成绩。根据这种假设，快速的呈现将排除被试者对刺激项目进行复述的可能性，所以，短时记忆容量方面的个别差异应该被消除。但是，从图 5-1 中我们可以看到，记忆广度大的一组被试者，无论在快速呈现（每秒呈现 3 个数字）条件下，或在慢速呈现（每秒 1 个数字）条件下，其直接回忆的成绩，都相应地比记忆广度小的一组被试者好。所以说，造成个体之间记忆容量大

小差异的原因，也很难说就在于有些人复述能力强而另一些人复述能力较弱的缘故。

那么，究竟什么是造成短时记忆容量差别的原因呢？有人又提出了另一种解释，认为儿童或记忆广度低的人，在刺激进入短时记忆之前的加工——如感觉记忆或模式识别——中，它已经受到阻抑。当然，关于这一方面的具体情况还不清楚，需要进一步加以探索。

三、短时记忆中的检索

当进行回忆时，必须要从记忆项目中从短时记忆中检索或提取出来。人们或许会认为，检索过程就是对存储空间进行检查，如果项目在那里的话，就报告这个项目。但实际上，事情并不是那么简单的。

（一）典型试验的事件及心理历程

关于短时记忆中的检索问题，主要是用斯顿伯格提出来的一种叫做“记忆扫描”（memory Scanning）的实验方法来进行研究的。被试者要参加一系列试验。在每一试验中，给被试者呈现一简短的记忆项目表。比如说，每张项目表上有一至五个数字不等，如2、4、7、3。记忆项目表上刺激的数量，要低于短时记忆的广度。被试者的任务是要记住这些刺激项目，但是，并不要求他从记忆中回忆全部项目。给被试者呈现刺激项目表后，紧接着呈现一单个刺激，即测试刺激（或叫检测刺激）。这个测试刺激可能是表中的一个成员，也可能不是。如果测试刺激与记忆中的一个项目相匹配，那么，被试者就反应“有”，否则，就反应“无”。被

试者通常利用按键的方式来作出“有”或“无”的反应。

在这种实验条件下，被试者往往可以完全正确地作出“有”或“无”的反应。实验者要搜集的材料，主要的不是反应正确或错误的百分比，而是被试者的“反应时”，即从呈现测试刺激到被试者反应“有”或“无”时所花的时间。

在理论上，我们可以认为，“反应时”为内部心理事件的发生花了多长时间提供一个客观的指标。所以，在这里，反应时包括被试者完成接收测试刺激，并把它与记忆集进行比较这样一些过程所需的时间。如果完成这些活动需较多的时间，那么，反应时就长，反之，反应时就会短一些。

那么，被试者在本实验任务的“反应时”期间，进行着什么样的加工活动呢？上面已经说过，当测试刺激出现时，被试者在短时记忆中已存储有一集记忆项目。因此，可以认为，随后的加工活动包括三个阶段。第一，被试者接受测试刺激并对它进行编码，使它成为相应的心理形式；第二，把

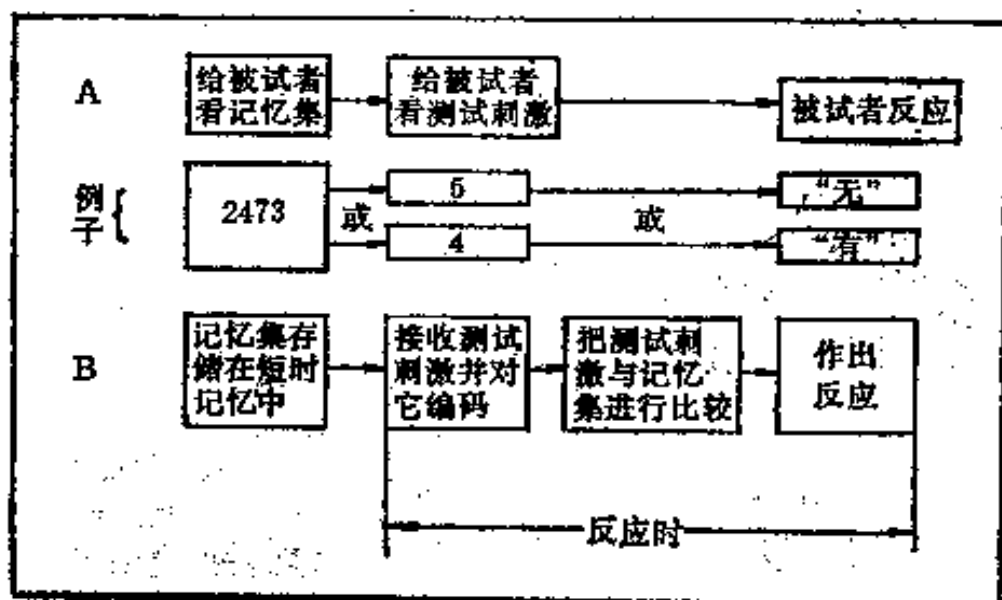


图 5-2 记忆扫描：(A)一个典型试验的条件；
(B)在试验期间的心理条件

测试刺激与记忆中的各个成员进行比较；最后，在这种比较的基础上作出反应。这三个阶段所花的时间的总和，就是反应时。图 5-2 给出了一个典型试验的事件及其被试者的心理历程。

根据反应时的变化，我们可以对被试者在试验任务的第二阶段所完成的比较过程作一些推论。在实验时，给被试者呈现的记忆项目数量是不等的。如果记忆项目数量多，毫无疑问，这意味着被试者要做更多的比较，因为测试刺激必须与记忆集中的每个成分进行比较。记忆集项目的多或少，对反应时是否有影响，这又取决于被试者以什么方式来完成比较任务。所以，弄清楚这些效应，就给我们一个工具去解释被试者是怎样加工这种信息的。

(二) 平行与串行的假设

首先，假定被试者在短时记忆中进行的比较过程是“平行”的，也就是说，被试者能在同一个时间里把测试刺激与短时记忆中的各个项目同时进行比较。那么，我们可以预计到，记忆项目数量的增加或减少，即记忆集的大小，对被试者的反应时不会产生什么影响。不管记忆项目是 2 个、3 个或 4 个，被试者完成这个比较过程所需的时间不会发生什么变化。这是因为，同时地把测试刺激与记忆集中的各个项目进行比较，并不比与一个项目作比较所花的时间多。这种平行模型预测的结果，如图 5-3 A 所示。

其次，假定发生的过程是“串行”的扫描，也就是说，被试者在一个时间里只检查或比较记忆集中的一个项目。在这种情况下，在记忆集中加上一个项目，就将增加完成这种扫描过程所需的时间。这样，我们可以预计到，随着记忆集增

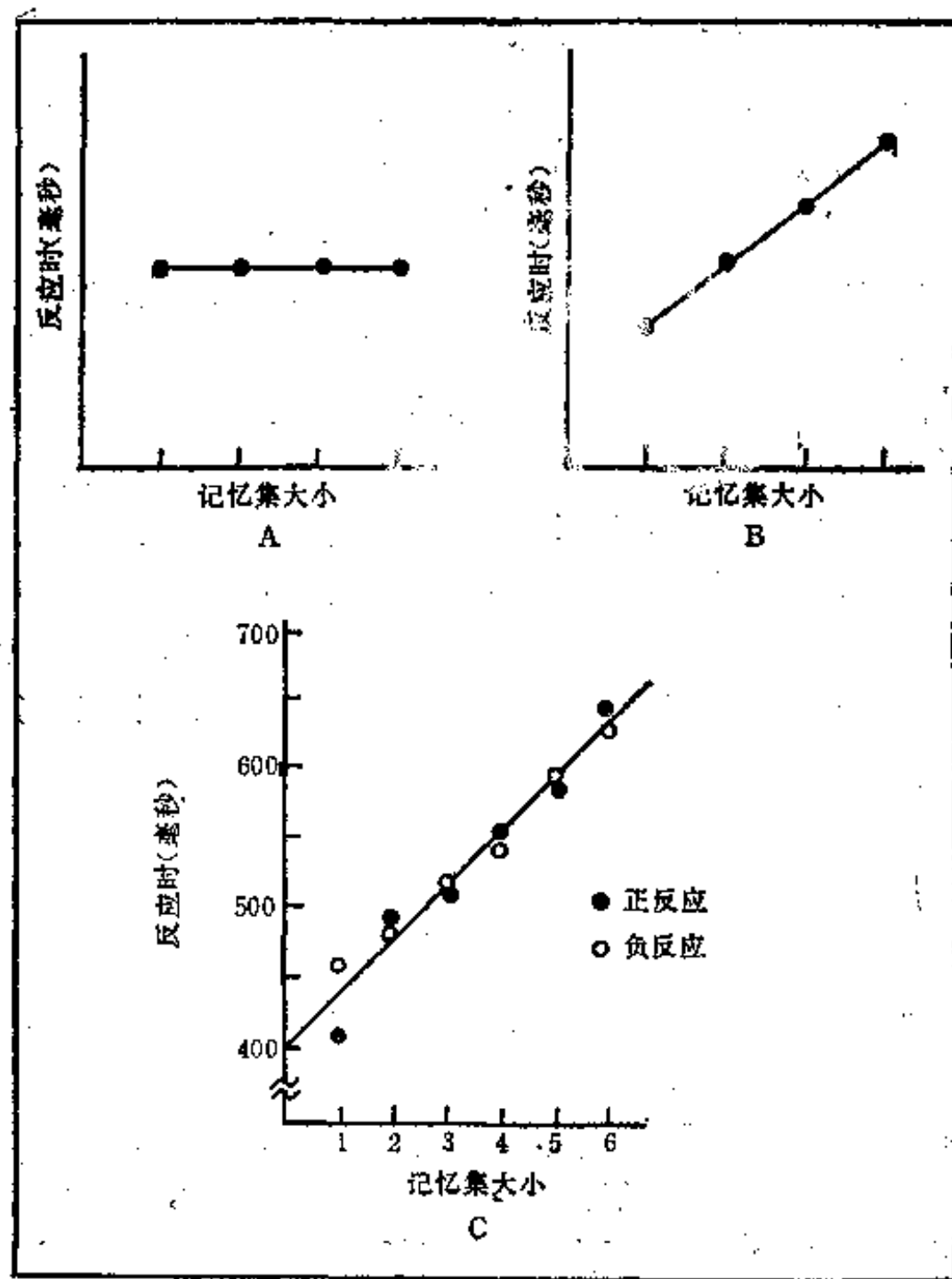


图 5-3 记忆扫描任务：(A)平行模型预测的结果；
(B)串行模型预测的结果；(C)扫描任务的实际结果

大，反应时也就增加。图 5-3B 给出了这种串行比较模型将得出的曲线。

(三) 自行终止与穷尽搜索的假设

再进一步分析,在串行扫描模型下面,还可能有二种方式。一种方式是“自行终止”,另一种是“穷尽搜索”。上面已经指出,在每个试验中,被试者的加工活动分成三个阶段,每个阶段都占用总反应时间量中的一部分。假定被试者对测试刺激进行编码需要花 e 毫秒时间;把记忆集中一个项目与测试刺激进行比较的过程需花 c 毫秒时间;而执行第三阶段(作出反应)则需花 r 毫秒。若记忆集中只有一个项目,那么,被试者可以在 $e+c+r$ 毫秒的时间里完成这个任务。这就是记忆集只有一个项目时的反应时。如果说记忆集中有五个项目,而且没有一个项目是与测试刺激相匹配的,那么,被试者要作出“无”的反应,必须花 $e+c+c+c+c+r$ 毫秒的时间。一般说来,在负反应(即反应为“无”)的情况下,被试者对测试刺激的“反应时”将是: $RT=e+N \times c+r$ 。在这里, N 是指记忆集的大小,即记忆集中包含的项目的数量。很显然,被试者只有把记忆集中的每个成员与测试刺激作比较之后,才能作出负反应。否则,他不能断定记忆项目中没有一个成员与测试刺激相匹配。所以,在这种情况下,记忆集中有多少个项目,就要比较多少次。

但是,在正反应(即记忆集中包含测试刺激,被试者的反应为“有”)的情况下,就比较复杂一些。这里可能会出现两种情况。一种可能是,被试者在发现某个记忆项目与测试刺激相匹配以后,加工过程就立即停止下来,不再去作所有其余的比较。这种搜索比较的方式称之为“自行终止”。自行终止的搜索意味着,每当被试者发现记忆集中某个项目与测试刺激相匹配时,他就终止扫描。第二种可能是,不管

是否发现测试刺激与记忆集中的某个项目相匹配，被试者在比较阶段总是要“穷尽”整个记忆集。也就是说，在比较阶段，不管他是否发现匹配，都要把记忆集中的每个项目与测试刺激进行比较。只有在“穷尽”整个记忆集之后，被试者才停止比较过程，作出反应。这种穷尽的比较，相当于负反应的加工过程。

通过实验上的安排，我们可以预计到，每当被试者发现测试刺激与记忆集中的某个项目相匹配时，平均来说，这个匹配项目将在检索该记忆集的中途发现。根据“自行终止”的理论，这就意味着，当作出一个“正”反应时，被试者只需搜索比较记忆集的一半，并在中途停下来（平均而言）；但是，在作出一个“负”反应时，被试者将要对记忆集的全部项目做完比较，走完整个路程。因此，如果被试者在进行搜索比较时，是“自行终止”的话，那么，对“正”反应来说，他只要进行 $(N+1)/2$ 次的比较。在这种情况下，他的反应时将是： $RT=e+r+(N+1)/2 \times c$ 。我们可以把这几项重新排列一下，使 RT 表示为 N 的函数，那么，我们就得到： $RT=(e+r+c/2)+(c/2)N$ 。从这里我们可以看出，反应时作为记忆集大小的函数，在“正”反应情况下，其斜度大约是“负”反应函数的斜度的一半。同时，根据“穷尽”理论，比较阶段对“正”反应和“负”反应来说，是没有差别的，两者都要作完所有可能的比较。因而，它们的反应时函数的斜度，也不存在上述这种差异。

（四）实验结果支持“串行穷尽”的假设

上面叙述了关于短时记忆中搜索比较的三种假设，它们是：平行比较、串行的自行终止的比较以及串行的穷尽式比

较。那么，究竟哪一种假设是符合于短时记忆中实际进行的搜索比较过程的呢？

为了检验这些假设，斯顿伯格做了一个实验。在实验中，他使用了若干种大小不等的记忆项目集，同时，使用了既有引起“正”反应的、也有引起“负”反应的测试刺激。这样，就可以得到每种试验类型（即“正”反应和“负”反应）和不同大小记忆集的平均反应时。所得的结果见图5-3c。从实验所得的结果来看，是支持“串行穷尽”扫描的理论的。从图5-3c中可以看出，根据“正”、“负”反应情境下获得的数据所画出的函数曲线，其斜率基本上是一致的。

这个结果是非常有趣的。“穷尽”式搜索比较似乎与我们的直觉背道而驰。这是因为，它意味着，在“正”反应的情况下，被试者在发现匹配以后，还要作许多不必要的比较。这是使人有点不可理解的。

当然，要解释为什么发生穷尽式的搜索比较过程，这也不是不可能的。我们可以把上述记忆扫描任务中的搜索比较过程分成二个组成部分。一是搜索比较行动本身；二是“决断”比较的结果是什么。如果被试者用来比较测试刺激和记忆集项目的时间非常快，而决断该比较的结果，即判定是匹配还是失匹配，却相对地慢，那么，串行穷尽的搜索可能是一种合理的加工方式。我们知道，若是自行终止的话，那么，穿越记忆集的过程将是：比较、决断、比较、决断……直至获得一个匹配或到整个记忆集被搜索完为止。另一方面，若被试者采用串行穷尽的加工方式，那么，其穿越记忆集的进程将可能是：比较、比较……决断。也就是说，在这种情况下，当记忆集被穷尽时才作出决断。因此，如果被试者能以极快的速度进行比较，而作出决断却相对地慢，那么，

穷尽的搜索将是一种方便而有效的加工方式，因为它只需要一次决断。

四、短时记忆中的编码

所谓信息在短时记忆中的编码问题，也就是信息以什么形式存储在短时记忆中的问题。

（一）听觉编码

有些实验表明，言语刺激经过模式识别以后，是以听觉编码的形式存储到短时记忆中去的。即使言语材料是以视觉的形式呈现时，也是以它们的读音而不是以它们的视觉形状来编码的。比如说，让被试者看 landlord 这个词，并要他记住。人们发现，他多半是记住了这个词的发音，而不是记住这个词的视觉形象。另外，从短时记忆再现时发生的错误中，也可以看到短时记忆存储是以听觉的形式进行的。例如，在实验中通过视觉给被试者呈现字母以后，让他把自己记住的字母写下来。他可能把一些字母写错。但从他所发生的错误中，我们可以看到一种很有趣的现象。比如说，他可以把C错写成T，而不是错写成视觉形象更接近的O；把F错写成M，而不是错写成E。这些结果表明，被试者所发生的错误主要是在声音上发生了混淆，而不是在视觉形状上发生混淆。

上述材料表明，识记的项目在短时记忆中存储时，是以听觉形式编码的。当被识记的项目回忆不起来时，他就根据还保存下来的部分听觉痕迹来重组遗忘了的项目。所以，他把C忘了，但还记得其中有〔i:〕的声音，于是把它重组成T；把F忘了，但还记得其中有〔e〕的声音，于是就把它重组

成M，如此等等。

但是，也有一些实验表明，信息可能以非听觉的编码形式存储在短时记忆中。所以，我们不能把短时记忆看作为只有一种特定类型的编码方式。

(二) 视觉编码

关于短时记忆中视觉编码的问题，波斯纳(M. I. Posner)曾做过一个有趣的实验。实验中，他要求一个被试者参加一系列试验。每一试验的时间很短。在每次试验中，给被试者看两个字母，并要求他指出，这两个字母是否有同样的名字(如A和a，B和b)，或者是否有不同的名字(如A和B)。被试者在执行这个任务时，是可以不犯任何错误的，所以，实验者要搜集的数据是被试者的反应时。在这里，“反应时”是由三个部分构成的：(1)知觉字母所需的时间；(2)把它们相互比较的时间；(3)决定它们的异同并按键反应的时间。从图5-4中可以看出，在两种情况下(即两个字母等同，或有同一名称)，被试者将作出“相同”的反应，这称之为“正”反应。除此之外，他将作出“不同”的负反应。被试者对这些条件的反应时是不等的。对等同匹配的反应时，差不多比名称匹配或负反应的反应时快0.1秒钟左右。这

试验类型	被试看到的刺激		被试的正确反应
等同匹配	A	A	“相同”
名称匹配	A	a	“相同”
负试验	A	b	“不同”

————— 反应时间间隔 —————

图 5-4 波斯纳字母匹配实验中的试验类型

就是说,被试者在执行这类任务时,其内部过程是有差别的。

为了找出差别究竟在什么地方,需要把被试者的活动进行分解,以便把造成反应时延长的那个成分孤立出来。被试者的活动可以分成:知觉字母、对这些字母命名、进行比较、作出反应等四个成分。在这里,知觉字母和按键作出反应这两者所需的时间,原则上应该是不变的。因此,造成反应时差别的原因,将在于命名过程或比较过程。事实上,实验者认为,反应时之所以出现差别,是因为当两个字母等同时,被试者根据它们的视觉形状,就可以作出“相同”的判断,没有必要对它们命名。只有当两个字母不等同时,才

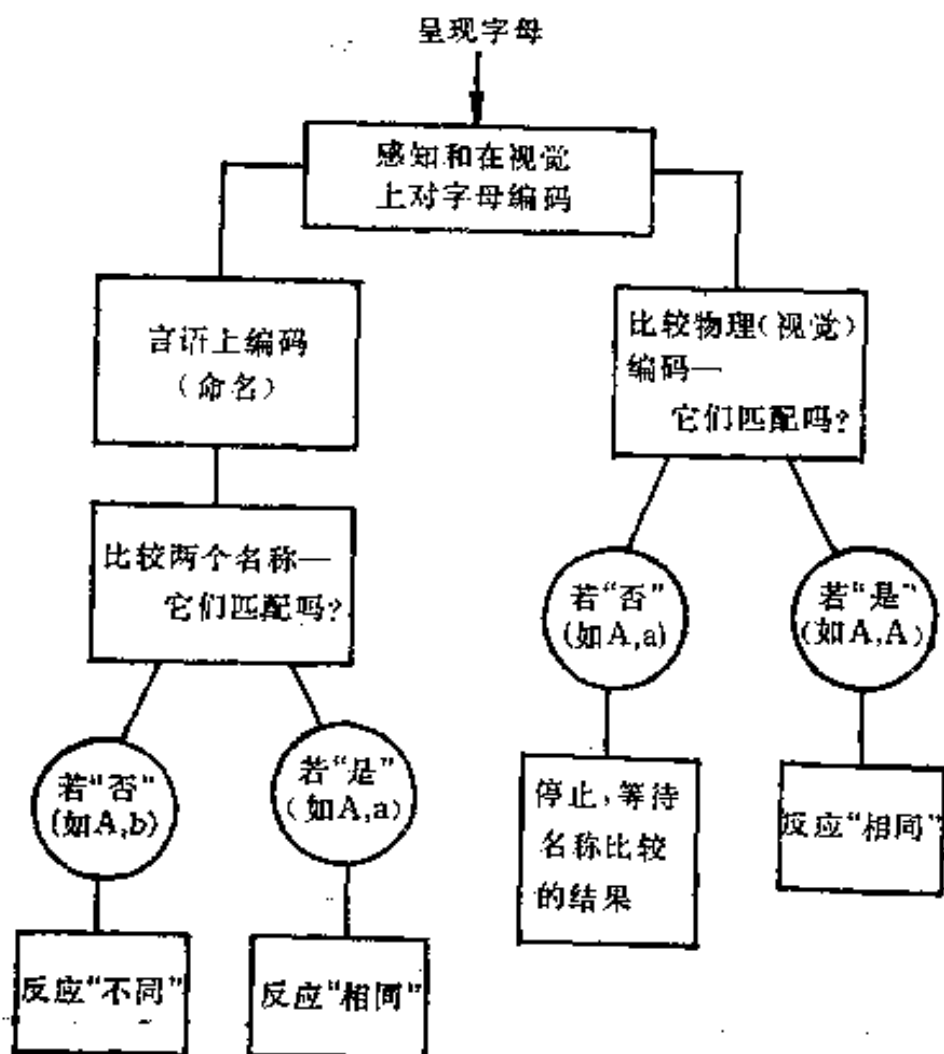


图 5-5 字母匹配任务中可能发生的加工过程

需要对它们命名并比较这两个名称。这里多了一个活动环节，从而导致反应时增加。简单地说，等同匹配是以视觉信息为基础的，而名称匹配或负反应，则是以言语命名为基础的。

图 5-5 给出了字母匹配任务中可能发生的加工过程。当然，在这个实验中，二个字母是同时呈现的，并且停留在视野中，直到被试者作出反应为止。

那么，在一个刺激从视野中消失以后进行的加工活动中，视觉信息是否仍然有效呢？

为了弄清楚这一点，可以把上面的实验修正一下，两个字母不是同时而是一先一后地呈现给被试者。首先，一个字

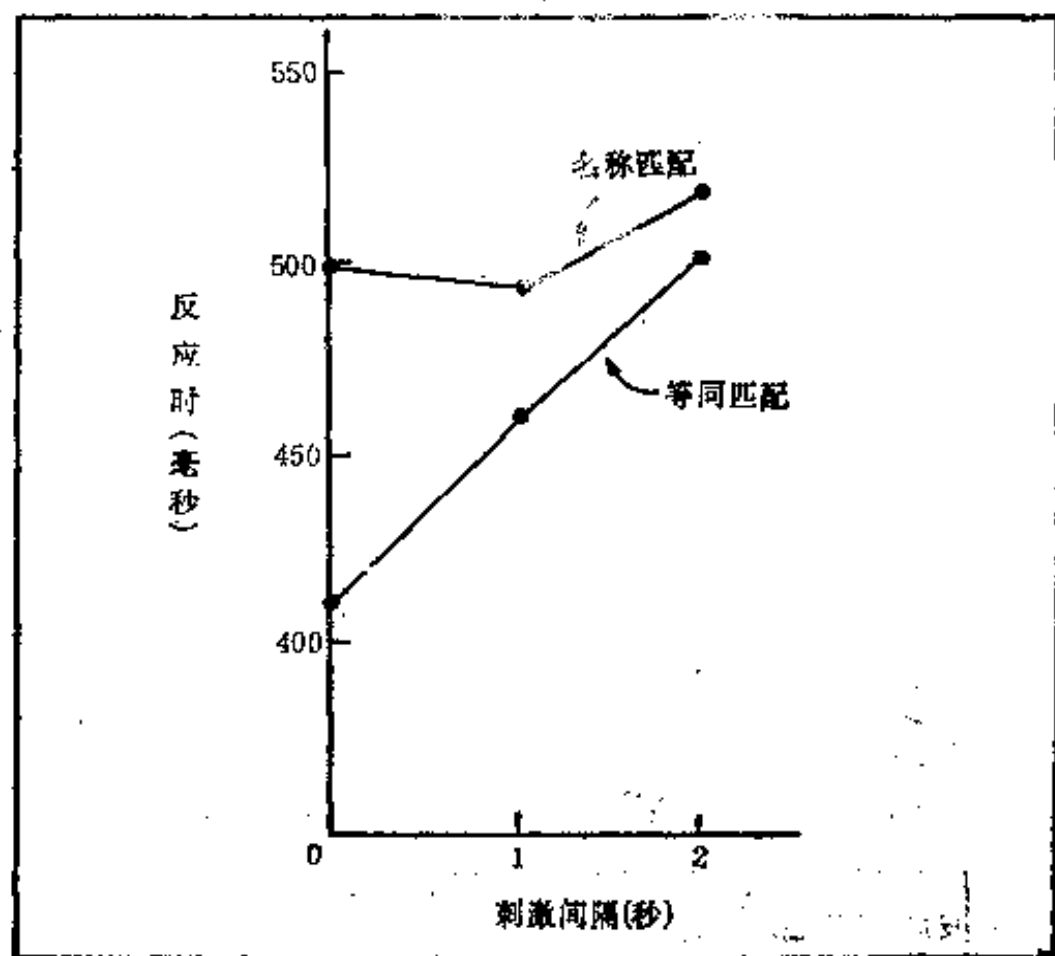


图 5-6 在先后呈现字母的匹配任务中，反应时作为刺激时间间隔的函数

母呈现半秒钟，然后让视野空白一会儿，再呈现第二个字母。第二个字母呈现到被试者作出反应这段时间，乃是被试者的反应时。在这种情况下，要比较两个字母，就必须利用记忆中的信息。那么，在这种情况下，视觉信息是否仍然起作用呢？换句话说，等同匹配所需的时间，是否仍然比名称匹配所需的时间短呢？回答是肯定的。如果两个刺激之间的时间间隔小于1秒的话，等同匹配所需的时间比名称匹配的时间少。但是，如果刺激间的时间间隔达到2秒以后，那么，两者的反应时的差别就不大了。图5-6给出了实验中观察到的这种现象。从这里我们可以看出，第一个字母的视觉信息在刺激消失以后，在短时记忆中可以保持2秒左右的时间。同时，由于视觉信息逐渐减弱，所以，等同匹配的反应时也就逐渐延长。

（三）语义编码

关于语义编码的问题，一部分证据也来自回忆时所发生的混淆错误的模式。也就是说，有许多混淆错误，是由于语义上的原因造成的。例如，给被试者呈现一张词表，紧接着呈现一个测试刺激，要求他回答词表中是否有该测试词。有些词表中有与该测试词“相同”的词，但另一些词表中只有“同样意思”（同义）的词。然而，在后一情况下，被试者也会回答“有”。这就是说，被试者错误地把词表中一个与测试刺激同义的词，看作是与测试刺激相同的词。这就是一种由于语义上的相似性而引起的混淆错误。被试者总是把词表中的一个同义词错误地与测试刺激混淆，不管这个词在词表的什么位置上。但是，他不会把一个无关的词与测试刺激混同起来。所以，在被试者的短时记忆中必定有关于这些词的语义编码。这种混淆正是由于刺激在短时记忆中的语义编

码内容的相同而产生的。

在短时记忆中，信息编码的形式不是单一的。而是多样的。它既有听觉的编码，也有视觉的和语义的编码。信息在短时记忆中究竟采取何种编码形式，或者同时采取几种编码形式，可能取决于刺激的性质和作为信息加工者的个人特点。

五、复述和信息向长时记忆转移

在打电话之前，我们查了电话号码并合上电话簿走到隔壁房间去打电话时，我们常常对自己不出声地默念这个电话号码，以免忘记。这种不出声的对材料的重复默诵过程，就叫做“复述”。为了不使信息从短时记忆中很快消失，这种有意识的反复默诵是必须的。复述实际上是一种隐藏的或不出声的言语过程。复述有两种不同的形式，一种叫做“机械的复述”（或叫“保持性复述”），它的主要作用是使短时记忆中的信息不断更新，强化；另一种叫做“精心的复习”（或叫“整合性复述”），它的作用是把被复述的材料整合到长时记忆的知识结构中去。图5-7表示这两种复述的不同作用。

在保持性复述时，被试者好象是把记忆项目对自己重新默念一遍，又把听到的项目再存储到短时记忆中去，从而使

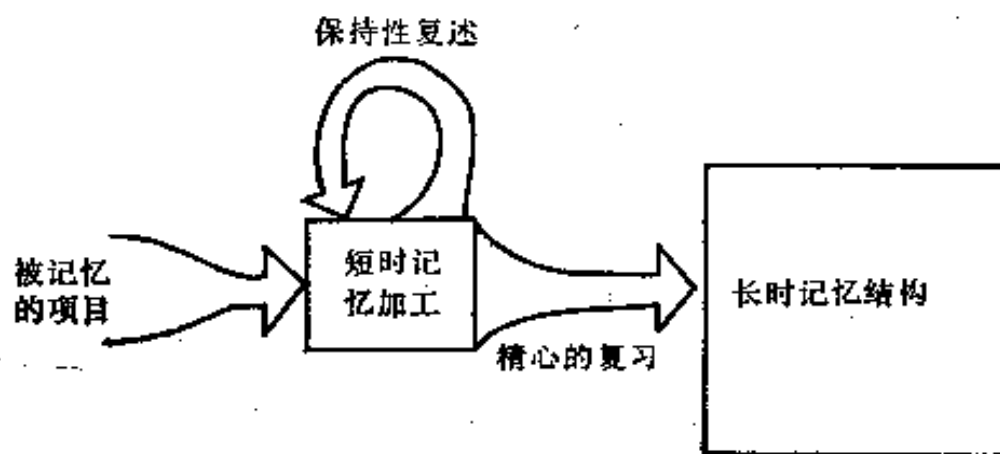


图 5-7 两种复述的不同作用

原来的记忆项目得到更新，保持原来的强度。通过保持性复述，能够使短时记忆中的信息不会变弱、衰退，但是，它并不能增加记忆系统的容量。如果需复述的项目太多，以致不能及时地对所有项目进行复述，那么，最后一些项目，在复述到它们之前，就将发生衰变，从而出现遗忘现象。

学习新材料的主要任务，是要把它们适当地整合到已经存在于长时记忆里的知识结构中去。这种整合过程究竟是如何进行的，现在还不太清楚。不过，保持性复述只是不加思索地、机械地重述某些项目，这样是不能把信息转入到长时记忆中去的。导致信息向长时记忆转移这一类复述，是一种精心的加工过程。在这种过程中，人要主动地作用于信息，比如说，在各记忆项目之间形成有意义的联系，尽力把它与已有的知识联系起来等等。这样才能使信息整合到长时记忆的知识结构中去。

有人曾做过一个实验，说明这两种复述的不同作用。他们通过听觉给被试者呈现七张词表，各个词表包含的单词数分别为：8，10，12，14，16，18，20。每2秒钟呈现一个单词。被试者分为A和B两组。每个被试者都有一个记录本。记录本每页的右侧有一竖列方格。A组被试者所用的记录本每页上的方格数与所呈现的单词数相等，而B组所用的记录本上每页的方格数均为30个。当呈现某一词表时，每呈现一个单词，就让被试者在方格中做一记号。这样，A组被试者在实验过程中可以知道哪些是最后呈现的几个单词，而B组就不可能知道。每个词表呈现完毕后，让被试者对词表上的单词进行立即回忆。回忆的时间因词表的长度不同而从60秒到90秒不等。做完第七个词表的立即回忆后，再让被试做1分钟的心算，然后，出乎意料地要求他们把七个词表上的单

词尽可能地回忆出来，时间为6分钟。实验结果表明，在立即回忆时，A组的被试者对最后呈现的几个单词的保存量高于B组被试者。但是，在延迟回忆时，则A组被试者的保存量低于B组被试者。人们认为，这是由于A组被试者知道哪几个单词是最后呈现的，对这些单词他们可以进行机械的复述。而B组被试者却不可能知道，所以，他们必须做精细的类似寻找意义联系的复习。这说明了机械的复述只对短时记忆的保持起作用，从而有利于立即回忆。但是，为了使信息转入长时记忆，则需要进行精心的复习。

平时，我们在学习新材料的时候，总是在进行着一种精心的、整合性的加工过程。比方说，过去学过网膜上有两种不同类型的细胞，它们的名字叫“锥体”和“棒体”。这名字本身不是一个新词。不过，在这里，它们是以一种新的方式被使用。我们可以利用它们的旧的意义来帮助自己。这些名字是指网膜上的细胞，但细胞的形状象个锥或棒，所以把它们叫做锥体和棒体。除了记名字以外，也应该记住它们位于什么部位，起什么作用。我们知道，锥体位于视网膜的中央，感受颜色刺激。这样，我们就把细胞的位置和功能与其名称联系起来了。那么，关于棒体的一些特性，又怎么记呢？你可以把棒体看作为锥体的“反面”。锥体在网膜中央，棒体则相反，它在网膜的边缘区域；锥体用于感受颜色刺激，棒体则对颜色刺激不能感受；棒体对光很敏感，而锥体却不敏感；锥体在白天强光条件下工作，而棒体却在晚上弱光条件下工作等等。所有这些学习的策略，体现了一种精心的复习过程。它在被学习的各材料项目之间，构成一种有意义的联系，并和已有的知识联结起来。这种精心的复习，将使信息有效地向长时记忆转移，整合到已有的知识结构中去。

第六章 长时记忆和知识表征

长时记忆 (Long-term memory, 简称LTM) 是一个庞大而复杂的信息库, 存储着我们具有的关于世界的一切知识。长时记忆中包含的知识, 使我们能够回忆各种事件, 认知各种模式, 进行各种推理, 解决各种问题。作为人类认知能力之基础的所有知识, 都存储在长时记忆中。

长时记忆中的信息量很大, 但是, 我们却能很方便地利用这些信息。例如, 当你企图回忆谁写了“子夜”这部小说时, 你不会感到麻烦就能从记忆中找到这个答案。看起来, 长时记忆的性质使我们能进行这种快速的检索。信息似乎是以一种非常有秩序的形式安排在长时记忆中的。

一、语义记忆和情节记忆

1972年, 塔尔文 (E. Tulving) 根据记忆的内容, 把长时记忆分成语义和情节记忆。这二种记忆都能长期地存储信息, 但是, 它们存储的信息的类型是不同的。塔尔文认为, 语义记忆是指我们关于词义的永久性知识和关于世界知识的记忆; 而情节记忆则是一种“个人经验及这些经验的时间关系”的记忆。按照这种看法, 关于“猫”这个词的意义的记忆, 以及关于大家都相信的哥伦布发现美洲新大陆这个事实的记忆, 都是语义记忆。而对于在心理学实验中呈现给我们的某个语句的记忆, 以及那些根据时空关系按其知觉属性对

某个事件的记忆，则均属于情节记忆。

如果我们考察一个典型的研究长时记忆的实验，那么，这种区分就是一目了然的。假定我们给被试者一张20个词的表，过一个小时以后，给他一个回忆测验。假定说，词表上有“桌子”这个词。当被试者记忆这张词表时，“桌子”这个词是否存储到长时记忆中去了呢？在词表呈现之前，在被试者的长时记忆中，肯定有关于“桌子”的知识。在实验中被编码和存储的是某种新的知识：“在一个实验里，我的词表上有‘桌子’这个词”。塔尔文认为，先前具有的关于“桌子”的知识，是语义知识；不管词表上是否有“桌子”这个词，关于桌子的概念在长时记中将依然是一样的。语义信息不依赖于这种情境的条件。另一方面，在词表上看到了“桌子”这种知识，则进入情节记忆，因为它是被试者生活中的一个“情节”，情节信息是一种“自传体的”知识。

由此可见，语义记忆和情节记忆的内容是不同的。语义记忆包含着我们使用语言时需要的全部信息。不仅包括词和它们的符号、意义、所指的对象（它们代表什么），而且也包括应用这些词的规则。语义记忆包含这样的内容，如语言的文法规则、化学公式、加减乘除的规则、冬去春来的知识——各种不依赖于特定的时间和地点的事实。与此相反，情节记忆包含的是一些临时编码的信息和事件，一件事怎么出现的和何时发生的知识。它是我们对个人历史的记忆，如“去年冬天我到上海去开了一次会”等等。此外，有人认为，它们还有另一个差别：情节记忆处于一种经常变动的状态；存储在它里面的信息常常被转换，或者变得不容易检索；而语义记忆的变动却很少，不大容易受影响，信息似乎始终是留在那里的。

塔尔文的区分两种长时记忆的见解，是有重要意义的。我们看到，先前有关利用词表来研究记忆的实验，实际上是在研究情节记忆而不是语义记忆。从艾宾浩斯开始到塔尔文为止，这段时间里用词表做了大量的关于记忆的实验，但是，语义记忆的研究却被忽略掉了。从七十年代开始，语义记忆已成为大部分研究的焦点，其中多数集中于长时记忆的结构问题，即信息在长时记忆中的存储方式或格式。人们也研究了长时记忆中信息加工的方式：关于概念意义的知识是怎样检索和使用的。

一个关于语义记忆的理论或模型，大体上说来，是要解决以下四个方面的问题：

第一，规定词义的心理表征方式，即对词义的各个成分及其结构，提供一个一般性的描述，并指明如何确定个别情况的意义细节。

第二，描述意义的组合过程，对如何从个别的语义单位组合成较大语义单位（如语句的意义）提供一般性的说明。

第三，把语义表征与现实世界的情境联系起来。也就是说，要在语义表征方式和知觉过程之间提供一个接口，以便我们能够说明人们如何根据有关的描述来认知客体或回答有关客体的问题。

最后，需说明人们根据句义作出的各种推理，或者根据组成句子的词的意义作出的各种推理。如“老张是单身汉”，可据此推出：“老张是男人”等等。

当然，这是就总的目标而言的。事实上，大多数研究，主要是探讨知识表征的方式问题，以及在理解语句时如何通过这种表征提取和利用知识的有关问题。

心理学家们也利用计算机模拟来发展有关人如何表征和

加工知识的理论。人工智能领域对这方面也表现出很大兴趣。人工智能要发展“智能的”计算机程序，能进行复杂的加工活动，如理解和使用语言。为了使计算机具有这种能力，就需要为它提供广泛的世界知识。存储这种知识的最佳方式以便有效地使用它们的问题，使得人工智能的工作者对知识表征的问题给予极大的关注。利用计算机来发展和检验记忆的模式，已成为认知心理学和人工智能共同的努力方向。

二、层次网络模型和标符搜索模型

（一）层次网络模型

柯林斯 (A. M. Collins) 和奎利恩 (M. R. Quillian) 认为，语义知识可以表征为一种由相互联接的概念组成的网。他们提出了“层次网络模型”，这是关于语义记忆的最早、也是最著名的一个理论。

图 6-1 是层次网络模型的片段。图中的圆圈称为“节

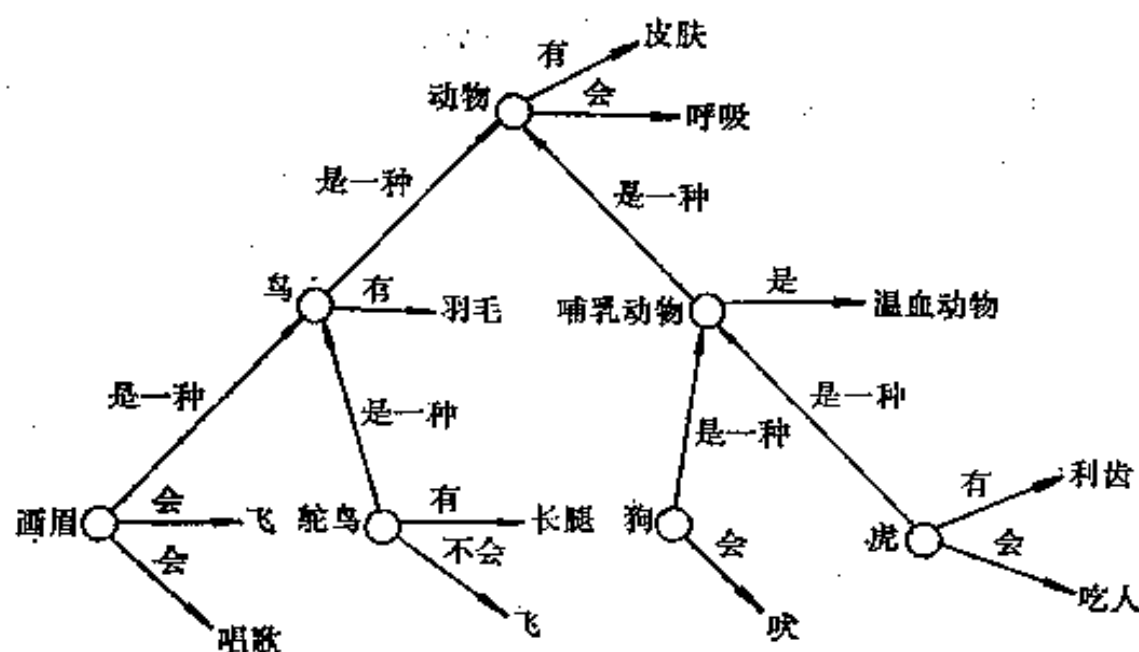


图 6-1 层次网络模型片段

点”，它们代表概念。节点之间带有箭头的连线表示概念之间的关系；连线上注明的标记或符号，说明概念之间关系的性质。连线和节点相互交织，构成一个错综复杂的网络。

这个模型把长时记忆中的语义知识描绘成一个相互联结的概念网。从图 6-1 中可以看出，此模型包含有四个重要的假定：

第一，各概念之间有两类不同的关系：（1）子集关系，用“是一种”来标示。例如，在“鸵鸟”与“鸟”节点之间带有箭头的连线上，标有“是一种”的标记。它表明“鸟”是“鸵鸟”的上属概念；而“鸵鸟”是“鸟”的下属概念，是“鸟”的一个子集；（2）属性关系，用“是”、“有”和“会”来标示。

第二，这是一个预存式的模型，它的某些子集关系及属性关系是直接表示出来的，也就是说预存的。因此，这个模型强调的是搜索过程，而不是比较过程。

第三，一个概念（如“虎”）的意义，是网络中该概念（节点）以及从它引出或通向它的标记关系这两者的复合物。一个概念的意义，也不局限于它的本质特征。

第四，在模型中，一个概念的意义，是用其他一些概念来表示的，没用一集语义基元来描述这些词或概念的意义。

模型对语句的加工，主要是通过搜索。当提出一个语句如“画眉是（一种）鸟”来要求验证时，人或计算机系统就从主项名词和宾项名词，即“画眉”和“鸟”的节点，进入语义记忆网络，并寻找把这两个节点联结起来的通道，核对通道上的标记是否与语句中所断言的关系相一致。应当指出，验证一个句子所需的时间可能是不同的，这主要是由搜索过程的差别所决定的。从上述例子而言，验证该语句所需的时

间，相对来说比较短。这是因为，在“画眉”和“鸟”之间，只有一条连线；这两个概念之间的子集关系，是直接存储在语义记忆中的。但是，如果给出的句子是“画眉是（一种）动物”，那么，验证该句所需的时间就要相对地长一些，因为这两个概念之间有两条连线。在这里，人们必需从已存储的概念关系“画眉是一种鸟”和“鸟是一种动物”中，推断出“画眉是一种动物”。

柯林斯和奎利恩的语义记忆模型，采用了层次结构的形式。名词概念在语义记忆中是一层一层地排列的。从图6-1中可以看到，“画眉”与“鸟”的距离，比“画眉”与“动物”的距离近。人们必须穿过“鸟”这个节点，才能把“画眉”和“动物”联系起来。另外，在这个模型中，把有关的属性存储在最高的或最一般的节点上。例如，在网络中，“会呼吸”这一属性，只附在“动物”这个节点上，而没有在较低的层次如“鸟”这一节点上附加这一属性。但是，“会飞”这个属性则附在“画眉”节点上，而不是附在“鸟”节点上。这是有道理的，因为并不是所有的鸟都是会飞的。例如，鸵鸟就不会飞。这种安排或处理属性的方法，人们称之为“认知经济”假设。

柯林斯和奎利恩提出的这个语义记忆模型，对人们如何表征和提取语义信息的问题，第一次提供了详细的描述。这个理论很快就被其他一些语义记忆的研究者所采纳，甚至被看作是语言理解的一般性理论。

但是，也有人提出了一些异议。首先是康拉德（C. E. H. Conrad）对认知经济提出了不同的看法。她认为，名词概念与其属性之间的关系紧密程度，与层次网络结构中把名词概念与其属性分隔开来的链的数量，是没有内在联系的，

而柯林斯和奎利恩却把这两者混为一谈了。例如，验证语句“鸟有羽毛”可能要比验证“鸟会吃”这个句子快一些。但那是因为在人们的头脑里，“鸟”与“羽毛”之间的联系比“鸟”与“吃”之间的联系更为紧密，而并不是它们在网络中所处位置的远近。名词概念与其属性之间的联想关系愈强，验证语句所需的时间就愈少。但是，验证语句所需的时间，与语义网络中所假设的连线的数量，没有必然的关系。

另外，对网络中名词概念是按层次组织起来的设想，也有不同的看法。在某些情况下，类属关系近的反反而比类属关系远的句子验证得慢些。例如，验证“狗是（一种）哺乳动物”，就比验证“狗是（一种）动物”要慢些。但是，在前一句子中，两个概念之间只有一条连线，而在后一句子中，两个概念之间却有二条连线。这也说明，验证句子所需的时间，取决于两个词概念之间联想的强度，而不是连线的数量。

还有一些问题，层次网络模型也很难解释清楚。比如说，有这样两个句子：“麻雀是（一种）鸟”和“鸡是（一种）鸟”。人验证前一句的速度快，验证后一句的速度则相对地慢一些。这就是说，范畴实例的典型性对验证句子的速度有影响。作为该范畴之实例的典型性高，则验证句子时所需的时间就少，反之，所需的时间就多。

当然，模型本身实质上是一种假设。任何一种心理学模型都不可能完美无缺。层次网络模型虽然存在这些问题，但它作为对人类语义知识表征的一个理论，具有重要的意义和价值。

（二）标符搜索模型

1975年，格拉斯（A. L. Glass）和霍利约克（K. J.

Holyoak) 提出了一个新的理论, 叫做“标符搜索模型”(Marker-search Model)。该模型有以下几个基本假设:

第一, 每一个概念都是由“标符”(即标记或符号)来表征的, 而大多数常用词代表的概念则仅与单个的定义性标符直接相联。例如, 鸟和麻雀是由标符“鸟类”和“麻雀”来表征的。在这里, 标符“麻雀”可以注释为“具有麻雀的本质属性”。

第二, 各标符之间是相互联系的, 因此, 一个标符可以支配或蕴含一集其他的标符。比如说, “麻雀”蕴含着“鸟类”, 后者又蕴含着“动物”。

第三, 这种蕴含的结构, 使标符搜索模型有点象层次网络模型。但是, 在标符搜索模型中, 有时在非邻接的标符之间可以有直接的通道, 这样, 多层次的联结就被缩短了。图 6-2 是该模型的图解。从图中我们可以看到, 在标符“鸡”和“动物”之间, 有一条注有记号d的捷径。这一点就不同于前述的层次网络模型, 或者说是层次网络模型的改进。

第四, 在标符搜索模型的网络中, 可以直接地表示出有关矛盾的信息。具体说来, 就是带有同一记号(图中连线上的字母)的二个通道, 在同一标符处汇合时出现的矛盾冲突。从图 6-2 中可以看到, “麻雀”和“鸡”同在“鸟类”这一点上汇合, 而且又带有相同的记号, 所以, 它们是相互排斥的。也就是说, 某个体不可能既是鸡, 又是麻雀。而“玩赏动物”和“鸟类”虽然也在同一点上汇合, 但它们带有不同的记号, 所以, 它们不是相互排斥的。这就是说, 某个体可以既是玩赏的动物, 又是鸟类。

标符搜索模型也是依靠对网络进行搜索来验证语句的。

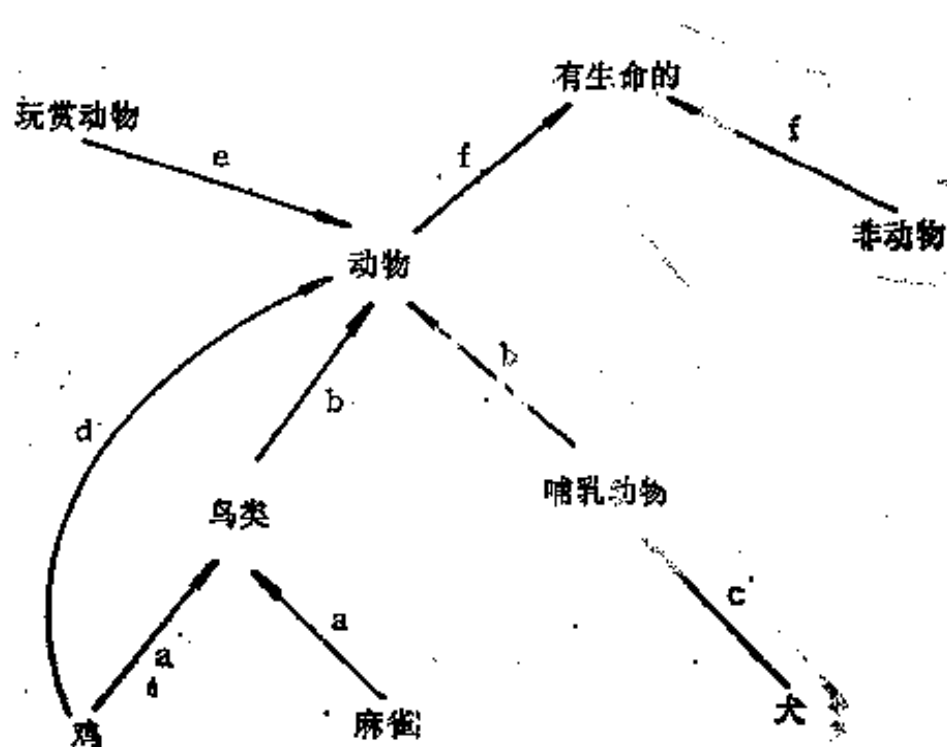


图 6-2 标符搜索模型图解

当呈现“所有鸟都是动物”这样一个语句时，被试者就去访问这两个名词的定义性标符，以及它们蕴含的或蕴含它们的一切标符，如“鸟类”、“动物”、“有生命的”和“麻雀”等，换句话说，就是对语义网络进行搜索。对于这样一个句子来说，当找到一条可接受的、把二个名词的标符连起来的通道时，搜索就终止。也就是说，当认定主项标符蕴含着宾项标符时，就停止搜索。这一点是和层次网络模型相象的。从图 6-2 中可以看到，从“麻雀”伸出来只有一条通道，走向“鸟类”，所以，验证“所有麻雀都是鸟”这个语句，相对来说比较快；而从“麻雀”到“动物”必须穿越二条通道，所以，验证“所有麻雀都是动物”这个语句时，需较长的时间。但是，从“鸡”走向“动物”，则有二条不同的通道。一条通道走向“鸟”，然后再从“鸟”伸出的通道走向“动物”；另一条通道则是直接走向“动物”的捷径。根据标符搜索模型的假定，凡有捷径时，通常是首先被搜索。所以，

验证“所有鸡都是动物”这个语句就比较快，而验证“所有鸡都是鸟”或“所有麻雀都是动物”这些语句则相对地比较慢。这样，“捷径”这个概念提供了解释联想强度和典型性效应的手段。这一点是不同于层次网络模型的。至于哪些标符之间能形成一条捷径，哪些通道首先被搜索，这是由某些标符在日常生活中同时出现的频率所决定的。

对假句的加工，其搜索过程基本上是相同的，但是，需要检查通道上的记号，看看这些记号是否提供了矛盾冲突的信息。例如，验证语句“麻雀是鸡”时，根据由“麻雀”和“鸡”各自伸出的通道上的记号，它们是相互冲突的，任何一个个体，它不可能既是麻雀，又是鸡。所以，可以断定这个语句是假。

本模型也是预存型的，它的许多子集关系由网络直接表示，如麻雀和鸟之间的蕴含关系。当然，另一些关系则是推导出来的，如麻雀和动物之间的关系。对语句的验证过程，就是对这些直接联结的搜索和根据这些联结作出推论。验证时间的长短决定于搜索的顺序和穿越通道的数量。这个模型依靠通道上的记号是否相同来决定二条通道是否相互排斥，但记号是否相同，只靠理论家的直觉判断，并无客观的标准。这是该模型潜在的缺点。

三、扩散激活模型

（一）扩散激活模型的基本思想

柯林斯和洛夫特斯 (E. R. Loftus) 在1975年提出了一个新的语义记忆模型，这可以说是对层次网络模型的修改，叫做“扩散激活模型”。在这个新的模型中，记忆网络不再

具有层次结构的形式，而是根据语义联系或语义距离构造起来的。图 6-3 给出了该记忆网络的一个小片段。

在扩散激活模型中，许多概念节点也是彼此联结在一起的，正如图 6-3 中那些椭圆形之间是用连线联结起来的那样。联结二个概念节点的连线愈短，表示二个概念之间的联系愈紧密。这样，“公共汽车”与“学校”或“学生”之间的联系，比与“救火车”或“事故”之间联系要紧密一些。从图中我们还可以看到，概念节点似乎是分群的。“花”属于一个概念群，而“车辆”则属于另一概念群，一群群相互联系的概念，实际上抓住了我们的直觉。

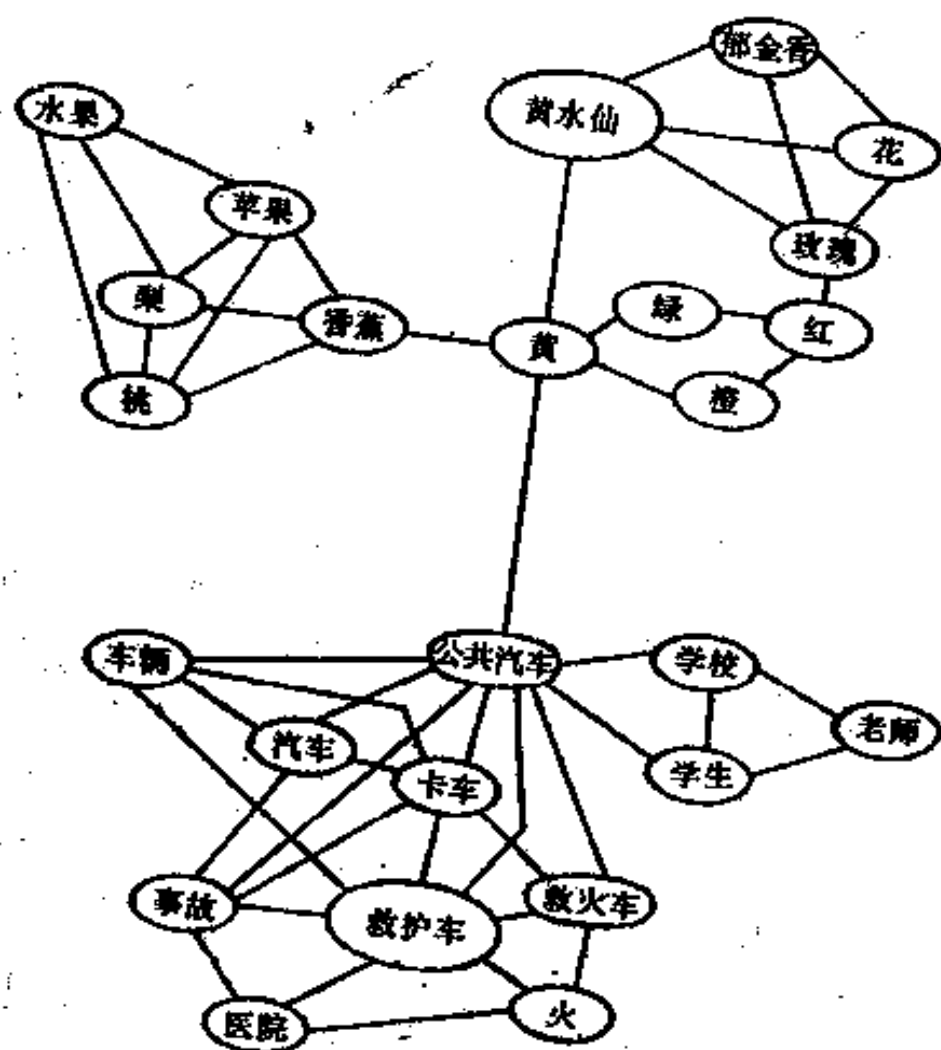


图 6-3 扩散激活模型提出的语义记忆网络片段

在模型中，不同概念之间相互联系的紧密程度，是由经验来确定的。确定概念之间联系程度的一种方法，是直接给被试者一对词，并要求他们根据量表来评定这两个概念之间联系的紧密程度。此外，也可以利用其他的尺度来评估概念之间的语义联系，如某些概念一起出现的频率的常模。当然，关于如何测量语义联系的问题，学者们并无一致的见解。如果能找到一种基本的尺度，可适用于评定联想的强度，那么，就会使语义记忆理论更为有力。这样一种发现，也必然会增强扩散激活模型的基础。

象层次网络模型一样，扩散激活模型中的连线，也标有“是（一种）”这类标记。该标记表示这一范畴隶属于另一范畴。例如，“金丝雀”和“鸟”之间，是由标有“是（一种）”的连线联结起来的。这个模型中，也包括标有“不是”的连线。例如，“蝙蝠”与“鸟”之间，是用标有“不是”的连线联结起来的。下面举个例子，说明模型中为什么要包括这类连线。

如果给被试者呈现下述句子：

“蝙蝠是鸟”。

根据层次网络模型，人们通过对记忆网络的搜索，确定该句子是假。他们循着从“蝙蝠”节点和“鸟”节点伸出的通道进行搜索，寻找“是（一种）”的关系。这种搜索过程是需要花时间的。所以，被试者要指出该句子是假，那应该会比较慢的。但是，事实上，人们确定“蝙蝠是鸟”这个句子为假，那是相当快的。如果在网络中用一条标有“不是”的连线来联结“蝙蝠”和“鸟”，那么，网络就可以适应上述观察到的事实。人们能很快地发现这条连线，从而迅速地作出“假”的反应。这样，关于蝙蝠不是鸟的知识，是直接存储

在网络中的，这片知识也称之为预存的知识。正如我们在前面已经谈到的，一般来说，网络模型假定我们的大多数知识是预存的。

扩散激活模型还包括若干有关加工的假定。第一个假定是说，网络中各连线在“易进入性”或强度上是有差别的：穿越一条强的通道时所花的时间比较短，而穿越一条弱的通道所需的时间则较长。一条连线的易进入性或强度，则依赖于某些变量，如该连线被使用的频率等。在我们的经验中，老师和学校这二个概念是经常在一起出现的，所以，这二个概念之间的连线是比较强的和容易进入的，相反，在我们的经验中，网球和火很少在一起出现过，所以，“网球”和“火”之间的连线是弱的，不容易进入的。

第二个关于加工的假定是说，当一个概念被加工的时候，激活就从该概念扩散到邻近的概念。这种激活的扩散，就象把一块石头投入水面平静的池塘里所发生的情形一样，波动从石头落水的那一点向四周扩散开去。波动的大小依赖于石头的轻重，依赖于离石头落水点的距离，同时，也依赖于石头落水后所经过的时间。在记忆网络中激活的扩散，大致同样地依赖于初始激活的强度、离初始激活点的距离以及激活开始后所经过的时间。此外，当从一个激活节点伸出的连线是强的时候，激活将扩散得比较快而远。但是，当连线是弱的时候，激活的扩散将比较慢而距离短。

这个模型也包括一集复杂的判定过程。例如，当我们读了“汽车是车辆”这个句子后，“汽车”和“车辆”的节点就被激活，而这种激活将从该二个节点扩散开去。激活的通道在网络的许多点上相交。判定过程的工作是评价各种不同相交的根据。某些根据是由联结这些节点的连线上的标记所

直接提供的。“汽车”和“车辆”之间的连线上标有“是(一种)”的标记，它表明，“车辆”是“汽车”的一个上属概念。这一标记为支持确证该句子提供了强的根据。同样，在评价句子“蝙蝠是鸟”时，判定过程会在“蝙蝠”和“鸟”之间发现一条标有“不是”的连线。这一信息将允许判定过程迅速地确定该句子为假。除了评价把两个概念节点联结起来的连线上的标记之外，判定过程也对特定节点上的激活强度进行评价。激活达到一个标准的水平时，才可以作出一个判定。从网络邻近点扩散来的激活累积起来，在某一节点上也可以达到激活的标准水平。

现在我们考察一下迈耶(D. E. Meyer)等人做的一个实验，看看这个模型是怎样操作的。实验中，他们每次给被试者呈现一张项目表，要求被试者尽可能快地回答，表上每一个项目是否是一个合法的词。这种实验叫做“词汇判定任务”。在这类实验中，不管一个项目是词或者是非词，判定它们所需的时间会受到刚才加工过的那个项目的影 响。例如，如果被试者刚才判定过“面包”这个词，那么，他们判定“黄油”这个词就较快，而判定“护士”这个词就较慢。这种影响称之为“语义起动效应”，因为对一个词的加工而使系统为加工下一个语义上有关的词进入准备并起动的状态。

激活扩散模型可以解释这种起动效应。被试者对“面包”是否是一个词的判定，激活了记忆网络中“面包”这个概念的节点。这种激活将扩散到与它联系紧密的有关概念节点，包括“黄油”这个概念的节点。因此，当“黄油”这个词实际上呈现给被试者时，“黄油”这个概念节点在一定程度上已经被激活。结果，在“黄油”这个词呈现后，相应的节点

很快就达到了对该词的判定所要求的激活的阈限水平。这样，被试者的反应时就较短。另一方面，“护士”与“黄油”之间的联系不紧密，它们之间不存在激活扩散的影响。因而，被试者对“护士”这个词的判定，不会减少随后判定“黄油”这个词所需的反应时。这个例子说明，在该模型中，激活扩散过程是一个重要的加工成分，它为语言上下文的效应提供了一种机制。

（二）表征和加工的关系

上面，我们从模型加工的角度讨论语义知识的表征。为什么在讨论知识表征的时候，总是要考察模型加工的问题呢？道理很简单，是因为我们不能直接观察知识的表征，只能从作业中推论出知识表征的特点。事实上，人的作业既依赖于知识的表征，也依赖于加工过程，例如，依赖于检索的操作。对知识表征的研究，实际上就是对知识的表征和加工这两者的研究。从模型加工这个角度去讨论知识表征，可以对各种不同的关于语义记忆的理论进行检验，同时可以对作业做出特定的预测。

对一个特定的模型来说，在关于表征和加工的假设之间，存在着内在的联系。我们知道，被试者验证句子“麻雀是鸟”所需的时间较少，而验证句子“鸡是鸟”所需的时间较多。根据扩散激活模型的假设，“麻雀”和“鸟”之间的通道，比之“鸡”和“鸟”之间的通道来说，比较容易进入和比较强。其可能的原因是人们经常使用“麻雀——鸟”这条连线，而人们在日常生活中却很少使用“鸡——鸟”这条连线。结果，在“麻雀”和“鸟”两个节点之间的激活扩散很快，人们能快速地验证“麻雀是鸟”这个句子。当然，也可

以从另一种机制来解释上述现象。例如，如果说所有的连线都是同等可进入的，但联系紧密的概念之间是由一条较短的连线联结起来的，而彼此联系不紧密的概念之间则是由较长的连线联结起来的。由于这个缘故，激活以同一速率沿着从一个节点伸出来的所有通道发散时，激活就将首先到达联系紧密的概念节点。从这里可以看到，一个特定的模型，或者借助于那些属于表征的概念，或者借助于那些属于加工的概念，能适应于特定的观察到的事实。

更进一步说，如果两个概念在网络中彼此之间的离距反映了它们之间联系的紧密程度，那么，是否就不必要再假定联结概念的连线在强度上存在的差别，也不必要再假定激活沿着某些通道扩散比较快，而沿着另一些通道扩散比较慢了呢？其实不然，关于表征和加工的观念，在模型中常常是交织在一起的，它们彼此是相互补充的，从而使模型能更好地适应观察到的各种事实材料。

（三）关于网络模型的评价

总的来说，扩散激活模型比层次网络模型提高了一步。扩散激活模型避免了记忆中的知识是有层次地并以最经济的方式组织起来的假定。此外，扩散激活的概念在说明上下文效应时，已证明是非常有用的。

扩散激活模型也是一个网络模型。一般来说，网络模型为人类知识的表征提供了有用的框架。人们往往把词看作孤立的心理学实体，认为它们的意义与其他的词是无关的。但是，网络模型认为，一个词的意义依赖于该词和记忆中它与之有关的大量概念之间的联系。在网络模型中，一个词的意义是由许多与它相联系的概念给出的，甚至可以包括那些构

成我们对世界的一般知识的概念。因此，网络模型具有较大的适应性。

网络模型具有潜在的发展余地和力量。有些网络模型能解释某些较为复杂的加工活动。象前面谈到的标符搜索模型，可以进行语义矛盾的加工（例如，处理“所有的麻雀都是鸡”这类句子）和语义异常的加工（例如，处理“所有的拖拉机都是海员”这种句子）。此外，网络模型也可以扩展，使其有可能适应对复杂句子和散文段落的理解和记忆。

四、特征比较模型

“特征比较模型”是由史密斯（E. E. Smith）等人提出来的。网络模型假定，我们的大多数知识是预存在记忆中的。而特征比较模型则假定，我们的大多数知识是从存储在记忆里的信息中计算出来的。

（一）特征比较模型的基本内容

与网络模型不同，特征比较模型假定，一个概念是由一集语义特征来规定的。这些特征描述了一个客体的属性。例如，“鸟”这个概念，具有“有翅膀”、“会飞”、“下蛋”等这样一些特征。一个概念的各种语义特征的结合，提供了这个概念的意义。

在特征比较模型中，用来定义一个概念的语义特征是预存在长时记忆中的（见图6-4）。在验证句子（如“麻雀是鸟”）时，是把主项（麻雀）的特征与宾项（鸟）的特征进行比较，如果宾项的特征是主项特征的一个子集，那么，该句子为“真”。所以说，在特征比较模型中，两个概念之间

的关系，是通过预存的特征计算或推导出来的。但是，这些语义特征所起的作用在程度上是不同的。有些语义特征对定义一个词概念来说是本质的，称之为“定义特征”；而另一些则只是该概念的一些特性。例如，“鸟”这个概念的定义特征是：有生命、有羽毛；而它的特性特征是：有一定的大小、与其他动物有捕食关系等等。特性特征对该概念来说，不是本质的，它只是使该概念更加具体化。

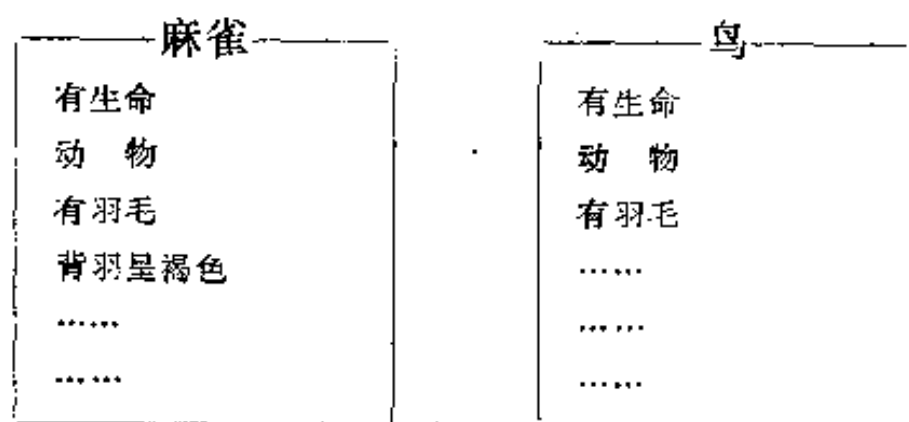


图 6-4 特征比较模型

(二) 模型的加工过程

特征比较模型对语句的加工是分二级进行的(见图6-5)。在第一级加工中,对主项和宾项的所有特征进行比较,看看这两项之间的特征的类似性程度。至于这些类似的特征是定义性特征还是特性特征,一概不去考虑它。如果两个概念的特征类似性程度非常高或非常低,人们就据此可以作出“真”或“假”的反应。更确切地说,如果两个概念之间的特征类似性大于一个高标准(如“麻雀”与“鸟”),人们就确定该语句为真;如果小于一个低标准(如“麻雀”与“铅笔”),就确定该语句为假。如果两个概念的特征类似性程度落在这两个标准之间(如“鸡”和“鸟”或“蝙蝠”和“鸟”),那么,

就要进入第二级加工。这时，就要把主项和宾项所包含的定义性特征分离出来进行比较。如果所有的定义性特征都匹配，那么，就确定该语句为“真”，否则，确定该语句为“假”。

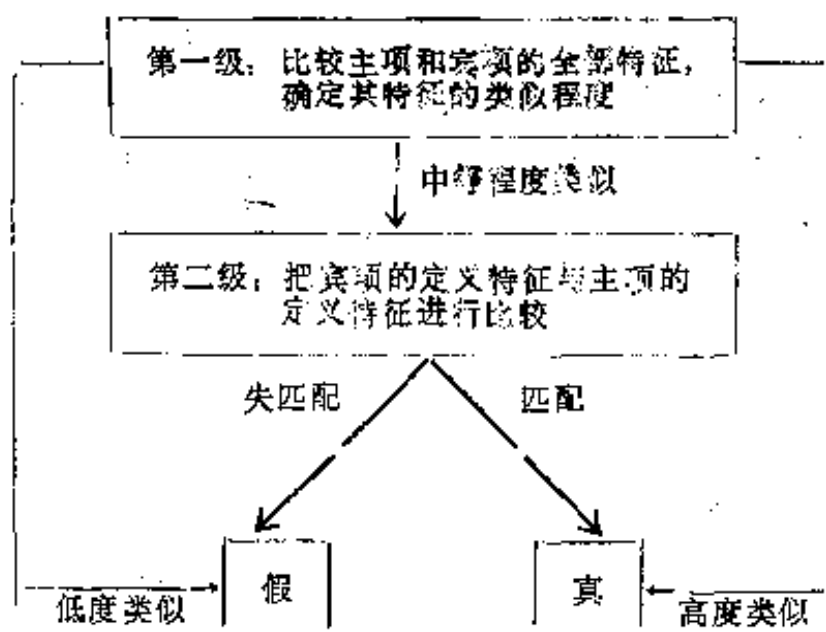


图 6-5 特征比较模型的二级加工

这个模型的特点，是在它的第一级加工中包含有一种试探式的推算程序，使它能解释典型性的联想效应。典型性反映着实例和范畴之间的特征的类似性程度。就真句而言，实例愈典型，它的特征与范畴特征之间的类似性程度愈高，从而超过类似性的高标准，在第一级加工时即可得到确证，因而加工的速度就快。反之，就假句而言，主项与宾项的特征愈类似（如“蝙蝠是鸟”）就不易低于类似性的低标准，这就需要第二级加工，因而要否定这种句子，就需较长的时间。当然，这个模型在加工过程中，有时也会发生错误，尤其是当许多类似性特征不是定义性特征，而是特性特征时（如蝙蝠和鸟），就容易产生这种现象。

特征比较模型和网络模型之间的差别，在于它们对范畴的描述方式不同。在网络模型中，所有的实例都以同样数量

的连线与它们的上属概念联结起来。而且，一个实例，或者属于某个范畴，或者不属于某个范畴，是很分明的。但是，事实上，并不是所有的实例都是同样明显地属于某个范畴的。有人认为，许多语义范畴是个“模糊”集，并没有清楚地规定了的成员。特征比较模型能处理概念的模糊性问题，因为它允许实例在一定范围里可以是变化的。鸡不是一种典型的鸟，因为它不会飞，而且体积大，但根据它具有的鸟的特征，它仍然是鸟。麻雀是一种非常典型的鸟，因为它具有所愿望的一切特征。

当然，特征比较模型也并不是处理范畴内各成员之间差异的唯一方法。柯林斯等人提出的扩散激活模型，也能处理这个问题。该模型的网络内，各节点之间的连线在强度上是不等的。每一条连线有一个强度值，它表示搜索通过此连线时有多快。强度值反映了两个概念之间的联系，一个范畴的较典型的实例，与该范畴之间的联结的强度就较大，经过此通道的搜索就较快，从而导致句子得到较快的验证。

五、命题表征

有人指出，知识的“基本单位”乃是一个事实，而不是一个概念，例如，“华盛顿是美国第一任总统”，它不单是“华盛顿”三个字的意思。因此，有人提出，人类知识（包括语义知识和情节知识）的基本单位是一些概念的结合，叫做“命题”。

（一）命题的含义

命题可能是知识的基本单位，但是，它的确切性质是很

难抓住的。安德森(J. R. Anderson)曾对命题的特性作过一些说明。(1) 一个命题是某种象句子一样的东西，但是它比较抽象；也就是说，它象句子的意义，而不是句子本身。两个非常不同的句子，可以由同一个命题来表征，如“张三看见了李四”和“李四被张三看见了”。因此，尽管我们可以用词来表征命题，但我们应该记住，它们实际上并不是词。

(2) 命题是根据一定的规则构成的。当然，各理论家所依据的规则有可能是不同的，但是，重要的是每个命题都是一种形式化的结构。

命题是一种符号表征，它表示两个或多个概念之间的联系。这种联系可以是真，也可以是假。例如，“碰撞(轿车，卡车)”是一命题，它可以用来表示句子“轿车碰撞卡车”的意思。这个命题中包括三个概念，即“碰撞”、“轿车”和“卡车”。在这里，“碰撞”称之为命题的关系项，因为它指明了“轿车”和“卡车”之间的关系。关系项写在命题的第一项，处于括号的外边，括号把二个由关系项联起来的概念圈起来，这两个概念称之为变元。这种类型的安排称之为命题格式。句子“女孩吃色拉”的意思，用命题格式可以表示为：“吃(女孩，色拉)”。

人们认为，理论家们之所以对命题表征很感兴趣，是因为命题能够包含那种复杂的知识，如理解散文、进行会话、形成映象等的知识。前面讨论的词概念的知识表征比较简单，往往不能用来描述我们日常生活中的许多任务。例如，人们阅读的课文段落的意义，就很难分解成语义网络，也不能用只包含简单的语义联系(例如“是”)的网络来加以表征。但是，许多研究者相信，命题表征能够包含人类具有的大多数知识。

(二) 命题和段落加工

许多事实表明，在加工句子和故事的时候，人们使用了命题表征。下面介绍麦科恩 (G. Mckoon) 和拉特克利夫 (R. Ratcliff) 的一个关于简单段落加工的实验。

一个段落的意义可以用一系列命题来表示，如表 6-1 所示。该段落包括 6 个句子，每个句子的意义都包含在一个命题中。p 1 (即命题 1) 表示第一个句子的意义；p 2 (即命题 2) 表示第二个句子的意义；如此等等。一个概念可以出现在两个命题中，这样，各命题也就彼此联系起来了。例如，p 1 和 p 2 二个命题都包含概念“昆虫”，而 p 5 和 p 6 都包含概念“乌鸦”。因为各命题联结成一直线或线性序列，

表 6-1 一个简单段落及其命题和线性意义结构

句 子	命 题
庄稼引来了昆虫。	P 1: 引来 (庄稼, 昆虫)
昆虫困扰农夫。	P 2: 困扰 (昆虫, 农夫)
农夫环视田野。	P 3: 环视 (农夫, 田野)
田野需要农药。	P 4: 需要 (田野, 农药)
农药毒死了乌鸦。	P 5: 毒死 (农药, 乌鸦)
乌鸦污染了农村。	P 6: 污染 (乌鸦, 农村)
命题联系:	
	P 1—P 2—P 3—P 4—P 5—P 6
名词之间的联系:	
	N 1—N 2—N 3—N 4—N 5—N 6—N 7
	庄稼 昆虫 农夫 田野 农药 乌鸦 农村

所以，我们说该段落具有线性的意义结构。邻接的命题在意义结构上较为紧密，而非邻接的命题在意义结构上相对来说就松一些。例如，p 1 和 p 2 的联系，较之 p 1 与 p 3 或 p 2 与

p 6 (之间的联系), 就较为紧密。

麦科恩和拉特克利夫假定, 阅读该段落的人觉察了表 6-1 中的命题和意义结构。为了检验这个假设, 他们利用了**在语义上有联系的词可以彼此起动(语义起动效应)**这一事实。根据语义起动效应, 他们预测, 在段落意义结构中邻近命题的词, 彼此的起动效应会较大。而在意义结构中彼此远离的命题中的词, 其起动效应将较小。这种预测与下述思想是一致的: 二个命题在意义结构上愈接近, 那么, 这些命题中的名词在心理学上的联系也就愈紧密。

在整个实验的每一次试验中, 被试者阅读两段彼此无关的段落, 然后作一次词识别的测验。在测验中, 呈现一个词, 被试者的任务是尽快地按两个键中的一个, 以指明该词是否曾在两个段落的任一段落中出现过。总共有 72 个试验, 所以被试者总共要阅读 144 个不同的段落。某些测试词跟着那些在段落意义结构中联系紧密的词之后呈现。例如, “农村”跟在“乌鸦”之后呈现。另一些测试词则跟在无关词之后呈现, 例如, 在呈现“汽车”之后立即呈现“庄稼”, 或者, 跟在那些意义结构上距离较远的词之后呈现, 例如, 在呈现“昆虫”之后跟着呈现“田野”。他们的预测是, 当测试词跟在段落意义结构中紧密地联在一起的词之后呈现时, 识别的速度将会最快。

结果, 这种预测得到了证实。二个名词在段落的意义结构中愈接近, 其起动效应也就愈大。例如, 当“农村”跟在“乌鸦”之后呈现时, 被试者的识别比较快, 而跟在“农药”之后呈现时, 被试者的识别就比较慢。这些结果表明, 在意义结构中词概念的接近程度决定了起动效应的大小。

但是, 从表 6-1 中我们看到, “乌鸦”和“农村”在段

落的物理空间位置上也是接近的。那么，前者对后者的起动效应会不会是由于这种物理位置上的接近所造成的呢？为了弄清楚这一点，他们又使用另外一些段落（见表6-2）进行了实验。在这类段落中，两个词或概念在意义结构上是紧密地联在一起的，但它们在段落中的物理位置上是彼此分开的。例如，“商人向侍者打手势”和“商人夸耀文件”在段落中是被三个句子隔开的。“侍者”和“文件”在意义结构上较接近，但在物理位置上则相隔较远。而“餐巾”和“顾客”在物理位置上的距离与前一对词概念是同样的，但在意义结构上则较远。结果表明，“侍者”对识别“文件”的起动效应大于“餐巾”对识别“顾客”的起动效应。这说明，意义结构上的联系决定了起动效应的大小。

表6-1 一个简单的段落和它的命题结构

句 子	命 题
商人向侍者打手势。	P1：向…打手势（商人，侍者）
侍者拿来咖啡。	P2：拿来（侍者，咖啡）
咖啡沾污了餐巾。	P3：沾污（咖啡，餐巾）
餐巾保护了桌布。	P4：保护（餐巾，桌布）
商人夸耀文件。	P5：夸耀（商人，文件）
文件解释了一个合同。	P6：解释（文件，合同）
合同满足了顾客。	P7：满足（合同，顾客）
命题联系：	
$P_1 \begin{cases} P_2 \rightarrow P_3 \rightarrow P_4 \\ P_5 \rightarrow P_6 \rightarrow P_7 \end{cases}$	
名词之间的联系：	
$N_2 - N_3 - N_4 - N_5$	
$N_1 \quad \text{侍者} \quad \text{咖啡} \quad \text{餐巾} \quad \text{桌布}$	
$\text{商人} \quad N_6 - N_7 - N_8$	
$\text{文件} \quad \text{合同} \quad \text{顾客}$	

总之，这些事实表明，在阅读一段故事时，被试者把它们的意义表征为命题的形式。在确定一个概念是否已在段落中出现过时，他们就访问命题表征。此外，也有人做过其他的实验。例如，让被试者阅读长度一样，但包含的命题数量不同的复杂句子。结果发现，句子包含的命题越多，被试者阅读这个句子的时间就越长。这一事实表明，理解的基本单位是命题，而不是单个的词，每增加一个额外的命题，就会增加理解的时间。因此，某些理论家利用命题作为知识的基本单位，构造较大的或综合性的关于理解和记忆的模式。HAM (Human Associative Memory 的缩写，即“人的联想记忆”) 是其中最具有影响的综合性模型之一。

(三) HAM模 型

HAM是由安德森和鲍尔 (G. H. Bower) 构造的一个模型。他们研制这个模型出自两个目的。第一，意图说明人们在执行阅读和记忆这类任务时所使用的知识。综合模型与那种简单的语义记忆模型不同，它涉及到的知识是人的复杂的认知活动的基础。第二，意图阐明人利用其知识的机制。

1. HAM中的特征 在HAM中，命题是表征知识的基本单位。命题是由五类联想组成的，如图6-6所示。上下文和事实构成第一类联想。“事实”指出发生了什么事件，而“上下文”则指出该事件的时间和地点。时间和地点之间是属于第二类联想，这是由第一类联想中的“上下文”分解而成的。主项和谓项之间构成第三类联想。主项指出事实的主题，而谓项则指出主项的特性，或者指出主项发生了什么。第四种联想是由关系和宾项构成的。这是由谓项分解而成的。概念和实例之间则属于第五种联想，例如，概念

“房子”和个别具体的房子或指该房子的词之间的联想。

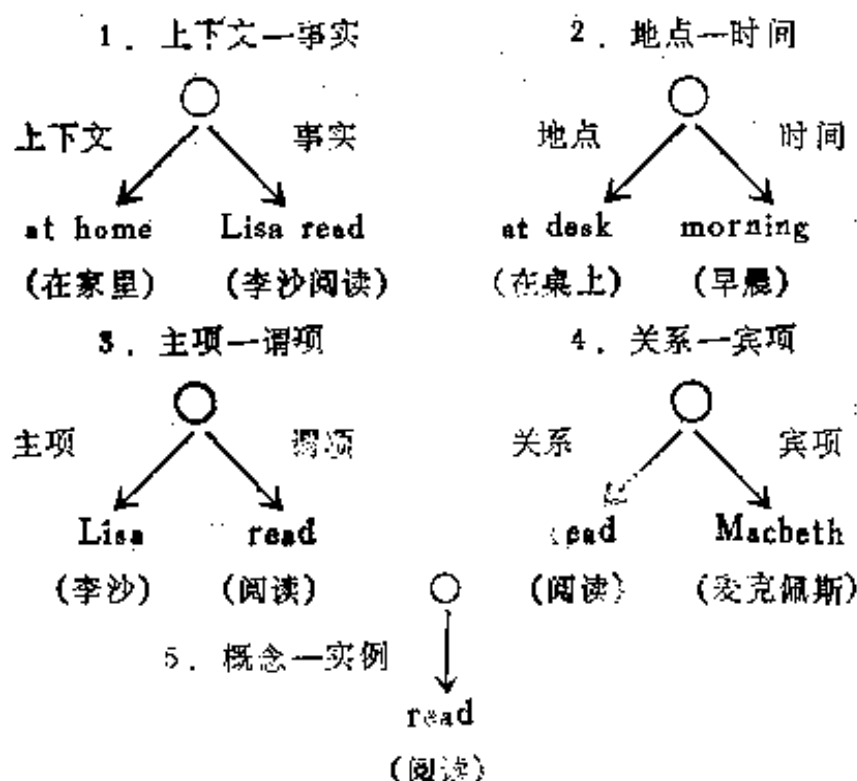


图 6-6 在HAM命题表征中出现的五类联想

注：“麦克佩斯”系莎士比亚的一个悲剧的剧名。

这五种联想结合起来，形成一种树状结构，如图 6-7 所示。这种树状结构是由节点和指针（带有箭头的连线）构成的，节点代表概念，指针代表联想。顶点A是该树的命题节点，它分成上下文节点（B）和事实节点（C）。上下文节点又依次分为地点节点（D）和时间节点（E）等等。整个树状图代表句子 At home Lisa read MACBETH（在家里，李沙阅读“麦克佩斯”）的命题结构。我们可以看到，HAM中使用的命题结构既可以表征语义记忆，也可以表征情节记忆。

HAM中的命题把句子表征为一种有层次的结构，从而说明了句子本身的内部结构，避免了那种传统的观点，即认

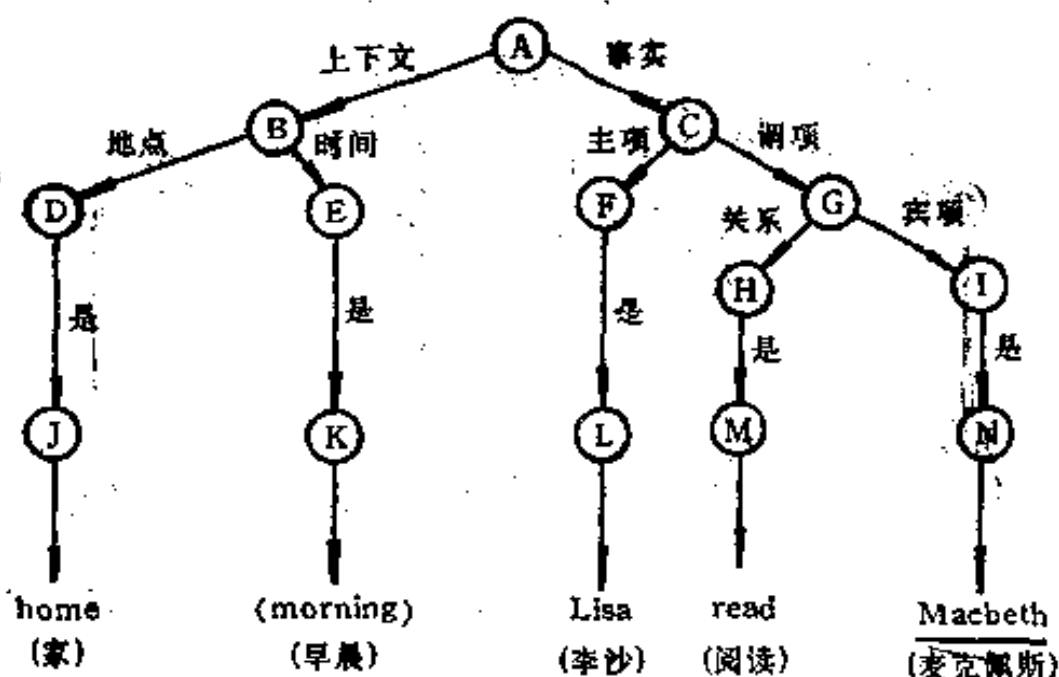


图 6-7 HAM命题表征中，句子“*At home Lisa read Macbeth*”
(在家里，李沙阅读“麦克佩斯”，)的命题结构

为句子是一条联想观念的链。这种有层次的表征，使得在一个命题中有可能嵌入另一个命题。例如，“*That Lisa had read MACBETH pleased the teacher*”（李沙阅读“麦克佩斯”，使老师感到高兴。）这个句子包含两个命题。这两个命题分别是“*Lisa read MACBETH*”（李沙阅读“麦克佩斯”和“*It pleased the teacher*”（这使老师感到高兴）。它们可以结合成一单个命题树，如图 6-8 所示。在该树状图中，主项是图 6-7 表示的命题。这种结合命题的方式，使我们有可能把许多命题有机地组合起来，从而构造一种复杂的表征。这一特点使该模型能够表征非常复杂的思想，显示出该模型的力量。

这种命题的表征形式对人们具有相当的吸引力，因为它把表征的力量和效率统一起来。这个模型作出了一个总括的、彻底的假设：人们具有的一切知识，都被表征为上述那

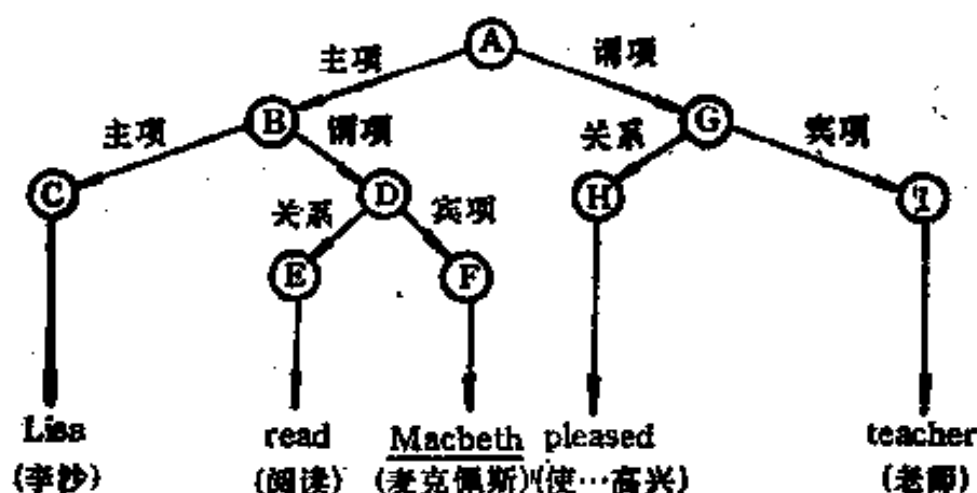


图 6-8 HAM 的命题嵌套,在此图中,“Lisa read MACBETH”是另一命题的主项(图略简化)。

样的命题树。也就是说,根据这个模型,我们的长时记忆是一个庞大的命题树网络。一个模型可以表征多样的信息,这是模型的一个重要性质,是它的力量的体现。由于命题可以表征我们的大多数知识,所以,人们认为它是一种非常有力的表征形式。

命题表征的效率也使它成为一种有用的知识表征方法。一个命题可以表征几种不同句子的意义。例如,图 6-7 中的命题树,既可以表征句子“Lisa read MACBETH at home”(李沙在家里阅读“麦克佩斯”),也可以表征句子“MACBETH was read by Lisa at home”(“麦克佩斯”被李沙在家里阅读)。因此,没有必要为加工过的每个句子都存储一个命题,同时也节省了存储空间。

HAM 中使用的命题表征,能有效地进行信息提取。命题树中标有标记的指针,把不同的节点联结起来,并且提供了对记忆进行直接搜索所需要的信息。例如,该系统对“What did Lisa read?”(李沙阅读什么?)这个句子的回答。HAM 首先针对该句子构成一命题树,如图 6-9

左边所示。然后，通过匹配或比较的过程，在记忆中寻找一棵相匹配的命题树。匹配过程中进行的搜索，不是漫无边际的。匹配过程利用标有标记的联想去指导搜索过程，限制待搜索的节点数量，从而提高了提取过程的效率。

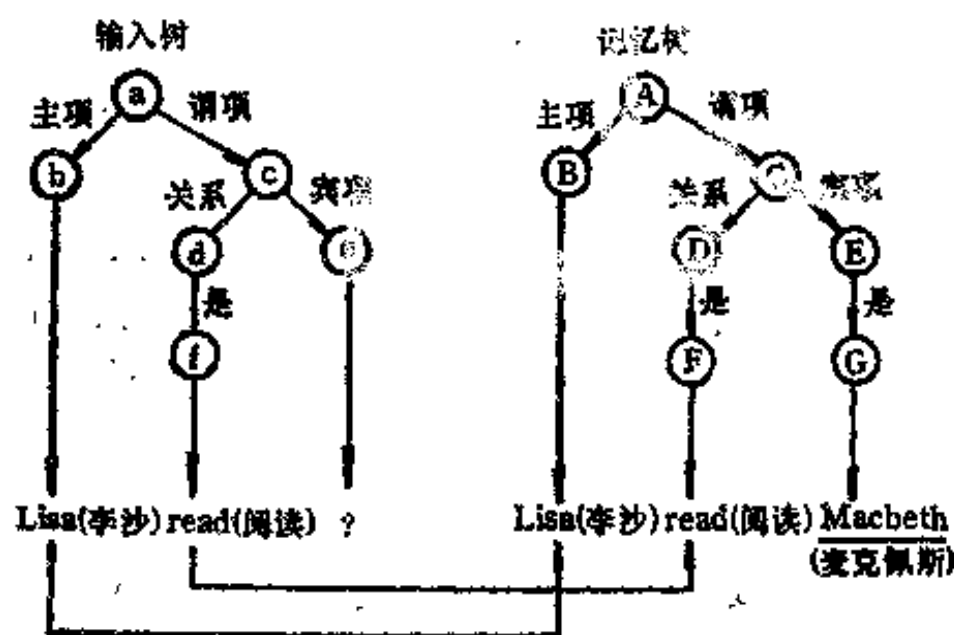


图 6-9 用于回答 “What did Lisa read?” 的匹配过程

在记忆中找到与输入树相匹配的命题树后，搜索与匹配过程就停止(如图 6-9)。由于找到了匹配的命题树，HAM就通过对谓项连线和宾项连线的搜索，并找到MACBETH的节点(E)。当然，如果它不能发现 MACEBTH 的节点，或者它不能从记忆提取相应的命题树，那么，它就不能回答这个问题，而只能答复说：“不知道”。

2. 对该模型的检验 一个模型的检验可以通过不同的方式。一种方式是吧模型嵌入一个计算机程序，然后考察这个模型是否能做一切期望它做的事情。当然，HAM 由安德森和鲍尔编成了计算机程序，能成功地用命题格式来表

征输入的新句子，并能正确地回答向它提出的许多问题。

但是，一个表征句子和回答问题的计算机模型，不一定是一个心理学模型。只有当一个模型能象人那样来加工信息，能预测一个特定任务中人将怎样作业时，它才是一个心理学模型。事实上，安德森和鲍尔在建造和评定HAM时，所依靠的主要是实验数据。

在HAM中，命题是加工的基本单位。从这一思想出发，可以预测，一个命题的重复会提高该命题内各成分的保持。安德森和鲍尔让被试者分别学一些句子。这些句子中有些词是重复的。这些重复的词又分二种情况，一是在同一命题中重复，另一是在不同命题中重复。接着，对被试者进行一次“提示回忆”测验（相当于填空的方式）。结果表明，一个词在同一命题内重复，会提高对这些词的回忆成绩，并且，重复的次数越多，其回忆的成绩就越高。相反，一个词在不同的命题中重复，不会影响被试者的回忆水平。这一实验以及其它许多实验观察，支持了HAM模型。

六、脚本和故事的结构

（一）脚本是知识表征的较高级的单位

“脚本”（Script）作为人类知识表征的一种方式，是由香克(R. Schank)和艾贝尔森(R. Abelson)于1977年提出来的。脚本描述的是事件，事件进行时发生的一系列动作，以及进入该事件的原因和有关的主要概念。因此，可以说，脚本是标准化了的事件序列，用来描述人类行为的某些模式，例如到饭馆吃饭或去医院看病等。香克和艾贝尔森认为，在人的长时记忆中，存储着许多各种各样的脚本。人们依靠这

些脚本，理解事件的各个方面，预测事件的进程，并据此作出推理等。下面是一个“餐馆脚本”所包含的内容：

表6-3餐馆脚本所包含的知识

脚本：餐馆	
人物：顾客、服务员、出纳员……	
道具：餐馆、餐桌、菜单、食物、饭钱……	
事件：	
1. 顾客走进餐馆；	2. 顾客在餐桌前坐下；
3. 服务员拿来菜单；	4. 顾客点食物；
5. 服务员端来食物；	6. 顾客吃食物；
7. 顾客付饭钱；	8. 顾客离开餐馆；

脚本的知识表征方式是对事件进行动态描述，用电影脚本的形式把一系列动作和有关内容联系起来。香克等人用这种脚本的知识表征模型，编制了一个理解人类自然语言的计算机程序。这个程序叫SAM (Script Applier Mechanism 的缩写，即“脚本应用机制”)，主要用来理解故事，进行推理和回答问题。比如说，告诉SAM系统下面一段故事：

“李利走进一个餐馆。他点了一份牛排。

他付了钱，并离开了餐馆。”

然后，就问SAM系统：

“李利吃了什么？”

在上述一小段故事中，实际上并没有清楚地说明李利吃了什么。要回答这个问题，不仅需要具有在餐馆里通常进行什么活动的知识，而且还要具备把这些知识应用于特定情境的能力。脚本提供了这种知识，可以用来进行常识推理。其推理机制是这样的：根据脚本中列出的事件1，系统可以认出“李利走进一个餐馆”这句话。同时，李利被指定为“顾客”的角色；根据事件4，可以认出“他点了一份牛排”这句话，

同时，把牛排指定为食物的角色；然后，又根据脚本中的事件 7 和事件 8，认出了故事中“他付了钱，并离开了餐馆”的话。这样，事件中的顾客就是李利，而食物就是牛排。所以，系统就可以把脚本中的事件 6 具体化为“李利吃了牛排”，并以此回答了上面提出的那个问题。

香克和艾贝尔森认为，人类所以能理解和进行各种各样的日常认知活动，这主要是因为他们利用了记忆中存储的各种脚本。

(二) 故 事 结 构

脚本是知识表征的较大的单位，它用来表示人们对日常生活事件记忆的高级知识结构。此外，罗迈尔哈 (D. E. Rumelhart) 提出，人们对故事也有一种知识结构。他设计了一套“故事文法”，这是一套规则，用来说明简单故事的一般框架。图 6-10A 给出了他的文法的最简单的形式。在他的文法中，利用了“重写规则”（在读这些规则时，把箭头解释为“可以重写为……”，即把左边的东西转换为右边的成分）。这些规则表明，象餐馆脚本一样，一个故事也可以分为一系列成分。故事的主要成分，象第一条规则所指出的那样，是一个“目标”（也就是一种所希望的状态），后面跟着一个或一系列情节，再跟着“解决”（相当于故事的最终结局的一种状态）。根据其他规则，其中每一个成分又可以再分解。例如，第二条规则指出，一个情节可以分成以下几个成分：一个子目标（在达到主要目标过程中完成的一种目标），该子目标上的一个或一系列企图，以及这些企图的结局。

图 6-10A 给出的规则，可以应用于下述简单的故事：

一个农夫有一头母牛，他希望母牛到牲口棚里去。他使劲拉这头母牛，但是母牛不动。所以，这个农夫要他的狗狂吠，把母牛吓到牲口棚里去。但是，这只狗拒绝狂吠，因为没有给它吃食物。所以，农夫走到房子里去，拿了些食物。他把食物给了狗。狗就狂吠，吓唬母牛，母牛就跑到牲口棚里去了。

图中的文法规则应用于这个故事的結果，是产生图 6-10B 的树状图。这个图表示重写规则成功地往下应用，使故事连续分解成各个成分。例如，第一条应用的规则是分解一个故事的规则，它所产生的各成分（目标、情节和解决），是从故事这个单位下面伸出来的枝。然后，又应用分解这些成分的规则，所产生的子成分是从它们底下伸出来的枝，如此等等，直到不能再进一步分解为止。其结果是，它把一般概念放在图的顶部，而把较特殊的概念放在一般概念的下面。最后，以农夫故事的特殊事件的分枝作为结束。当然，这个图解主要是用来说明由文法规则所描述的故事的结构，如何适用于农夫这个特定的故事。

有人根据罗迈尔哈的故事文法，进行了一系列的研究。其结果表明，当人们理解故事的时候，他们实际上的确是利用这样一种知识结构的。比如说，在一个实验中，给被试者呈现一个故事，然后，要求被试者对这个故事进行回忆。实验者把原初的故事写成各种不同的版本，每个版本都有不同的结构。目的是为了控制故事符合于结构规则的程度。一种版本把故事的要点（把母牛放到牲口棚里的目标，即上述农夫故事的第一句）放在故事的末尾；另一种版本把故事要点省略了，第三种版本把故事作为一集事实加以简要地讲述，而没有把原因及其结果联系起来；还有一种版本则是一些任

意次序的句子。结果，实验者发现，所呈现的故事越是符合于规则所描述的正规的形式，被试者对故事的回忆就越好。很明显，当被试者理解故事的时候，他们的理解过程涉及到使故事适合于结构规则。故事适合于结构规则的程度，则影响他们后来回忆故事的能力。

在另一个实验中，实验者让被试者挨次学习两个故事。这第二个故事或者重复第一个故事的结构，但使用新的人物；或者重复第一个故事的人物，但用新的结构。我们可能直觉地认为，由于人物重复，这就意味着两个故事中有许多词是同样的，所以，人物重复将导致对第二个故事的回忆成绩较好。但是，事实并不是这样的。如果第二个故事是在一种新的结构中使用第一个故事中的那些人物，那么，对第二个故事的回忆成绩就较差；如果第二个故事在任何方面都不重复第一个故事，其回忆成绩反而较好。当然，如果第二个故事重复了第一个故事的结构，其回忆的成绩就更好一些。这些结果表明，故事的结构，是被试者学习的对象之一。事实上，故事的结构，似乎是独立于填满这种结构的特定人物或动作而由被试者记住的。所以，这种结构能够被应用于新的人物和动作。这种情况表明，故事的一般结构，是我们长时记忆中存储的知识的一部分。

人们认为，脚本或故事的结构规则，把许多构成命题的信息结合在一起；并且，它们是以高度有组织的方式来综合这些信息的。此外，它们本身又可以结合起来，构成更为复杂的知识表征。例如，许多脚本可以结合起来，构成更复杂、更高级的表征，从而产生用来描述“城市之夜”（把剧院脚本和餐馆脚本等结合起来）那样的成批使用的脚本。在长时记忆中，知识的宽度和复杂性，看起来实质上是无限的。因

- A: 故事→目标+情节+解决
 情节→子目标+企图+结局
 企图→事件或情节
 结局→事件或状态
 解决→事件或状态
 子目标、目标→所希望的状态
- 意指该项可以重复任意次数

B:

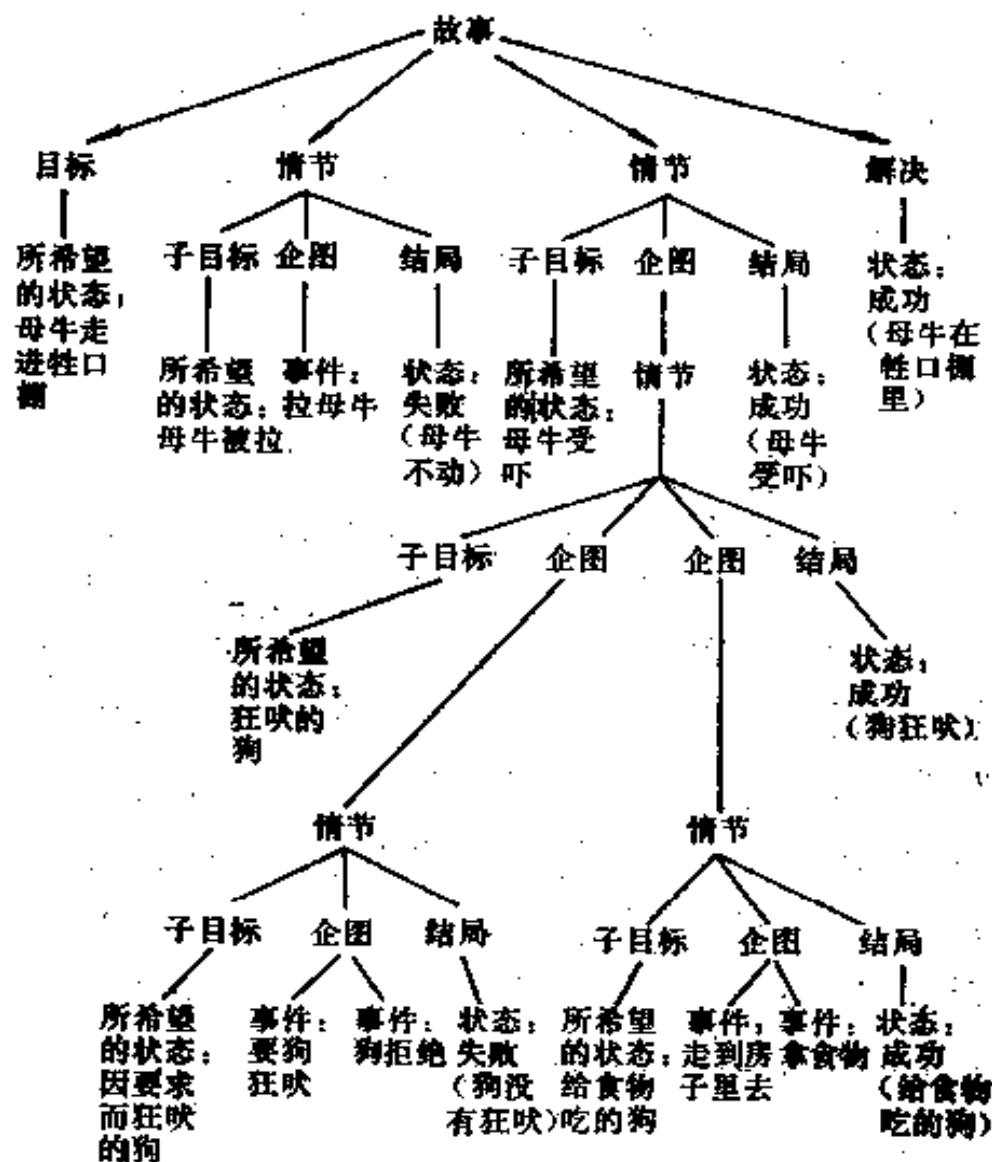


图 6-10 (A) 故事结构文法的规则

(B) 根据规则表示农夫故事的结构树状图

此, 知识的表征问题, 非但是一个问题, 而且是一个复杂的、远没有真正解决的问题。

第七章 语言和理解

人们常说,语言*是人类智力活动的窗口。这句话并不夸张。我们有理由认为,语言的理解和生成,或许是人类所做的最复杂的认知活动。所以,关于语言加工的研究,可以说是人类信息加工研究的基础,在认知心理学中占有极其重要的地位。

一、语言和交往

语言是人们表达思想,进行交往的工具。进行交往则是一个复杂的过程,它不等于把一些词按照文法规则串在一起。讲出来的话所传递的意思,有时候与字面上解释的意思是不同的。说话的声调、姿态和表情,都会给所说的话以特定的含义。所以,一件事情怎么说,与我们说什么事同样地重要。因此,要成为一个语言的有效的使用者,就必须要考虑到你与其进行交往的人们的特征,例如,他们的社会背景,他们的知识,以及他们参加交往的目的等等。事实上,微妙的社会习俗支配着及人们之间的交往。

• 注:确切地说应该是“言语”。语言和言语是两个相互联系而又有区别的概念。语言是指词的符号系统(如英语、汉语等),它们是交际的工具。言语则是人凭借某种语言来进行交际的过程。显然两者是不同的。但是,另一方面,两者又是不可分割的。离开了语言,就不能进行言语活动;而没有言语活动的要求,也就不可能产生语言。在本章中为了叙述上方便起见,对语言和言语这两个概念,不再加以严格的区分。

(一) 注意倾听

在人们彼此之间进行交往的时候,要注意倾听对方的话,这可以说是进行语言交往的一条基本规则。在交往中,如果违反了这条规则,表现出心不在焉或思想开小差,那将被人们认为是一种不礼貌的行为,甚至于会使语言交往不能正常地进行下去。比如说,一位老师正在给某个班讲课。老师的任务是把有关的知识传授给学生,而学生的任务则是从老师那里吸收新的知识,并把它们存储到长时记忆这个庞大的知识库中去。但是,如果有一位同学没有听清楚老师所讲的东西,并且站起来说:“老师,我思想开小差了,没有听清楚,请再讲一遍。”他这样说话,将被老师和同学们看作是一种不合适行为。然而,事实上,他说的或许是实际情况。思想开小差就是注意力分散了,这是每个人都可能发生的事情。但是,社会习俗是不允许我们在交往中思想开小差的,强烈的社会规范支配着人们之间的交往。要注意倾听对方的话,这是语言交往的一条基本规则。谁违反这条规则,将被认为是一种不合适行为。

(二) 默认的会话规则

在人们彼此交往的过程中,有一些不成文的会话规则。它们反映了谈话者之间对支配会话的社会习俗有一种共同的默契。从一个集团到另一个集团,从一种文化到另一种文化,这种谈话的规则是不同的,它们甚至依赖于正在说话的人。所以,支配两个儿童之间谈话的规则,不同于支配教授和学生之间谈话的规则。然而,有些规则是普遍的,例如:“一个时间里只有一个人说话”。尽管是皇帝与臣民之间的谈话,这条规则也

是适用的。但是，其他的规则往往是不同的，例如，表示希望说话的方式，在不同的集团或人们之间，是各不相同的。在社会同辈之间，只要对方停止说话有足够长的时间，一个人就可以开口说话。如果一个人不希望另一个人开始说话，那么，他就必须用一些无意义的声音，如“啊”、“嗯”等来填满这些间歇，以给对方一种本人需要连续说下去的信号。或者故意用“并且”、“但是”作为每句话的结尾，示意对方自己的话还没有说完。但是，在社会非同辈之间，社会地位高的人有时会用话来打断别人的谈话，在这种情况下，后者通常就会停止说话，转而倾听前者的话。

（三）传递知识

语言是传递知识的工具。人们通过彼此交往，要在他们的知识结构中增加新的、但与原有结构有联系的信息。在交往过程中，单纯重复某些听者已经知道的东西，那将导致使人厌烦；或者，陈述那些全然是新的、不熟悉的、以致听者还没有办法把它们插入自己的知识结构中去的观念，那将导致迷惑不解。交往是踏着这两者之间的道路进行的。教师们发现，这种道路是狭窄的。因为，在通常的情况下，一个老师在教一班参差不齐的学生时，总是或者有些学生感到厌烦，或者有些学生感到迷惑不解，要维持合适的水平是不容易的。即使在两个人谈话时，情况也往往是如此。

假定说，在一次偶然的交往中，甲问乙：伦敦大桥在哪里？（假设乙是知道的，而甲确实不知道）。对于这个问题有以下几种可能的回答：它是在哈瓦苏湖；它是在亚利桑那州；它在美国西部；它在美国。或者，甚至可以回答说，它不在哪里；它已经不在伦敦了；它已被某个古怪的美国人买

走了。这些回答中究竟哪一个最合适，这完全取决于甲为什么要问这个问题。所以，乙若要适当地回答上述问题，就要了解甲为什么问这个问题，甲对这件事了解多少。在这里，交往的规则是：在通常的情况下，不是去跟别人说那些他们已经知道的东西。这条规则的意思是指，说话者要考虑听者的知识结构，并在听者的知识结构中插入新的、不是多余的东西。如果甲和乙都在伦敦，那么，最后一个回答或许是正确的。为了达到以一种给予信息的方式来回答别人的问题，那么，就必须知道问者在寻找什么信息，他已拥有什么知识。如果你告诉某个人他已经知道的东西，那么，肯定不会对他有多大帮助。就上例而言，若甲是在亚里桑那州寻找伦敦大桥而问伦敦大桥在那里，假若乙不考虑甲的知识结构，不考虑甲在寻找什么信息，而回答说：它已经不在伦敦了，它已经被某个古怪的美国人买走了，那么，这肯定对甲不会有什么帮助。

交往中所遵循的谈话规则，是随具体情境的不同而变化的。在某些礼貌性的会话中，其目的是把一些声音来回传送，而不打算传递什么知识结构。此外，当一个学生回答老师的问题时，其谈话的目的就不同了，所以规则也是不同的。如果教师问学生一个问题：巴黎公社的领导人是谁？那么，学生最好不要去管前面提到的规则（即“不是跟别人说那些他们已经知道的东西”），而是要遵循这样的规则：你必须把老师们确实知道的东西告诉他们。

二、语言的表面结构和意义结构

在说话和写作的时候，我们说出的或写出的句子本身叫

做语言的表面结构。这种表面结构是人类语言行为可以看到或听到的东西，也是可以测量的东西。但是，由于这些句子传递的意义却是看不见摸不着的，而是要通过参加交往的人的记忆和思维来加以理解的。我们把句子传送的意义叫做句子的意义结构。所以，用来传递思想的词和思想本身是不同的，我们必须加以恰当的分。

(一) 同一意义结构可以有不同的表面结构

对于给定的意义结构来说，可以用各种不同的句子（即表面结构）来描述它。而且，在某种条件下，传送这种意义结构不必用完整的句子，信息的接受者通过上下文提供的信息或记忆中存储的知识结构，也完全可以理解它的意思。下面的例子表明，同一意义结构可以有不同的表面结构：

“小宝宝吃了苹果。”

“苹果被小宝宝吃了。”

从上面两个句子来看，语态和词序都是有差别的。所以说，它们的表面结构是不同的。但是，这二个句子传送的意义（即信息）是相同的。也就是说，它们具有或属于同一种意义结

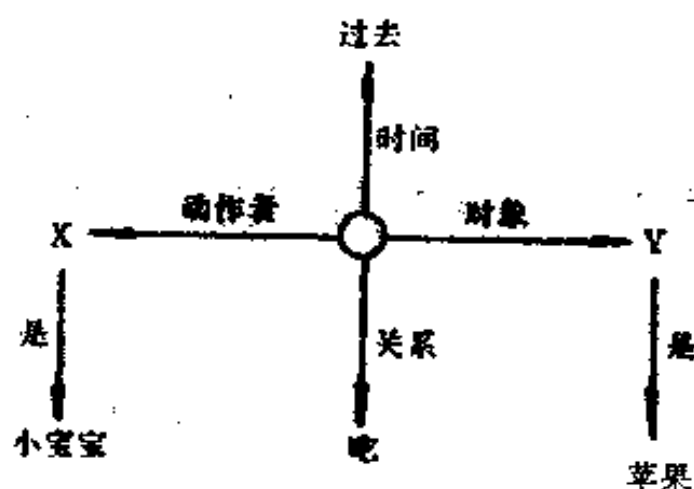


图 7-1 “小宝宝吃了苹果”和“苹果被小宝宝吃了”二个句子的意义结构

构。我们可以用图把它们的意义结构表示如下（图7-1）。

图7-1的意思是说，“小宝宝”与“苹果”之间存在着一种“吃”的关系。那么，究竟谁“吃”谁呢？从图中可以看出，“吃”这个动作的执行者是X，而“小宝宝”是X的一个实例。“吃”这个动作的对象是Y，而Y的实例乃是“苹果”。所以，自然是“小宝宝”吃了“苹果”，而决不是“苹果”吃了“小宝宝”。同时，图中也表示小宝宝吃苹果这个动作是发生在过去。这个例子说明，同一意义结构，可以有不同的表面结构。换句话说，同一个意思，可以用不同的句子来表达，这正体现了语言的丰富多采和灵活性。

（二）同一表面结构可以有不同意义的义结构

另一方面，正如不同的表面结构可以表示一个基本的意义结构一样，同一个表面结构也可以有不同的意义结构。比如说，有这样一个句子：

“财务组要进行清查”。

这个句子可以有两种不同的理解，或者说有两种不同的意思。它既可以理解为：（1）财务组要对某件事进行清查；也可以理解为：（2）要对财务组进行清查。这是有关一种句法结构对应多种语义关系的问题。所以说，“财务组要进行清查”

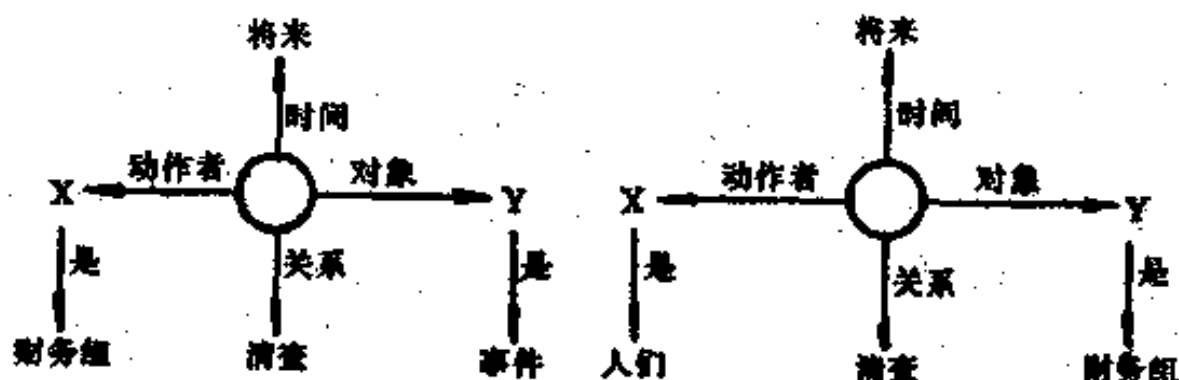


图7-2 “财务组要进行清查”的两种意义结构

查”这句话，有以下两种不同的意义结构图（图7-2）：

同一个句子可能有几种不同的理解或意思，这种情况也就是我们常说的句子的歧义问题。有歧义的句子在人类的语言中是经常可以碰到的，我们可以再举一些例子。如：

（1）我看这本书很合适。

（2）这是反对老张的意见。

（3）小王已经写上了。

（4）他知道不要紧。

对于句（1）来说，它有两种可能的含义：一是这本书适合于说话人阅读；二是说话人认为这本书适合于某种用途。句（3）也是同样的情况，既可以理解为“小王已经写上了某些字句”，也可以理解为“已经把小王的名字写在上面了”。但是，这些有歧义的句子，在一定的上下文和背景知识的情况下，它们是有确定的意义的。比如说，如果有人问某本书给孩子看行吗，那么，说者所说的“我看这本书很合适”，其意义显然是明确的。在这里，这句话传送的信息是：说话人认为，这本书适合于孩子们阅读。同样，如果人们知道财务组的帐目十分混乱，大家的意见很大，那么，“财务组要进行清查”这句话，显然表示说者认为要对财务的帐目进行清查的意思。

歧义的句子在其它语种中也是常见的。

句子的歧义是语言加工中面临的一个重要的问题。人类借助于丰富的生活知识和对母语的熟练的掌握，在语言交际中可以正确地把握句子的意义，排除大量的句子的歧义，大大地缩小了误解的可能性。上面我们所指出的例子就说明了这一点。但是，如果用计算机来处理语言，那么，对一句多义的问题就很难正确地加以处理了。这是因为，人类的知识范围很广，计算机不可能把一个人的知识全部存储下来。科

学家们研制的自然语言处理的计算机系统，都是范围有限的。它所容纳的词汇、句法和语义知识都是很少的，尤其是背景知识，更是微乎其微。计算机处理自然语言的能力与人类处理语言的能力相比，是差得很远的。所以，在碰到有歧义的句子时，对句子产生误解或不解的可能性就会大大地增加。

人们在说话或写作的时候，必须要仔细思考他们自己所要表达的意义结构，并且根据有关的语法规则来构成合法的表面结构（即符合文法的句子）。但是，在听或读的时候，这个加工的过程恰好是相反。人们必须根据一定的规则和知识去分解句子的表面结构，构造（或恢复）原来的意义结构。这种关于表面结构和意义结构之间进行转译的认知过程，是心理语言学所要加以研究的问题。在某些特定的情况下，信息的传送是否遵循正确的语法规则，似乎并不是决定这种转译过程是否成功的一种可靠因素。所以，有些句子看起来似乎具有不规整的、含糊的表面结构，但实际上，它的意义是清楚的、也是容易理解的，因为构成句子的这些词之间，只有一种可能的关系。或者说，只有一种可能的意义。例如，“射张三老虎”是不符合语法规则的，但是，如果需要的话，是容易加以转译的，也就是说，是可以理解的。在这句话里，只有一个词即“张三”是可以作为“射”这个动作的执行者，而另一词“老虎”，却是“射”这个动作的合适的对象。从这里可以看出，人具有关于世界的知识，在这种转译的认知过程中是非常重要的。

三、句子加工的心理语言学模型

语言是交流思想的工具，而句子则是表达思想的基本单

位。所以，人们在探讨语言理解过程时，往往从句子的理解上入手。下面，我们谈谈人在加工一个句子时，在执行语义水平上的认知操作时可能涉及到什么样的心理过程。

（一）句子和图形的内部表征和验证任务

理解一个句子，当然会涉及到人怎么对句子所包含的信息作内部表征的问题。不少人认为，句子信息的内部表征，很可能是一种命题的格式，即是由“谓词”（即关系项）和“变元”（或叫做“自变量”）组成的一种形式化的结构。至于这种习惯的表示法具体怎么表征一个句子，各理论家之间可能有所不同。

根据卡彭特(P. A. Carpenter)等人的见解，如果我们对一个简单的陈述句，如“圆点是红色的”，进行分析，那么，这个句子的内部表征可以是：

[AFF (红色, 圆点)]

这里，“红色”和“圆点”称之为变元，AFF(affirmative的简写，即“肯定”的意思)则称之为谓词。这种表征方式的意思是说，“红色”是对“圆点”的预测，而这种预测可以是肯定的，也可能是否定的。而谓词AFF则表示，“圆点”和“红色”这二个变元之间的关系是一种“肯定”的关系。

在研究句子理解过程时，人们经常采用“句子——图形”验证方法，即先给被试者呈现一个图形，接着再呈现一个句子，并要求被试者就图形与句子的意思是否一致立即作出“真”或“假”的判断。相对于某一特定的图形而言，句子可以分成四种类型，即“真肯”、“假肯”、“真否”和“假否”。比如说，给被试者呈现的图形是一些红色的圆点，那么，图形的内部表征格式应该是：

[AFF (红色, 圆点)]

确定了图形的内容以后,我们就知道,上述例句属于

“真肯”,因为它与图形的内容是一致的,而且是个肯定句。

同时,我们可以列出其他三种类型句子的表征格式。“假肯”句(即“圆点是黑色的”)的内部表征,应该是:

[AFF (黑色, 圆点)]

而“真否”句(即“圆点不是黑色的”)的内部表征乃是:

[NFG (黑色, 圆点)]

在这里,“NEG”这个符号表示后面二个变元之间存在着一种否定的关系。最后,“假否”句(即“圆点不是红色的”)应作如下的表征:

[NEG (红色, 圆点)]

所以称它为“假否”句,是因为句子本身是否定的,而句子的意思与图形的内容又是不一致的。

当然,各种类型句子的详细的表征形式,到现在为止,可以说还没有很好地从经验上加以验证。比如说,语言学因素怎么决定上述这种心理学的“谓词—变元”结构的表征形式,这类研究还很少。上面介绍的内部表征方式,也只是许多表征方式中的一种。另外,同样一些句子,在不同的情境下,有时可以作不同的表征。这是因为,表征中所包含的信息,是人们从句子中抽取出来的。一个人从句子中抽取什么信息,这要依赖于前面的那些句子,依赖于句子所处的情境,也依赖于听者(或阅者)以前的知识。所以说,心理学上关于句子的内部表征的概念,与语言学上深层结构概念是有区别的。

(二) 心理操作过程的模型

下面介绍由卡彭特提出的理解句子的一个心理学模型。这个模型是以上述“谓词——变元”的表征方式为基础的，用来说明人在执行“句子——图形”验证任务时所实现的心理操作。所以，在比较句子和图形时模型实现的操作过程，是我们要讨论的重点。

这个模型叫做“成分比较模型”。它有个基本的假定，即在执行“句子——图形”的验证任务时，要对两种表征（即句子和图形的内部表征）中的相应成分进行成对的提取和比较。同时假定，这种提取和比较的操作数量，是验证句子时反应时长短的主要决定因素。因为上述命题的表征结构，为我们提供了各个成分之间的一种有序的关系。这种次序决定了各成分进行比较的先后顺序：命题内部的变元先进行比较，而外部的谓词则稍后进行比较。一般来说，图形总是被编码成一种肯定的形式。“发现和比较”过程，是一种串行地进行的重复操作，可以作用于用多重括弧括起来的表征形式。这种重复性的操作，使得模型能一般化而不需要额外的假设。

模型的一个最重要的假定是，每当来自句子和图形表征的两个相应的成分彼此失匹配（或不匹配）的时候，那么，整个比较过程就要从新开始。为了防止比较过程在失匹配的成分上无限地循环下去，所以，这个模型又假定，第一次发现两个成分失匹配时，就对这两个成分记上标记，说明已经比较过了，在以后再次作比较时，这两个成分就被默认为是匹配的。图7-3是“成分比较”模型的操作流程图。

既然两个相应成分失匹配的情况将引起比较过程从头开

始,那么,比较操作的整个数量,因而也就是验证句子的反应时,会随着失匹配的次数而增加。而且,在比较过程中,外层成分的失匹配要比里层成分失匹配引起更多的重复比较。所以,整个反应时,实际上是失匹配的次数和成分在表征中的位置这两者的函数。另外,图7-3中列出的“反应变址”(记录反应的地址),监控着成分之间匹配或失匹配的情况。反应变址有两种可能的状态,或者说它有二个不同的值,即“真”或“假”。它表示人在验证句子时的反应本身。每一次失匹配使它从当前的状态改变为另一种状态,即从“真”改变为“假”,或者相反。

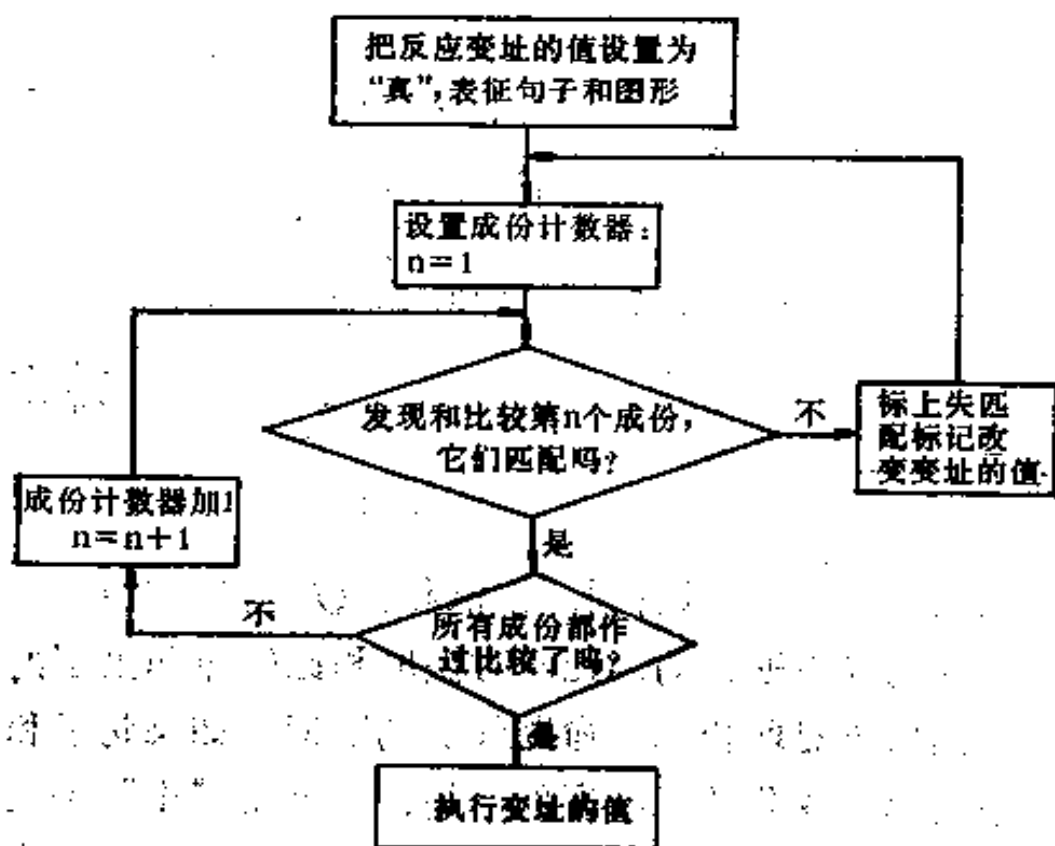


图 7-3 成分比较模型的流程图

当我们用这个模型说明上述四种句子类型的反应时间的模式时,就可以得到表7-1中列出的具体过程和结果。

表7-2四种句子类型反应时的预测

句 型	刺 激	表 征 和 比 较
真 肯	句子: 圆点是红色的 图形: 红色的圆点	$\begin{array}{cc} \{AFF(\text{红色}, \text{圆点})\} & \\ \{AFF(\text{红色}, \text{圆点})\} & \\ + & + \\ \text{反应} \rightarrow \text{真} & \\ \text{比较次数} = K & \end{array}$
假 肯	句子: 圆点是黑色的 图形: 红色的圆点	$\begin{array}{cc} \{AFF(\text{黑色}, \text{圆点})\} & \\ \{AFF(\text{红色}, \text{圆点})\} & \\ + & - \\ \text{反应} \rightarrow \text{假} & \\ \text{比较次数} = K + 1 & \end{array}$ 变址 → 假
假 否	句子: 圆点不是红色的 的 图形: 红色的圆点	$\begin{array}{cc} \{NEG(\text{红色}, \text{圆点})\} & \\ \{AFF(\text{红色}, \text{圆点})\} & \\ - & + \\ + & + \\ \text{反应} \rightarrow \text{假} & \\ \text{比较次数} = K + 2 & \end{array}$ 变址 → 假
真 否	句子: 圆点不是黑色 的 图形: 红色的圆点	$\begin{array}{cc} \{NEG(\text{黑色}, \text{圆点})\} & \\ \{AFF(\text{红色}, \text{圆点})\} & \\ - & - \\ + & + \\ \text{反应} \rightarrow \text{真} & \\ \text{比较次数} = K + 3 & \end{array}$ 变址 → 假 变址 → 真

从表7-1中可以看到, 提取和比较是自右至左或者说自里向外地逐步进行的。先同时提取和比较命题里层的二个变元, 然后再提取和比较命题外层的谓词。如果句子和图形表征中的相应成分是相匹配的, 那于, 就用“+”来表示; 如果两个相应成分失匹配, 那么就用“-”来表示。在发生失匹配时, 比较过程就从头开始, 但另起一行。从表7-1中可以看到, 在“真肯”句条件下, 句子和图形的表征之间, 没有失匹配发生。第一次比较(命题里层的变元之间的比较)

是匹配的；第二次（谓词之间的）比较也是匹配的。这样，经过二次比较，完成了句子—图形的验证任务，其反应变址的值仍然为“真”，所以，给出的反应为“真”。在验证“真肯”句的条件下，模型假设成分比较的次数为 K 。在“假肯”句条件下，句子的命题表征中的变元与图形的表征失匹配，而失匹配将重新发动比较过程，并改变反应变址的值。这样，在“假肯”句条件下，比较的次数为 $K+1$ ，而给出的反应为“假”。在“假否”句条件下，第二个成分发生失匹配，引起比较过程从头开始，因而使比较的次数增加到 $K+2$ 。但是，失匹配也只发生了一次，所以，它的反应变址的值也是“假”。对于“真否”句的条件来说，由于命题里层的变元和谓词都失匹配，因而使比较次数增加到 $K+3$ 。同时，由于发生了二次失匹配，反应变址的值改变二次，因而其反应仍然为“真”。从这个模型中我们可以看到，从真肯句、假肯句、假否句到真否句，它们的比较次数是依次增加的。这就意味着，验证这些类型的句子所需的时间，也是线性地增加的，它们的反应时的时间模式应该是：

真肯 < 假肯 < 假否 < 真否

当然，模型是用来描述和说明实际心理操作过程的，模型并不是实际过程本身，两者不是一回事。但是，它能够说明我们通过实验而观察到的现象。上面所谈的“成分比较模型”所作出的预测，与卡彭特所做的“句子——图形”验证的心理学实验所得到的数据，是相当吻合的。同时，它对国内外其它有关这类句子理解的实验结果，也能作出比较合适的解释。另外，有了这样的模型以后，我们还可以根据它来编制计算机程序，用计算机来模拟理解句子的实际过程。当然，上述“成分比较模型”还是有局限性的。例如，为什么

假设在失匹配的情况下整个比较过程要从头开始，这一点似乎很难得到解释。

四、生成转换文法

“生成转换文法”是由乔姆斯基 (N. Chomsky) 提出来的语言学理论。乔姆斯基曾强调指出，尽管他的理论不是通过心理学实验研究提出来的，但它具有心理学的现实。人们普遍认为，乔姆斯基的理论是语言学中的一次革命性的变革，给传统的语言学一个巨大的冲击。同时，它也是现代认知心理学发展的巨大动力。

根据乔姆斯基的理论，在句子的生成过程中，涉及到两种类型的句法规则。一种叫做“短语结构文法”，另一种叫做“转换文法”。

(一) 短语结构文法

短语结构文法也可称为“短语结构规则”，它是乔姆斯基的理论基石。短语结构文法把组成句子的各个元素(词)在不同的层次上分成各种短语和子短语，并且表示出这些短语之间的关系。人们可以根据短语结构本身和一系列规则的运用，来分析一个句子或解释一个句子的生成过程。为了说明这个问题，下面举一个具体的例子。假如说有这样一个句子：

The tall boy saved a dying woman.

(这位高个子的男孩救了一个快死的妇女。) 我们可以根据短语结构文法对这个句子加以分解。首先，“The tall boy”构成一个名词短语(用 NP 来表示)；而“saved a dying woman”则构成一个动词短语(用 VP 来表示)。

同时，动词短语是由动词saved（用V表示）和名词短语“a dying woman”组成的。因此，如果用S来代表该句子的话，那么可以得到：

(1) $S \rightarrow NP + VP$

(2) $VP \rightarrow V + NP$

规则（1）是说，一个句子可以分为一个名词短语和一个动词短语；规则（2）是说，一个动词短语可以分为一个动词和一个名词短语。对名词短语还可以作进一步的分解。上面两个名词短语具有同样的结构，它们都是由冠词(the或a)加形容词(tall, dying)再加上名词(boy, woman)构成的。这样，就可以把它表述为：

(3) $NP \rightarrow Art + NP_i$

(4) $NP_i \rightarrow ADJs + N$

(5) $ADJs \rightarrow ADJ + ADJs \mid E$

在这里，Art代表冠词，NP_i代表名词子短语，ADJ代表形容词，ADJs表示可以有多个形容词。所以，规则（3）的意思是说，一个名词短语可以由一个冠词加上另一个名词子短语组成的；而规则（4）告诉我们，名词子短语可以继续分解为形容词和一个名词；规则（5）则表示，在名词之前可以允许有几个形容词来修饰它，但也可能是一个形容词也没有。式中“E”表示在形容词的位置上是空的。这些规则是基本的，在实际使用时也可以任意取舍。例如，我们还可能碰到“very old man”这样的名词短语。由于这种名词短语不包含冠词，因此，我们必须再加上一条规则：

(6) $NP \rightarrow NP_i$

如果再添上某些步骤，把由上述规则分解所得的结果（V、N、Art等）与相应的词联系起来，那么，我们就完成

了该句子的分析过程，如表 7-2 所示。

表 7-2 使用短语结构文法分析句子

"The tall boy saved a dying woman"

$S \longrightarrow NP + VP$	
NP	VP
$NP \longrightarrow Art + NP_i$	$VP \longrightarrow V + NP$
$NP_i \longrightarrow ADJ_s + N$	$NP \longrightarrow Art + NP_i$
$ADJ_s \longrightarrow ADJ + ADJ_s \mid E$	$NP_i \longrightarrow ADJ_s + N$
$Art \longrightarrow the$	$ADJ_s \longrightarrow ADJ + ADJ_s \mid E$
$N \longrightarrow boy$	$V \longrightarrow saved$
$ADJ \longrightarrow tall$	$Art \longrightarrow a$
	$N \longrightarrow woman$
	$ADJ \longrightarrow dying$

同时，我们还可以利用树状图来描述上述这些短语结构规则的应用过程，或者说对该句子指派的结构。这种图称为句子的“分析树”，如图 7-4 所示。

许多句子都具有差不多同样的结构，例如，“A wolf ate the Little lamb”（一只狼吃了这只小羊羔）这样的句子，用上述这些规则也可以完成对它的分析。但是，实际上，一个句子除了有名词短语和动词短语以外，还可能有介词短语。介词短语又可以分解成一个介词和一个名词短语。因此，为了分析和生成某些句子，短语结构规则（有时也叫做“重写规则”）集是可以扩大的。乔姆斯基曾强调指出，提出短语结构文法的目的，是要能够分析和生成一个语言的全部语句。换句话说，短语结构规则应该能够分析和生成所有可能的英语句子。但实际上，要分析和生成一个语言的所有可能的句子，它需要的规则集将是很大的。所以，乔姆斯

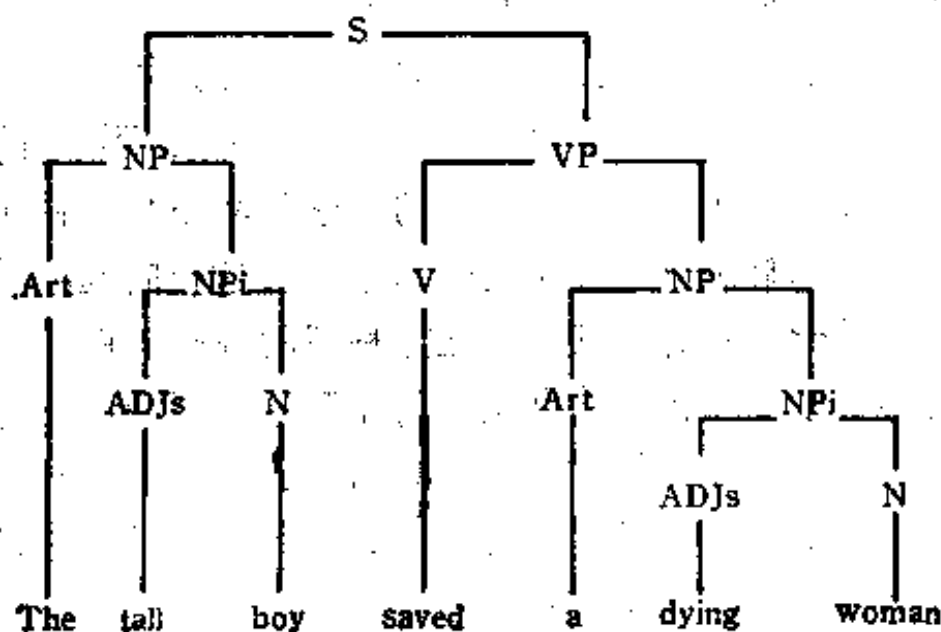


图 7-4 句子的分析树

基本人也未能写出这样一个完满的规则集来。

(二) 转换文法

对一个语言的各种句子来说，除了句子内部的短语结构及它们之间的联系之外，在句子之间也存在着某种转换的关系。短语结构文法没有考虑或说明这种转换关系。比如说，它没有考虑到英语中主动语态和被动语态之间的关系；它没有提供一种手段，去识别那些具有不同的表层结构但其意义基本相同的句子。例如，下面几个句子：

(1) The tall boy saved a dying woman.

(这位高个子的男孩救了一个快死的妇女。)

(2) A dying woman was saved by the tall boy.

(一个快死的妇女被这位高个子的男孩救了。)

(3) Did the tall boy save a dying woman?

(这位高个子的男孩救了一个快死的妇女吗?)

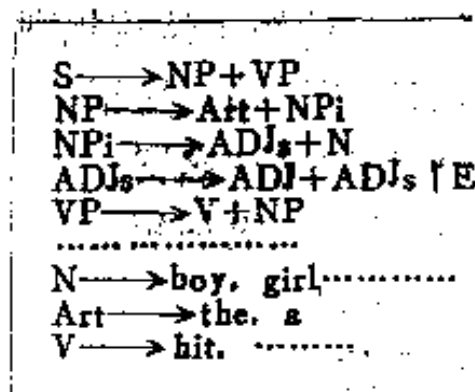
(4) Was a dying woman saved by the tall boy?

(一个快死的妇女被这位高个子的男孩救了吗?)

这些句子在形式上,或者说在短语结构上是不同的。但是,它们之间在句法上是有联系的。这种联系是一种转换的关系。怎么来解决这种关系呢?为了解决这个问题,乔姆斯基提出了另一集新的规则。他把这种新的规则叫做“转换规则”。转换规则与短语结构文法中的重写规则不同,它的目的并不是用来生成什么句子,而是应用于已经生成的一串词,并对该串词中的某些元素进行重排,对另一些元素进行删除或替代等等操作。因此,乔姆斯基的转换生成文法在分析和生成句子的过程中实际上分成了两步:第一,运用短语结构文法中的重写规则,生成基本的词串;第二,运用转换规则对该词串进行改组或重排。基本词串是一个简单的、主动的、肯定的和陈述的句子,人们有时把它叫做核心句。由核心句经过转换规则的应用,从而产生出被动的,否定的或疑问的句子等等,这个过程我们把它叫做句法的“转换”。图7-5举例说明了短语结构文法与转换文法之间的关系,以及转换规则的使用。

后来,乔姆斯基对他的理论作了某些修改,建立了所谓的“标准理论”。修改有两个方面:一是用深层结构代替核心句;二是在词汇的表达中插入了语义信息。在其标准的理论中,短语结构文法中的重写规则,生成抽象的符号串,这叫做句子的深层结构。深层结构不包括词。通过规则的进一步应用,深层结构则变成实际的语句。其中,某些规则把词插入深层结构中,另一些规则把深层结构重新排序。总之,转换规则使深层结构转化为实际的言语表达(即实际的句子)。这种实际的言语表达,就叫做句子的表层结构。另外,在标

用写规则生成一个
基本词串或核心句



该核心句作为输入
送交转换规则处理

S = the + boy + hit + a + girl

转换规则

NP₁ + V + NP₂ → NP₂ + be + V + by + NP₁

通过转换规则生
成的转换句

a + girl + be + hit + by + the + boy
→ a girl was hit by the boy

图 7-5 句子的生成和转换过程

准理论中,词汇是利用一部独立的词典来表达的。每个词根据其句法类别(如名词、动词等)和某些语义信息分别加以列出。这些信息支配着该词与其它词相结合的方式。由于加上了这种词典,乔姆斯基就在他的句法理论中引进了语义信息。这是一个重要的变化。

短语结构规则在人的语言理解和记忆中是起作用的。短语结构规则很自然地把句子中的一系列词分成若干组(或组块),这部分地决定了一个句子的知觉和记忆。一些心理学的实验结果表明,被试在知觉和记忆一个句子时,每次加工的不是单个的词,而是几个词。用短语结构规则对句子进行划分而产生的语段,是与人的知觉单位相适应的。

(三) 汉语与生成转换文法

生成转换文法在一定程度上也适用于汉语。我们可以举一个例子来说明这个问题。比如说有这样一句话:『优秀的

学生学习心理学”。它的短语结构规则应该是：

〈句子〉→〈名词短语〉〈动词短语〉

〈名词短语〉→〈形容词〉〈名词〉

〈动词短语〉→〈动词〉〈名词短语〉

〈名词短语〉→〈名词〉

〈形容词〉→优秀的

〈名词〉→学生，心理学

〈动词〉→学习

这样一个句子的生成过程，同样也可以用树状结构图来加以描述。另外，转换文法在一定范围内也适用于汉语，也就是说，汉语也可以借助于转换文法，进行句式之间的变换。但是，汉语有它自身的特点，在用生成转换文法来处理汉语时，可能会遇到不少问题。首先，汉语有自己的短语结构，不能简单地用“ $S \rightarrow NP + VP$ ”这种两分法的方式来表示。这种两分法很难概括汉语表层结构的多种句式。一个汉语句子，既可以没有〈名词短语〉，也可以没有〈动词短语〉；既可以由单项〈名词短语〉和单项〈动词短语〉组成，也可以由单项〈名词短语〉和多项〈动词短语〉组成；或者由多项〈名词短语〉和单项〈动词短语〉组成，等等，例如：

(1) 我的钢笔呢？（没有VP）

(2) 从北京到大连没有特快。（没有NP）

(3) 交通问题我们负责联系解决。（多项NP，多项VP）

(4) 苹果园坐地铁直达。（单项NP，多项VP）

其次，英语的形态比汉语丰富，词形和句型的变化也较有规律，所以，有可能找出一套进行转换的规则来。而汉语则不同，例如汉语表示时态方式就不是很规则的，而且，时态之间可以产生歧义，即同一句式可以理解成两种不同的时

态。象这样一个句子：

(5) 他们来了吗？

句(5)是一个“完成时”问句，问的是他们是否已经到了这里；但有时也可以理解为是一个“进行时”问句，意指“是否正在来”的意思。总之，汉语句式的这种灵活性，使转换生成文法的适用性受到一定的限制。

但是，不管怎样，生成转换文法为语言的结构提供一种特定的、形式化的描述手段，这一点是很重要的，因为它给计算机理解自然语言提供了有效的途径。有一些理解自然语言的计算机系统，就是以这种理论为基础而建立起来的。这种短语结构文法，稍加修改，还可以用来描述和识别图象模式。这叫句法（结构）模式识别，是近年才提出来的。

五、扩展的转移网络

“扩展的转移网络”(Augmented transition networks, 简称ATN)是由伍兹(W. A. Woods)提出来的。严格说来，扩展的转移网络自身并不包含任何语言的文法。它只是提供一种用来定义和使用某种文法的机制。

(一) ATN的结构

ATN是由状态和弧二部分组成的。在弧上可以标上某种记号，用来指导对句子的分析过程。例如，在弧上可以标上某个“推向动作”，把对句子的分析过程推向某个子网络，或者标上某种语义知识，用于检验某个词的语义特征。图7-6给出了用于分析英语句子的ATN网络片断。

图7-6以图解的方式提供了ATN的一个例子。图中状

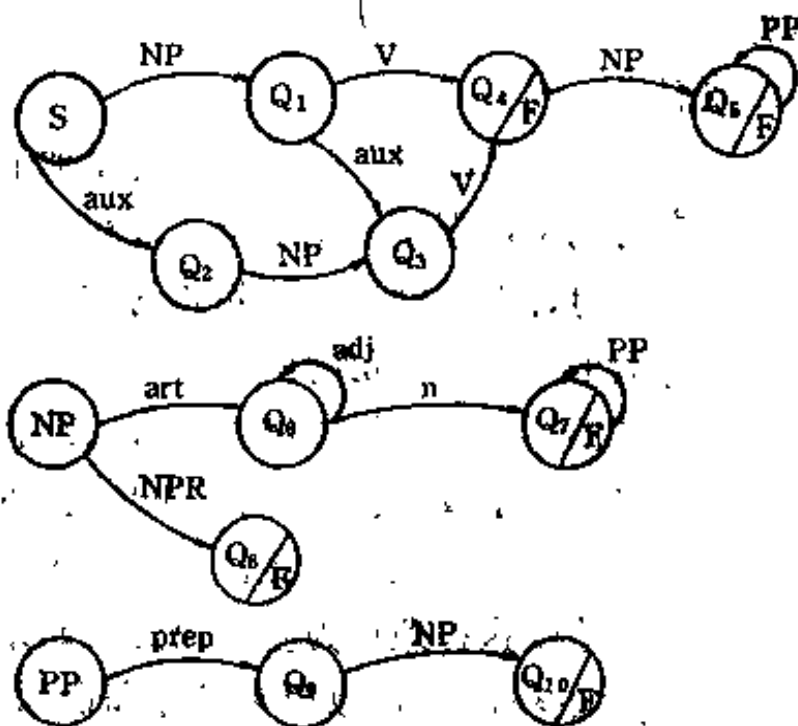


图 7-6 分析英语句子的ATN网络片断

注：s表示句子； $Q_1, Q_2, \dots$ 表示状态；状态 Q_i F表示某一特定的分析过程在此状态可以结束；PP代表介词短语；aux代表助词；Prep代表介词

态S和NP分别表示主网络和子网络；弧上的NP和adj分别表示一个下推动作和一个测试或检验动作，下面看一看ATN在分析句子时是如何工作的。

(二) ATN对句子的分析过程

假设有这样一个句子：

The brown dog has gone.

ATN对这个句子的分析过程是如何进行的呢？首先进入主网络，从状态S开始。根据从状态S伸出的弧上的标记，分析过程下推到NP（名词短语）子网络。然后，根据从状态NP伸出的弧上的标记，做一次范畴检验，看看句子的第一个词是否是一个冠词。这一检验的结果是肯定的，因此，

把冠词存储起来,并走向状态Q6。根据从该状态伸出的弧上的标记,需要做一次检验,看看句子的第二个词brown是不是一个形容词。这次检验的结果也是肯定的。于是,就把该形容词存储起来,并继续检验第三个词。看它是否是形容词。这次检验的结果是否定的,因为它不是一个形容词。因此,分析过程就沿着另一条弧进行,检验句子中的dog一词是否为名词。这一检验的结果是肯定的,所以把dog存储起来,并走向状态Q7。

如果句子已经完了,分析过程就将到此终止。Q7状态中标有“F”,说明该状态是一个可结束的状态。但事实上dog后面还有词存在,所以分析过程根据由该状态伸出的弧上的标记,下推到PP(介词短语)子网络,并检验下一个词是否是一个介词。这次检验的结果是否定的,因此,分析过程返回到状态Q7。但是,从该状态来看,已没有什么事可以做了。所以,分析过程将返回到主网络,处于状态Q1。这时,通过对句子的第一个阶段的分析,已经获得一个名词短语的结构,并初步知道该句子的类型是陈述句,并把它们分别存储起来。

接着,分析过程从状态Q1出发,检验下一个词是否系一动词。此次检验的结果也是肯定的,所以,把has放在动词寄存器,并走向状态Q4。根据从该状态伸出的弧上的标记,分析过程将下推到NP子网络。但是,下一个词gone既不是一个冠词,也不是一个合适的名词,所以,分析过程将上托并回到状态Q4。在状态Q4除了停止分析过程以外,已没有什么事可做了。但句子并没有完,所以,分析过程将返回到状态Q1,沿着由该状态伸出的另一条弧,检验has是否是一个助动词。检验的结果是肯定的,这样,就把has存

到助动词寄存器,清掉原来动词寄存器中的内容,并走向状态Q3。根据由状态Q3伸出的弧上的标记,将检验下一个词是否系一动词。在这里, *gone* 是一个动词,检验的结果是肯定的。因此,分析过程将走向状态Q4。因为整个句子已经完了,而Q4又是一个合法的结束状态,所以,对该句子的分析过程就在状态Q4结束。对该句子分析的结果是:

(S DCL)NP (DEF the) (ADJ brown)(N dog))
(VP (Aux has)
(V gone)

其中, DCL 表示该句型为陈述句, DEF 表示定冠词, Aux 表示助动词。

ATN为分析自然语言提供了一种非常有用的机制。这种机制是相当灵活而有适应力的,因此,在理解自然语言中,它或许算得上是当今最成功的分析策略之一。不少计算机理解自然语言的系统,都采用了这种机制。

至于这种分析句子的策略,是否有其心理学的现实,这个问题似乎很少有人去探讨。

六、格 文 法

格文法(case grammars)是由菲尔莫尔(C. Fillmore)首先提出来的,它为句法的和语义的分析相互结合起来的问题提供了一种不同的途径。

(一) 格文法的含意

在格文法分析中,一个句子是以动词为中心的。其他的词,都与动词处于某种关系中,这种关系称之为“格”关系。

格关系表示其它词与动词之间的语义关系。下面举二个例子来说明格文法的特点。

图 7-7 给出了二个句子的句法分析树的简化形式。在这

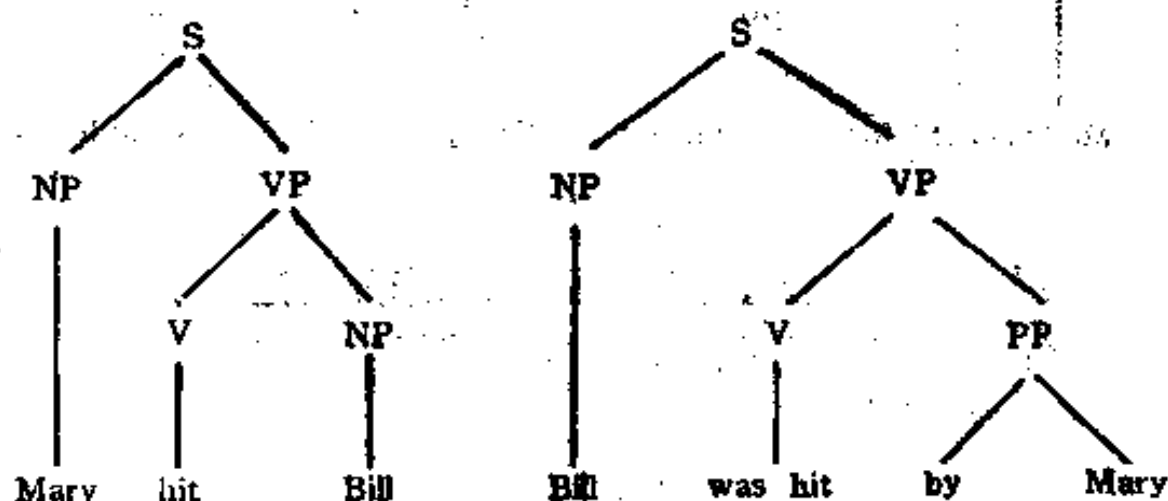


图 7-7 主动句和被动句的句法分析树

二个句子中，尽管Mary和Bill的语义角色是相同的，但是，它们的句法角色却是相反的。从图中我们可以看出，Mary分别在一个句子中是主语，在另一个句子中却是宾语；而Bill却相反。但是，如果用格文法来分析，这两个句子的解释将是同样的：

(Hit(施事 Mary)

(与格 Bill))

现在我们再看另一个例子。图7-8给出了“Mother baked for three hours”（母亲烤了三个小时）和“The Pie baked for three hours”（馅饼烤了三个小时），这二个句子的句法分析树。我们可以看到，这二个句子的句法结构几乎是相等的。在一种情况下，Mother，（母亲）是baked（烤）的主语；而在另一情况下，pie（馅饼）是主语。但是，从语义角度上看，“母亲”和“烤”之间的关系，完全不同于“馅饼”

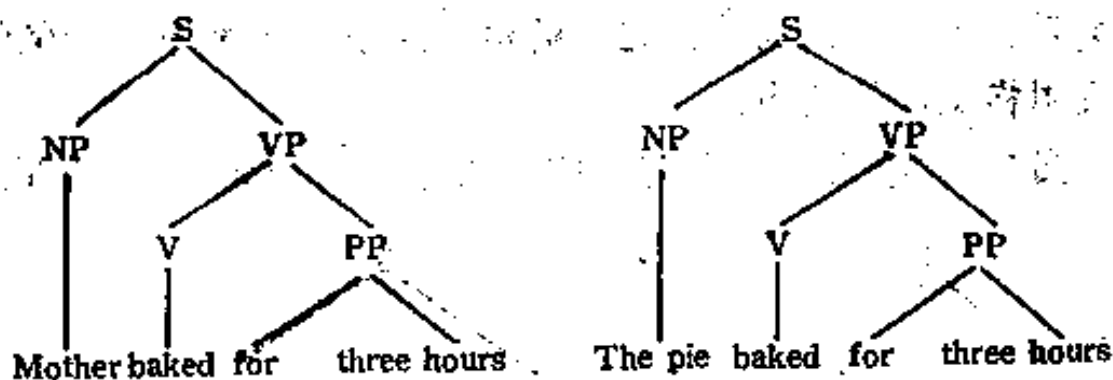


图 7-8 二个类似的句子的句法结构分析

和“烤”之间的关系。如果用格文法来分析这两个句子，那么，就能反映出这种差别。格文法将把第一句解释为：

(baked (施事 Mother)

(时间 three hours))

第二句将被解释为：

(baked (客体 the pie)

(时间 three hours))

从这里我们可以看到，“母亲”和“馅饼”的语义角色分清楚了。它们在句子中所担任的语义角色是不同的。前者是“施事”，即动作的执行者，而后着乃是“客体”，即被动作影响或改变的对象。

格文法所使用的格，描述了动词与它的变元之间的关系。这与通常的表层格的文法概念是不同的。例如，在英语中，*I*属于主格，而*me*属于宾格，但这乃是一种句法格或表层格。格文法中所涉及到的格，是一种语义的或深层的格。表层格和深层格这两者是不同的。当然，某个给定的句法格或表层格，有时可以预示各种不同的语义格或深层格。

(二) 格文法中“格”的种类

格文法究竟包含哪些深层格，还没有十分一致的看法。

大致说来，主要的深层格有如下几个：

- 1、施事格 (Agent): 典型是生物，是动作的执行者，简称A格；
- 2、工具格 (Instrument)：引起某个事件时用到的东西或某事件产生的原因，简称I格；
- 3、与格 (Dative)：受动作影响的实体，是指生物，简称D格；
- 4、使成格 (Factitive)：由事件或动作引起的产物或存在，例如：I wrote a poem. (我写了一首诗)；
- 5、处所格 (Locative)：事件发生的地点；
- 6、源点格 (Source)：某种东西开始移动的地方；
- 7、终点格 (Goal)：某种东西移向的地方；
- 8、受益格 (Beneficiary)：事件的受益者，指生物；
- 9、时间格 (Time)：事件发生的时间，简称T格；
- 10、客体格 (Object)：被作用或变化的实体，指非生物，简称O格，例如：A boy broke the window (一个男孩打碎了这扇窗户)。

根据格文法理论，任何一个动词，都要求特定的格，这些格形成一个有序集。这种格的有序集，称之为“格框架”。这样，每个动词都有自己特定的格框架。下面看一下几个动词的格框架及有关的例句。

open. (开) [——O (I) (A)]

The door opened. [O]

(门开了。)

Bill opened the door. [O, A]

(皮尔开了门。)

Bill opened the door with a Key.

[O I A]

(皮尔用钥匙开了门)。

Kill(杀死) [—D (I) A]

Bill killed a pig. [D A]

(皮尔杀死了一头猪。)

Bill killed a pig with a knife. [D I A]

(皮尔用刀杀死了一头猪。)

die (死) [—D]

Bill died. [D]

(皮尔死了。)

在每个动词后面的方括弧中给出了该动词的格框架，例如，动词Open的格框架为：[O (I) (A)]。其中，用圆括弧括起来的I格和A格是供选用的。也就是说，在以动词open组成的句子中，可以出现I格和A格，但也可以不出现这二种格。但是，O格却是在句子中必需出现的。

这样，格文法就为每个动词的使用环境提供了一个框架。在理解句子时，一旦把动词确定下来以后，就可以根据动词的格框架，预期句子中将会出现什么成分。因此，使用格文法的语言理解系统，是一种“期望驱动”的系统。

格文法反映了人类加工句子的特点。人类在理解和记忆诸如“李利打了王英”或“王英被李利打了”这些句子时，同样地主要是确定和抓住“谁打谁”这一句子的深层关系；至于其表层形式，往往是被忽略掉的。因此，有些人认为，格文法可以看作是语言理解的心理学模型。

格文法强调的是句子各成分之间的语义格关系，较少受不同种语言特点的制约。因此，我们可以预料，格文法是适用于汉语处理的。

第八章 概念和概念形成

一、概念形成的含义

概念形成、概念掌握或概念归纳，都是同义语，它们均是指同一种学习过程。所谓“概念形成”，简单地说，就是对刺激物进行分类，把具有共同的本质特征或属性的东西归入一个范畴，而把特征不同的东西归入另外不同的范畴。或者说，概念形成就是把一个名称和一类事物联系起来。

世界上无论哪个民族，不管他们使用何种语言，都把现实世界的各种各样事物分成不同的概念范畴。我们把某些事物称作家具，而把另一些东西叫做狗，等等；在日常生活中，我们也借助于有关概念的知识去回答各种问题。比如说，海豚是鱼还是哺乳动物？如果我们以前没有掌握关于鱼、海豚和哺乳动物的概念，那么，我们就无法回答这个问题。在模式识别中，有关概念知识的重要性也是非常明显的。我们知道，模式识别的过程，实际上就是把输入刺激指派给某个现存的概念范畴。总之，我们关于客观世界的知识，是由概念和概念之间的联系构成的。概念是我们具有的关于客观世界的知识——不管是语义知识还是情节知识——的重要成份，也是人类许多认知活动的基础。

概念形成的过程或者说对事物的分类过程，必然会涉及到抽象、概括和辨别这几种活动。我们根据狗的一些共同属性，把某些动物纳入“狗”这个概念的范畴。在这里，我们首先

要抽取不同狗的共同特征，同时，又从一只狗概括到另一只狗。可以这样说，如果我们能对新的未见过的狗作出正确的反应，那么也就掌握了关于狗的概念；从而，我们也就注意和辨别了狗和其他动物之间的差别。反过来说，如果某个小孩把猴子和马都叫做狗，那么，证明他还没有掌握“狗”这个概念。

总之，抽象，概括和辨别在人类认知世界的活动中是非常重要的。抽象能使我们发现一类事物的共同属性；概括使我们能够把具有共同属性的事物归入一个范畴，从而把客观世界的无穷多样化进行了整理和缩减，使人易于认识和掌握。辨别则使我们能够对不同类的刺激作出重要的区分。如果我们不能辨别羊和狼，不能辨别美味的食物和毒药，那么，我们或许就无法生存下来，更谈不上去讨论这个问题。

因为我们的概念是在抽象、概括和辨别的基础上形成的，所以，概念传送着大量的信息。例如，如果有人告诉我们，外面停着一辆“上海桑塔纳”轿车，那么我们就知道，这辆轿车是用西德技术在上海制造的，汽油发动机为 1.8 公升，采用前轮驱动技术；而且知道它不是拖拉机，不是一匹马等等。所以说，告诉我们一个客体属于什么概念范畴时，也就传送给我们有关该客体是什么和不是什么的许多信息。总之，概念反映了事物的共同属性，携带着一类事物的重要信息；它使我们能有条理地去认知和把握世界的多样性，使我们能够应用现有的知识去处理新的情境。

二、概念的结构

所谓概念的结构是指概念的内部组织。就某个特定的概

念来说，其外延可能是大的，也可能是小的；有可能是规定严格的，也可能是规定不严的、其边界是模糊的。另外，只有当我们在某一特定的时间里检验一个概念时，才觉得它是稳定的；而事实上，在人认识世界的整个过程中，概念是动态的，经历着由不精确到精确的变化。例如，3岁的儿童有时把一切四条腿的动物都称之为狗，这就犯了扩大外延的错误。但是通过经验的积累，儿童的概念发生变化，逐渐与成人的概念相一致。当然成人的概念也是随着人的认识活动的发展而变化着的。比如说，我们的“致癌物”这个概念，随着医学研究的发展，其外延在不断地变化。

（一）规定概念的属性和规则

概念可以包括数量几乎是无限的成员。例如，概念“鱼”不仅包括我们已经看过见的鱼，或者书上读到的鱼，而且也包括大量我们从未见到或听到过的鱼。因此，我们仅依靠列举某个概念的成员，是不可能适当地描述这个概念的。当然，一个概念所代表的许多成员有各种共同的属性。例如，大多数鸟都有羽毛和钩形嘴，而大多数狗都有毛，会吠。因此，有理由认为，在描述我们的认知概念时，客体的属性是重要的。事实上，人类对客观事物的识别和编码，必然涉及到对它们的特征或属性的分析。

但是，我们关于概念的知识，不仅包括各种属性，而且也包括这些属性之间的联系。这些联系可以用规则来描述。这些规则是一种逻辑的陈述，它们详细地说明各种属性如何结合。因此，大多数心理学家认为，一个概念的定义，不仅包括一集刺激的属性，而且还包括把这些属性联系起来的规则。也就是说，大多数心理学家都是借助于那些根据规则而

结合起来的属性来定义概念的。比如说，我们可以用“黄色羽毛”、“会唱歌”和“家养的玩赏的鸟”这样一集特征来定义“金丝雀”这个概念，其规则是要求某个实例必须具备所有这些属性，它才能被确认为是一只金丝雀。又如，“水”这个概念，我们是用“无色”、“无嗅”、“无味”和“液体”这些属性的同时存在（或称“合取”）来定义的。某种东西在具备了所有这些特征时，我们才会承认它是水。

属性和规则规定一个概念的主张，可以表述为以下的公式：

$C=R(x, y, \dots)$ 在这里， C 代表一个概念； x 和 y 是它的有关属性； R 是一条把这些属性联系起来的规则。

为了进一步说明这一点，下面我们来看一看在一些实验室中研究人工概念形成所使用过的一集人工刺激，以及用来描述刺激的各种特征之间的联系的规则。图8-1中给出的刺激，包括三种维度的变化：明度、大小和形状。每个维度有

规 则	□	□	■	■	△	△	▲	▲
1、肯定（如，若暗，则+）	-	-	+	+	-	-	+	+
2、合取（如，若亮且方块，则+）	+	+	-	-	-	-	-	-
3、可兼的析取（如，若亮且/或方块，则+）	+	+	+	+	+	+	-	-
4、条件（如，若三角形，则暗的，若方块，则亮的）	+	+	-	-	-	-	+	+
5、共否定（若不是小的，且不是三角形，则+）	-	+	-	+	-	-	-	-

图 8-1 可以定义概念的五条规则，每一概念的正例用“+”表示，负例用“-”号表示

两种属性。方块和三角形是形状这个维度的属性；亮和暗是明度的属性；大和小是大小这个维度的属性。属于该概念的实例称为“正例”，不属于该概念的实例称为“负例”。下面谈一下图中列出的五条规则的意思。

1. **“肯定”规则** 此规则指定一个特定的、独自规定该概念的属性。在图 8—1 中，肯定规则指定“暗”这个属性为定义属性。所有暗的刺激，都是该概念的正例，而亮的刺激则属于负例。在这个例子中，亮度是与定义一个概念有关的维度；而大小和形状是与定义概念无关的维度，它们并不参与规定概念的正例。

2. **“合取”规则** 人类大多数概念是较复杂的，不能单用“肯定”规则来描述。例如，“父亲”这个概念，包含“男性”和“有孩子”这两个属性。由几个属性的同时存在来定义的概念，就需要用“合取”规则来加以描述。在图 8—1 中，“亮方块”这个概念意味着，它的正例必须既是亮的，又是方块。这两个属性必须同时存在。

3. **“可兼的析取”规则** 这条规则是说，概念的所有正例或者包括一个属性，或者包括另一个属性，或者两个属性都具备。例如，概念“亮和/或方块”就是这种情况。亮的刺激物或方块形状的刺激，或者即是亮的又是方块的刺激物，都是该概念的正例。我们可以看到，“可兼的析取”规则与“合取”规则是不同的。后者要求必须同时具备两个属性才能作为该概念的正例，而对可兼的析取规则来说，刺激物只需具备一个特征就可以成为该概念的正例。将此规则应用于“坏歌”这个概念时，歌词坏的歌或曲子坏的歌，或者歌词和曲子都坏的歌，都是该概念的正例，都被包括在该概念之内。

4. **“条件”规则** 图8-1中对该规则所举的例子是说，如果刺激是暗的并且是三角形，或者，如果刺激是亮的并且是方块，那么，它们是概念的正例。比如，“好学生”这个概念是由条件规则来规定的。如果一个学生学习努力，积极锻炼身体，热心为同学服务，但是，只有在不违犯校规时才是一个好学生。条件规则的变式是“双条件”规则。例如，当三角形亮的时候，并且只有三角形亮的时候，它们才是概念的实例。这就是双条件规则的一个例子。

5. **“共否定”规则** 顾名思义，这条规则是通过几个特征的共同否定来定义一个概念的。在日常生活中，“和平”这个概念的定义可以作为这条规则的一个例证。“和平”部分地可以规定为“没有打仗”、“没有恐怖”；也就是说，这个概念是由“打仗”和“恐怖”这样一些特征的共同否定来加以定义的。

其实，用属性和规则可以恰当地描述概念的观点，不仅适用于人工概念，也适用于某些自然概念，比如说，“负数”这个概念是用“小于零”这个属性来定义的。因为这是一个逻辑概念，所以这个概念的成员（-1，-129，-1/12）和非成员（2，1/3）之间的区别是清楚的、明显的。一般说来，数学家、逻辑学家和科学家们提出的许多概念，是在逻辑上加以定义的。这些概念可以用属性和规则来加以规定。然而，有许多自然概念，是不适于使用严格的逻辑定义的。

（二）概念的边界是模糊和可变的

上面谈到，对一个概念的规定，涉及到有关的属性和把这些属性结合起来的规则。这种观点强调，概念是根据逻辑规则构造起来的。这意味着，一个概念的所有实例同样地都

是该概念的良好实例，而不同概念之间的区别，是界线分明的和符合逻辑的。在实验室里关于人工概念学习的研究中，通常都采纳和强调这一见解。

然而，许多自然概念的边界，是模糊不清的。“生”和“死”、“植物”和“动物”、“溪”和“河”之间的差别究竟是什么，往往很难说清楚，常常会引起一些争论。这一事实说明，自然概念包括的成员，不是一个“全或无”的问题，而是一个“度”的问题。这就是概念边界的模糊性。由于概念边界的模糊性，所以我们平时常用“大致地说”或“差不多是”这类短语来修饰自己的陈述。比如说，当我们不能肯定一条水流是河还是溪时，我们往往会说：“大体上，说这是一条河”。这类语言学的修饰方式，常使我们说出的话是“模棱两可的”。

许多概念的边界不仅是模糊的，而且是变动的。例如，拉博夫（W·Labov）曾做过一个实验，证明“杯子”这个概念的边界，明显地依赖于上下文变动。实验时，拉博夫给被试者呈现20张素描（如图8-2a），要求他们对所画的客体命名。被试者在四种不同的条件下去完成这一分类的任务。这四种不同条件是：

（1）上下文是中性的。告诉被试者，想象这个客体是在某个人的手里。

（2）上下文是咖啡。要求被试者想象某个人握着这个客体，并正在喝里边的咖啡。

（3）上下文是食物。要求被试者想象客体里面装满了土豆糊，并坐在餐桌旁。

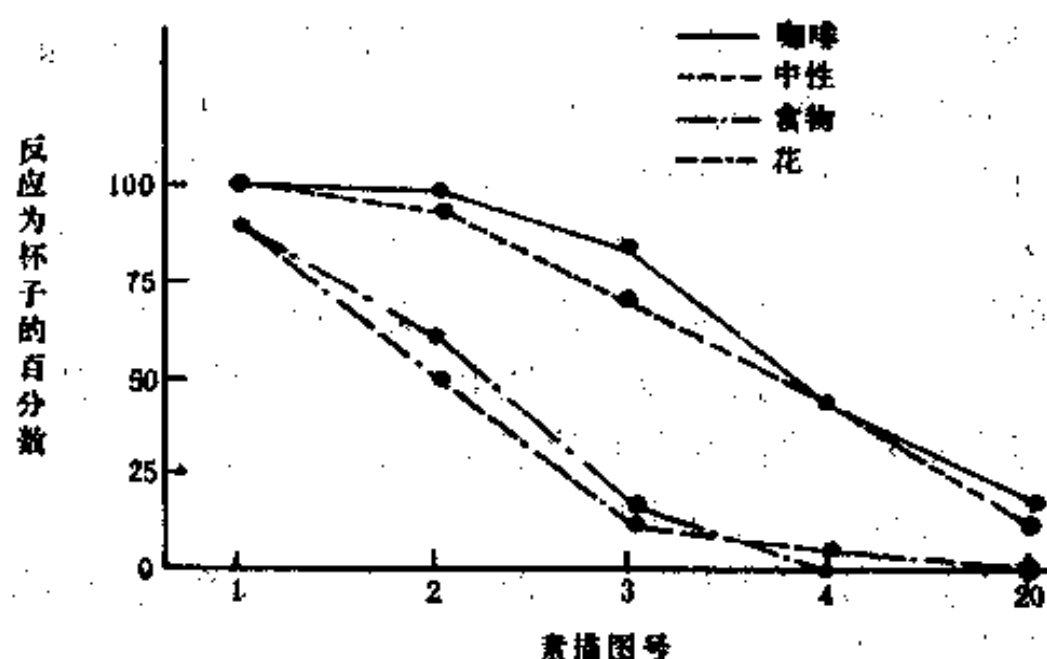
（4）上下文是花。要求被试者想象客体在架子上，里边插着修剪过的花。

实验结果表明，被试者对素描图怎样进行分类，是受上

下文所影响的（图 8-2 b）。在上下文是咖啡和中性的条件下，几乎每个被试者都把 2 号素描图说成是一只杯子；大多数被试者把 3 号素描图也叫做杯子。但是，在上下文是食物的条件下，被试者倾向于把它看成是一只碗而不是一只杯子。当上下文是花的条件下，被试者很少会把 3 号素描图看



(a)



(b)

图 8-2 “杯子”这个概念的边界(a) 被试者在不同的上下文里对之进行分类的 4 张素描图；(b) 把客体称之为杯子的被试者的百分数。作为杯子，而大多数被试者把它看成一个花瓶。这一实验结果说明，“杯子”这个概念的边界是变动的，也是连续的、模糊的。

从日常生活经验来看，概念边界对情境的依赖性也是常

见的。“高”和“矮”这两个概念的边界就是连续的，依情境而变动的。一个人身高1.7米，以前人们认为他属于高个子，而现在人们或许会把他看成矮个子了。

（三）概念的原型和内部结构

在传统上，人们是用一个概念所包括的全体成员所共有的那些属性，来描述和规定该概念的。当然，事实上，一个概念包含的所有成员或样例，在心理学上彼此并不是相等的。就拿鸟作为一个例子来说，我们觉得麻雀、燕子或喜鹊是典型的鸟；而鸡、鸭或企鹅就显得不很象鸟，尽管实际上它们均为同一个概念“鸟”的成员。这就是说，概念有它的结构，某些成员处于概念的中心，而另外一些成员则可能处于它的边缘。

1、**典型性** 罗希（E·Rosch）曾多次强调，许多自然概念都有其内部的结构，也就是说，概念所包括的成员在典型性的程度上是有差别的，其中某些成员的典型性程度高，而另一些成员的典型性程度可能较差。罗希曾做过一个实验，给被试者呈现一些日常概念范畴的成员名称（如表8-1所示），并要求被试者针对每个项目进行等级评定，指出它作为本概念的典型实例处于什么等级上。被试者对2个概念范畴的一些成员作了评定，其结果见表8-1。当两个项目，如椅子和沙发，都得同一的等级分数时，说明它们在典型性程度上处于同一等级。从表中我们可以看出，胡萝卜与西葫芦比较起来，被试者认为前者是一种较为典型的蔬菜；同样，被试者认为，椅子较之镜台，是一种较为典型的家俱。一个概念范畴的最典型的实例，一般可叫做“原型”。从概念的内部结构来看，原型处于概念范畴的中心，而那些非典型的成员，则分布在概念范畴的边缘区域。

表8-1概念范畴的不同成员的典型性的等级评定

概念：家具		概念：蔬菜	
实例	典型性等级	实例	典型性等级
椅子	1.5	豌豆	1
沙发	1.5	胡萝卜	2
睡椅	3.5	青蚕豆	3
桌子	3.5	菜豆	4
安乐椅	5	菠菜	5
梳妆台	6.5	芦笋	7
摇椅	6.5	玉米	8
咖啡桌	8	花椰菜	9
双人坐椅	10	球菜	10
柜子	11	西葫芦	11
书桌	12	莴苣	12
床	13	芹菜	13
镜台	14	黄瓜	14
茶几	15.5	甜菜	15

从表8-1来看，被试者认为豌豆是蔬菜的典型实例，胡萝卜、青蚕豆次之；椅子和沙发是家具的典型实例，而桌子和睡椅则次之。这里需要指出的是，这种评定结果，有可能受文化、习惯和传统的影响。如果我们对东方人进行这种实验或调查，其结果与西方被试者所给出的结果不会完全相同。所以，罗希的实验结果只是说明，概念的不同成员其典型程度是不同的。

罗希等人还对属性怎么分布在各概念成员之间的问题进

行了探讨。他们选定了6个概念（如家俱等），从每个概念范畴中选出20个普通客体的名称，要求被试者列举出客体的所有属性。也就是说，凡是被试者能想到的客体的属性，都把它们列举出来。结果发现，被试者列举出的对某概念范畴的所有成员来说是共同的属性是很少的。平均来说，被试者列举的属性中，有10个属性是该概念的两个成员所共有的；大约4个属性是该概念的10个成员所共同的。

某一特定概念成员与其他概念成员共有多少属性，也是随概念成员的典型程度不同而有较大的差别。一般来说，非典型的概念成员与其他概念成员共有的属性是很少的；相反，高度典型的概念成员却具有许多与其他成员共有的属性。例如，实验结果表明，在“家俱”这个概念范畴中，五个很典型的成员有13个属性是共同的。但是，在该概念的典型性差的5个成员中间，只有2个属性是彼此共有的。由此可见，概念的一个成员有与其他概念成员共同的属性愈多，被试者就会判断这个成员的典型性愈高。换句话说，概念成员的典型性水平，依赖于该成员与其他成员共有属性的多少。此外，罗希等人还发现，一特定概念成员的典型性愈高，则该成员与其他概念范畴的成员共有的属性就愈少。总而言之，概念成员的典型性高，那么，它与同一概念范畴的其他成员将会共有不少属性，而与不同概念范畴的成员则很少有共同的属性。

当我们听到一个概念名称的时候，我们往往会倾向于想起作为该概念的原型，或者想起该概念的其他高度典型的成员。这正是因为它们具有许多与其他成员共同的属性，成了该概念的一种具有代表性的成员。

2. 原型加变换 不少心理学家相信，概念范畴最好是

用原型来描述。上面已指出，一个概念的原型，就是该概念的最典型的成员。它具有许多与其他典型成员共同的属性。但是，这些属性并不是该概念所指的全体成员都共有的。所以，概念的不同成员之间存在着一个“度”的问题。一个概念成员如果有许多属性是与原型共有的，那么，它就被认为是比较典型的；如果与原型共有的属性很少，它就被认为是非典型的。这种描述概念范畴的方式有它的优点，它抓住了概念的内部结构。避免把各种属性看作是概念的所有成员都共同的。

有人把这一思想进一步具体化。认为概念可以由原型加上一集变换规则来加以描述。这些规则规定了概念的边界，并且表示了一特定的概念成员与其原型之间的类似性程度。

弗兰克斯 (J. J. Franks) 和布兰斯福德 (J. D. Bransford) 的实验，例证了原型加变换的思想，他们制作了一张卡片，上面有 4 个图形：方块、菱形、红桃和三角。这一卡片作为原型（见图 8-3 顶部）。然后，他们又确定了若干变换原型的规则，把这些规则应用于原型，从而产生一集与原型卡片有所不同的“目标”卡片。例如，一条规则是把原型卡片左右两边的图互换位置；另一条规则是从原型卡片中删去一个图，如此等等（见图 8-3）。在实验时，先呈现各种目标卡片给被试者学习，但不给他们看原型卡片。然后，对被试者进行一次“再认”测验，在再认测验中，他们给被试者看一系列卡片，其中包括原型卡片和目标卡片，以及一些新的没有给被试者呈现过的卡片。被试者的任务是要从众多的卡片中认出那些已给他呈现过的目标卡片，并且利用五点量表来评价他对自己的决定有多少可信度。

结果表明，被试者不但认为看见过原型卡片，而且其可

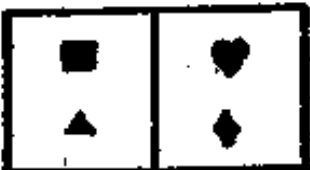
原型	
左右位置变换	
上下位置变换	
删除	
使用三条变换规则的卡片	

图 8-3 弗兰克斯和布兰斯福德实验中使用的原型及其变换

信度是最高的。事实上，实验者并没有给被试者看过原型卡片。另外，目标卡片与原型愈类似，它得到的可信度等级就愈高。从图 8-4 中我们可以看出，变换的次数愈少，则被试者对自己以前看见过这张目标卡片的可信度愈高；由原型经过多次变换而产生的目标卡片，被试者则对自己以前曾看见过这张目标卡片的可信度极低。因此，这一结果有力地支持这样的观点，即人们形成的概念范畴是可以由原型加转换规则来加以描述的。

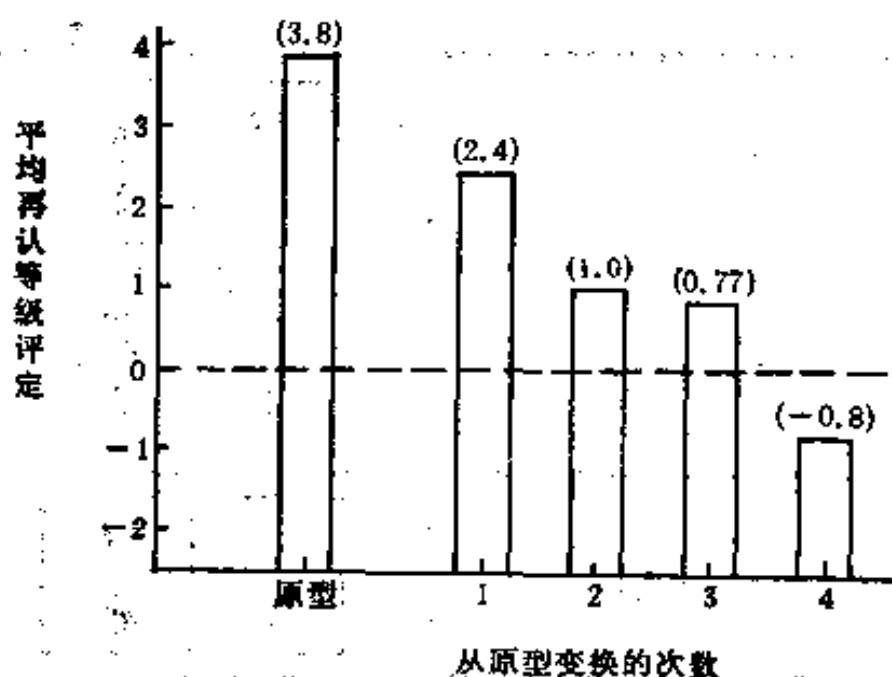


图 8-4 弗兰克斯和布兰斯福德的实验结果

(四) 概念的层次和基级概念

现在，再谈一下概念之间的联系问题。概念范畴之间联系的一种重要形式叫做“类包含”。例如，“动物”是一个上属概念，它包含着“脊椎动物”和“无脊椎动物”等下属概念。脊椎动物又分为“爬行动物”、“鱼”、“鸟”和“哺乳动物”等子概念。这些概念又可以进一步分下去，结果是形成了概念范畴的一种层次结构的排列。许多自然概念，都具有这种类包含的关系。例如，“水果”是一个上属概念，它包含着“苹果”、“梨”等下属概念。而后者又可分为“国光”、“红香蕉”和“红玉”或“鸭梨”、“京白梨”等等。

在概念的层次结构的排列中，有一个层次是特殊的。在这个层次上出现的概念，叫做“基级概念”，它具有特殊的性质。基级概念包含的属性是大多数成员所共有的，而在其

他概念中几乎是不可能出现的。有人做过一个实验，给被试者看不同层次的概念名称，要求被试者列出它们各自包含的所有特征。例如，“家俱”是一个上属的概念，而“椅子”是一个基级概念，“扶手椅”则是一个下属概念。

结果表明，被试者对上属概念只能列举出很少的属性，但对基级概念和下属概念能列举出许多属性。因此，基级概念似乎是最包含的概念，也就是说，它具有的属性对许多成员来说都是共同的。另一方面，各种基级概念之间彼此的差别是很大的，它们几乎没有什么属性是彼此共同的，就象椅子和书桌那样，彼此的共同属性是很少的。但是，各个不同的下属概念之间，却有许多属性是共同的。

基级概念提供的信息是丰富的，因为它们具有许多其成员所共同的属性。如果某个朋友告诉我们他新近买了一把椅子，那么，我们就会知道许多有关的属性。但是，如果他只告诉我们他买了一件家俱，那么，我们得到的有关该客体的信息就很少了。由于基级概念传递着丰富的信息，所以，我们往往是以基级概念来划分世界的。当然，认知经济是通过形成高度包含的上属概念来获得的。但是，我们知道，这种经济上的获得被信息上的损失低销了。基级概念恰好在认知经济和提供信息这两个要求之间到达了一种平衡。

其次，基级概念反映了我们周围环境的结构。在我们感知的世界中，各种属性是倾向于以“组”或“群”的形式出现的，而不是胡乱地分布的。例如，“有翅膀”、“有钩鼻”、“会飞”和“有羽毛”这些属性，是倾向于在一起出现的。人类的基级概念，恰好抓住了在我们周围环境中出现的各种属性之间重要的相互联系。

三、概念形成过程

人们形成概念的过程，是一个积极活动的过程，涉及到提出假设和测验假设，以及使用各种复杂而巧妙的策略。

(一) 设计概念形成实验的要求

关于人类概念形成的实验，虽有各种不同的具体形式，但有一个共同的模式。西蒙教授曾提出，设计一个概念形成的实验，需要考虑和确定以下四个方面：

1. **定出刺激的维度和值** 确定实验使用的刺激有哪几个方面的特征，每个特征又有哪几个不同的值。例如，“形状”就是一个维度；而形状又分为三角形、圆形或方块等，它们就是形状这个维度的不同的值。“颜色”也是一个维度，而红、绿、黄、蓝等，又是颜色这一维度的不同的值。“大小”又是另一个维度，在这个维度上，大、中、小是不同的值。我们设计一个概念形成的实验，可以编出不同的维度和不同的值。

2. **找出样例** 设计概念形成的实验，要找出适当的样例。这里所说的“样例”，就是有具体形状、颜色或其他维度及其不同值的东西。它们是各种维度的各个值的具体组合。

3. **确定样例的选择** 即由被试者自己选择下一个呈现的样例，还是由实验者决定呈现什么样例。西蒙教授认为，采用前者被试者要学得快，因为由实验者提样例时，被试者能得到的新信息较少，不一定能解决其不确定性的问题，所以，被试者掌握概念的过程就慢了。

4. **确定概念的不同性质** 最简单的人工概念只有一个

维度的一个值，如“红的颜色”；另一种概念的性质是相加的、在逻辑上称为“合取”的概念，如“红色方块”；第三种概念的性质是“分离的”，在逻辑上称为“析取”的概念；第四种概念的性质更为复杂，是“条件”性的概念；等等。本章的第二部份对此已有较为详细的叙述。人工概念的性质从简单到复杂是不同的。现在的大多数研究都是关于前两种概念。

当然，上述几点只是一般要求，具体的还要根据实验的任务和目的来考虑。

（二）假设测验模型

1. 假设测验模型及其实验证明 概念形成涉及提出假设和测验假设的过程，但是 有时候人们并不意识到他们形成概念时经历的这些过程。因此，实验者需要设计一种特殊的程序，使这些过程外化成可以观察到的材料。

莱文 (M. Levine) 曾设计了一种特殊的方法，叫做“插入试验法。”这种方法包括先做一个学习试验，后面跟着进行 4 个插入试验，并且重复循环。在学习试验中呈现的刺激，可以是一张卡片，卡片的左边有一个大而黑的字母 X，右边是一个小而白的字母 T，如图 8-5 中间部分所示。字母在颜色、大小和位置方面可能是不同的。实验者设定的概念由单个属性（如“白”）所构成。在每一学习试验中，被试者或者选择字母 X，或者选择字母 T；同时，实验者告诉被试者这种选择是否正确，也即是否符合实验者设定的概念。我们可以设想，被试者会利用这种反馈信息，就什么属性规定了这个概念作出某种假设。为了确定被试者随后是否在测验假设，所以，在每一学习试验以后跟着进行 4 个插入试验。

在每一插入试验中，呈现一个刺激（一张卡片），而被试者则要选择其中一个字母。但是，实验者不给被试者反馈信息。通过 4 个插入试验中被试者作出的反应模式，我们可以看出被试者正在测验什么样的假设。例如，假定说，在 4 个插入试验中呈现的卡片是图 8-5 给出的那几张。如果被试者提出的假设为：“黑色是正确的”，那么，他就会跟踪那个黑字母，其反应模式就是图 8-5 最左边一列所给出的那样。同样，如果被试者的假设是：“小是正确的”，那么，在插入试验中，被试者会始终选择小的字母，他的反应模式是：右、左、左、右。若被试者在插入试验中并没有测验某一特

假设				刺激				假设			
黑	X	左	大					小	右	T	白
•	•	•	•	X	T	•	•	•	•	•	•
•	•	•	•	X	T	•	•	•	•	•	•
•	•	•	•	T	X	•	•	•	•	•	•
•	•	•	•	T	X	•	•	•	•	•	•

图 8-5 用在插入试验中的刺激。各种可能出现的假设列在刺激的左、右两边，每一列表示一个选择模式，该模式说明被试者已测验过的假设。例如，若被试者的选择模式是：左、右、右、左，那么，他测验的假设是“大”。

定的假设，那么，他的选择模式就会背离图 8-5 给出的情况。莱文通过学习试验和插入试验数次交替循环的办法，就可以观察到被试者在解决问题的过程中怎样修改他的假设。

研究表明，成年人在超过95%的试验中采用假设测验的办法。此外，成年人倾向于采取一种“胜则保持、败则转移”的策略。只要某个假设还能起作用，他们就坚持这个假设。若一个假设在某次反馈试验中得到肯定，那么，在随后的整个一组无反馈的插入试验中，有95%的时间他们仍然维持这个假设。但是，若一个假设在反馈试验中没有得到印证，那么，在随后的无反馈的插入试验中，他们就不会再使用这个假设。

总之，假设测验是一种盛行的概念形成的认知模型。这种模型强调，在概念形成过程中，概念的学习者是一个积极的问题解决者，不断地提出假设，又不断地测验假设和提出新的假设。因此，假设—测验一再假设一再测验，直至成功，这就是假设测验模型的基本模式。

2. 概念形成是渐进的还是突变的 在许多日常情境下，学习似乎采取一种渐进、连续的形式。概念学习或许也是一种渐进的过程。

当然，也有一种可能性，即人们掌握概念的过程是一种非连续的、突变的过程，是以“全或无”的方式进行的。早期的关于假设测验的模型，强调概念形成是一个突变的、全或无的过程。根据这种观点，被试者学习一个概念的时候，不是一个逐渐改善的过程，而是一个突然变化的过程。也就是说，被试者对某一个试验的作业成绩，会突然达到高而稳定的水平。

莱文等人利用插入试验法进行的实验获得一些证据，支持这种突变的观点。他们发现，被试者的作业成绩，在问题获得解决之前进行的那些试验中，并没有什么改善。被试者在解决问题之前，难得把正确的刺激维度作为他们的假设的基

础。例如，若实验者设定的概念为“黑色”，在被试者找到正确答案前的那些试验中，看不出他有把“黑”和“白”作为假设的明显倾向。而一旦被试者使用了正确的假设，在以后的试验中就不会犯错误。因此，被试者是在一个试验中实现了从受错误支配的作业到不受错误支配的作业的转移。这些材料表明，概念形成是一个全或无的突变过程，而不是一种渐进过程。

近年，国内有人（陈家麟）采用汉字假设测验模型研究了概念形成过程。结果表明，在三种不同条件下（“是”、“非”和“交替”三种概念的形成），被试者形成概念的学习曲线的形状大致是相同的。这些曲线的形状表明，在未形成概念前，有相当长时间的高原现象，尔后，高原现象消失，所犯错误的单元数迅速减少，突然形成概念，表现出一条突变曲线。这一实验同样说明，概念形成似乎是一个全或无的突变过程，而不是一个逐渐改善的渐进过程。

当然，我们也应该注意到，莱文及其他许多人所做的实验，实际上是关于概念确认的研究。被试者的任务是确认一个概念，如“白”；这个概念本身他以前已经学过了。如果要求被试者学习一个全新的概念，这个概念是由一些他先前从未见过的属性规定的，那么，这种学习过程有可能是逐步提高的渐进过程。当然，对这个问题很难作出完美的令人满意的回答，需要进一步研究和讨论。

3. 测验多个假设 在一个时间里，被试者是测验一个假设，还是同时测验多个假设，对这个问题，莱文也曾利用插入试验法进行过探讨。如果被试者在一个时间里只测验一个假设，那么我们可以想象，若产生的假设是“白”，他们就选取白色字母；若这种选择是不正确的，那么，他们就会

仅仅拒绝“白是正确的”这个假设。但是，被试者也许会同时测验几个假设，以使自己从一个不正确的反应中学到更多的东西。比如说，若选择“小的、白的T”乃是一个不正确的反应，那么它就会使这些假设：“T”、“白”、“小”和“右”都无效。所以，在这种情况下，一个反馈试验也许能消除一半候选的假设。

莱文根据若干实验观察的材料，对被试者在一个反馈试验中测验的假设的数量问题作了分析和估计。莱文认为，如果被试者在一个时间里只测验一个假设，并且不需要去记住哪个假设已经测验过，那么，可以预料，其结果将如图8-6顶部的那条线（标有“无记忆”）所示；图8-6底部的那条线（标有“完美的加工”）表示，被试者在每个反馈试验中尽可能多地消除掉假设，并且相当好地记住以前的假设和

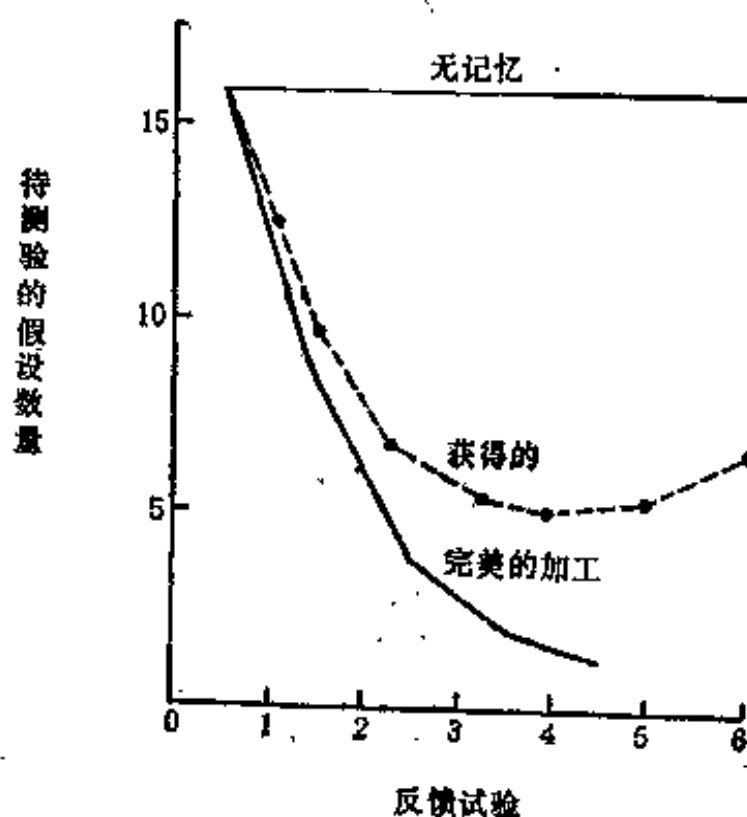


图 8-6 被试者在相继的反馈试验期间测验的假设数量

测验的结果。图 8-6 中间的那条曲线表示，被试者并不以最大的效率进行作业，但是，根据一个特定的反馈试验，被试者消除掉的假设确实不止一个。大多数被试者的作业属于后者。

这些结果与某些早期的概念形成理论是有矛盾的。以前，人们认为，被试者在学习概念时，并不记住前面已测验过的假设。而要象莱文的被试者那样有效地进行作业，他们就必须记住哪些假设已经测验过，哪些假设的候选者还没有测验过。其他的一些实验结果也表明，被试者在一系列试验过程中确实记住一些特定的假设。此外，戴尔(J. C. Dyer)和迈耶(P. A. Meyer)已经证明，若训练被试者利用形象记忆术去记忆概念的一些实例(正例)。可以促进被试者的概念学习。因此，看起来，记忆确实对概念形成的进程是有积极影响的。当然，记忆的重要性是依任务的复杂程度而异的。

人的认知是相当灵活的，所以，我们在谈论人的信息加工活动时，切忌简单地用非此即彼的方式来加以描述。下面，从人们学习概念时往往使用各种不同的策略中，我们也可以看到人类信息加工活动的灵活性。

四、概念形成中的策略

人们怎样形成概念，有各种不同的看法或理论，概括起来有三种，(1)通过联想来学习概念；(2)通过假设测验来形成概念，也就是前面讨论的假设测验模型；(3)通过实例记忆来学习概念。

联想主义理论把概念学习看作为是刺激特征和反应联结

起来的一种过程。因为联想主义者相信，联想是通过刺激一反应的接近而自动地形成的。但是，心理学家们早就指出，人们是通过积极地计划一种方法或策略来学习概念的，这些方法或策略指导着他们的作业。这些策略超出了联想主义的范围，鼓励了积极的信息加工理论的发展。

这里有两种基本类型的策略：分析和非分析的策略。

(一) 分析的策略

“分析策略”涉及特征、规则或原型的抽取。前面提到的莱文的被试者，采用的是分析策略，因为他们从呈现的刺激中抽取各种属性，并就这些属性是否规定这个概念加以测验。而在关于原型的实验中，被试者似乎也抽取了概念的原型或最佳代表。

布鲁纳 (J. S. Bruner) 等人所做的关于概念形成的实验是一个经典性的研究，说明了成人使用的某些分析策略。实验的步骤是：给被试者看一组卡片，卡片上有各种图形，如方块、圆、十字形等。各卡片上的图形类型、图形的数量和颜色，以及边的数量是不同的；首先，让被试者选一张卡片，接着实验者告诉被试者，这张卡片是否属于概念的一个实例；如果被试者错了，就让他们再选另外的卡片，如此重复下去，最后要求被试者认出这个概念。这种方法叫做“选择法”，由被试者选取每个试验中的刺激。这种方法有它的优点，它允许实验者能观察被试者所做的一系列不受限制的选择。而这种选择的系列，则可以用来推论被试者使用了什么策略。

1. “稳妥集中”策略 布鲁纳推论出来的一种分析策略叫做“稳妥集中”策略。这种策略包括：认出一个正例，形

成一个假设，这个假设包括该实例中含有的全部属性；然后，系统地测验每个属性的关联性。假定说，被试者选择的第一张正确的卡片上有1个红色十字和3条边。因为这张卡片是属于概念的正例，所以使被试者能排除掉这样一些属性，如绿、2个图或1条边等。在这种情况下，一个采用“稳妥集中”策略的被试者，会注意到哪些属性是无关的和哪些属性是有关的，并提出一个全局性的假设。这个假设包括所有潜在地有关的属性，即1个图、红色、3条边。然后，被试者将会对其中每一个属性进行检验。例如，被试者或许会选一张上面有1个绿十字形和3条边的卡片来检验属性“红”。如果这张卡片也属于概念的正例，那么，被试者就会从那些潜在地有关的属性中排除掉“红”这个属性。下一步，被试者可能挑一张有1个红十字形和2条边的卡片。如果这张卡片不属于概念的正例，那么，3条边这一属性将作为有关属性而保留下来。这样，一次检验一个属性，不断地做下去，被试者就会掌握这个概念。

所以，“稳妥集中”策略的特点是，被试者在测验假设过程中一次检验一个属性。

2. “冒险集中”策略 这是一种冒风险的策略。这种策略与“稳妥集中”策略不同，它每次要改变2个或更多的属性。仍以上述为例，如果被试者选的第一张正确的卡片是1个红十字形和3条边，采用“冒险集中”策略的被试者，下一步就可能选那张有1个红方块和2条边的卡片。这样，被试者同时改变了两个属性即十字形和边的数量。

这种策略与“稳妥集中”策略相比，带有较大的潜在风险，也有较大的潜在得益。其风险是，一旦选择错误，被试者往往不能明确地断定问题在何处。也就是说，若被试者一次

改变了 2 个属性，而且是错了，被试者很难确定究竟哪个属性是无关的。“冒险集中”策略的优点是，如果被试者改变了 2 个属性，并且选择的卡片仍然是正确的，那么，2 个无关的属性，并且，或许所有的属性，都能在一次选择中识别出来。这样，“冒险集中”策略，既可以使人很快地解决问题，也可以使人很慢地解决问题。当然，究竟发生哪一种情况，很大程度上依赖于机遇。

3. “扫描”策略 这种策略与集中策略不同，被试者测验的假设，并不包括所有可能有关的属性。比如说，如果第一张被选的卡片上有 1 个红十字形和 3 条边，而且属于概念的正例，那么，被试者可能选定属性“红”来加以系统的检验，但不管其他的潜在有关的属性。因此，被试者接着就可能选择一张有 2 个绿圆和 1 条边的卡片。但是，在这里，他并不是在测验形状和边的数量这些属性，而只注意到颜色。这种策略称作“逐次扫描”，被试者一次只测验 1 个属性。这种方法的效率相对来说比较低，使用这种策略的被试者，并没有充分利用哪些属性是有关的和哪些属性是无关的一切信息。

除了“逐次扫描”策略以外，还有一种叫“同时扫描”。使用这种策略的被试者，一次检验多个属性。但是，它与“冒险集中”策略之间是有差别的。被试者在使用这种策略时，他的假设并不包括全部有关的属性。例如，若第一张属于正例的卡片是：1 个红十字形和 3 条边，那么，被试者接着就可能测验两个属性：“红”和“3 条边”。这个策略的主要问题是记忆的要求很高，因为被试者必须记住哪些属性已检验过了，哪些属性是有关的或无关的。

在布鲁纳的实验中，许多成年被试者都使用了“稳妥集

中”策略，并且，他们解决问题常常要比使用其他策略的人快。“稳妥集中”策略有一个主要的优点，即它对记忆的要求不高。使用这种策略的被试，不需要记住哪些卡片已经选过了，哪些是正例或负例。只有当前那个包含全部潜在地有关的属性的假设是必须记住的。但是，对于形成概念来说，“稳妥集中”策略并不是一个理想的策略。对那些复杂的概念来说，记住哪些卡片已经选择过，哪些是正例和哪些是负例以及其它等等，这是很有帮助的。因此，“稳妥集中”策略只是任务的逻辑要求和记忆要求这两者之间的一种折衷的策略。

（二）非分析的策略

非分析策略的特点是强调实例记忆，不涉及从概念的正例中进行属性、规则的抽取。被试者在使用非分析策略时，首先记住了某概念范畴的特定正例，然后，根据该正例与新刺激进行比较，以确定新刺激是否属于该概念范畴。例如，用非分析策略去形成“奔驰”汽车这个概念的人们，首先要记住特定的“奔驰”汽车，以后碰到别的新汽车时，就把它与先前记住的特定的正例进行比较，以确定它是否属于“奔驰”这个概念范畴。如果新汽车与记得的“奔驰”之间有高度的复盖，那么，新汽车就被指定为该概念范畴的一个实例。

根据样例记忆的概念形成理论，人们最初知道的是某个概念范畴的一些实例。概念形成的非分析策略与基于样例记忆的概念形成理论是一致的。

布鲁克斯 (L. R. Brooks) 的研究提供了被试者使用非分析策略的实验证据。他在实验中使用的刺激是二集字母

系列（A和B），这些字母系列是由英文字母X、V、M、R和T组成的。每一集字母系列都是通过特定的序列规则而构成的。例如， $X \rightarrow M$ 、 $X \rightarrow R$ 算是两条规则。它们分别表示，字母X后面可以跟着字母M或者字母R。当然，A集字母系列的规则和B集字母系列的规则是不同的。被试者分别在两种条件下进行试验。

第一，一部分被试者在熟记条件下，先学习一些对子-联想表。这种联想表中有两种类型的对子。其中的一半是用A类字母系列作为刺激，以城市的名称作为反应；另一半对子是用B类字母系列作为刺激，以动物的名称作为反应。被试者都不知道这里有两种不同类型的字母系列，而且，在表上两类对子是彼此混合的。实验者告诉被试者，在每个刺激呈现时，要立即说出与之相联的名词。这种对子-联想的学习程序，意图在于让被试者掌握这些特定的字母系列，一直继续到被试者正确地通过30个不同的对子为止。

第二，另一部分被试者在规则学习的条件下看同样的30个字母系列。实验者告诉被试者，字母系列是根据两集不同的规则生成的，他们的任务是对这些字母系列进行分类，把每一个字母系列或者归入A类，或者归入B类。在被试者每次作出反应后，实验者给予一个肯定或否定的反馈。这样，直至被试者能对全部30个字母系列进行正确的分类为止。这一步骤的意图在于鼓励被试者从各字母系列中寻找和抽取规定这两类字母系列的不同的特性和规则。

在两组被试者完成上述各自的任务后，就要求所有的被试者对一些新的字母系列进行分类，把它们或者归入A类，或者归入B类，或者指出既不属于A类也不属于B类。前面提到，熟记条件下的被试者并不知道联想表上的字母系列属于

两个不同的范畴，而且，这些被试者也说，他们并没有注意到这里有两个不同的范畴。在要求这些被试者对新的字母系列进行分类时，实验者告诉他们，他们先前学过的字母系列刺激，分属于A和B两个不同的范畴。这些被试者对他们的实验任务有些茫然，他们诉说不知道怎样去对那些新字母系列进行分类，但同意试试看。

实验结果却是令人惊奇的。规则学习条件下的被试者只有对46%的新字母系列进行了正确的分类，他们的成绩，几乎与那些没有参加过对子-联想试验和规则学习试验的被试者一样。相反，熟记条件下的被试者能对60%的新字母系列进行正确的分类。由于他们没有表现出对两类不同范畴的字母系列的特性有什么知识，所以，他们必定是使用了非分析的策略。布鲁克斯认为，这些被试者在对新字母系列进行分类时，把它们与那些特定的、他们在对子-联想学习试验中记住的字母系列进行比较，是根据这种比较来进行分类的。记忆中的实例是被试者对新字母系列作分类的依据。从这里，我们也可以得出结论说，非分析的策略有时比分析的策略更为有效。

在大多数情况下，成年人混合地使用分析的和非分析的策略。例如，有人观察到，某些参加象上述那种对子-联想试验的被试者，对用来产生字母系列的规则表现出有一定的知识；这些被试者也利用特定刺激实例的知识。所以，成年人两种策略都用，有时把它们结合起来使用。至于哪种策略占优势，将取决于任务的类型及其许多情境因素。

（三）利于使用不同策略的若干因素

人们在什么情况下使用分析策略或非分析策略，这受许

多因素影响。

1、有利于使用分析策略的因素 (1) 记忆因素可以促使人们使用分析的策略。若任务要求对许多刺激进行加工,那么,记住所有这些正例是困难的。用抽取正例的特征这种办法,就可以减轻记忆的负担; (2) 实验使用的刺激和步骤这些因素,也可以促使人们去使用分析的策略。当刺激具有非常突出、引人注目的属性时,例如,红翅膀的黑鸟,那么,人们倾向于抽取这种突出的属性。同样,当规定一个概念范畴的那些规则是突出的时候,人们也倾向于使用分析的策略。某些实验安排方面的因素,如告诉被试者去寻找规律性的东西,可以导致被试使用分析的策略; (3) 先前的学校训练也可以导致人们去使用分析的策略。在整个正规教学中,我们就是学习给许多概念下定义。在学习给一个概念下定义时,老师总是告诉我们要说出规定该概念的那些属性。这种训练,使我们在学习新概念时,去寻找那些可以规定该概念的属性或其他规律性的东西。这种训练是很有价值的。因为对特定事例的抽象是一种高度有用的认知活动。在某种意义上说,对客体和事件之特性和联系的抽象,是科学探索的目标。因为抽取出共同的特性和关系,可帮助我们去注意和利用我们周围世界的规律,所以,在日常生活中,我们经常使用分析的策略,这是毫不奇怪的。

2、导致使用非分析策略的因素 (1) 如果有关的属性不显著,或者,我们只看到复杂概念的一、二个实例,这种情况会促使我们使用非分析的策略。在对非典型的概念范畴的成员进行分类时,我们往往也采用非分析策略; (2) 当人们对特定的实例记得非常好,或者,必须要作快速分类时,他们倾向于使用非分析的策略。

因此，在不同的情境下人们会使用不同的策略。同时，这两种不同类型的策略是相互补充的，两者均可用于某一特定的情境。人们利用非分析策略，可以获取有关某个概念范畴的许多实例的具体知识；这种知识又可以用来作为抽取各个实例共有属性的基础，使人们随后可以使用分析的策略，事实上，在概念学习时，当被试者既注意属性和规则，又注意特定实例时，概念形成就最快。

五、影响概念形成的有关因素

概念形成的速度受一些基本因素的制约。例如，概念范畴的内部结构及人的活动特点等对概念形成都会有一定的影响。

（一）规则类型

根据我们的直觉，某些概念似乎比另一些概念要难一些。概念的复杂性是不同的，部分是因为定义概念的那些规则的复杂性不同。有人根据概念规则包含多少逻辑操作而对概念规则作过分类。最简单的规则只是规定：一个特定的属性在所有正例中必须存在或必须不存在；比较复杂的规则，如合取规则和共否定规则，规定在所有的正例中必须存在或必须不存在两个特定的属性。这样，他们根据规则所要求的逻辑操作的数量，把规则的难度分为几个层次，并且，让没有经验的被试者在不同的难度层次上学习概念，也就是让被试者学习由不同难度的规则规定的概念。结果表明，一般来说，规则的逻辑复杂性愈大，被试者学习概念所需的时间就愈长。

那么，某些概念规则比另一些概念规则难一些，这是否是规则本身固有的特点呢？对此问题也有人作过实验研究。他们让不同的被试者学习由不同类型的规则所规定的概念。（实验中使用了4条规则，它们是：合取规则、析取规则、条件规则和双条件规则。）每组被试者必须连续地学习由一种规则规定的几个概念。为了让被试者集中注意力去学习这些规则，而不是去发现有关的属性，所以，实验者把那些与问题有关的属性告诉了被试者。

实验结果见图8-7。从图中我们可以看出，各组被试者在解决第一个问题时所作的试验次数，彼此有显著的差别。合取和析取规则学得很快，而条件和双条件规则学得相

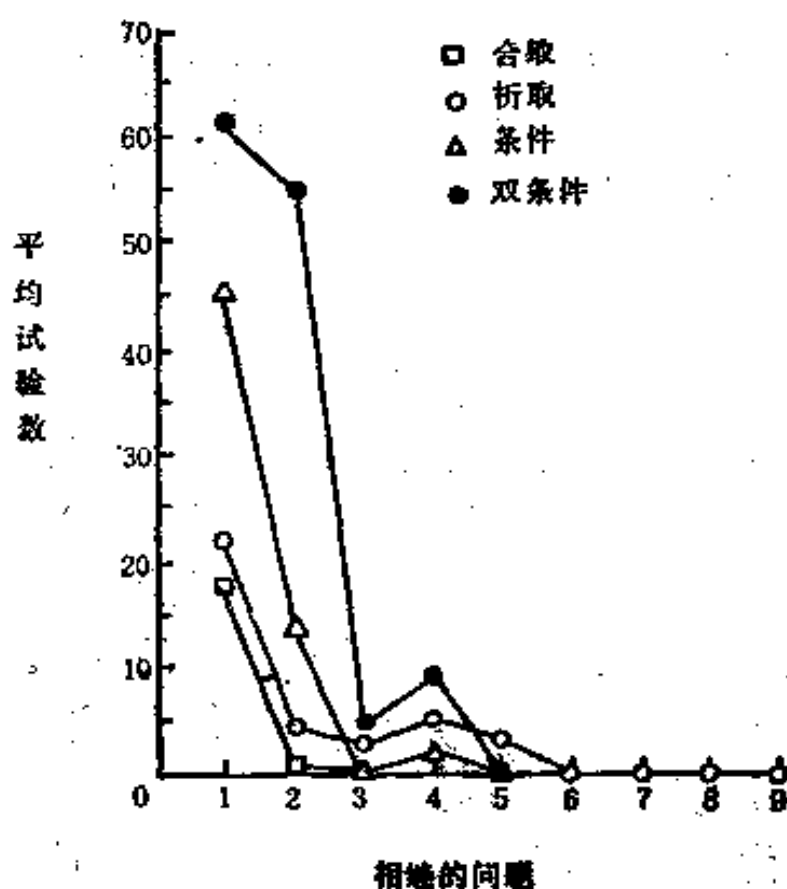


图 8-7 学习各种类型概念规则所需的试验数量

当慢。但是，随着实际练习的增加，学习速度上的差别就消

失了。这一事实说明，原初的学习速度上的差别，并不是由于规则固有难度上的差别造成的。大多数人或许对双条件规则没有多少实践，而对合取规则有相当的经验，因为日常生活的概念大多数是合取的。所以，学习不同类型规则时起初出现的差异，可能是由于先前的经验造成的，而不是由于各种规则固有难度的差异所造成的。

另外，在学习由合取规则规定的概念方面的练习，也可以使人们在学习由析取规则规定的概念时提高速度。这一事实说明，对不同类型规则的练习所获得的经验，有助于人们提高整个逻辑规则系统的知识。

（二）注意和概念形成

在概念形成过程中，注意是一个重要的因素。

皮亚杰和他的同事们已经证明儿童在掌握“守恒”这个概念时出现的差异是和注意这个因素有关的。例如，给儿童看两只同样的玻璃杯（甲杯和乙杯），里面装有等量的水，并且问儿童，两只杯子中的水是否一样多。大多数儿童都说“是”。然后，把甲杯的水倒入另一较细长的丙杯，并让儿童看着这样做的过程。自然，丙杯子的水平面要高于乙杯子。接着，问儿童，丙和乙两个杯子中的水是否一样多。如果儿童说两个杯子中的水量是相等的，并能恰当地证明他的回答；那么，就可以说，他理解了量守恒的概念。结果发现，对这类守恒问题，大多数学前儿童不能作出正确的回答，但学龄儿童的回答很正确。

皮亚杰认为，当儿童认知发展达到具体运演阶段时，在他们的思维中就表现出“可逆性”。简单一点说，年纪较大的儿童认识到，一个操作——把水从一个容器倒入另一容器

——的后果，可以通过做相反的操作把它反转过来。可以假定，还没有达到具体运演阶段的幼小儿童，他们不能理解可逆性原则，因此，他们不能理解和解决上述这类守恒问题。

这种解释当然是有道理的。但是，另一方面，有人已证明，这种现象部分地起因于对无关的刺激属性的注意问题。在日常生活中，容器中液体的高度是液体量的一个指标，然而，在上述守恒问题中，它是一个无关的属性。年幼儿童往往只注意到杯子中水的高度，从而造成水不等量的结论。

另一个是棍子长度守恒问题。两根等长的棍子并排放着，它们的头成一线。然后，把一根棍子往右移一点，使它的头伸出在另一根棍子的前面。当两根棍子排成一线即对齐的时候，幼儿往往说它们的长度是相等的；但是，当一根棍子的头伸到另一根棍子前面时，幼儿倾向于说，前者较长。这种现象也与幼儿注意无关的刺激属性有关。

有人做过一个实验，实验中把儿童分成二组即训练组和控制组；给训练组的幼儿一些奇特的问题要求他们去解决：每个问题涉及三根棍子，其中两根棍子是完全相同的，但彼此放在不同的角度上；幼儿的任务是从三根棍子中选出两根相同的棍子。在某些问题中，有两根不相同的棍子具有一个共同的无关属性（如“红色”），解决这种问题时，就要求儿童忽略掉该无关属性。对控制组的幼儿，只让他们玩呈现给训练组的那些棍子，但不要求他们去完成选择相同棍子的任务。然后，让两组幼儿重新解决水和棍子长度的守恒问题。

结果发现，训练组幼儿对守恒问题的作业成绩大大提高，而控制组的作业成绩却没有什么改善。这并不是因为训练组幼儿通过少量练习掌握了守恒这个概念，而是他们学会了不去注意那些无关的刺激属性。

(三) 信息和概念形成

刺激所包含的信息如果丰富，就会减少我们的不确定性。例如，如果对方告诉我们说，“某人姓赵，名字是二个字”，那么，要比说“某人不姓张”所传送的信息多得多。前者排除了许多可能的姓；后者只排除了一个姓，而听者仍然不知道某人姓什么。

早期的实验告诉我们，人们通过正例来学习概念比较容易，而通过负例来学习概念较为难一些，因为正例提供的信息较为丰富。回想一下我们怎么学会许多自然概念，这也会告诉我们，从正例中学习概念是最好的。对自然概念来说，负例在数量上远远超出正例，其结果是负例的信息量少于正例。事实上，我们是很少会根据负例来学习概念的。

信息这个变量是概念形成的基本的决定的因素。例如，原型被认为是一个概念范畴中最富信息的成员。原型具有的与本概念范畴其他成员共同的属性最多，而与其他概念范畴成员共同的属性又最少，也就是说，它携带的有关本概念范畴的信息是最丰富的；原型反映了客体的结构，因此，如果我们知道了某个概念的原型，我们也就掌握了许多有关概念范畴的信息。所以，在概念形成中，原型处于重要的地位，这是可以理解的。

当然，从宏观上考虑，或许还有其他许多因素，如文化和语言，对概念的形成也会有一定的影响，这里就不细述了。

第九章 推理和判断

根据事实作出各种判断，得出结论，做出选择或决定，都是人的认知行为的重要方面。这些活动表现的形式是多种多样的。其中，有些可能属于相对自动的加工过程，有些则属于高度意识的加工过程。比如说，给一个人观察一个微弱信号（如一个光点）的任务，这里就牵涉到作出判断和决定的问题。当然，对于这种简单的信号觉察任务来说，它所涉及的作出判断和决定的过程，是一种属于相对自动的、快速的过程。但是，在一些语义知识的加工过程中，或根据一些事实提供的信息而抽取出规则的学习过程以及其他许多认知活动中，这里所涉及到的比较、判断和选择，则是一种“深思熟虑”的加工过程。自然，本章将要讨论的是那些属于“深思熟虑”的推理和判断过程。

我们知道，形式逻辑这门学科所涉及的是达到正确结论的系统方法。在传统上说，形式逻辑要研究各种陈述的真或伪，提供那些把各个陈述彼此联系起来的规则。而心理学在这方面的主要兴趣，则是研究“普通人”怎样处理和判断的任务。大致地说，本章讨论的内容包括三个方面的问题：（1）首先将考察人们在评定论据（一集陈述）的内部一致性以及把证据与一个陈述的真值（真假）状态（truth status）联系起来时涉及的推理过程；（2）有关人们的回答与逻辑结论之间的一致和矛盾的问题，以及人们在获得他们的回答时所涉及的一些加工过程；（3）人们如何处理不确定信息的

问题，即人们如何评估事件的可能性及事件之间的联系等问题。

一、三段论式的推理

研究推理的一个传统的任务，是研究三段论式的推理。

（一）三段论推理的形式及其加工过程

三段论式的推理（或称三段论推理）包括一个论据（由二个或二个以上的前提组成）和一个结论。前提假定为真，人们的任务是要确定在此条件下三段论的结论是否必然地能推导出来。例如，有这样一个三段论推理：

前提 1：所有的猫都是动物。

前提 2：有些动物有尾巴。

结论：有些猫有尾巴。

如果我们给予人们一个三段论推理，我们并不是要求他们去确定前提 1 或前提 2 在事实上是否为真，而是要求他们假定这些前提是真的；同样，我们也不是要求他们去确定结论在事实上是否精确；而只是要求他们确定，从这些前提中是否必然能产生这个结论。为了着重说明三段推理的内部倾向性，再看一看下面的例子：

前提 1：所有的狗都是飞机。

前提 2：有些飞机有翅膀。

结论：有些狗有翅膀。

根据我们对世界的一般知识，在第一个例子中的陈述（猫是动物等）应该说是合理的。但是，第二个例子中的陈述看起来显然是错误的，是稀奇古怪的。然而，从逻辑的角

度来看，这二个例子是相等的，都符合下面这种抽象的一般形式：

前提 1：所有的A都是B。

前提 2：有些B有C。

结论：有些A有C。

如果有人发现自己对这两个三段论推论的例子的结论有不同的评价，对其中一个的结论是接受的，而对另一个结论却是拒绝的，那么，他并没有根据形式逻辑的要求来处理这两个三段论推理。从逻辑上说，这二个例子是相等的，因为在形式逻辑中，内容是可以不管的，而其前提却被假定是正确的。但是，一个三段论推理中各个陈述的内容在心理学上却是重要的。

表 9-1 提供了三段论推理的各种样本。这些样本是通过改变前提中所给出的联系类型而产生出来的。比如说，一个前提把A和B联系起来，另一个前提把B和C联系起来，而结论则把A和C联系起来。表中指出了每个三段论推理的逻辑结论。上面所举的这二个例子，符合表 9-1 中第二个三段论推理形式。其中的陈述“有些飞机有翅膀”，可以写成“有些飞机是一种有翅膀的东西”，因此，它符合于“有些B是C”的形式。在这里我们可以看到，从前面两个例子的结论中，是得不出逻辑结论来的。

图 9-1 表示形式逻辑中用量词来限定的各种陈述的意义。从图中我们可以看到，“所有的A都是B”这个陈述，可以指A集和B集之间的二种关系中的任何一种。A集可以是包括在更大的B集中，或者，A集与B集是相同的。前一种关系即A在B内，叫做类包含关系，如“所有的狗都是动物”；后一种关系则代表二者是等价的，如“所有的母亲都是女性”。

表9-1 三段论推理的样本

前 提	逻辑结论	前 提	逻辑结论
(1) 所有的A都是B 所有的B都是C	所有的A都是C	(11) 有些A不是B 所有的B都是C	无结论
(2) 所有的A都是B 有些B是C	无结论	(12) 有些B是A 没有B是C	有些A不是C
(3) 没有A是B 所有的C都是B	没有A是C	(13) 所有的B都是A 所有的B都是C	有些A是C
(4) 有些B不是A 所有的B都是C	无结论	(14) 有些A不是B 有些B不是C	无结论
(5) 所有的A都是B 没有C是B	没有A是C	(15) 有些B是A 有些C不是B	无结论
(6) 没有B是A 有些B不是C	无结论	(16) 所有的A都是B 所有的C都是B	无结论
(7) 有些A是B 所有的B都是C	有些A是C	(17) 没有B是A 有些B是C	无结论
(8) 所有的B都是A 所有的C都是B	有些A是C	(18) 有些B不是A 有些C是B	无结论
(9) 没有A是B 没有B是C	无结论	(19) 所有的B都是A 没有B是C	有些A不是C
(10) 有些A是B 有些B是C	无结论	(20) 所有的A都是B 没有B是C	没有A是C

注：“无结论”的意思，是指从前提中不能得出把A和C联系起来的特定陈述。

家长”。“没有A是B”是唯一没有歧义的陈述，它只能意指在A集与B集之间没有重叠。使用“有些”或“有些/不”的那种陈述，是高度歧义的，因为每一类都可能有几种不同的意义。从图中可以看出，“有些A是B”的最后二种意思，也就是“所有的A都是B”的二种可能的意思；而“有些A不是B”的最后一种意思，也就是“没有A是B”的意思。在逻辑中，“有些”的意思是指：至少是有一些，或许是全部即“所有

的。”这与我们日常的语言是很不同的。在日常的语言中，当我们要表示“全部”或“所有的”这个意思时，是不会去用“有些”这个词的。如果我们用“有些是”这个词组时，这就暗指“有些不是”，反之亦然。因此，当人们处理那些涉及到“有些”的陈述时，可能会得不到逻辑结论。

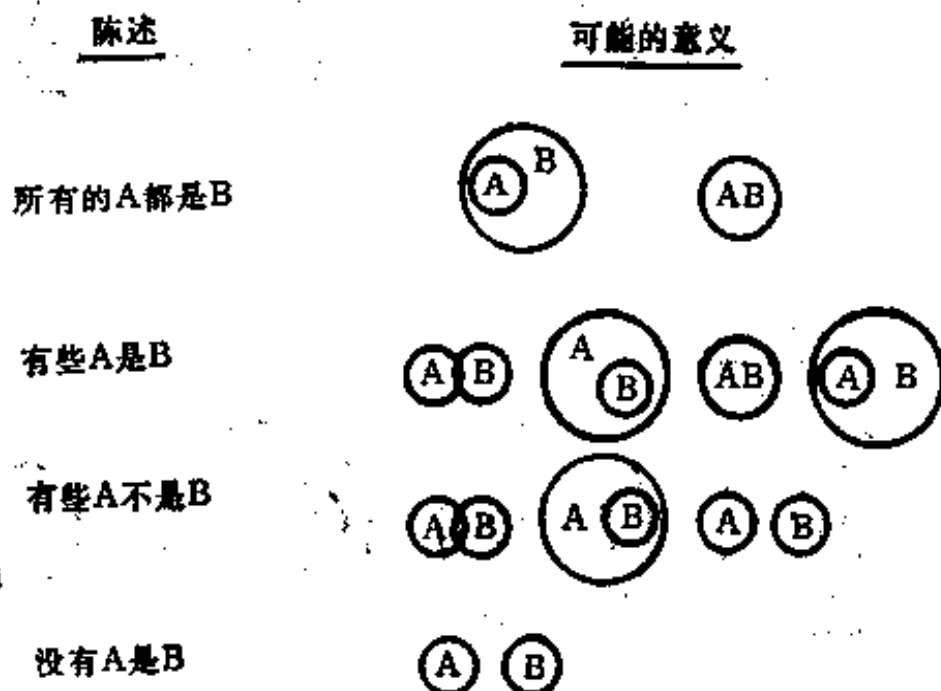


图 9-1 表示形式逻辑中以量词限定的各种陈述的意义

在逻辑中，只有当结论适用于前提的一切可能的意义组合时，这个结论才能成立。这就意味着，进行逻辑推理需要大量的信息加工。人们必须考虑每个前提的一切可能的意义；构造关于这些意义的一切可能的组合；确定一个结论是否与生成的一切意义组合相一致。

我们也可以另用一种方式来说明逻辑的要求。如果它有可能生成一种其结论对之不能适用的前提意义的组合，那么，就不能得出这个结论。图9-2给出了表9-1中第8个三段论推理的逻辑分析。以该图为例，如果结论是“所有的A都是C”，那么，它只适合于第4种前提意义的组合，但不适合于

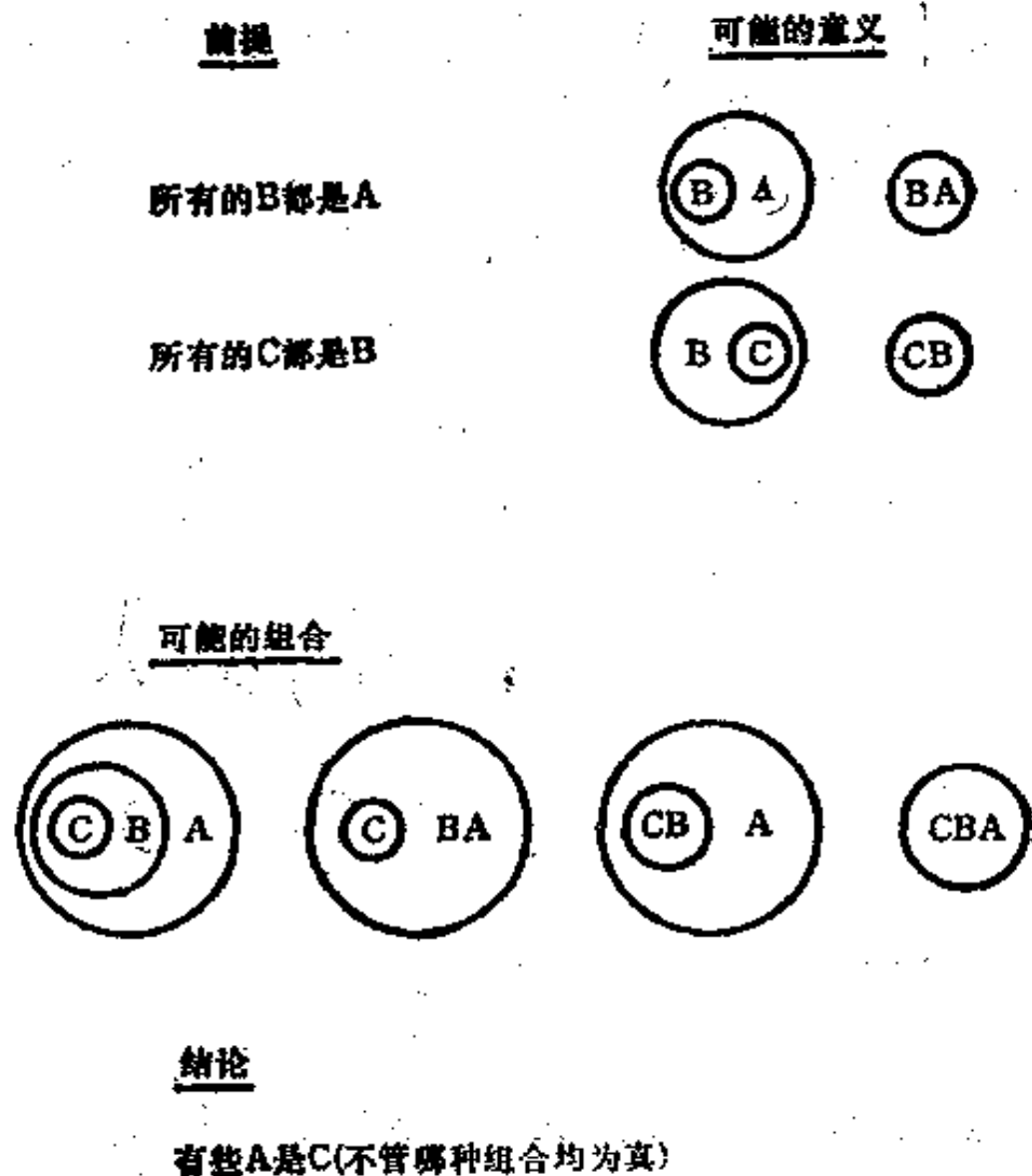


图 9-2 一个三段论推理的逻辑分析

前面三种意义组合，因此，它不是一个合适的逻辑结论。表 9-1 中有许多三段论推理，其逻辑反应是“无结论”，即不允许得出任何特定的结论；这是因为，没有一个陈述可以复盖一切可能的前提意义的组合。由于从前提中推不出什么结论，所以这个三段论推理是“无效的”。一些研究发现，人们对无效的三段论推理进行加工时，常常感到非常困难。例如，被试者对有效的三段论推理的作业，可以获得 77% 的正

确结论（可以从前提中符合逻辑地得出的结论）；但是，对无效的三段论推理的作业，只能获得33%的正确结论（即“无结论”）。这些结果表明，人们在三段论推理过程中所犯的普遍错误，是接受那种不能从前提中符合逻辑地推出来的结论。

三段论推理的加工，主要分三步：

- （1）解释前提；
- （2）生成前提的意义的各种组合；
- （3）把结论与所生成的各种意义组合进行比较。

每个人在进行逻辑推理时，都要完成这几步。当然，由于人的信息加工的能力有一定的限度，所以，有的人在进行逻辑推理作业时，时常不能获得正确的逻辑结论。这种情况，并不表明他们的行为与逻辑联系不起来，而只是说他们不能完成逻辑所要求的全部加工过程。有关的研究表明，在一些逻辑上没有受过训练的大学生中间，大约只有10%的人能够以彻底的逻辑方式进行三段论推理的作业。

（二）气氛假设

这是许多年前提出的观点。这一假设是说，人们往往根据前提造成的气氛或全局印象来得出各种结论。例如，包含“有些”的前提，往往造成一种气氛，使人容易接受那种包含“有些”的结论。气氛假设可以用下述三点来说明：

（1）当一个或一个以上的前提是“否定”的陈述时，人们将易于接受一个否定的结论；

（2）当一个或一个以上的前提是“特称”的陈述（包含“有些”）的时候，人们往往就会接受一个特称的结论。

（3）当前提不包含“否定”的和“特称”的陈述时，

人们往往就会选择一个“全称”（“所有的”）肯定的结论。

由此可见，气氛假设实际上就是预测人们就任何一对前提将会接受或认可什么样的结论。例如，一个包含“有些”的前提和一个包含“没有”的前提，将会导致人们去接受一个“有些不”的结论。

人们在认可无效的三段论推理的结论中所犯的错误，是与气氛假设的预测相当一致的。表9-2提供的数据证明了这一点。例如，表9-2中的第一项可以具体化为，

所有的A都是B

所有的C都是B

从逻辑上讲，这样一对前提是没有结论的。但是，一些被试者却错误地承认它是有结论的。在犯错误的被试者中，绝大部分是接受一个“全称”肯定的结论，即“所有的A都是C”。关于表9-2有几点需要加以说明。第一，这些数据

表9-2 在无效三段论推理中所犯的错误

观察到的选择频率

前提对偶	所有的	有些	没有	有些不
所有的，所有的	22*	3	1	1
所有的，有些	6	91*	2	19
所有的，有些不	3	34	6	147*
有些，有些	2	89*	4	40
有些，没有	3	15	34	72*
所有的，没有	1	3	61*	4
有些不，没有	8	76	44	98*
没有，没有	15	15	47*	15

注：从一行到另一行频率的变化，部分是因为在研究中使用的每一类三段论推理的数量不同。“*”（星号）标示的数据，是气氛假设的预测。

只是来源于那些无效的三段论推理；第二，那些三段论推理只包括抽象的项，例如“所有的X都是Y”；第三，在这里，气氛假设的预测和数据之间并不完全一致，并且，对各种不同形式的三段论推理来说，两者一致性的程度是不同的。例如，对包括“没有A是B”这种前提形式的三段论推理的反应，和其他形式的三段论推理相比，与气氛假设预测的一致性较差一些。

表9-2中列出的这些数据在一定程度上是与气氛假设的预测相吻合的，这是事实。但是，现在人们却并不认为气氛假设是对人类推理的一种很好的说明，尽管在有限的情境下气氛可以影响结论的选择。这里有几个问题。首先，气氛不能解释为什么人们对有效三段论推理要做得好些，而对无效三段论推理的作业要差一些。就气氛假设来说，有效三段论推理和无效三段论推理之间并没有差别。其次，气氛也不能解释为什么在前提中填入一定意义的内容会产生不同的效应，而根据气氛假设，内容应是无关系的。一般来说，与气氛假设一致的那些数据，也可以用其他的理论观点来解释。最后，当人们面对一个三段论推理并有较充裕的时间去解决问题时，气氛假设看起来往往并不是对人的作业的一种有力的说明。但是，当人们只有很少时间去考虑一个论据时，气氛是会影响人的推理的。例如，在听一个政治演说时，人们往往没有多少时间去对演说中包含的论据进行非常认真的加工。在这种情境下，气氛是有重要影响的。

（三）前提错解

许多事实表明，人们经常错解三段论推理的前提。人们对前提所包含的信息进行加工时，往往不是去考虑一个前提

可能有的全部意义，而是倾向于只是对一种意义进行编码。这里，最普遍的错误，是推想当前的陈述和它的反式是真。例如，对“所有的A都是B”这个前提，人们在解释它的意义时，倾向于认为“所有的B都是A”也是真。这是一种叫做“前提换位”的错误。图9-3提供了这种错误的若干情况。

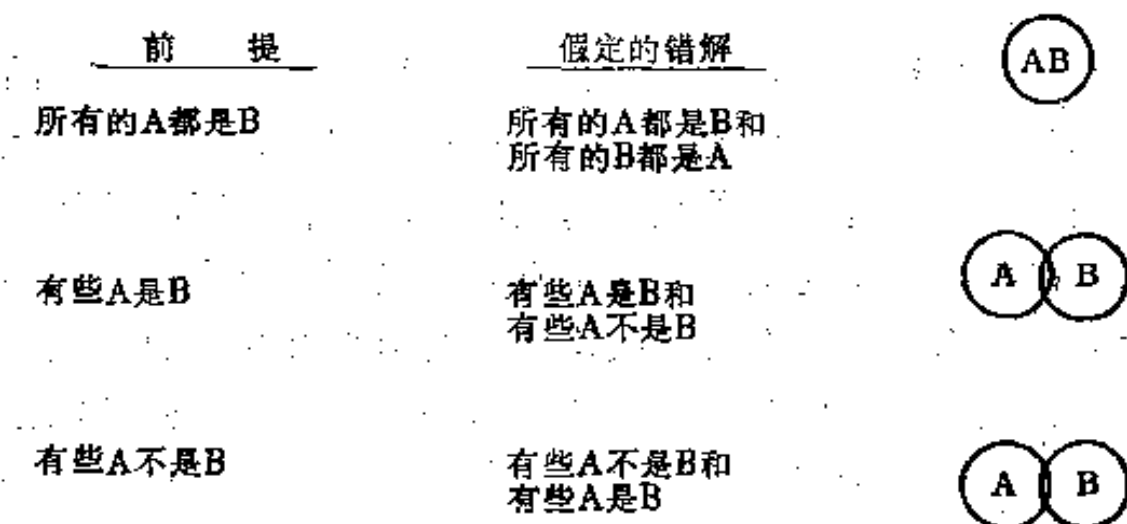


图 9-3 前提错解的若干情况

人们一方面对前提作了错误的解释，另一方面，又把他们组合起来去得出一个结论。前提换位可能会导致结论的逻辑错误，当然，也可能不会引起这种错误，如图9-4所示。这就是说，对某些三段论推理来说，对前提作有限的解释，也会得出一个有效的结论，就象对前提作了完全的逻辑分析一样。但对另一些三段论推理来说，前提换位将导致人们认可在逻辑上有错误的结论。

前提换位是人们进行三段论推理时发生错误的一个重要的根源，这一点已经通过某些方式得到证明。通常，实验者设法堵住被试者进行前提换位的可能性，给被试者一个较长的、没有歧义的前提，如“所有的A都是B，但有些B不是

A”。对于这种没有歧义的前提，人们几乎总是会得出逻辑上正确的结论。还有，实验者给被试者呈现一些有意义的、经过精心挑选的前提，以使被试者所具有的关于世界的知识能防止前提换位。例如，给被试者这样的前提：“所有的狗都是动物”。一般来说，被试者不会去设想它的反式“所有的动物都是狗”也是正确的。所以，被试者在进行推理时不

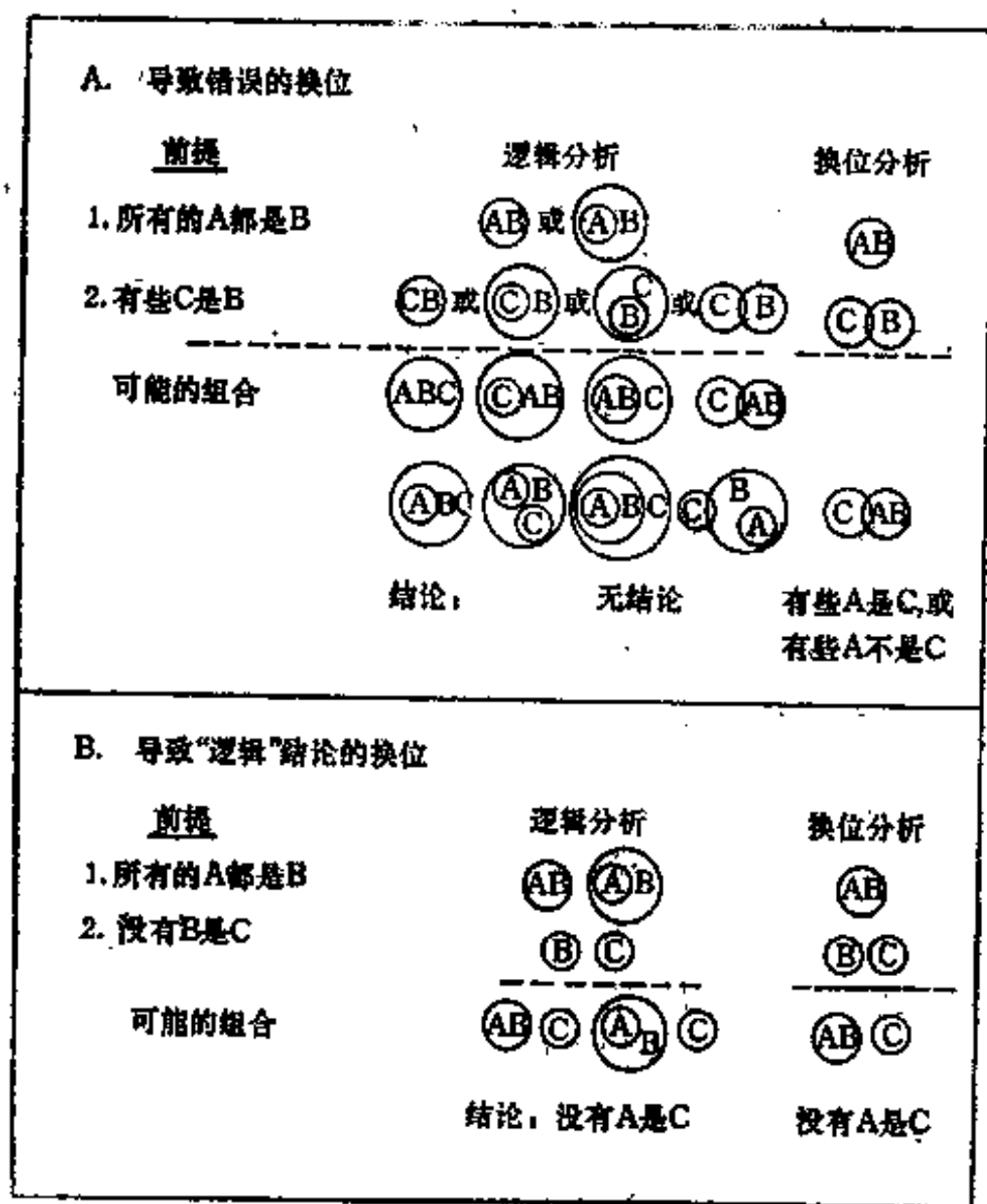


图 9-4 前提换位对推理的影响

会犯换位驱动的错误。

(四) 部分组合

前提错解并不是人们在加工三段论推理时所犯的唯一错误。前面已经指出过，对有些三段论推理，人们的正确反应乃是“无结论”。这种推理实际上是最困难的。例如，根据前提“有些A是B”和“有些B是C”，不管我们是考虑每个前提的所有可能的意义，还是仅考虑由前提换位假定的有限的意义（部分重叠），都得不到逻辑结论。假如A和B部分重叠，并且B和C部分重叠，那么，A和C之间的关系可以有三种情况：（1）完全重叠；（2）部分重叠；（3）完全不重叠。因此，即使就有限的前提意义来说，也不是必然能得出结论的。再进一步说，象“没有A是B”这种形式的前提，是无歧义的，没有人会错解的。但人们在处理这种前提时，也会犯错误。前提错解不能解释这种情况，这里必然还牵涉到另一种过程。

许多事实证明，人们在对三段论推理进行加工时，可能只生成某些（而不是全部）前提意义的组合。人的短时记忆的容量是有限度的，因而往往使人们只是选择几种前提意义对偶。这种不完全组合的策略，会导致各种各样的错误。图9-5给出的前提对偶可以用来说明这种情况。有人把图中的一对前提呈现给被试者，并要求被试者从以下几个陈述中选出一个在逻辑上正确的回答：

- （1）“所有的A都是C”，
- （2）“有些A是C”，
- （3）“没有A是C”，
- （4）“无结论”。

实验结果表明，选择（1）作为正确回答的被试者占5%，选择（2）的被试者占10%，选择（3）的被试者占35%，而选择（4）的被试者占50%。图9-5也指出了在不同情况下可能得出的结论。例如，若被试者刚好考虑组合1，那么，他就可能接受“所有的A都是C”这个结论。

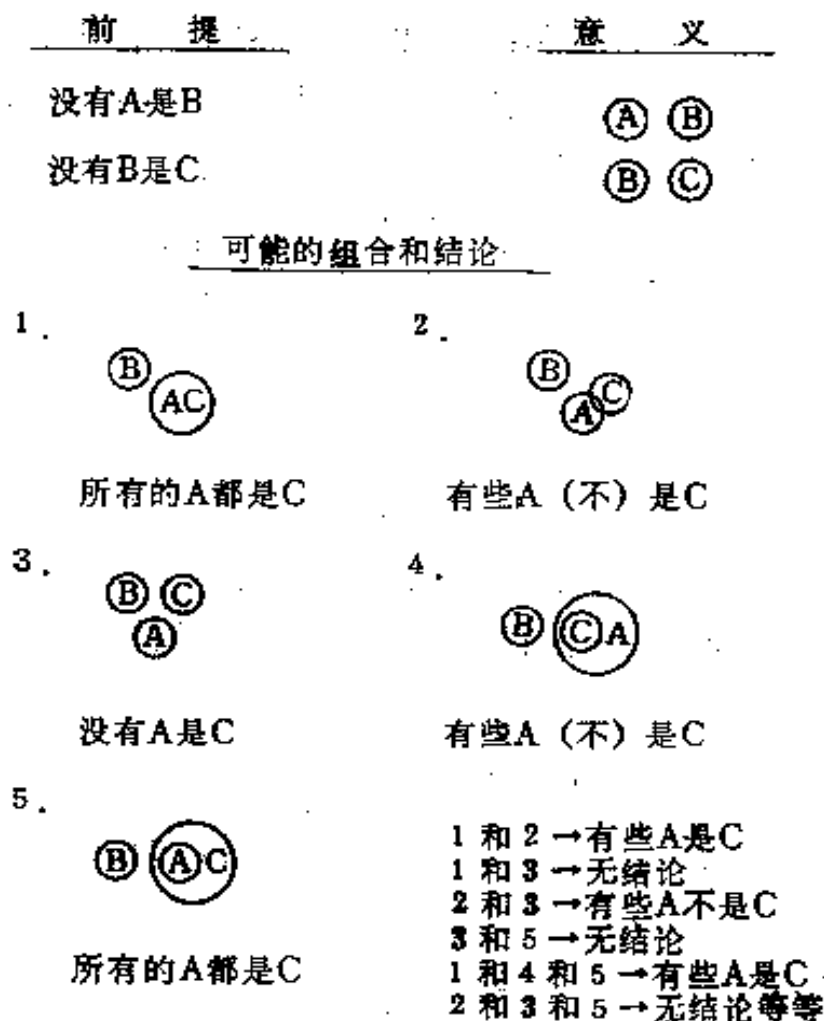


图 9-5 前提意义的部分组合

不完全组合策略既可以导致人们把只是“可能真”的结论看作为“必然真”，也可以导致人们把“可能真”的结论看作为“不可能的”而加以拒绝。例如，恰好考虑组合2和3的人可能会说，“所有的A都是C”这个结论是不合适的。总之，前提意义的部分组合，是人在处理前提对偶时所犯错误的主要方面。因此，为得出一个在逻辑上正确的结论所要

求的前提组合愈多，人们能够做到的可能性似乎就愈少，从而在推理过程中犯错误的机遇就愈大。

事实上，人们往往不能生成一切可能的前提意义组合，这是由于人的信息加工能力的限度所造成的。比如说，A和B部分重叠，B是C的一个子集（B完全包含在C内），在这里，A和C之间就不可能没有重叠。然而，有少数被试者把这样一种联系看作为“可能真”。另外，我们还可以看到，在某些情况下，人们只是构成一些错误的、不合逻辑的组合。也就是说，他们会把一种在逻辑上不能由前提构成的组合看作为是“可能真”或甚至是“必然真”。

（五）评估结论和信仰偏见效应

不管对前提是否作了错误的解释，也不管对前提意义已经进行了完全的或是不完全的、或者甚至是不合逻辑的组合，人们最终必须要选择一个结论，或者针对已经考虑过的前提对偶核对一个给定的结论。根据形式逻辑的要求，一个结论只有当它与所有可能的前提组合一致时才能成立。在评估结论时，至少有二种错误是可能发生的。这二种错误是：

（1）如果一个结论适合于人们考虑过的某些（但不是全部）前提组合，那么，他们可能会把这一结论看作为“必然真”。实际上，人们有时不能理解“逻辑上必然的”意义；

（2）人们有时犯一些与前提换位非常类似的错误；有人把这种倾向称之为“结论换位”。有些前提在逻辑上可能只允许一个方向的结论，如“有些A是C”。但是，人们或许会错误地得出结论，认为其反式“有些C是A”也是可以接受的。

一般认为，人们在进行三段论推理时，首先要解释每个前提；然后，构造或生成出部分或全部前提的意义组合；最后，针对所生成的组合来评估其结论。这是人们关于三段论推理加工过程的一种看法。此外，约翰逊-莱尔德（Johnson-Laird）等人还提出了另一种看法。他们认为，人们或许会根据一种偏爱的或似乎最有理的前提对偶，得出一个似乎有道理的（试验性的）结论。为了获得一个正确的结论，就必须对这个试验性的结论进行检验，看看是否有任何一个前提组合会使其失效。当然，这种检验有时可能是不合适的或不正确的。

在这里，信仰偏见起着重要的影响。我们知道，三段论推理可以用抽象的符号，也可以用具体的有意义的材料。当构成三段论推理的那些陈述包含有意义的、要求人的世界知识的材料时，常常会出现二种效应。第一种效应，就是前面已经提到的，有意义的陈述可以防止前提换位。例如，人们不会认为，“所有的猫都是动物”也意味着“所有的动物都是猫”。人们关于世界的知识防止了这种换位的错误。其第二种影响就是信仰偏见效应。信仰偏见效应是说，人们可能会根据他们的信仰而不是通过逻辑分析来评估有意义的逻辑推理的结论。换句话说，人们往往会不顾结论的逻辑状况而易于接受那些符合他们的信仰的结论。

埃文斯（Evans）等人曾对上述问题进行过研究。他们给被试者一些完全的三段论推理（两个前提和一个结论），要求被试者确定，其结论是否可以从其前提中得出来。对给定的前提来说，这种结论或者是十分可信的（例如：“有些优秀的滑冰者不是职业冰球运动员”），或者是相当不可信的（如“有些职业冰球运动员不是优秀的滑冰者”）；或者

是有效的结论，或者是无效的结论。实验结果表明，信仰偏见对无效的三段论推理有较大的影响。这就是说，人们不仅是拒绝一种不可信的但有效的结论，而且，似乎更容易接受一种可信的、但无效的结论。埃文斯等人认为，这种结果反映了一个事实，即当人们在加工三段论推理时，他们面临着逻辑和信仰之间的冲突。所谓有效的但不可信的结论是指，该结论应该是被接受的，但结论本身是难于相信的。所谓可信的但无效的结论是说，该结论尽管不是必然从前提中推异出来的，但也没有被前提所否定。这种结论是“可能真”，但不必然是真；并且，既然它是一个非常可信的陈述，所以，人们就倾向于接受它。实验中，当被试者对三段论推理进行作业时，还要求他们“大声地思维”；同时，实验者把被试者的话语记录下来，进行分析。这些材料表明，当被试者力图从前提到结论进行推理时，这种解释是最适用的。特别是对那些有效的但不可信的结论，常常使人们产生信仰与逻辑的冲突。另外，实验获得的材料表明，当人们把注意力主要集中于结论本身时，他们的反应将在更大的程度上取决于他们的信仰——可信的结论被接受，而不可信的结论被拒绝。总之，信仰偏见效应是存在的，尽管其形式是不同的。这就是说，人们或许会面对逻辑的和他们为世界所信奉的东西这二者之间的冲突。而这种冲突，将会使利用有意义材料的推理过程不同于用抽象符号的推理过程。

二、检验陈述的真假

三段论推理的任务是强调其论据的内部一致性。当要求人们确定一个结论在逻辑上是否能从前提中推导出来的时

候，首先要求他们假定其前提是真的。这里，还有一个既与此相联系但又不相同的问题，即人们会如何使用证据去确定一个特定的陈述是真还是假呢？这个问题很有意思。

（一）检验陈述真假的实例

有关这个问题的大多数研究，都采用“条件句”（如果……，那么……。）和选择任务的形式。在选择任务中，要求人们详细规定，为了发现命题是真还是假，哪些项需要加以检验。图 9-6 是一个选择任务的样本。

为了理解对这个问题的回答，先考察一下图 9-6 中给出的这条待检验规则的意思，以及它给各种可能性指定的真假值（truth value）。简单地说，这条规则是：“如果元音

检验命题的选择任务

“假定下面每一个方块代表放在桌子上的一张卡片，每张卡片的一面有一个字母，另一面有一个数字。因此，有数字 5 或 4 的卡片，另一面有某个字母；而有字母 A 或 K 的卡片，另一面就有一个数字。在这里，有一条规则：如果一张卡片上有一个元音字母，那么，它的另一面就有一个偶数。你的任务是指出，为了发现这条规则是真还是假，你需要把哪几张卡片翻转过来，你会选哪一张或哪几张卡片？”

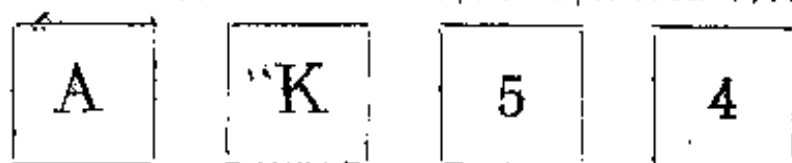


图 9-6 检验命题的选择任务

字母，则偶数”。这是条件语句的一般形式“如果 P，则 Q”的一个特定例子。表 9-3 根据形式逻辑给出了条件陈述的真值表。从表中我们可以看出，一个条件陈述只有一类实例是假的（即不存在的），这就是 $P\bar{Q}$ 或者说是一面有一个元音字母和另一面有一个奇数的卡片。在这里，P 或 $\bar{P}$ （非 P）与

表9-3 一个条件陈述的真值表

一般形式, “如果P, 则Q”

例子, “如果元音字母, 则偶数”

可能性		真 值
一般形式	例子	
PQ	元音字母, 偶数	真
P $\bar{Q}$	元音字母, 非偶数 (元音字母, 奇数)	假
$\bar{P}Q$	非元音字母, 偶数 (辅音字母, 偶数)	真
$\bar{P}\bar{Q}$	非元音字母, 非偶数 (辅音字母, 奇数)	真

Q或 $\bar{Q}$ (非Q) 的任何其他的组合, 都是和规则一致的, 即不违反规则的。

现在我们就这条规则从逻辑的观点来考察图 9-6 中的那 4 张卡片。有字母 A 的卡片是必须要检查的。如果卡片的另一面有一个偶数, 这是符合规则的。但是, 如果它的另一面是一个奇数, 那就反证 (即反驳) 了这条规则。而上面有字母 K 的卡片能提供的信息, 在这里可以说是无用的。该卡片的另一面究竟是奇数还是偶数, 这对证实或反证这条规则来说没有什么关系 (K 相当于 P, 而 PQ 和 $\bar{P}Q$ 在逻辑上都是真)。同样, 上面有数字 4 的卡片也是无关紧要的。“4” 相当于 Q, PQ 和 $\bar{P}Q$ 都不违反这条规则。相反, 上面有数字 5 的卡片却是至关重要的。如果它的另一面有一个元音字母, 那它就可以证明这条规则不能成立。因此, 在逻辑上, 对这个问题的正确的回答, 应该是上面有字母 A 和数字 5 的这两张卡片是必须翻过来加以检查的。

可是, 当给予被试者这种抽象材料的选择任务时, 大概

只有10%以下的人能够作出正确的选择。大约有50%的被试者说，他们要翻查上面有字母A的卡片和数字4的卡片。35%的被试者仅选择上面有字母A的卡片。其余的被试者则有其它各种不同的选择。从这里我们可以看到，很少有人提出来要翻查上面有数字5的那张卡片，而事实上，这张卡片是必须要加以检查的，这种现象是很奇怪的。

这种情况，或许与人们对条件句的理解有关。被试者可能把条件句作了换位，认为“如果P，则Q”也意味着“如果Q，则P”。这与三段论推理中的前提换位（“所有的A都是B”也意味着“所有的B都是A”）性质上是一样的。另外，人们或许没有正确理解条件规则和任务的选择，而只是把规则中提到的术语指定为他们的“选择”的目标。在使用抽象材料对被试者进行实验时，这种情况是常见的。但是，当条件陈述具有意义内容时，陈述内容的性质，会对人们如何解释条件规则产生一定的影响。可以说，影响被试者作业的因素是多方面的。

假定说，我们正在验证一个“如果P，则Q”的陈述，从逻辑上说，没有必要去检查“非P”的实例。并且， PQ 和 $P\bar{Q}$ 实质上都是无关的。在这里，最关键的实例是 PQ 和 $P\bar{Q}$ 。 PQ （元音字母与偶数）确证这条规则，而 $P\bar{Q}$ （元音字母与奇数）将否定这条规则。从逻辑的观点来看，通过观察有限数量的 PQ 实例来证明“如果P，则Q”的规则为真，这是不可能的。但是，只要一个 $P\bar{Q}$ 实例，就能证明这条规则是假的。例如，假定说你检验这样一个陈述：

“如果雌鸟的羽毛上有许多斑点，那么它的蛋壳上也有斑点。”

从逻辑的观点来看，不管你已经看到过多少羽毛上有斑点的

鸟下出来的蛋也有斑点,这并不能证明这条规则必然是真的。因为,在某个你至今尚未观察过的地方,仍然有可能存在着一种羽毛上有斑点的鸟,但它下出来的蛋却没有斑点。然而,如果你恰好看到一只羽毛上有斑点的鸟,它的蛋壳上并没有斑点,这就可以证明这条规则是假的。因此,从逻辑上说,那种能证明一个陈述为假的证据是十分重要的。联系这一点,我们再来考察一下上述选择任务的作业。

上面给出的规则是“如果元音字母,则偶数”,那么,一面有元音字母和另一面有偶数的卡片(PQ)属于确证的实例;一面有元音字母和另一面有奇数的卡片则属于反证的实例,它将反驳这条规则。强调反证的重要性,就意味着,适当的行为应是去寻找那种可以反证规则的实例。换句话说,核对元音字母(P)的卡片,是因为它可能是 $P\bar{Q}$,而核对奇数($\bar{Q}$)的卡片,也是出于同一缘由,即它可能是一个 $P\bar{Q}$ 的实例。但是,如前面已指出的,在选择任务中,很少有人去核对 $\bar{Q}$ 卡片,而人们往往去核查P卡片和Q卡片。怎么解释人们的这种行为呢?大致可以说,人们倾向于寻找那种能确证规则的证据,而没有看到那种能反证规则的证据的重要性。实际上,人们从P卡片或Q卡片只能获得一个确证规则的实例,但 $\bar{Q}$ 卡片却能够提供一个唯一的反证该规则的实例。

(二) 具体性的效应

如果用具体的而不是抽象的材料,人们的推理作业是否会好一些呢?有一种观点认为,由于实验中使用了抽象的材料,从而妨碍了人们的推理活动。如果这种观点是正确的,那么,使用现实的、具体的材料来进行实验,被试者将会有

较好的作业成绩，能够表现出人们的推理能力的“真实水平”。沃森（Wason）等人曾利用一种具体的材料对英国被试者作了实验。实验中使用的每张卡片都涉及一次旅行，卡片的一面是一个城市，另一面是交通工具。要求被试者检验的规则是怎么旅行，例如，“每次去伦敦，我都乘火车”。呈现给被试者的4张卡片上分别有“曼彻斯特”、“伦敦”、“火车”和“汽车”。结果发现，62%的被试者都作了正确的选择。这就是说，他们为了验证上述这条规则是否成立，选择了上面有“伦敦”及“汽车”的这两张卡片。这种选择是唯一正确的。可见，使用具体的材料来进行实验，被试者的作业成绩就较好。

但是，后来又有一些研究表明，这种看法是有问题的。一些研究者发现，不管实验中使用的是抽象材料或现实的、具体的材料，被试者的作业水平却很少或几乎没有什么差异。

于是，格里格斯（Griggs）等人提出了“记忆提示假设”以解释这种实验结果的矛盾现象。这个假设认为，在下述情况下，人们将易于作出正确的选择：

- （1）被试者对被检验的规则是熟悉的；
- （2）他们对规则的反例具有一定的经验。

这个假设的含义是说，人们所以对那种使用现实的具体材料的选择任务会有较好的作业成绩，这并不是说因为他比较全面地分析了条件陈述和每张卡片的可能的含义，而是他们已经知道要检验什么，是记忆中的知识在起作用。有些实验所用的材料虽然是具体的，但被试者对它们并无足够的知识经验，因此，他们的作业成绩并没有明显的改善。美国大学生对上述关于旅行和交通工具的验证任务，其作业成绩

较差，但对“饮酒年龄规则”的验证任务却有较好的作业水平。后者在他们国家中是一条法律，实际上所有的被试者对此都是熟悉的，知道在什么情况下将违反这条规则。

由此可见，简单地说人们对现实的具体材料能较好地进行推理，这不一定是很正确的。

(三) 验证的偏见

我们已经看到，人们在前述的验证任务中，表现出一种较为普遍的倾向性。他们不是去检验Q卡片，而是去检验Q卡片。事实上，前者是必须要加以检验的，而后者却是不必要加以检验的。所以出现这种倾向性，是因为人们希望寻找那种确证的证据去证实这条规则，而没有看到反证的证据在验证规则时的重要价值。人们在验证任务中表现出来的这种倾向性，一般被称之为“验证偏见”。

比如说，在三段论推理作业过程中，人们时常未能去考虑所有的前提意义的组合，其中一个可能的原因是“验证偏见”的影响。在产生了一个与前提一致的结论时，人们就可能把加工过程停止下来，接受这个可能的但在逻辑上不是必然的结论。人们似乎没有想到应对一个给定的结论去寻找可能的反例。在概念学习任务中，当人们检验自己的假设时，也时常出现这种情况。人们倾向于认为他们目前的假设似乎是正确的，在少量的验证以后，就想把它作为正确的结论加以接受。即使还有其他假设也是站得住的，但往往被忽视了。

这里再举一个说明“验证偏见”的实验例子。实验中给被试者一个任务，要求他发现一条适用于由三个数字构成数字系列的规则。首先，由实验者给被试者呈现由2、4、6

三个数字构成的数字系列，并告诉他这是符合待发现的规则的。然后，由被试者向实验者提出另外一些数字系列，实验者则立即指出被试者提出的数字系列是否符合待发现的规则，给被试者适当的反馈信息。事实上，这条待发现的规则是很一般的，即“上升序列的任意3个数”。但是，这里却有许多比较特殊的规则可以生成正例。例如，其增长值为2的任何三个数字的系列；或者，第二个数是第一个数的二倍，而第三个数是第二个数的1.5倍，等等。

被试者接受这个实验性任务后，他们往往会构成一条比较特殊的规则来作为假设，如“增长值为2的数字系列”。接着，他们就提出一些数字系列，意在证实自己的假设。这些数字系列都是符合他们所假设的规则，例如4、6、8、10、12、14、24、26、28、7、9、11等数字系列。它们不但符合被试者所假设的“增长值为2的数字系列”这一较特殊的规则，也都符合实验者头脑中确定的、有待于被试者发现的较一般的规则。因此，实验者告诉被试者，每个数字系列是符合待发现的规则的。结果，造成被试者越来越相信，他的局部假设是正确的。于是，被试者就宣布了自己假设的规则，并认为这就是待发现的规则。

实验结果表明，大约有80%的被试者宣布的规则是不正确的。对这样一个显然很简单的任务，为什么被试者的作业水平这么差呢？其中一个较为重要的原因是，他们在处理自己假设的规则时，单纯地采用了一种“证实策略”。这种策略在于搜索那些能直接证实自己所假设的规则的证据，而不去搜索那些会把自己的假设证明为假的证据。这种做法显然是不正确的。事实上，被试者需要提出一个与其局部假设不一致的数字系列，检验一下自己假设的规则是否是错误的。

比如说，提出 5、6、7 数字系列。这与“增长值为 2 的数字系列”是不一致的。如果发现这也符合待发现的规则，那么，就可以知道，“增长值为 2 的数字系列”并不是待发现的那条规则。也就是说，被试者通过搜索那种能够证明当前假设为假的证据，可以找到待发现的一般规则；然而，靠搜索那种能确证当前假设的证据，却不能做到这一点。但是，人们却往往倾向于去搜索那种能证实自己的假设的证据，说明这里存在着一种“验证偏见”。

被试者的这种行为模式也不能说是一定不合适的。事实上，我们不但需要知道什么是错误的，而且需要知道什么可能是正确的。从这一观点出发，被试者搜索那种能够证实所提出的规则的证据是必要的，它能使该规则成为值得考虑的结论。而且，通过寻找那种能够证实规则的证据，也有可能发现某种能证明该规则为假的证据。当然，“验证偏见”是有其危险性的。

三、关于不确定性的判断

三段论推理和有关陈述的验证，都有确定的标准。一个结论，或者根据逻辑可以从前提中得出来，或者根据逻辑不能从前提中推导出来；一个陈述必须或者被判断为真或者被判断为假。但是，我们平时处理的许多信息是不确定的，并不是绝对地正确或绝对地错误的，而是或多或少是正确的。概率论和统计学是用来处理不确定信息的工具。也可以说，对一个涉及不确定性的问题的正确回答，是借助数学公式来确定的。在这里，认知心理学家关注的是下面二个问题：

(1) 人类的判断在什么程度上符合于从数学理论和公

式推导出来的回答：

(2) 人类是怎样作出关于不确定性的判断的。

早期的一些研究曾认为，人是非常好的“直觉的统计学家”，但是，近来的许多研究表明，人类的判断经常明显地偏离形式推导的结果。

(一) 概率和概率判断

概率是从 0 到 1 的一些数字。当一个事件的概率是 0 时 ($P=0.00$)，我们可以肯定地说，这个事件是决不会出现的。抛一次硬币，要得到正反两面都朝上的结果，其概率是 0。也就是说，这种情况是决不会出现的。对概率为 1 的事件 ($P=1.00$)，我们可以肯定地说，该事件必然会出现的。因为每个人都会死，所以死这个事件的概率为 1。在这两端之间的各种概率，则意味着事件多少是要出现的。也就是说，可以预料，该事件在一个较长的时间内，它有一定的百分比的时间将出现。把一个硬币往上抛 1000 次，预料它将有 50% 的次数是正面朝上。

客观的概率是以对事件根源的详细观察和分析为依据的。通常，人们可以通过分析产生这些事件的手段的自然特性来作出概率分配。比如说，骰子是 6 面的，而且，骰子落下来时各个面朝上的机会是均等的。因此，可以说，骰子落下来时任何一个数朝上的概率为 $1/6$ 或 .17。假定说，一副扑克牌共 52 张，其中 13 张是黑桃，那么，从这付扑克牌中随机挑出一张黑桃的概率是 $13/52$ 或 .25。这里所用的“随机”这个术语，意指一副纸牌中的每一张牌在挑选过程中都有同等的被挑出来的机会。有时候，人们对作为事件出现之物理基础的机制不大了解，所以通过观察事件在过去出现的相

对频率来确定它的概率。如果有人声称，贵阳在某个特定的日子下雨的概率是0.40，那么，这就是根据这个城市过去下雨的记录来确定的。这些是基于一定情境的物理特性的客观概率。当然，我们个人自己对一定事件有多大可能出现的估计，不一定与这种客观的概率相一致。

人们对实际的客观概率作出的符合他们对事件可能性的主观知觉的判断，称之为主观概率。换句话说，人对事件概率的判断就是主观概率。如果给某人一对骰子，他想会有多大可能掷出一个7的数来？骰子这种东西是没有记忆的。它每一次的滚动完全不依赖于上一次的滚动。在最近的连续100次的滚动中，从未出现过7这个数是完全可能的。那么下一次会出现什么呢？也许有人会打赌说，下次将出现7这个数，因为在最近100次的投掷中从未出现过这个数，这一次它一定会出现。这种打赌来自对概率的主观判断，这就是主观概率。

人的概率判断或主观概率与客观概率之间的关系，是使一些研究者感兴趣的问题。

(二) 概率与频率

用形式术语来说，概率是一种成比例的频率。随机从一副扑克牌中挑出一张黑桃的概率，等于扑克牌中黑桃的数量除以牌的总数，即 $13/52$ ，等于.25。如果教室里有24个女生和18个男生，那么，随机挑出一个女生的概率是：

$$P(\text{女生}) = \frac{24}{24+18} = \frac{24}{42} = 0.57$$

但是，有证据表明，在某些情况下，人们作概率判断时，不是根据概率，而是根据频率。例如，埃斯赫斯（W.

K. Estes) 曾做过一个实验, 他让被试者接受一系列想象的民意测验结果的信息。每一民意测验均涉及二种候选者的比较(如二种产品或二个经理候选人等)。其中一个获得民意测验的胜利。实验者一次一个地给被试者呈现各种不同比较内容的民意测验结果。然后, 要求被试者对随后的民意测验结果作出预测。预测的对象是那些彼此还没有被比较过的、但具有不同胜败历史的候选者。例如, 要求被试者预测产品A和E在民意测验中哪一种产品将获胜。产品A曾参与过6次民意测验, 获胜5次; 而产品E曾在18次民意测验中获胜了9次。但被试者从未看见过一次比较A和E的民意测验。

如果被试者根据产品获胜的概率来进行比较, 那么, 他们就会预测产品A将会获胜, 因为产品A获胜的比率是0.83($5/6=.83$); 而产品E获胜的比率只有0.50($9/18=.50$)。但是, 实验结果却表明, 大多数被试者预测产品E将获胜。这说明, 人们通常是根据候选者获胜的频率而不是根据候选者获胜的概率来作预测的。尽管实验者可以引导被试者去注意候选者的失败历史, 但结果还是那样。

(三) 代表性

概率论有一些把概率结合起来的规则。比如说, 有两个事件A和B, 它们的概率分别为 $p(A)$ 和 $p(B)$ 。A和B事件一起出现的概率 $p(A和B)$, 不可能大于两者之中较小的那个概率。这条规则称为“合取”规则。例如, 在某个地区居民中, 少数民族占30%, 汉族占70%。因此, 作为少数民族的概率是0.30。该地区居民中, 已婚者占45%, 其概率为0.45。这样, 在该地区碰到一个既是少数民族又是结婚的居民的概率, 不可能大于0.30。但是, 在某些情况下, 人们经

常要犯“合取错误”。也就是说，人们往往认为，二个事件合在一起时的概率，比单个事件的概率要大。

卡尼曼 (Kahneman) 等人做过一个实验，给被试者看下面一段话：

“林达今年31岁，单身、坦率并且非常聪明。她学的专业是哲学。当她是一位大学生时，她非常关心歧视和社会正义这些问题，也参加过反核示威运动。”

然后，问被试者：林达是位银行职员，或者，既是一位银行职员又是一位积极的男女平等主义者，究竟哪种可能性更大一些？结果发现，超过80%的被试者选择后者（合取），认为后者的可能性更大。事实上，这种回答不可能是正确的。为什么人们会犯这种错误呢？很显然，这是因为上面那段描述林达个性的话，使人产生一个印象，她似乎更象一个男女平等主义的银行职员，而不仅仅是一个银行职员。

这个例子说明，被试者在作概率判断时使用了“代表性”策略，即藉助于证据和结果具有类似特征的程度来判断“概率”。

在日常生活中，我们经常使用“代表性”策略来判断事件的概率。但是，这种策略往往会使我们产生错觉，从而使我们在判断事件概率时发生错误。比如说，我们用计算器加一系列数，得出的总数恰好是12345。我们会觉得，这个数似乎不大可能，怎么会那样凑巧得出如此一个整齐的数值，感到它缺乏随机代表性。所以，我们又把一系列数重加一遍，结果总数还是12345。事实上，这个特殊的数与另外一个数（如21357）具有同等的出现概率。

概率论指出，一个样本愈大，它就愈可能象它从中抽出

来的总体的性质。相反，人们却倾向于期待甚至非常小的样本是总体的非常好的代表。这样，就导致他们错误地过分解释小样本的信息。关于小样本代表性的这种错误的信仰，在日常生活中或许常常是无害的。但是，它们是“赌徒错误”的基础，即相信在非常短的时间里“事情将会自我纠正”。举一个很简单的例子：如果一个铜币向上抛2次，并且落地后每次都是反面向上，那么，赌徒的错误就在于相信，当第三次把铜币往上抛时，铜币落地后正面向上的机会将超过50%。

我们再看另一个例子。假定说，一个班有18个男生，12个女生；必须从中选出三个学生，让他们去参加社会服务或得到某种他们所希望的东西。教师告诉这个班说，现在已选出了三个学生，他们都是随机地选出来的，并且宣布他们是2个女生和1个男生。然而，整个班里有60%是男生，选出来的样本却是女生多。于是，相信小样本的代表性的同学可能会认为，这几个学生不是随机地选出来的，是教师使用了某种有偏见的方法，并且，教师对学生说了慌。事实上，从18个男生和12个女生的“总体”中随机抽取出来的3个人的样本，在1/3的次数里会出现2个女生和一个男生。所以，上述老师宣布的结果，几乎是不奇怪的。但是，要对那些错误地相信小样本代表性的人解释这件事，有时或许是困难的。

（四）可利用性

人们除了使用“代表性”策略外，也常常依据“可利用性”来进行概率判断。“可利用性”即有关信息从记忆中是否比较容易提取出来，对人的概率判断有重要的影响。

梯厄斯基（A. Tversky）等人曾问被试者，第一个字母

是K的英文单词多，还是第三个字母是K的英文单词多？大约有2/3的被试者回答说，字母K出现在第一个位置上的英语单词，比出现在第三个位置上的英语单词多。事实上，这种判断是错误的。情况恰好相反，字母K出现在第三个位置上的英语单词比较多。这种判断上的错误主要是由于人们对词知识的提取难易程度不同而造成的。我们知道，根据第一个字母来回忆一个词比较容易，而根据第三个字母来回忆一个词就比较困难。这就是说，它们的“可利用性”是不同的。当梯厄斯基问被试者这个问题时，他们就使劲地回忆以字母K为起头的英语单词以及在第三个位置上出现字母K的英语单词。结果，他们能比较容易地回忆起以字母K起头的英语单词。所以，他们判断，在英语单词里，字母K出现在第一个位置上的可能性比出现在第三个位置上的可能性大。这个例子说明，人的记忆组织的性质，影响记忆材料的可利用性，从而也影响人的概率判断。原则上可以说，许多其他的能影响人对信息记忆力的因素，也同样会影响信息的可利用性，从而将会影响人的概率判断。

在另外一项实验中，要求被试者判断在美国各种原因的死亡率。实验结果同样表明，人的判断有偏向性，主要受这样一些因素的影响：特殊实例的可回忆性，想象一个人死于某种特定原因的容易程度等。比如说，让被试者判断由龙卷风造成的死亡率大，还是由气喘病造成的死亡率大。大多数被试者认为，由龙卷风造成的死亡率大。事实上，气喘病造成的死亡是龙卷风的20倍。

当人们具有那些可利用的、与进行概率判断有关的信息时，就很难摆脱这些信息对他们的判断发生的影响。例如，一旦人们知道某一特定结果已经出现，他们就倾向于相信自己

也预计到这种结果的发生。有人曾作过一个研究，给被试者有关一个历史事件（二个军队之间的战斗）的背景知识，并要求他们评价各种结果的可能性。例如，哪方将取胜，或者，军事上的僵持是否可能会取得政治解决。那些仅仅获得背景信息的被试者估计，根据背景信息，A方只有三分之一的机会能获胜；而那些既获得同一背景信息，又同时被告知是A方获胜了的被试者则判断，根据背景信息，可以预测A方获胜的机会大于二分之一。即使实验者告诉这些被试者，要撇开有关结果的信息，独立地根据实验所提供的背景信息来进行判断，上述效应却依然存在。

从上面的叙述中可以看出，“代表性”和“可利用性”对人们判断事件的概率有明显的影响。这种影响，常常使人的关于事件概率的判断偏离事件本身的客观概率。同时，有人认为，人的概率判断或主观概率，具有下述三个倾向性的特点：

（1）人们倾向于高估低概率事件的出现，而低估高概率事件的出现。

（2）人们倾向于预测暂时未出现的一个事件很可能在最近的将来就要出现。

（3）人们倾向于高估对他们有利的事件的真实概率，低估对他们不利的事件的真实概率。

第十章 问题解决和决策

当某个人面临一项任务而又不知道如何去完成它时，他就面对一个问题，一旦找到了完成这项任务的方法或策略，这意味着他找到了该问题的解。问题解决就是导致某个问题获得解决的思维活动。决策也是人类重要的思维活动，不过，决策的特点是从一集可供选择的对象中作出抉择。

一、问题和问题行为图

(一) 两类问题和四个步骤

人面临的问题，大致可以分成两大类。一是规定良好的问题；二是规定不良的问题。所谓规定良好的问题，是指该问题具有清楚地说明了的起点（或起始状态）和终点（或目标状态）。例如：

“从天安门怎么乘公共汽车到颐和园？”

“摄氏20°等于华氏多少度？”

这些问题的出发点和目标是明确的。解决这类问题，是要求问题解决者在正确的方向上采取适当的手段去获得问题的解。

另外，在生活中也经常碰到许多规定不良的问题。它们不象有明确的起点或终点，甚至两者均是含糊的；从而没有一个客观的尺度去衡量是否已成功地达到了问题的目标状态。例如：

“怎样把物价上涨率限制在群众可接受的程度？”

“你能修好我的破汽车吗？”

当然，不管是规定良好的问题还是规定不良的问题，问题解决的过程基本上都可以归纳为四个步骤。

1. **识别和理解问题** 它要求认出问题的起始状态和目标状态。在解决含糊规定的问题时，这一步是较为困难的。

2. **设计或规定一种解决问题的策略、方法或计划** 人们在解决问题时，使用各种不同的策略，包括最简单的“尝试和错误”法到复杂的类比推理。规划一种有效的策略就要求利用各种有关的知识，包括问题类型的知识、某特定策略取得成功的可能性的知识等。例如，在下象棋时，棋手究竟采取什么策略，这要看他有多少象棋知识和某种策略取得成功的历史及对手的性格和思维习惯等。

3. **执行策略** 如果在解决一个简单而且规定清楚的问题时，这一步往往是容易的。但是，在解决某些比较复杂的问题时，一种策略的执行也决不是一种简单的过程。

4. **评价问题解决的进展** 这要求就是否出现了向目标状态的进展，问题是否已获得解决以及当前的策略是继续执行还是把它放弃等问题作出评定。例如，若问题解决者原先采取一种“做任何一个行动都必须使问题更接近目标状态”的策略，但这种策略很可能行不通；问题解决者不得不放弃这种策略，改而使用一种“退一步进二步”的策略。

当然，在解决问题的过程中，上述四个步骤决不是绝然分割的，而是彼此相互联系和相互影响的。

心理学家对问题解决的研究，集中在那些清楚地规定的问题上。通过对这类问题的研究，心理学家企图了解人们关于问题的内部表征；解决问题时应用什么策略，遵循何种法则；以及，随着问题一步步向解决的方向前进，人们是怎么评价

其进展的。

为了开始下面的讨论，这里举一个具体的例子。有人给一位大学生提出了这样一道密码算术题：

$$\begin{array}{r} \text{DONALD} \quad D=5 \\ + \text{GERALD} \\ \hline \text{ROBERT} \end{array}$$

这个算式里有十个字母，每个字母代表一个不同的数字。给被试者提出的任务，是要发现每个字母被指派的数字，使得以相应数字代替这些字母时，其等式的计算是正确的。这里已给出的一个条件是 $D=5$ 。

(二) 原始素材

我们的目的是要观察人在解决问题时的行为。但是，人在解决这类问题时，他是默默地进行其内部的心理操作的。对这种心理操作，我们不能直接观察到。所以，实验者要求被试者“大声地”思维，也就是说，被试者一边试图解决问题，一边把他们正在做的那些心理操作出声地说出来。实验者把被试者解题时的口语记录下来，结果就是被试者的言语化的思维过程的素材（原始记录）。它们提供了有关问题解决时所包含的思维过程的有用的最初信息。当然，在这里，其主要的目的是要记录、分析和发现人们在解决问题时做了什么样的心理操作，至于是否能把问题真的解决了，那是无关紧要的。（你也可以用三十分钟的时间，试解上述密码算术题。如果你有录音机的话，也采用“大声地”思维的方式，把你解题时的言语记录并整理出来。）

上述那位大学生在解这道密码算术题的20多分钟的时间里，其口语记录约有2200字。下面我们列举和分析其中的一

小部分：

“每个字母有一个并且仅仅有一个数值？（实验者回答：‘一个数值’）。有10个不同的字母，并且其中每一个都有一个数值。所以，我可以考虑两个D——每个D都是5，因此，T是零。我把问题先写下来，5，5等于零。现在，我还有没有另外的T？没有。但我还有另一个D。就是说，我在那边有个5。现在，我分别有两个A和两个L，而R则有三个。两个L等于一个R。当然，我现在可以进1，这意味着R必然是一个奇数。因为两个L——两个任何数字相加必定得出偶数，而我将有一个奇数。所以，R可能是1、3、而不是5、7或9。（有一个长的停顿）现在看G，既然R必然是奇数，而D是5，G必定是一个偶数。我看这道题的左边，那里说D+G。哦，可能还要加另一个数，如果我必须从E+O进1的话。……可能，做这道题的最好的方式是尝试各种可能的解法。我不能确信这是不是最容易的方式。”

从上述的原始记录中，我们可以看到，这位问题解决者没有以一直向前的、直接的方式来行动。他通过一些假设来积累信息，看看这些假设将把他引向何处。他常碰到死胡同，然后倒回来，再作另外的尝试。他一开始很有把握地发现 $T=0$ （即零）。然后，他想发现他所知道的 $T=0$ 和 $D=5$ 在题内的其他地方能否用得上。他先寻找T，结果落空了，他又回过头来寻找D。结果他发现还有一个D。但是，他马上又在问题中发现了看来有希望的另一一点，这就是两个L。他推断出R必定是个奇数，而且把这个推论重做一遍，好象在核查推理。当然，这时他将他的推理进行得更前进了一步，明确地列举出一些可能的数字，即R可能是1或3，而不是5、7或9。

可是，在长时间停顿以后，他放弃了这条路子。看来，他没有办法从上述列举的数字中为R选取一个特定的数值。所以，他再次回到R是个奇数这一基本思想上来，看看它能否提供有关G信息。结果，他从R是奇数，D是5这一事实出发，推断出G必定是一个偶数。

上面是有关问题解决者行为的最初几分钟的原始记录及其分析。从中我们可以看到这位被试者的问题解决行为的一般模式。被试者知道他正试图达到的总的目的，即问题的目标状态。为了达到问题的目标状态，他首先把它分成许多小的步骤或子目标。然后，他连续地使用各种简单的策略，想寻找一些新的信息。有些策略行得通，问题解决就前进了一步；有些行不通，被试者就退回来，并改用另外的探索方式。

(三) 问题行为图

口语的原始记录用起来是不方便的。为了详细地考察问题解决的过程，人们制作了一种关于问题解决时进行的操作序列图，称之为“问题行为图”。它用来描述被试者在整个问题解决过程中的行为及其进程。

被试在任一时刻所知道的关于问题的信息，称作他的“认识状态。”每当他应用某个操作于问题时，就会积累问题的新信息，从而也就改变了他的认识状态。问题行为图的目的，是要用图来勾画出被试者这种认识状态的进程。可以用一个方块来表示一个认识状态，以箭头来表示被试者从一个“认识状态转入另一个认识状态时的操作，如图10-1所示。这样，言语的原始记录就可以转译成许多方块和箭头，从而可以更清楚地看出被试者在经历连续的认识状态的转换时所

走的路线。

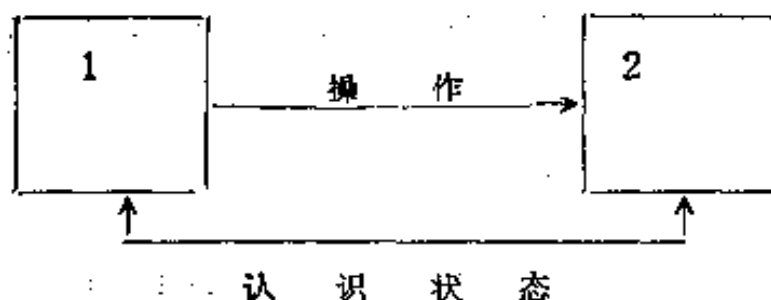


图10-1 认识状态转换的表示方式

被试者在解上述密码算术题时，最初的陈述不过是在证实他理解作业的规则。当他考虑到D是5，因此T是零时，才出现第一个实际的推理。这时，他显然是在加工D+D=T这一列的信息。我们暂将这个操作称作“加工第一列”

（将该题的六列自右至左地编号）。这个操作使被试者从起始的认识状态（D=5）进到一个新的状态。在这第二个认识状态时，他已知道T=0，然后，他就尝试去取一个新列，或者是有T的一列，或者是有D的一列。第一个操作（取T列）没有成功，但寻找有D的新列的操作成功了。因此，问题的行为图又前进了一步。到此时，问题的行为图表示了三个认识状态的序列，如图10-2所示。

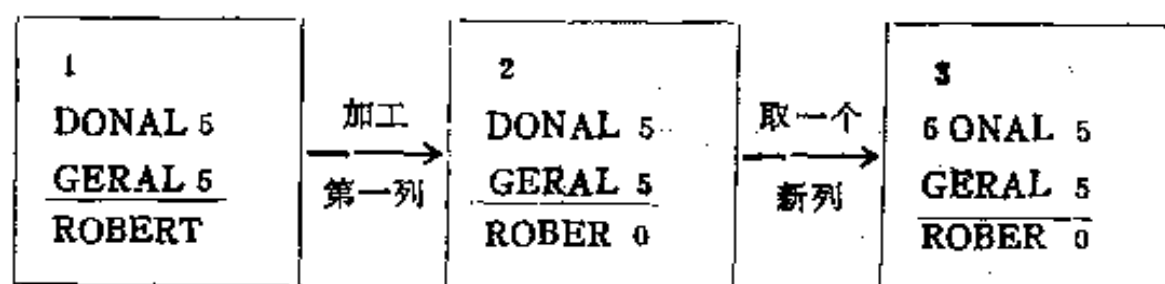


图10-2 表示三个连续认识状态的问题行为图

这时被试者又再次试图取一新列，先试图取第三列，但他马上发现，R有三个，使他明白是值得去加工第二列的。因

此，他取了第二列来加工，结果，推动他从第四个状态进到第五个认识状态。在这里，他作出了R是奇数的结论，见图10-3。

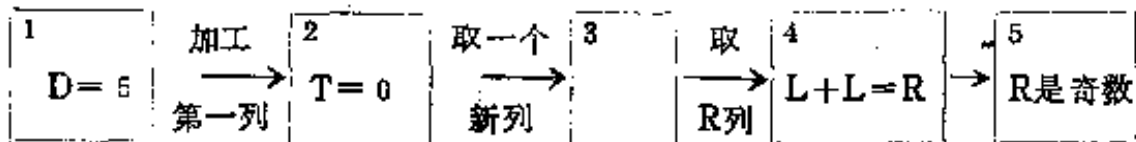


图10-3 表示五个连续认识状态的简写行为图

可是，在第五个认识状态时，他又决定提出确定数字的实际可能性。为了达到这一点，他又退回到第四个状态，把原来的推论检验一遍，并在此基础上提出了R可能是1或3，而不是5，7，或9的假设。用图来说明倒退过程时，我们把下一个状态（即第六个状态）画成发自第四个状态（如图10-4）。在图10-4中，从第四个状态画一个垂直箭头指向第六个状态，以此来表示这一过渡。第六个状态与第四个状态是完全相同的。在第七个状态他又重新提出R是奇数，但是，在第八个状态，他井井有条地推敲了全部奇数。

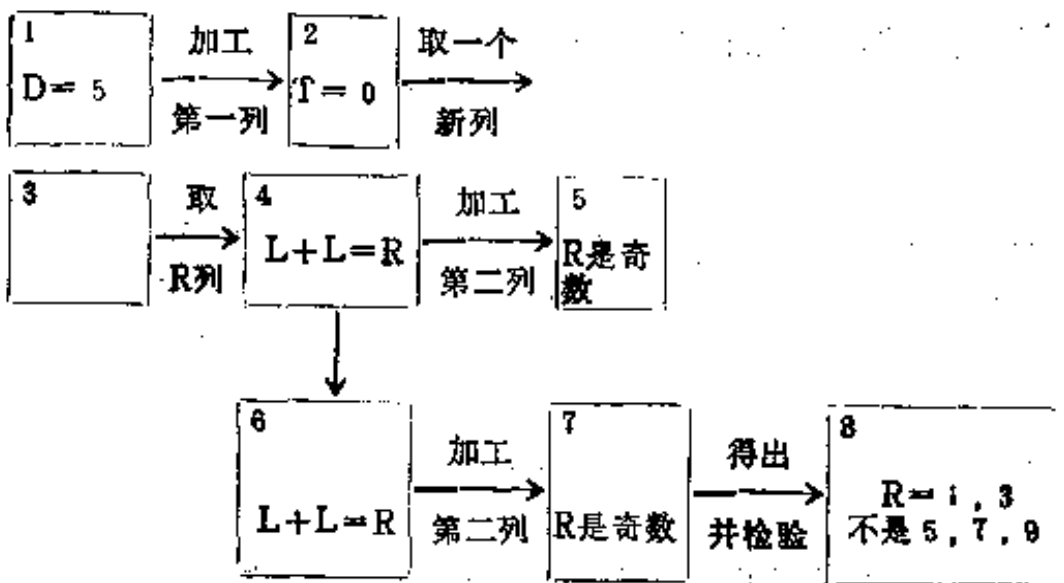


图10-4 表示倒退过程的问题行为图

当被试者在为R提出候选的数字时，他还是比较仔细的。他排除了已经用了的数值。例如，他否定了 $R=5$ 的可能性，因为D已经是5了。从原始记录中，我们没有取到关于如何确定R的可能数值的证据。相反，被试者再次倒退回来，带着 $D=5$ 和R是奇数的结论转向第六列，并作出G必然是偶数的结论。这导致被试者进入第十个认识状态。当然，这个推理是不正确的，但是，它乃是被试者的一个现实的认识状态。从原始记录中我们看到，G不一定是偶数这一事实，很快就被他认识到了，因为他看到 $E+O$ 有可能进一。这表示被试者加工第六列的另一个开端，它终止于第十二个状态（考虑到进位），然后又再次倒退回来，并忘记给G的数值。前面引用的原始记录的整个片断，可以用图10-5来表示。

问题行为图可以用来仔细地分析问题解决的过程，并把这个过程分解为一系列小步骤。它用图解来表示问题解决作业中出现的成功和失败的混合情况。当然，对某个特定的问题来说，被试者所应用的特殊规则，依赖于该问题的特殊性质。但是，他的问题解决总的结构是非常相似的。比如说，被试者把总的问题改组成一系列较简单的小目标或小问题；通过应用某种操作，从一个认识状态前进到下一个认识状态；解决问题的进程是不稳定的，其中包含着不断的尝试和错误，考验着各种操作的有效性；当一系列操作引到死胡同时，就倒退回来，在某一点上重新开始等等。所以，问题行为图可以应用于许多不同的问题情境。

（四）算术密码题的解法

前面谈了算术密码题 $DONALD+GERALD=ROBERT$ 的原始素材和问题行为图。许多人在解这个密码算术题时，

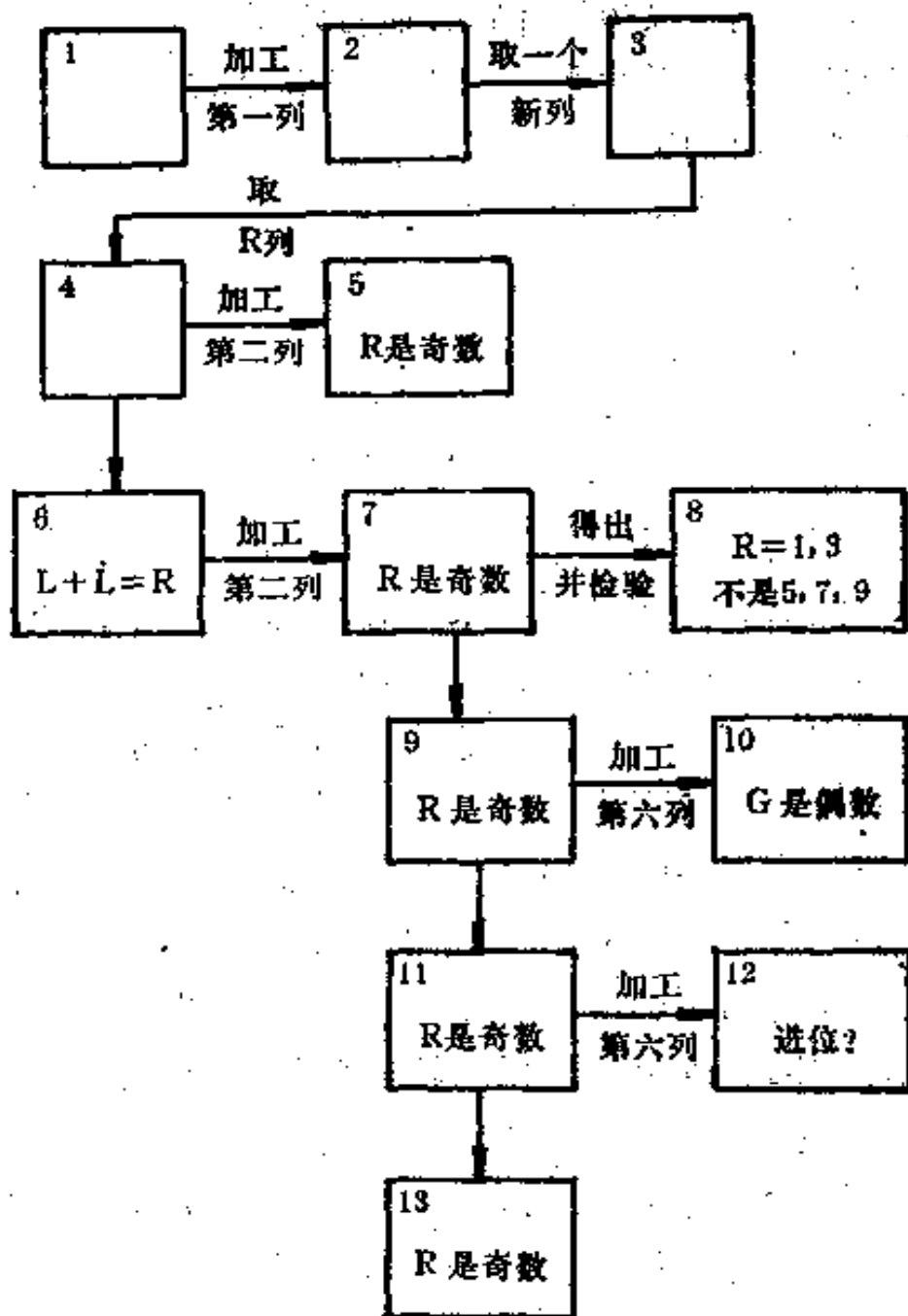


图 10-5 相应于原始记录片断的问题行为图

都经历着类似于上述这位被试者那样的曲折道路。但实际上，这个密码题有一个解法，看起来还是一个比较简捷的方法。现把它述说如下。

将该题的六列自右至左地编号。

$$\begin{array}{rcccccc}
 & 6 & 5 & 4 & 3 & 2 & 1 \\
 & D & O & N & A & L & D \\
 + & G & E & R & A & L & D \\
 \hline
 R & O & B & E & R & T &
 \end{array}
 \quad D=5$$

(1) 因为 $D=5$ ，所以， T 必然是 0 (并向第 2 列进位)。

(2) 看第 5 列， $O+E=O$ 。只有当 0 或 10 与 0 相加时，才会出现这种情况，因此， E 必定是 9 (加上进 1) 或 0。但是，我们知道， T 已经是 0 了，所以， E 必定是 9。

(3) 再看第 3 列。如果 E 是 9，那么， A 必定是 4 或 9 (都需要加上进 1)。但由于 E 已经是 9 了，所以， A 必定是 4。

(4) 在第 2 列里， $L+L$ 再加上进 1，等于 R 并向第 3 列进 1，所以， R 必定是奇数。现在，奇数只剩下 1、3 和 7。从第 6 列中知道， $5+G=R$ ，所以， R 必定大于 5 (因为后面已没有进 1 的情况了)。因此， R 必然是 7。这样，就可以知道， $L=8$ ，而 $G=1$ 。

(5) 再看第 4 列， $N+7$ 等于 B 加上进位。因此， N 是大于或等于 3。现在剩下来的数字只有 2、3 和 6，所以， N 是 3 或 6。但是，如果 N 是 3，那么， B 应是 0；而 T 已经是 0 了，所以， N 必定是 6；结果， B 必然等于 3。

(6) 剩下来的只有字母 O 和数字 2，所以， O 只能是等于 2。

因此，其结果是：

$$\begin{array}{rcccccc}
 & 5 & 2 & 6 & 4 & 8 & 5 \\
 + & 1 & 9 & 7 & 4 & 8 & 5 \\
 \hline
 & 7 & 2 & 3 & 9 & 7 & 0
 \end{array}$$

这个解决办法，看起来确实是非常直接和有效的。但是，

由于各种心理因素的制约，人们在解这个密码算术题时，往往并不是按照上述这些有效的步骤来进行的。

(五) 问 题 空 间

问题空间是由一切可能的关于问题的认识状态构成的。这些认识状态是通过把有效算子应用于问题的某种状态而生成的。信息加工理论把问题解决看作为穿越问题空间搜索一条通往问题目标状态的路径。对于不同的人来说，一个给定问题（特别是那些含糊地规定的问题）的问题空间，可能是不同的；对于那些清楚地规定的问题来说，不同个体的问题空间有可能是类似的。前面所描述的问题行为图，代表了该密码算术题的问题空间的一部分状态。被试者或许会以另一种方式来处理这个问题，从而生成其他许多可能的状态。事实上，对于那些较为复杂的问题来说，其可能的认识状态的数量是相当大的。

问题空间也可以用树形图来表示。举一个简单的例子：一把锁有二个转盘；每个转盘上有英文字母A、B、C；现在锁的设置是AA；要求找出打开这把锁的字母组合。图10-6表示组合锁可能产生的一切状态。为了发现打开这把锁的字母组合，我们必须从AA（节点1）开始搜索。这里有二种可能的操作：可以转动1号盘一个位置，产生BA，即图中节点2显示的位置；我们也可以转动2号转盘，使之产生AB组合，即图中节点3表示的位置。同样，对节点2和节点3来说，也各有二种可能的操作，结果将产生节点4至节点7的不同状态。

这样，一个树形图表示在问题解决过程中的一集可能的操作以及这些操作所导致的各种状态。为了解决这个问题，

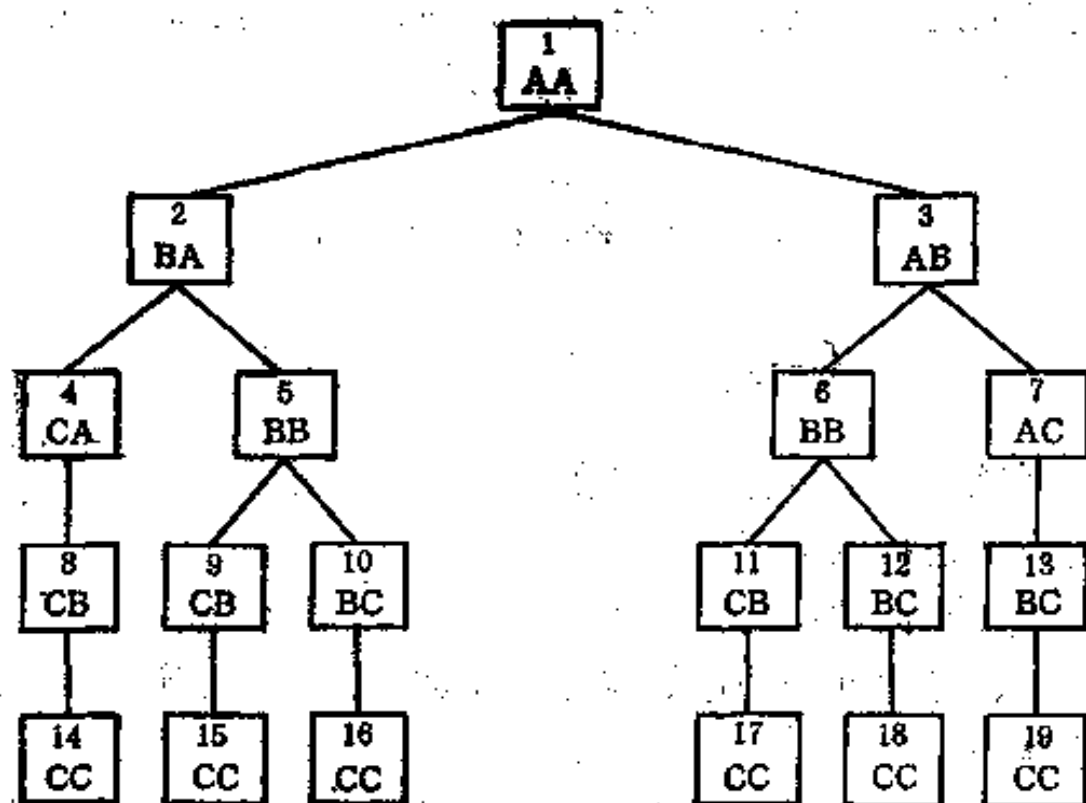


图10-6 二个转盘（每个转盘有3个字母）
组合锁可能产生的状态树形图

就要找到一条路径，穿越这树形图所表示的问题空间，走向目标状态——可以把锁打开的字母组合，如果字母BC组合可以打开锁，那么，达到目标状态的路径可能是1—3—7—13节点。在这里，我们可以看到，当问题空间以这种形式来表示的时候，一个特定的状态可以出现几次，而且可以通过不同的路径到达该状态。事实上，对大多数问题来说，可以通过一条以上的路径来达到问题的解决。问题解决的过程就是对问题空间的搜索过程。

总之，问题空间是对某个特定问题的一切可能的认识状态的集合。穿越该问题空间的路径，可以用问题行为图来表示。尽管我们不能肯定地说问题行为图完全抓住了问题解决

者的思维过程，但不管怎样说，这种分析还是比较客观的，尤其适用于问题解决过程的计算机模拟。

二、问题的表征

(一) 问题表征的基本因素

问题可以用多种方式来加以表征，但是，一般而言，问题的表征包括四个因素。

(1) 问题的起始状态：作为一个问题解决者，必须先要认识、理解问题的起始状态，也就是说，要了解当前处在什么状态。

(2) 问题的目标状态：即问题得到解决时的状态，也就是说，问题要达到什么样的状态才算获得了解决。

(3) 算子：使问题从起始状态转变成目标状态所需要采取的“步”或“步骤”。算子可以用来改变问题的状态。

(4) 对算子的约束：在问题解决过程中，算子的应用往往有一定的约束。也就是说，算子的应用不是无条件的，而是有条件限制的。

因此，问题表征涉及上述四类信息。

问题的表征是重要的，因为它影响和决定着问题解决者所考虑的一系列合法的、适当的“步”。这些“步”的实现构成了从问题起始状态通向问题目标状态的路径。其中一条（或一条以上）的路径将把问题的起始状态与问题的目标状态连接起来，结果产生问题的解。所以，前面谈到的问题空间，也可以说是由问题解决者对一个问题所考虑的一集可能的“步”来规定的。因此，问题空间的大小和性质，取决于问题

解决者如何理解该问题，而在这里，问题的表征是十分重要的。

（二）表征的适当性对问题解决的影响

对一个问题的表征有各种不同的方式，而问题表征的适当性将影响问题解决的难易程度。例如，请考虑下面这个“和尚爬山”的问题：

一天早晨，刚刚日出，一个和尚开始沿着盘旋的山路爬一座高山，到山顶的一座庙里去。和尚向上爬时，速度时快时慢，而且在路上多次停下来休息和吃随身带的东西。他到这这座庙时，刚好日落。在庙里住了几天后，他开始沿着同样的路线下山，也是在日出时起程，行走速度也是时快时慢，而且沿路有多次停歇。当然，他下山的平均速度比上山的平均速度要快一些。即使这样，当他到达山脚下时，也将日落。请问，和尚在往返的路上，是否曾在一天中的同一时刻通过沿路的同一点。

如果说我们觉得这个问题的解答不明显，那可能是由于我们对这个问题的表征方式的缘故。在考虑这个问题时，我们往往只是从一个和尚在不同的日子里上下山这个角度出发，所以，该和尚是否会在一天中的同一时刻通过同一点是难以肯定的。但是，如果我们把情境想象为有两个和尚，在同一天里，一个上山一个下山，那么，问题就容易解答了。很显然，他们必然会在一天中的同一个时刻在路上的同一个地点会面。同样，若用图10-7来表示这个问题，我们就会发现，问题

的解答是明显的。

另外一些例子，如残缺的棋盘问题和媒人问题，也说明表征在问题解决中所起的作用。先看一下残缺的棋盘问题。一个棋盘有64个方块，32个白的和32个黑的；还有一付32张长方形的骨牌，每一张骨牌盖住棋盘上二个方块。所以，我们可以用这些骨牌把棋盘全部盖住。但是，如果从棋盘的二个对角切掉二个黑方块（如图10-8所示），那么，用31张

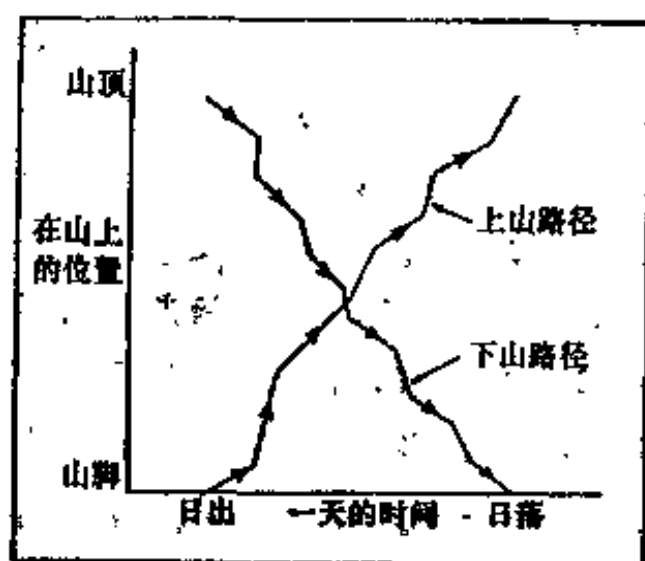


图10-7 和尚爬山问题的图形表示

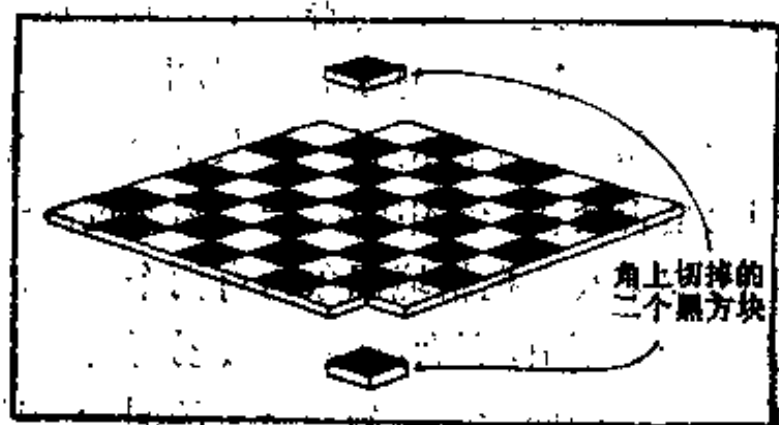


图10-8 残缺棋盘问题的情境

骨牌是否也能恰好盖住棋盘上剩下的62个方块呢？

对这个问题的回答是否定的。31张骨牌不可能盖住这个残缺的棋盘。这是因为，不管骨牌是横着放或是竖着放，每张骨牌只能盖住一个白方块和一个黑方块。31张骨牌只能盖住31个白方块和31个黑方块。但是，现在我们已经从棋盘的二个角上切掉了二个黑方块，留下来的是32个白方块和30个黑方块。所以，用31张骨牌就没有办法把这个残缺的棋盘恰好盖住。

现在我们再来看一下媒人问题。

在一个城市附近的小村庄里，住着32个年轻健壮的单身汉和32位未婚女子。一些爱管闲事的人，经过巨大的努力，安排了32对和谐的婚姻。整个村庄是幸福的。然而，在一个疯狂的星期六晚上，刮起了暴风雪，气温降到零下40℃。二个年轻的男子进行耐力考验，喝了数加伦的豌豆汤，结果死了。这些媒人是否能在62个幸存者中间安排好31对满意的婚姻？

这个问题比残缺的棋盘问题容易得多了。因为每对婚姻必须包括一个男的和一个女的，现在少了二个男的，就不可能安排出31对婚姻。

上面所举的例子，清楚地表明，我们表征一个问题的方式，对问题解决的难易有显著的影响。

当人们构造问题的表征时，有两种过程似乎是起作用的。第一，人们可以有选择地注意呈现给他们的信息。我们知道，并不是问题中的所有信息对问题解决来说都是必需的。事实

上，问题中有些信息对问题解决来说可能不是必需的，有时甚至会把问题解决者的思路引向不适当的方向。第二，在解决新问题时，人们可以利用先前的知识，包括有关特殊的问题类型的知识。有关问题类型的知识，可以帮助人们识别出新问题的类型，从而有助于获得问题的解决。识别出一个熟悉的问题类型，不仅可以使人们知道要注意什么样的信息，而且也指导人们从记忆中提取有用的解决问题的策略。

三、解决问题的策略

(一) 算法和启发法

概括地讲，成人解决问题的策略可分为二种：一是算法 (algorithm)，另一是启发法 (heuristic)。算法是解题方法的一种精确的描述。它用来求解问题的特点是，如果问题至少有一个解的话，那么，它保证以有限的步数获得问题的解，不管问题解决者是否理解了该方法的含义。比如，天气预报说今日气温为摄氏 30° ，那么华氏表的温度是多少呢？为了获得华氏表的温度，我们可以把摄氏表温度乘以1.8，再加上32度。这就是一个很简单的算法例子。

相反，启发法是一种经验规则（有人把它叫做“大拇指规则”）它可以用来帮助解决问题，但是，它并不保证使问题一定获得解决。比如说，下棋时人们常用的一种启发法，是控制棋盘的中央，这种启发法并不是一种获胜的保证，但是，利用启发法能使人增加获胜的机会。

既然算法能保证使问题得到正确的解，为什么人们在解决问题时还要使用启发法呢？这里有几方面的原因：（1）

对某些问题来说并没有一个可用的、有效的算法；（2）虽然有算法可用，但为了节省时间和精力，就靠经验来解决问题。例如，一个保险柜有四个转盘，每个盘上有10个号码。打开保险柜的密码丢了，并且记不起来。怎么办？四个转盘，每个盘有10个号码，这里可能的组合就是 $10^4=10000$ 种。如果采用算法，系统地对每种组合进行尝试，以期找到打开保险柜的密码，那就需要好几天的时间。但是，若根据经验指导，把寻找号码组合的尝试集中在最有希望的范围内，那么，也许不用太多的时间和力气，就能发现这个密码；（3）有些问题虽然有算法，但由于某些原因，这种算法事实上是无法执行的。例如，下棋比赛，如果比赛一开始就考虑所有可能的棋步和对手可能采取的反应，以及你可能给对手的反应，对手的再反应等等，并在此基础上选择那能导致获胜的一步棋，那么，尽管这种办法在理论上将保证你获胜，但实际上是不可能实现的。塞缪尔(A. L. Samuel)曾计算过，用这种方法来下跳棋，将涉及到 10^{40} 可能的棋步，若以每毫秒考虑3步计算，也仍然需要 10^{21} 个世纪的时间；而对国际象棋来说，他估算其可能的棋步数量是 10^{120} 。因此，若利用算法来下棋，那么，太阳系消失了，棋或许还没有下完。

应该指出，人们使用什么样策略，这是问题解决的一个决定因素。我们平时碰到的问题，在复杂程度上自然是有差别的。但是，问题解决的难易和可能性，却往往取决于问题解决者使用的策略类型。另一方面，某种策略是否合适，这不但依赖于问题的性质，也取决于问题解决者的能力和技巧。如果某种策略既符合于问题的性质，又适合于问题解决者的特点，这样的策略将是一种有效的策略。

(二) 手段-目的分析

有许多问题被称之为“转换的问题”。在解决这类问题时人们常使用的一种启发法,叫做“手段-目的分析”策略。这种策略的要点在于,识别起始状态和目标状态之间的差别,并采取行动(使用算子)去缩小这种差别,使问题一步一步地接近和达到目标状态。所以,从本质上说,手段-目的分析策略乃是一种缩小差别的策略。

这种策略通常用来解决那些一般性的问题,例如,冬天在户外怎么保持温暖的问题。这个问题的起始状态是感觉到冷,其目标状态是要达到产生暖的感觉。问题解决者在使用这种手段-目的分析策略时,无论采取什么步骤,都是有利于产生暖热的,如增加衣服,燃烧一堆篝火,寻找一个避风的草棚等。当然,对于某个人来说,在没有策略性的计划和目标意识时,这些动作或许也会进行的。但是,在大多数情况下,人们是自觉地利用手段-目的分析法,去努力达到这些目标的。

从解决河内塔问题的作业中,我们可以看到手段-目的分析法的优越性。图10-9表示只有三个圆盘的河内塔问题

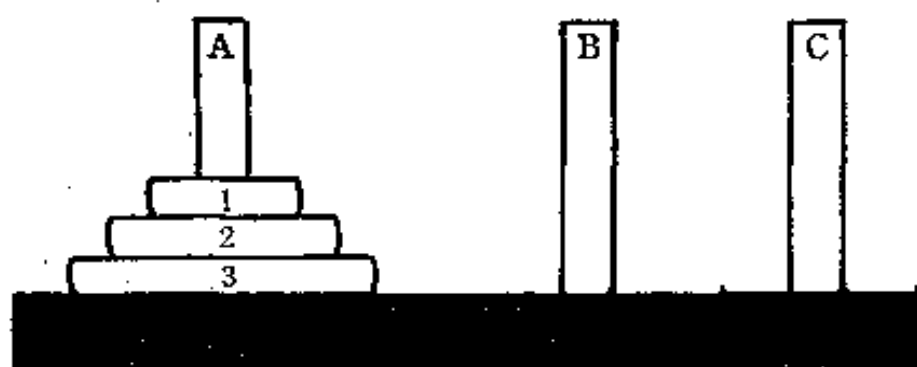


图10-9 河内塔问题

的初始状态。被试者的任务是，以最少的步子，把三个圆盘从A柱移到C柱。同时，要遵守二条规则：（1）一次只能从一列圆盘的顶部移动一个圆盘；（2）小的圆盘只能位于大的圆盘之上，不允许把较大的圆盘放在较小圆盘的上面。当圆盘以1、2、3的次序排列在C柱上时，问题解决者就算取得了成功。传说河内附近的和尚曾企图解决一个很复杂的、包括64个圆盘的河内塔问题。根据这个传说，当和尚们完成这一任务时，世界的末日也就到了。当然，我们不用惊慌，因为要实现这个问题的解决，需要很长很长的时间。据估算，即使僧侣们每秒钟能移动一个圆盘，而且都是正确的，他们也需要大约一万亿年的时间，才能根据上述二条规则，把64个圆盘从A柱移到C柱。

河内塔问题的目标结构是可以清楚地规定的。它的总目标是把三个圆盘从A柱移到C柱。这个总目标可以分成若干子目标。一个较明显的子目标，是首先要 把圆盘3移到C柱。而这个目标本身，再可以分成若干属于第二个层次的目标，如把圆盘1从A柱移到C柱，把圆盘2从A柱移到B柱，然后再把圆盘1从C柱移到B柱。也就是说，把圆盘3移到C柱，需要完成下列几个第二层次的目标：

圆盘1从A柱移到C柱，

圆盘2从A柱移到B柱，

圆盘1从C柱移到B柱，

圆盘3从A柱移到C柱。

在实现了把圆盘3移到C柱这个子目标后，问题解决者就着手完成下一个子目标，即把圆盘2移到C柱。实现这个子目标比较容易，只需二步就可以了：

圆盘1从B柱移到A柱，

圆盘 2 从B柱移到C柱。

手段-目的分析法是一种有用的策略。其中一个原因，是它对加工能力的要求相对地轻一些。它把总目标分解为若干子目标，并一个个地去解决子目标，从而减轻了问题解决者的加工负荷。

手段-目的分析策略在解决那些没有明显子目标的问题时，同样是有用的。例如，有一个“传教士和野人过河”的问题：

三个传教士和三个野人在河的左岸。他们必须用一只小船渡到河的右岸，但小船最多只能载二人。如果在任何一边的岸上，野人的数量超过了传教士的数量，野人就会吃掉传教士。怎样才能使三个传教士和三个野人从河的左岸渡到河的右岸，同时又不至于发生野人吃掉传教士的现象呢？

解决这个问题的方法可分以下几步（图10-10）：

- （1）二个野人先乘船过河；
- （2）一个野人划船返回左岸；
- （3）再由二个野人乘船到河的右岸；
- （4）一个野人又划船返回河的左岸；
- （5）二个传教士一起划船过河；

（6）一个传教士和一个野人划船返回左岸。对于解决这个问题来说，这是关键的一步。想到了这一步，以后几步就比较容易了。

实验结果表明，在解决这个问题时，大多数人采取一种平衡策略。他们力图使传教士和野人的数量在任何一边岸上

都均等。他们之所以使用这种策略，或许是出于问题中要求防止野人吃掉传教士的缘故。实际上，这是不必要的。我们知道，若所有的传教士在河的一边岸上而所有的野人在河的另一边岸上时，是合乎规则的，不会发生野人吃掉传教士的现象。图10-10表明，为了使问题获得解决，这种排列是必然会出现的。西蒙等人发现，被试者走了几步以后，大多数放弃了平衡策略，并转而采用手段-目的分析法。被试者在使用手段-目的分析策略时，倾向于把大数量的生物渡到

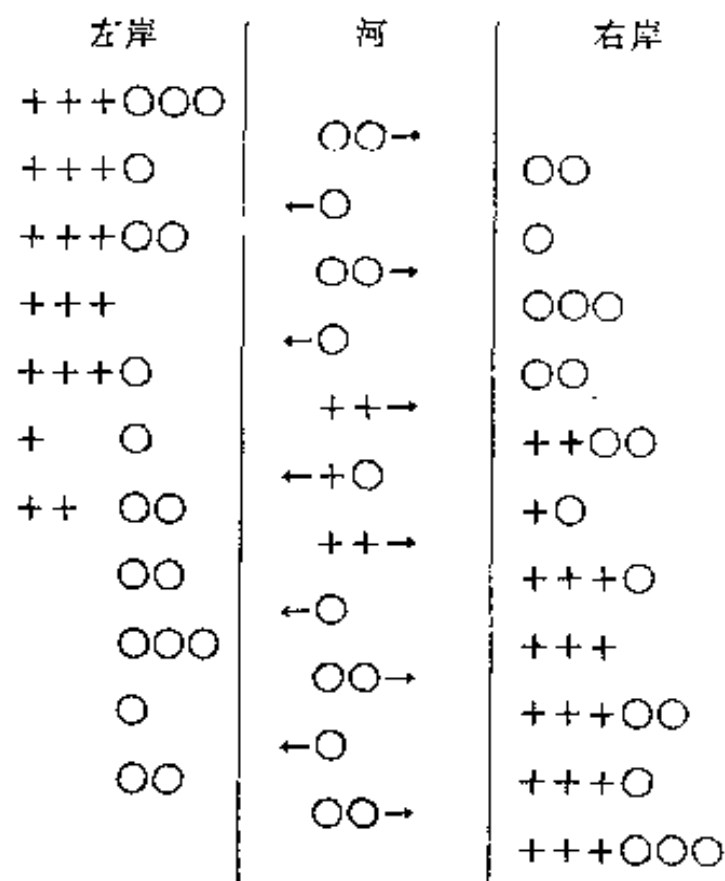


图10-10 解决传教士和野人过河问题的步骤。“+”代表传教士，“○”代表野人，箭头表示船的走向，箭头后部的标记表示每次横渡的乘客。

河的右岸去，而把小数量的生物带回河的左岸来。显然，被试者这样做的目的，是企图尽可能地缩小起始状态和目标状态之间的差别。他们所采取的步骤，往往使传教士和野人在河两岸是不等量的。显然，这与平衡策略的出发点是不同

的。

手段-目的分析法有时也会导致错误。例如，被试者使用这种策略时，在最初几次往返旅程中，往往只涉及一个生物。然后，有的被试者或许倾向于固定下来了，以致在随后所有的往返旅程中，都只涉及一个生物。显然，这样做的结果是注定要失败的。

事实上，为了解决一个问题，有时候扩大当前状态和目标状态之间的差别，是完全必要的。

当然，人们在解决问题时，很少严格地依附于一种纯粹的策略。在大多数情况下，人们总是要使他们的方法适应任务的要求，同时依靠多种策略来进行作业。

(三) 搜索的策略

搜索是问题解决活动的一个普遍特征。就拿执行手段-目的分析策略来说，就涉及搜索。这种搜索过程在于寻找一个适当的步，以缩小当前状态和目标状态之间的差别。同样，在解决那些字谜游戏问题时，也涉及到搜索过程。例如，根据字母串TERALBAY构成一个合法的词这样一个问题，就要求搜索各种不同的字母组合，并找到一个合法的、真正的词。一个人如果企图在图书馆的许多架子上寻找一本参考书，或在词典中查找一个不认识的单词，搜索活动是必不可少的。

搜索可以由两种不同的方式进行，一种是宽度优先搜索，另一种是深度优先搜索。宽度优先搜索是先沿着水平方向进行，然后指向重直方向。深度优先搜索的情况恰好相反，首先沿着垂直方向进行搜索，然后再指向水平方向。这两种不同类型的搜索方式适用于不同的目的。当然，人们在

解决实际问题时，并不是单独地采用宽度优先搜索或深度优先搜索的，而往往是交替地使用这两种搜索方法。

1. 生成-测试策略 并不是所有的搜索策略都象宽度优先搜索或深度优先搜索那样有系统。事实上，人们在解决问题时往往利用一种简单的叫做“生成-测试”的策略来进行搜索。这种策略的意思是说：生成一个问题的各种可能的解答作为候选者，同时对每一候选者进行测试，看看它是否属于该问题的一个合理的解答。比如说，用字母串YADL造一个词这个问题，人们可以生成这些字母的一切可能的组合，同时去确定每种组合是否是一个词。这样，可以使问题获得解决。

生成-测试策略用起来比较简便。但是，这种策略在很多情况下效率是很低的，或者甚至是不能用来工作的。首先，这种策略在问题只有少量可能的解需要考虑时才是可以使用的。但是，大多数问题具有大量可能的解答，因而无法用生成-测试的方法来有效地进行搜索。前面提到的用字母串TERALAY来构成一个合法的词这个问题，其可能的字母组合数量相当大（有40320个可能的字母组合）。因此，该问题的实际解在一个合适的时间内或许不可能生成。也就是说，这个问题的搜索空间，即问题解决者意欲考虑的该问题的所有可能解答的集合，是比较大的，所以，生成-测试策略的效率就比较低。一般来说，搜索空间越大，生成-测试策略的效率就越低，因为它对一切候选者都同样对待，不管它们在多大程度上有可能成为问题的实际解。而对有些问题来说，由于其搜索空间很大，以致于在短时间内根本不可能生成实际的解答。所以，我们需要一种办法来缩小选择的范围，也就是缩小搜索空间，以便只考虑那些较好的候选者。

生成-测试策略还有另一局限性。它涉及到生成一切可能的解，因此，它可以成功地用于一些规定清楚的问题，如简单的字谜游戏问题；而那些规定含糊的问题，其解答的性质是不知道的，就很难使这种策略获得成功。解决这类问题的有效途径，是在于识别那些有助于指导随后搜索的部分解。总之，从本质上说，生成-测试方法不能把搜索空间限制到一个易于处理的范围。

2. 限制搜索空间的记忆策略 借助于各种记忆策略，可以限制搜索空间。最简单的一种记忆策略，只是回忆和搜索过去经历过的事情的历程。比如说，如果问题的要求是构成一个以字母Q开头的词，那么，我们就可以开始搜索我们以前见到过的以Q起头的那些词。这种策略尽管并没有什么创见，但在解决那些属于“重现”的问题时是很有用的。

另一种比较复杂一点的记忆策略，涉及到类比推理。这种策略有一定的优点，它使我们能应用自己具有的有关以往事态的老知识去解决那些新的类似的问题，从而减少了为解决新问题我们所必需的搜索数量。比如说，为了做一个人工髋关节，需要寻找一种结实而轻巧的材料。这个问题的搜索空间应该说是相当大的，如果我们采用没有指向的搜索过程，那么，就会耗费很长的时间，或许永远也不会找到一种合适的材料。但是，如果我们注意到，这个问题与制造飞机的问题有类似之处，那么，搜索就有了指向，解决这个问题搜索空间就会大大缩小。科学家就可以把搜索活动集中在那些构造飞机的材料上。事实上，科学家们经常利用这种策略，使已知的东西去影响或帮助解决那些未知的、但类似的问题。

3. 建构性搜索 这种搜索策略是高度指向性的，在这

种搜索中,问题解决者利用自己的语义知识和世界知识,对问题产生一个不完全的解。这个不完全的解会缩小仍待搜索的项目的数量。为了理解什么是建构性搜索,我们可回忆一下前面提到的密码算术题。

在解决这个密码算术题时,人们常常针对同列中字母进行思考,以此力图识别出一个特定字母的数值。例如,大多数人一开始就利用他们的算术知识去进行推理;既然 $D=5$,那么 $T=0$,因为 $5+5=10$ 。被试者知道, $L+L$ 的和必定是偶数,同时,第2列还有进位1,所以,他推论出 R 必然是一个奇数。既然 $D+G=R$ (第6列),那么,就可以得出 R 必定是大于5,等等,以这种方式继续下去,许多被试者都解决了这个问题。如果使用一种没有指向的搜索来解决这个问题,比如说,任意地用一个数字来代替一个字母,再计算一下看看行不行。那肯定是要失败的,因为在这个问题中,用数字代替字母,有几十万种可能的方式。在数量这么大的组合中去搜索一个正确的答案,对一个人来说几乎是不可能的。

建构性搜索策略经常用来解决一些日常问题。例如,打字员经常要辨认一些不规则的或潦草的词。他们会利用语言规则和上下文的知识去确定这些词的可能的文法类别和意义。他们有指向性地在语义记忆中搜索,而不是采用那种效率很低的生成-测试策略,任意地生成一些词,并对每个词进行检验看看是否合适。又如,侦探在缉查某个杀人犯时,他们往往把罪犯的作案工具或留下的脚印作为线索,并根据他们具有的关于这方面的知识,指导他们的搜索过程。

人们在使用建构性搜索策略来解决问题时,常会导致子目标的形成。例如,在上述例子中,侦探首先力图发现谁最

可能拥有这样的作案工具或鞋子，就这一方面而言，建构性搜索策略是启发式的“手段-目的”分析法的一种补充。

一个效率高的问题解决者的技巧，在于他能够熟练地以一种新的、有力的方式把不同的策略结合起来，去解决各种不同的问题。

四、问题解决的计算机模拟

现代认知心理学的主导思想，可以简要地用一句话来加以表述：人类的认知系统可以看作是一个信息加工系统。这个思想构成了当代认知科学的基础，它在“计算机模拟”这个领域中得到了最充分的发展。这个研究领域的主要目的，是要构造各种计算机系统，模拟人类外显的行为和作为这些行为之基础的认知过程。一个模拟问题解决的计算机系统，不但要能够解决人类所解决的问题，而且要利用人类在解决问题时所使用的那些表征和策略。

认知科学家对构造计算机模拟系统的兴趣出于多种的原因。第一，他们希望检验他们提出的理论的可靠性。认知心理学家认为，如果人们认识了某方面问题解决过程的规律，并形成了一定的问题解决的心理学理论，那么，就可以根据这个理论来编写计算机程序，并使计算机程序能以类似于人的方式去解决问题和达到类似的结果。这样，这个理论就得到了证实。否则的话，就可以发现这个理论的不足及其存在的问题。第二，认知科学家们希望对某些重要的术语，如“表征”、“理解”等，尽可能清楚地加以规定。理论的明确性不是一个无足轻重的问题。事实上，只有理论是明确的，它们才是可以检验的。第三，计算机模拟是名符其实的

心理学理论，它有助于指导我们的研究和揭示有用的概念。可以说，这是认知科学家对计算机模拟感兴趣的最重要的原因。计算机和人在物理构成成分上是完全不同的，但两者的功能是类似的：它们的加工原则，即信息加工原则，却是一致的。两者都是信息加工系统，都是符号信息加工系统。它们都操作各种各样的符号，遵循逻辑上类似的步骤来完成特定的动作。所以，人的认知活动可以用计算机来模拟。从精确地模拟人类作业这个方面来说，计算机模拟是一种合适的心理学理论。

在设计模拟模型时，认知心理学家常常要利用其他学科如计算机科学的信息。比如说，若认知心理学家还没有搞清楚人们所使用的各种知识表征的类型，那么，他们就会从计算机科学那里去寻找有关知识表征的各种思想。事实上，这些思想或概念常常被证明在心理学理论中是非常有用的。认知心理学家确实引进了不少新的、强有力的概念，使计算机模拟获得成功。这些都是真正的理论成就。

（一）通用问题解决者

通用问题解决者（General Problem Solver，简称GPS）是由纽厄尔（A. Newell）和西蒙等人研制成功的。它是一个用来模拟人类问题解决行为的计算机程序。由于它适用于多种类型问题的解决，所以才称之为“通用”。

GPS包括一个长时记忆，存储各种语文知识有关的世界知识，并存有使用这些知识的各种算子；另外，还包括一个容量有限的短时记忆，以实现和信息进行串行的加工。短时记忆是一个工作空间，各种信息在短时记忆中进行比较，并对下一步的行动作出决策。

GPS系统是通过问题空间的搜索和考察来进行工作的。前面我们已经提到，问题空间是由该问题的一切合法形态组成的。所谓合法状态，是指当系统对一个问题进行作业时系统可能会进入的问题状态。GPS利用“手段-目的”分析的启发式方法，对问题的初始状态和目标状态进行比较，然后选取一种操作（算子），以缩小当前状态和目标状态之间的差别。

GPS内部的知识，主要是以产生式规则的形式来表征的。产生式规则分二个部分，一是条件部分，二是行动部分。所以，产生式规则也叫条件-行动规则。如果某一规则的条件部分得到满足，那么，该规则就指导系统去执行由规则指定的某个特定的行动。这里举一些简单的产生式的例子：

如果是红灯，那么就停车。

如果上课的铃声响了，那么就走进教室。

如果等号两边都有一个常数，那么两边都减去等号左边的那个常数。

从这里我们可以看出，每条产生式规则都是由一个“情境-行动”对构成的。它把系统在什么样的特定条件下采取什么样的行动这一小块知识结合在一起。一个模拟人类认知活动的计算机系统，往往包含有成百上千条产生式规则。

为了设计GPS系统，使它象人类那样地去解决问题，就得首先了解人类在解决问题时所经历的各种步骤。因此，他们首先采用“大声地思维”的方法，对人类被试者解决问题的过程进行了实验研究。然后，对被试者的“口语素材”进行分析，以了解和揭示问题解决者所使用的问题空间和问题解决者当前所处的知识状态，以及问题解决者把问题推向目标状态时所使用的各种操作。可以说，这些关于人类解决问

题过程的实验材料，是他们设计GPS系统的基础。同时，它们也是用来评定GPS是否象人类那样的方式去解决问题的依据；是评定人类和GPS两者所使用的问题空间、知识状态以及各种操作有多大程度类似性的依据。事实上，在人类心理过程的计算机模拟中，口语素材或原始素材分析这种方法，已证明是一种有效的方法。

对那些规定良好的问题，如前面提到的传教士和野人过河的问题，GPS系统的作业与人的作业是相类似的。当然，在某些方面，与人的作业还是有一定的差别，比如说，GPS的作业是严格按照串行的方式进行的，而且，与人类被试者相比，GPS更多地依赖于手段-目的分析方法。但是，总的来说，它取得的成就是卓著的，尤其是因为在解决各种任务时，它只使用了少量基本的加工过程，如在短时记忆中对二个项目进行比较，选择恰当的算子等。另外，GPS系统研制成功这一事实本身就具有重要的意义，它不但表明人的认知过程是可以用来计算机来模拟的，同时，它的成功也标志着一门新的学科——人工智能的真正诞生。

（二）计算机模拟遇到的问题

用计算机来模拟人类的认知过程，是心理学研究中的一个重大的进展，这是无可非议的。但是，这种努力也面临着一些障碍，包括在方法论上和理论上存在的问题。

第一，评定标准问题。怎么去评定计算机的作业在细节上是否完全模拟了人类的作业，这个问题现在还不太清楚。这是属于方法论上的问题。一般都承认，通过口语素材的分析，认知心理学家能相当好地揭示人在解决问题时所采取的有意识的加工步骤；利用口语素材分析，我们可以对计算机

采取的加工步骤和人所采取的有意识的步骤之间的类似性作出评定。但是，在人类解决问题过程中，某种无意识的加工也许在起作用；另外，有意识的加工步骤在问题解决者的口语素材中是否全部得到反映，这也是没有把握的。所以，把口语素材分析的数据与有关计算机作业的数据进行比较，并不能十分正确地告诉我们，人和计算机在执行同一特定任务时是否完全使用了类似的基本过程。这就是说，通过口语素材分析数据与计算作业过程的比较，我们只能断定两者在完成某一特定任务时所采取的全局性的、宏观的步骤是相同的。

第二，人的认知活动是由大量的背景知识指导的。例如，一位老练的棋手，在下棋时，并不是以串行的方式有意识地对每一可能的棋步逐一深思后才决定怎么走的。相反，他们利用自己关于象棋的背景知识，自动地“瞄准”最强有力的一步。这种背景知识并不直接进入问题的表征，但是，它可以影响和指导问题的表征，是问题表征形成的基础。既然问题表征的形式影响问题的解决，因此，必须在理论上说明背景知识的结构和性质以及人们是怎么利用背景知识的。有人认为，人类具有的背景知识本质上是整体性的，不可能用当前计算机中所使用的形式规则来加以恰当地描述。另外，人类具有的背景知识是通过学习获得的，并且是易于变化的。我们可以使计算机模拟人类学习的某些方面，如学习解代数题或学习概念，但是，对计算机学习问题的研究，尚处于婴幼儿时期。计算机是否能象人类那样的方式去学习背景知识，这个问题现在还很难回答。

第三，人的心理过程是极其复杂的，要精细地模拟人类解决问题的过程，计算机程序就会显得非常大而复杂，以致难于理解。即使是最好的程序编写者，也不会那么幸运地有

一个仔细写出来的程序，在第一次运行时就能正确地工作。通常，程序编写者总会在程序中发现有一些缺陷或“故障”。编程者必须要对这个错综复杂的程序进行修改、调试，必然会把大量微小的步子插入系统中去。结果，有时候会出现一种令人啼笑皆非的现象：程序或许可以工作了，而编程者却并不完全清楚地理解它为什么会工作和工作的详细过程。另外，当一个系统装入了过多的微小细节时，就有可能掩盖这个系统的主要的理论主张。

第四，现在，我们不能说计算机也有情绪。因此，当我们用计算机来模拟人类认知过程时，必须假定认知过程和情绪过程彼此是独立无关的。但是，事实上，这二种过程彼此之间是相互影响的。大量的事实表明，人类的决策既取决于认知因素，也依赖于情绪的作用。情绪因素运转起来比认知因素快，因而情绪因素有时可能首先起作用。所以，它们能够影响或甚至指导认知加工过程。例如，如果我们对某人产生一种赞赏的最初反应，那么，我们在认知上似乎就会倾向于他的正面特性；如果最初的反应是否定的，那么，我们就会倾向于寻找这个人的某些负的性质。这样，情绪活动在影响着人类的认知活动。然而，至少到目前为止，我们还没有理由说，情绪能够体现为那种指导计算机工作的符号结构；因此，也没有足够的理由说，计算机可以模拟人类的情绪的活动。

上面谈了计算机模拟中存在的四个问题。这并没有给我们任何理由去低估或忽视用计算机模拟人类认知活动的重要意义及其真正的价值。这只是说，在计算机模拟人类认知活动这个方面，还有一些重要的问题需要我们从理论上和实践上去进行探讨和解决。

五、决策及其策略

决策或抉择是人们日常生活中一种最普遍的活动，从一般意义上说，每当人们要从一集候选者中选取某个候选者时，人们实际上就是在进行抉择或决策。人们在作抉择或决策时，不一定都清醒地意识到自己的选择过程。比如说，象早晨起床这样简单的动作，也包括大量的抉择。它不仅仅涉及到起来或不起来的问题，而且也涉及到先睁开右眼还是左眼，或是二只眼同时睁开；先抬起头还是先移动腿；是移动右腿还是左腿等等的抉择。从这个意义上来说，我们每天或许要作出数以万计的抉择或决策。一般而言，我们把决策看作是一种有意识的行为。尽管人们不必定意识到决策是怎样作出来的，或者说不一定意识到决策过程本身，但是，决策确实是指向一定的任务的。当然，这并不不是说，一切抉择或决策都是经过深思熟虑而作出来的。我们可以想象，如果一个人作出起床的最初抉择时，对大量微不足道的、但与起床有关的抉择都要进行深思熟虑，那么，很可能一天的时间已经过去，而这个人却依然还在床上呢。

事实上，人们在生活中形成的习惯，是一种相对自动化的动作。它解放了人类信息加工的能力，使人有可能去从事更重要的活动。然而，尽管自动化的或出于习惯的行为具有这种不必深思熟虑的优点，但它同样需要一定的代价。这种隐蔽的自动的决策活动，是一个有意义的、也是重要的问题。当然，大多数研究者的注意力，都集中在人的那些比较深思熟虑的决策行为上。

人们在作出决定时，需要对各种候选者进行某种形式的

比较。每一候选者都具有多种属性或特性。针对这些属性或特性，在各候选者之间进行比较，这是进行决策的基础。

(一) 预期-效用模型

几十年来，预期效用模型一直是决策的主要理论。这个理论的最简单的形式，使用了预期金钱值的概念，并且可以通过对各种赌博的抉择来具体地加以说明。金钱赌博的预期值的计算方法，是用每个候选者的值乘以它的概率，并把它们的积加在一起。我们看一下表10-1，例如，对赌博1来说，它的预期金钱值是：

$$(.5 \times 12 \text{元}) + (.5 \times 8 \text{元}) = 10 \text{元}$$

表10-1 你愿选哪种赌博来打赌10元？

赌博1：50%的机会获得12元和50%的机会获得8元。	(预期值=10元)
赌博2：75%的机会获得12元和25%的机会获得4元。	(预期值=10元)
赌博3：10%的机会获得28元和90%的机会获得8元。	(预期值=10元)

在这里，决策者碰到的最重要的问题，是求确定每种赌博的预期值。并且选择那个具有最高预期值的赌博。根据这种简单的预期效用理论，决策者对表10-1中列出的这三种赌博，将会同等看待。因为从表中可以看出，这三种赌博都具有同样的预期金钱值。

预期金钱值的理论，其局限性比较大。所以，人们把它的一般方法从不同的角度进行了修正，从而产生了预期效用理论的各种版本。所以作这些修正，主要是考虑到以下三个事实：

(1) 金钱值和对决策者的价值并不是以一对一的方式联系在一起的。简单地说，在大多数情况下，对于一位决策

者来说，1元和2元之间的差别不同于100元和101元之间的差别。

(2) 决策不仅仅涉及客观概率，而且常常受主观概率的影响。

(3) 在大多数的决策活动中，候选者的特征并不具有金钱价值，但它们对决策者来说确实具有某种价值。换句话说，不是所有的价值都可以转换成金钱的。

预期效用模型虽然有几种不同的变式，但它们具有共同的核心思想。这些思想包括：

(1) 一个候选者的预期效用是被预测的愉快和痛苦的一种加权的平均数；

(2) 权数（或概率）和值是各自独立的，但是，在确定预期效用时，则把它们合在一起相乘；

(3) 抉择完全是根据预期效用来确定的。

由此可见，根据这种理论，对每个候选者的评估，是依据每个候选者的好的和坏的方面的某种加权的组合来进行的。所以，每个候选者有一种预期的效用，这种预期的效用既考虑了它的正的特性，又考虑了它的负的特性。一旦所有候选者的预期效用确定下来以后，再作抉择就比较简单了。这就是说，不管在任何情况下，我们作抉择时，自然要选取那个具有最高效用的候选者。现在，我们举一个具体的例子来说明这个问题。比如说，有人考虑选择一个工作。假定说，与决策者有关的一些工作的特性，包括以下几项：收入、提升的机会、工作的兴趣水平、退休金和休假旅行等。对决策者来说，其中每一个特性都有一定的效用或价值。但是，对不同的决策者来说，它们的效用或价值可能是有变化的。因此，对某位决策者来说，每种工作将有某种（包括一

切的、平均的)预期效用。决策者要考虑的所有工作,可以根据预期效用从最高到最低地排列起来。根据这种排列的次序,决策者就可以确定作何种抉择。表10-2列出了三种不同的工作,每一工作的值都是假想的。为了简便起见,我们假定所有工作的特性都同等地加权。这样,一个候选者的预期效用,就是它的5个维度值的简单平均。

表10-2 假设三种工作的预期效用

工作	收入	提升机会	兴趣水平	退休金	旅行	预期效用
售货员	9	5	2	2	3	4.2
会计	7	4	4	6	5	5.2
教师	5	3	7	6	8	5.8

注:数字是任意的单位,反映一个特性对决策者的价值。数字越大,表示价值越高。

表10-2中给出的任何一个特性的值,既反映了抉择者实际察觉到的该工作的特性,也反映了抉择者本人对该特性的主观值。例如,收入项下的7这个值,代表候选者本人对该工作所预计的收入假定值。从表中我们可以看出,各个项目下的值是不同的,有的高,有的低。根据预期效用模型,对这些值要通盘考虑。所以,它们彼此之间可以相互平衡和补偿。例如,售货员这个工作的兴趣水平较低,但它的收入较高,后者可以补偿前者。当然,决策者作抉择的依据,不是个别项目下的个别的值,而是那个总的值,也就是预期效用的值。根据表中所列出的数据,这位假想的决策者应该选择教师这一工作,因为教师工作具有最大的预期效用。

一般认为,预期效用模型是人们进行正确决策的标准,但也有人对此提出一些问题。例如,对预期效用模型是否充

分地描述了实际的决策过程，人们是有不同看法的。有不少的证据表明，人们作出的决策，常常与预期效用模型作出的预测是不一致的。这种不一致性的产生可能有许多原因。例如，根据预期效用理论，象收入和休假旅行的时间这样一些特性，可以放在一个尺度上来进行比较。但是，对人来说，要这样做将会是有困难的。前面谈到的对各种赌博作出抉择这个例子，如果人们因为赌博了可以赢的数量相对要大一点（不管赢的概率有多大）而选择赌博3，那么，他们这样的考虑与预期效用的原则是不一致的。因此，研究者们已经看到，人们事实上还利用其他一些不同的策略来进行决策。

（二）非补偿的策略

预期效用模型有一个主要的特征：它对一个候选者进行评价时，要考察候选者的所有特性；而在确定总效用时，这些特性又是可以彼此补偿的。然而，在日常生活中，人们还使用另一些抉择的策略。这些策略不是采用各种特性彼此补偿的办法，或者说，它们只根据候选者的一集特性中的部分特性来进行抉择。所以，这种策略通称为非补偿的策略。

1. **“最大”策略** 决策者在使用这种策略时，对不同候选者的最佳特性进行比较，如果某个候选者具有最强的最佳特性，那么，就把它作为选择的对象。例如，若把这个策略用于上面所列举的选择工作的例子，人们就会选择售货员的工作，因为这个工作的最佳特性（收入）是这些工作中最强的最佳特性。

2. **“最小”策略** 这种策略的特点，是要求对各候选者的最差特性进行比较，并选择那个具有最强的最差特性的候选者。如果人们利用这种策略来对上述三种工作进行选

择，那么，他们就会选择会计工作。因为会计工作的最差特性在这三种工作中是最强的。

3. **“合取”策略** 决策者在使用这种策略时，对每一维度设定某个最低的可接受值。如果某一候选者的所有维度上的值都超过了这最低的值，那么，它就成为决策者选取的对象。因而，在使用这种策略时，决策者可接受的候选者或许会多于一个；当然，也可能发生这样的情况，即没有一个候选者是可以被接受的。这完全依赖于决策者确立的标准。例如，把这种策略应用于上述选择工作的例子时，若决策者设定的标准为“每个特性的值至少是3”，那么，教师和会计这二种工作都是可以接受的；若决策者设定的标准是“每个特性的值至少是5”，那么，其结果将是没有一个候选者是可以接受的。

4. **“方面排除”策略** 这里所说的方面，可以是候选者的任何种类的特性。如录音机是否有电脑控制、集体旅游是否包括膳食、年薪是否超过5000元等等。人们在进行抉择的时候，根据所知道的重要性或爱好，以一定的次序对候选者的各个方面加以考察，一次比较一个特性。决策者把第一个方面选出来，针对各候选者进行比较，并把所有缺乏这一特性的候选者排除掉，以后不再考虑它。如果剩下来的候选者还在一个以上，就再考虑第二个方面，并用同样的方式来排除候选者。这种排除过程一直要继续到只剩下一个候选者为止。这最后一个候选者就是决策者选择的对象。在用这种方法或其他的一个接着一个地考察特性的策略时，考察特性的次序有可能影响抉择过程的结果。我们再以前面列举的假想的工作选择为例，如果决策者用规定为“至少是4的判断值”首先来考查“退休金”这一方面，那么，

售货员这项工作就会被排除掉，而剩下的实际上就是要对会计和教师这二种工作进行抉择。相反，如果决策者用“至少是4的判断值”首先来考察提升的机会，那么，教师这个工作就会被排除掉，剩下的只是要对售货员和会计这二种工作做出抉择。由此可见，在利用这类策略时，各种特性被考察的顺序，对选择的结果是有影响的。

5. “方面比较”策略 这种策略要求把各候选者直接进行比较，一次比较一个方面的特性。若某个候选者多数方面的特性优于其他候选者，那么，它就成为选择的对象。比如说，候选者甲有三个方面的特性优于候选者乙，一个方面的特性与候选者乙相等，还有一个方面的特性却劣于候选者乙。根据“方面比较”策略，应选择在多数方面占优势的候选者，因此，决策者将会选择候选者甲。

从这里我们也可以看到，在使用这样一些非补偿策略的时候，选出来的候选者不一定具有最高的总的预期价值或预期效用。如果我们把“合适的抉择”规定为具有最大的预期效用，那么，使用非补偿策略进行抉择的结果，显然与这一规定是不一致的。我们可以把这种不一致看作决策中的错误。但是，为什么许多人在进行决策时，不使用那种补偿的预期效用理论呢？这里有几种可能的原因。从根本上说，预期效用模型要求对各种不同结果的概率进行评估，要求对各种各样特性的值转换成可进行比较的效用尺度，只有这样才能得出每个候选者的预期效用。我们已经指出过，人们关于概率的直觉亦即主观概率常常与客观概率不一致，而且，要把一些非常不同的特性转换成共同的效用尺度，也是很困难的。事实上，人们往往喜欢那些用起来简易方便的决策策略。这种策略是易于理解的，并且能够提供明确的抉择，而又

不需要进行复杂的计算。非补偿策略倾向于能满足这些要求。此外，在许多情况下，利用非补偿策略也可以作出与预期效用理论相同的抉择。这就是说，基于非补偿策略的决策和基于补偿策略的决策，常常是非常类似的。所以，人们往往偏爱于使用非补偿的策略来进行抉择。

当然，我们把注意力集中在待选择的候选者时，也不能忽视另一个潜在的重要因素，这就是决策过程的代价问题。

（三）抉择过程的代价

决策理论似乎集中考虑决策的结果，也就是集中考虑被选择的和被拒绝的候选者的性质。但是，一个决策者不仅需要处理候选者，而且也要处理抉择动作本身。有人指出，决策者作出的抉择是要指导其将来的行动的，而选择本身即意味着要承担一定的义务。所以，这就意味着，选择本身涉及代价问题。人们在进行决策时，常常限制用于选择的时间和努力程度，其目的就是力图减少这种代价。

大多数人或许会相信，投入决策的时间和努力愈多，所作出的抉择将会愈好，也就是说，这种抉择将受益愈多。但是，事实上并不完全是这样的。决策的投资与决策的质量（预期的效益）之间的关系，到一定程度后，将出现报酬递减的现象。这就是说，在一定量的时间、努力或金钱已经投入决策的情况下，再增加决策的投资，其最后抉择的质量的提高，将出现逐渐减小的现象。在这种情况下，我们可以认为，最优的决策策略是扩大选择结果的效用与决策代价之间的距离。也就是说，我们应该力求为进行决策所花的代价要尽可能地缩小，而所作的抉择的效益要尽可能地大。运用这种策略来进行决策，将产生一个比不顾一切代价寻求的最佳

候选者要差一点的抉择。比如说，你要买一架录音机，你可用的录音机根据其特性可以排列成最好到最差的一系列等级。录音机A是最佳的，录音机B则次之，等等。你的任务是收集有关录音机的信息，并以某种方式对这些信息进行评价，以便选择一架对你合适的录音机。假定说，你若用10个小时的时间来收集和评价信息，将导致你购买录音机B，而要花100个小时的代价，你才会决定购买录音机A。如果只看最后的结果，那么，花10个小时的决策其结果不是最优的，因为你没有选择最佳的录音机A。但是，如果考虑一下决策性代价，考虑一下所付出的时间和心理上的努力，那么，10个小时的决策所作出的决定，完全有可能是一个最佳的决定。

著名认知科学家西蒙教授，对决策理论曾进行了创造性的研究。他从经济决策的角度提出了“比较满意”的决策原则。他认为，这是实际生活中进行决策的一个普遍的、指导性的策略或原则。我们知道，在经济、政治、科学等实践领域中，情况是相当复杂的。当人们在进行决策的时候，常常会发现很难找到一个最优的决定。但是，寻求一个比较满意的决定，看起来就省力得多了。问题的最优的解和问题的比较满意的解，这两者是不同的。如果我们不是去追求一个最优的解，而是限制在只是找一个令人满意的解，那么，我们就会发现，整个决策过程就大大简化，投入决策过程的时间和心理努力就会大大减少。一句话，决策的代价就会大大降低。西蒙教授曾举了一个很形象的例子来说明这一思想。一个人中午饿了，到地里去找一个玉米棒子吃。如果他不是力图去找一个最大、最嫩的玉米棒子，而是只找一个够他吃的、比较嫩的就行了，那么，问题的解决就容易多了，其所花的代价比找一个最大、最嫩的玉米棒子要小得多。反过来

说，若他非要找一个最大最嫩的玉米棒子来吃，那么，他所要花的时间就和玉米地的面积大小成正比。也就是说，玉米地的面积越大，他需花的时间就越多，其代价也就越大，因为他必须找遍整个玉米地后才知道哪个玉米棒子是最大最嫩的。但是，如果他想找的不是一个最大最嫩的玉米棒子，而只要达到使自己比较满意的程度就可以了，那么，寻找的时间就与玉米地面积的大小无关，其抉择的代价就会大大降低。

事实上，人们在决策时确实是考虑决策的代价问题的。比如说，在选择决策的策略时，人们会考虑到时间耗费的问题。

除了抉择过程本身的基本冲突以外，还有二个因素影响决策的难度。一是需要加以评估的候选者的特性的数量；二是所有特性在各候选者之间的变异性。当一个候选者的各种特性对其他候选者来说拥有压倒的优势时，人们要作出抉择是比较容易的。反之，抉择就比较困难。同时，一般来说，需要加以评估的候选者的特性愈多，各候选者之间要加以比较的特性愈不一致，那么，作出抉择的过程就会愈困难，其代价就会愈大。非补偿策略往往能够降低对候选者进行多方面比较的需要，所以，它可以缩小决策的代价。

不管怎么说，在给定的情境下，选择“最佳候选者”的收益，需要与作出一个决定所花的代价相平衡。这可以说是决策者进行决策时需要考虑的一个要点。

第十一章 人工智能和认知心理学

五十年代中期，由于世界科学技术的迅猛发展，出现了一门新兴的学科，叫做“人工智能”（Artificial Intelligence，简称AI）。人工智能科学在整个科学园地中处于一种令人瞩目的地位。有人把人工智能与空间探索和生命奥秘的探索一起誉为当今科学研究的三大前沿阵地。人们认为，人工智能的发展将对人类的生活方式以及人类认识自己的方式产生巨大而深刻的影响。

一、人的智能和人工智能

“智能”这个术语，也象大多数心理学术语一样，是从日常语言词汇当中吸取来的。科学家们对“智能”这个概念并无一致公认的定义。有的人强调，智能是一种学习的能力；有的人把智能与抽象思维的能力联系在一起；也有人认为，判断和推理的能力代表人的智能；或者是人利用自己的知识去应付新情境、解决新问题的能力；或者是人预见新问题和创造新的相互关系的能力等等。当然，这些对“智能”的不同规定，虽然都不够完善，但它们是彼此相互补充的，而不是相互矛盾的。

总而言之，人的智能是指人认识客观事物、获取新知识的能力以及运用已有知识去解决实际问题的能力，它表现为人的观察力、记忆力、理解力、想象力、判断力、思考力和创

造力等。从这里我们可以看出，智能不是指知识，它们是二个不同的概念。知识是存储在头脑中的有用的信息，而智能则是以某种有效的方式获取、保存和应用这些知识的能力。一个人如果接触了许多事物，但什么也记不牢，或者，一个人若记住了许多知识，但不会运用这些知识去解决问题，应付新的情境，那么，就不能说他有正常的智能。当然，从另一方面讲，它们又是相互联系和不可分割的。可以说，没有智能就不可能有知识和知识的利用；同样，没有知识也不可能进行智能活动。

人工智能则是相对于人的自然智能而言的。顾名思义，人工智能是一种人造的智能，是把人的某些智能赋予机器，让机器在功能上模仿并代替执行人的某些智能活动。下面引用几位著名学者对“人工智能”这一学科的定义：

温斯顿 (P. H. Winston) 在他所著的“人工智能”一书中说，“人工智能是研究如何使计算机去做过去只有人才能做的智能工作”，而“人工智能的主要目标，是要使计算机更加有用，以及理解那些使智能行为成为可能的原则。”

麦卡锡 (J. McCarthy) 在为“科利尔百科全书” (Collier's Encyclopedia) 写的“人工智能”词条中认为，人工智能是计算机科学的一个分支，从事于使机器行为智能化的研究。最终，为了使计算机程序具有人类水平的或较高的一般性智能，人工智能研究者希望充分地理解人的智能的实质。

尼尔森 (N. J. Nilsson) 则强调，人工智能是关于知识的科学——怎样表征知识以及怎样获得知识和使用知识的科学。（见“信息加工74”中“人工智能”一文）。

从这里我们可以看出，人工智能研究的是如何使计算机

具有人类的智能，使计算机象人类那样智能地工作，去完成那些需要人的智能才可以完成的工作，从而使计算机更加聪明、更加有用。从另一个角度也可以说，人工智能研究如何使人的智能机械化。为了实现这一宗旨，就需要深入了解人类智能的实质和机制。而人的智能这个问题，一向是心理学研究的重要内容，甚至可以说，在过去，它主要是由心理学来研究的问题。所以，从人工智能本身来看，它内在地与心理学有着密切的关系。

那么，计算机的工作怎样才算是一种智能行为呢？人工智能的先驱者英国数学家图灵（A. M. Turing）从工程的观点出发，对机器的行为是否属于智能行为提出了一种规定。他认为，如果机器在某些规定的条件下，能够非常好地模仿人回答问题，以致使讯问者在一定的时间内误认为它不是机器，那么，就可以认为机器的行为是属于智能的，或者说机器是能思维的。这种考核机器智能的规定，普遍地称之为“图灵试验”，如图11-1所示。

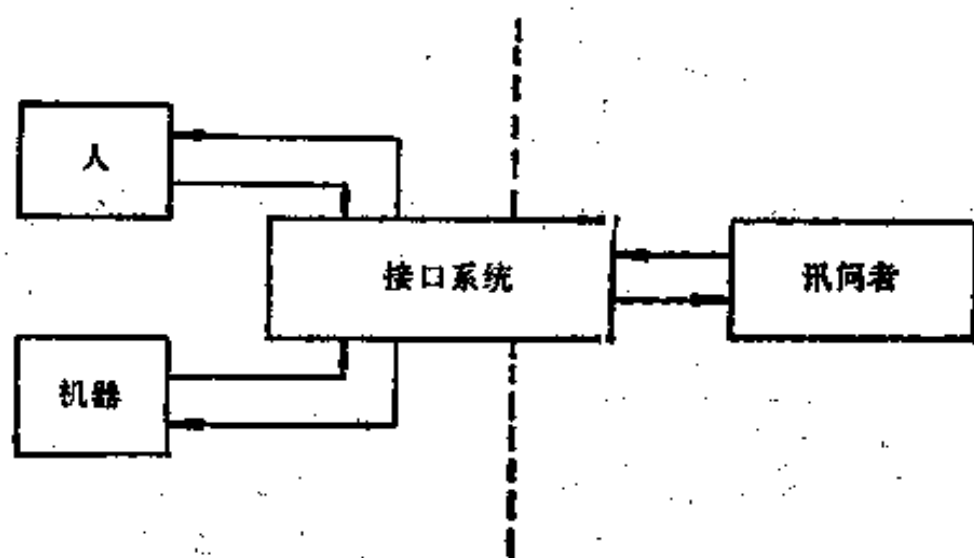


图11-1 图灵试验

在图灵试验中,问讯者坐在一个房间里,一台机器和一位问题回答者隐蔽在隔壁另一个房间里。问讯者向他们提出问题并要求回答。如果问讯者分辨不出是人还是机器在回答他的问题,那么,就可以认为,机器具有了人那样的智能。图灵试验的思想是他在1950年提出来的。这种思想强调从功能上考察计算机是否具有人的智能,根据机器工作的最后结果或外部行为,评定计算机能否思维。它并不考虑人如何思维和计算机是否象人类那样思维。因此,这里所涉及的是计算机模仿人的思维的外在表现,而不是模拟人的思维的内在过程。图灵曾预言,在50年内,计算机可以模拟人的思维行为达到这样的程度,使70%的问讯者在5分钟之内,分辨不出回答他问题的是人还是计算机。当然,问讯者为分辨谁是机器谁是人所需的时间越长,说明计算机的智能水平越高。

不管怎样,人们普遍认为,人工智能和人类智能的研究是紧密相关的。人工智能研究的性质决定了它与认知心理学的密切关系。由纽厄尔和西蒙在五十年代发展起来的信息加工心理学与行为主义的“刺激—反应”理论不同,认为在机体的输入和输出之间的内部过程是相当复杂的,提倡用编写计算机程序这种手段来模拟人的复杂的解题行为。人工智能和信息加工心理学都必须确定,人在解决各种类型的问题时,需要什么样的智能机制。

二、人工智能的诞生及其研究范围

(一) 人工智能的诞生

英国数学家图灵在1950年写了一篇文章,名字叫“机器

能思维吗？”。这可以说是第一篇关于人工智能的文章。而第一个职业性的关于人工智能研究的小组，是由纽厄尔和西蒙于1954年在卡纳基·梅隆大学建立起来的。这二件事为人工智能的诞生作了理论上和组织上的准备。

但是，人们一般认为，人工智能是在1956年正式诞生的。

这一年的夏天，由美国数学家麦卡锡发起，邀集了一些有志于使机器行为智能化的学者，在美国达特默思学院开了一次讨论会。参加这个会议的有数学家、计算机科学家、心理学家、神经生理学家及信息论等方面的专家。这次讨论会历时二个月，其任务是，从不同学科的角度来探讨使机器行为智能化的可行性，探讨人类各种认知活动的特点及其基础，并研究如何在原理上进行精确的描述，以便于计算机模拟。这是一次人工智能的创始性的讨论会。在会议给有关部门的文件里，第一次使用了“人工智能”这个术语。从此，这门学科有了自己特有的名称。目前，美国已形成的人工智能研究的三个中心，即麻省理工学院、卡纳基·梅隆大学和斯坦福大学，主要是由当时的与会者们组建起来的。

同时，纽厄尔、肖（J. C. Shaw）和西蒙，也在这一年研制成功了能证明符号逻辑定理的计算机程序，叫做“逻辑理论家”（Logic Theorist，简称LT）。这个程序模拟了人类被试者证明符号逻辑定理的思维活动过程，第一次就把怀特海德（A. N. Whitehead）的数学名著“数学原理”第二章中的52条定理证明了32条，后来又作了进一步修改，证明了第二章中的全部定理。这个程序是计算机模拟人类高级心理过程的第一个收获，也是第一个成功的实验性人工智能系统。因此，人们把1956年看作是人工智能的真正开端，是有道理的。

LT程序的研制成功，有以下几点重要的贡献。第一，它首次成功地模拟了人类解决复杂问题的思维活动，使计算机象一个没有学过数理逻辑的大学生那样来证明数理逻辑定理；第二，建立了第一个启发式程序，这对后来的人工智能的研究有重要的影响。启发式 (heuristics) 的原意，是指发明的艺术，或帮助发明和发现的技巧。它不是一种理论推导，而是从生活经验中得到的一些解决问题的规则、技巧或窍门。把启发式规则编进计算机程序，这种程序就叫做启发式程序；第三，研制并应用了符号信息处理语言 IPL (Information Processing Language)。这一语言是现在人工智能研究中普遍应用的LISP语言的前身。

LT也是现代认知心理学的真正开端。它是信息加工理论在心理学研究中的具体化；它表明，启发式解题程序的作业与人类解决问题的作业是一致的；它为用计算机模拟本身可以作为一种重要的心理学理论形式树立了一个范例，是心理学研究方法论上的一个突破。所以，我们可以说，现代认知心理学和人工智能是同时产生的。

(二) 研 究 范 围

人工智能是一门新兴的正在发展中的学科，关于它的具体研究内容究竟包括哪些方面，还不能说有一种统一的看法。尼尔森在“人工智能”一文中描述了人工智能研究的主要内容及它与其他学科的联系（见表11—1）。他把人工智能研究分成二个部分，一是属于人工智能本身的方法学和技术的基础性研究；二是属于应用性的研究。基础性研究包括四个方面：1、启发式搜索；2、常识推理、演绎和问题求解；3、知识的模型化和表征；4、人工智能的系统和语言。

应用性研究包括八个方面：1、博弈（如下棋）；2、计算机辅助设计；3、数学定理证明；4、自动程序设计；5、机器人研制；6、机器视觉；7、自然语言的理解系统；8、信息加工心理学。

从这里我们可以看到，心理学和人工智能有着极其密切的关系。尼尔森甚至把信息加工心理学纳入了人工智能研究的范围。当然，尼尔森对人工智能研究内容的看法，不能说是大家一致公认的，但大多数学者基本上是接受的。

巴尔（A. Barr）和费肯鲍姆在由他们主编的《人工智能手册》（1981年）中，把大多数人工智能研究的课题归纳为以下几个方面。

1. 问题解决 包括能够解谜题和下棋之类的计算机程序。人工智能研究的第一个巨大成就正是在这一方面取得的。现在，这方面的程序，已能下竞赛水平的西洋跳棋和十五子棋，也能下非常好的象棋。

2. 逻辑推理 已经建立了一些程序，它们能够通过事实的数据库的处理，证明某些论断。只要最初的事实正确，那么，程序就能证明根据这些事实得出的所有定理。证明定理是逻辑演绎的过程，机器证明定理就是把人证明定理的过程通过一套符号体系加以形式化，变成一系列能在计算机上实现的符号运算过程，把推理演绎过程机械化。

3. 自然语言处理 这个领域始终吸引着人们的兴趣。现在已经研制成功了一些程序，它们能够回答用人类自然语言提出的问题，能够把一种语言的句子翻译成另一种语言，能够领会以自然语言给予的指令去执行某些操作。有些系统能理解人说出的自然语言指令。尽管这些自然语言的处理系统离开人的水平有相当的距离，但是，它们可以满足某种范。

有围限的特殊用途。

表11-1 人工智能研究的范围

研 究 范 围		有 关 的 学 科
基 础 性 研 究	启发式搜索	图论、现代控制论、运筹学
	常识推理、演绎和问题求解	心理学、逻辑学
	知识的模型化和表征	心理学、认识论
	人工智能的系统和语言	计算机语言、系统程序编制
应 用 性 研 究	博 奕（如下棋）	
	计算机辅助设计	管理科学、符号运算、图示学、有关的数学和科学分支
	数学定理证明	逻辑学、有关的数学分支
	自动程序设计	逻辑学、系统程序设计、算法分析、计算理论
	机器人研制	控制论、空间探索、工业自动化
	机器视觉	光学、心理学、图示学、模式识别
	自然语言的理解系统	心理学、语言学、语音学、声学
	信息加工心理学	心理学、控制论、语言学、逻辑学

4. 自动程序设计 它研制的系统能够根据有关某个程序的目的是各种说明, 写出相应的计算机程序, 对自动程序设计的研究, 既可以产生半自动的软件研究系统, 也可以产生

那种通过修正它们自身的编码来达到学习的人工智能程序。

5. 计算机学习 虽说人类智能的一个突出的和显著的特点是学习的能力,但我们对人的这种认知活动了解得很少,所以,在人工智能研究中这方面取得的进展也最小。当然,现在也有一些程序具有学习的能力,包括“例中学习”的程序和“做中学习”的程序,以及根据给予的指令来学习的程序。

6. 专家系统或知识工程 这是人工智能研究中一个重要的子领域,是人工智能显示其应用价值的主要途径,它给人工智能的发展以巨大的推动力。专家系统是把某个领域的专家的专业知识及解决专业问题的经验编进计算机程序,构成一个能在一定程度上替代专家解决专业问题的计算机系统。人们在与专家系统打交道时,就象与某个领域的专家打交道一样,能够从它那里得到专业性的咨询和指导。当前已建立的实验系统,包括化学和地质学的数据分析以及医疗诊断等方面的专家系统,在咨询服务方面,已达到相当高的水平。

7. 机器人学和视觉 现在工业生产中使用的成千上万个机器人,依然是一些简单的机械装置。它们中有些是由程序控制去执行某种重复性的任务,叫做“重复型”机器人;还有一些属于“远距离操纵型”机器人。它们在恶劣或危险的环境下进行作业(如海底作业和清除核废料等),但由人在安全的地点进行远距离操纵。它们没有感知觉。而人工智能要研究的是“智能机器人”,这是一种具有感觉(视觉、甚至触觉)、识别和规划、决策能力的机器人。当然,关于机器人学的研究,是一种综合性的研究,它几乎涉及人工智能研究的各个方面。对视觉信息进行加工也是人工智能中的一个重要领域,现在已研究成一些程序,可以识别各种客体和分辨图象的微小变化。

8. **系统和语言**：这是属于人工智能本身的新工具的研究。

当然，这是他们根据以往大多数人工智能的研究课题归纳出来的。因此，这不等于说人工智能研究的内容就是或者只能是这八个方面。

三、研究的途径和理论前提

(一) 三 条 途 径

上面介绍了人工智能研究的内容，从这些研究内容中，可以看到人工智能与认知心理学之间的密切关系。另一方面，从实际情况来看，围绕着上述内容，人工智能的研究是沿着不同的途径来展开的。

人工智能是一门综合性的边缘科学。它是由计算机科学、数理逻辑、控制论、信息论、心理学、神经生理学、语言学等许多学科相互渗透而形成的。这些领域的学者们在从事人工智能研究时，由于其自身的专业知识背景不一，强调的侧面各异，所以，他们沿着不同的途径来进行研究。具体说，人工智能研究中实际上存在着三条不同的途径，即仿生学途径、心理学途径和工程技术途径。

1. **仿生学途径** 沿着仿生学途径进行研究的人们，主张从生物学的角度来直接模拟动物或人的大脑及各种感觉器官的结构和功能，力图寻找一种可以描述自然神经系统的方法，建立神经生理学模型。许多学者沿着这条途径做了不少工作，在人工智能发展的前期取得了一定的成绩。

数理逻辑学家麦卡洛克 (W. McCulloch) 和神经生理

学家皮茨 (W. Pitts) 在1943年曾提出第一个神经元模型。他们认为, 任何计算的功能都可以由一个组织适当的理想神经元的网络来实现。这种理想神经元可以是一种阈限逻辑元件, 但它与实际神经元的性质是相似的。也就是说, 它的活动遵循全或无定律; 神经元之间有二种联系方式, 即兴奋性联系和抑制性联系, 后者具有“否决权”; 每一神经元以一定的正整数值作为阈值, 当输入超过这一阈值时, 它就产生一个恒定的输出, 等等。问题是需要网络有一个合理的重新组织的原则, 以使本来是随机联接的理想神经元由它们自己组成一个计算装置。这种重新组织的原则是一个学习过程。

。后来, 1948年, 加拿大心理学家赫布 (D. O. Hebb) 又提出了“细胞团”学说。他认为, 皮质细胞本来是一些随机单元, 当动物或人与环境接触以后, 外界特定刺激 (如一个图形) 可以引起皮质一定部位的一团神经细胞的活动。这一图形刺激的反复出现, 使这一细胞团逐渐形成一种固定的网络, 这样就达到了学习的效果。细胞团学说是一种神经学的学习理论。

1958年, 罗森布拉特 (F. Rosenblatt) 在这二个理论的基础上提出了“感知机” (perceptron)。感知机是脑的信息存储和组织模型, 实际上是一个关于学习的神经模型。它模仿了三层细胞, 即S层 (感受细胞)、A层 (即联络细胞) 和R层 (即反应细胞)。它的实验系统包括一个主机 (感知机), 一个刺激环境和一个强化控制装置三大部分。主机本身又分为三个部分, 即感受单元、联络单元和反应单元。感受单元把外界的某一物理量变成相应的信号; 联络单元则对该信号进行加权兑和及阈值运算以巩固正确反应, 淘汰错误反应; 反应单元则对结果进行分类显示。它通过学

习，可以分辨不同的简单图形，或者，通过几小时的学习，可以分辨潜艇和海豚发出的声响。

关于感知机的研究，在60年代中期以前，曾盛极一时。据说，全世界曾有一百多个组织围绕感知机进行研究。这些研究，不仅在图形辨认方面有一定的贡献，而且对学习机的研究工作也有所促进。但是，人们逐渐发现，这类机器对简单图形的学习辨认效果还可以，但对复杂的图形却无能为力，因为图形中包含的信息量是巨大的，若不考虑预处理工作，特征抽取的加工以及各特征之间的“关系”等问题，是很难有正确的学习和辨认的。现在，感知机的研究已经无人提倡了，但它在人工智能的发展史上，特别是在图形辨认研究的初期，起了很大的鼓舞作用。

60年代后期，日本结合脑的功能也展开了一些研究。中野在1969年提出了模拟人的联想功能的“联想机”，它可以根据颜色和形状认出物体，也可以反过来根据物体读出颜色和形状。福岛根据哺乳动物视觉系统的结构和原理，建立了图形特征抽取的数学模型，并在计算机上进行了模拟。后来，他又设计了一个自组织（即系统根据环境情况并按照一定的指令或规则，对内部结构和功能进行调整）的多层网络模型，称为“认知机”。

但是，人脑是一个异常复杂的组织。人类各种形式的复杂的心智和行为，都是人脑活动的产物。大脑两半球皮层的厚度只有2—5毫米，但却包含着大约140亿个神经元。神经元的类型很多，但有组织地分成六层，通过树突和轴突组成极其复杂的神经网络。它们是人类智能活动的中心，一切科学发现和创造都是在这里孕育出来的。然而，大脑的任何一种功能，既不是可以由个别神经元的功能来实现的，也不

是许多神经元的功能的简单相加，而是由成千上万个神经元按特定方式组织起来的整体的功能。人脑是如此的复杂，在现今的科学技术条件下，近期内要搞清它的结构和活动的机理是不可能的。所以，沿着仿生学途径进行的人工智能研究，遇到了一定的困难。

2. 心理学途径 主张从这一途径来研究人工智能的不仅是心理学家，也有不少计算机科学家。他们相信，“世界上最聪明的是人”，要使机器的行为具有人的智能特征，最好的办法是了解、研究人，然后模拟人的智能行为。但是，由于现代神经生理学能提供给我们的关于人脑及其机能的知识还很少，从神经元的活动到逻辑思维这中间有很长的一段是不清楚的，所以他们主张暂时撇开脑的内部结构和机制，采用心理学模型。沿着心理学途径来研究人工智能的人们，应用实验心理学的方法，搜集和考察人在各种情境下解决各类问题时的言语素材和行为表现。然后，他们对这些实验材料进行认真的分析研究，了解人在知觉、记忆、理解或解决各种问题时采用什么样的步骤和策略，有什么窍门或捷径，怎样规划自己的行动，怎样进行推理等等。在总结出一些规律的基础上，提出心理学模型，然后用计算机进行模拟，达到使机器表现出智能行为的目的。这种模型最初可能是粗略的，计算机模拟的结果可能和人的行为并不很一致，但经过模拟，修改，再模拟，再修改，可以使计算机的作业与人的行为逐渐接近。心理学途径有时也叫“启发式途径”。通过这种途径设计的计算机程序，就是启发式程序。前面所介绍的LT程序，就属于这一类。又如，纽厄尔和西蒙研制的GPS程序，也是首先请大学生当被试者做了心理学实验后，才在计算机上进行模拟的。

心理学途径不仅在人工智能的开创时期起了奠基性的作用，而且，在国际人工智能研究中，现在仍然是颇有生气，很受人重视的。

3. 工程技术途径 主张这一途径的人主要是电子技术和计算机科学的专家。他们认为，只要创造出某种方法，能使计算机解决那些需要人的智能才能解决的任务，这就达到了使机器行为智能化的目的。至于机器的作业过程与人在解决同样任务时的心理操作过程是否一致，暂时可以不去考虑。例如，在机器模式识别的研究中，人们利用计算机处理来自人造卫星的遥感信息，识别地球表面的景物乃至各种资源，利用计算机识别细胞是否发生癌变，识别指纹、文字、图形、语音等等。他们在研制这类计算机系统时，不一定考虑人的识别过程，事实上，其中有些系统的识别过程，显然与人的识别过程是不一致的。但它们的外部行为却表现出人那样的智能。这一途径所受的限制比较少，灵活性比较大；沿着这一途径进行人工智能研究的队伍相对来说也比较庞大；因此从目前来看，采用这一途径的研究工作还是比较活跃的，也取得了不少的成就。

（二）物理符号系统

三十年前，纽厄尔和西蒙等人提出了“物理符号系统”的假设。他们认为，人是一个信息加工系统，信息加工系统也就是“符号操作系统”，或者叫做“物理符号系统”。物理符号系统的基本功能，就是对各种符号进行操作，也就是对各种符号及符号之间的联系进行比较、辨认等等。他们认为，一个完善的物理符号系统具有以下六种功能。

1. 输入符号 符号通过各种途径输入系统中去。例如，

人能够通过眼睛看、耳朵听，用手触摸物体等不同方式输入符号；计算机也能输入符号，用纸带打孔输入或用键盘打字输入等。

2. 输出符号 人说话、写字、面部表情乃至躯体某些部分的动作，实际上都属于符号的输出。计算机的输出是通过显示器显示出来或通过打印机打印出来的。

3. 存储符号 人类可以把各种输入符号保存在头脑里，人的记忆就是存储各种符号及符号结构的仓库，而计算机则可以用各种不同的方式存储信息，或者存储在磁带和磁盘上，即存储在外存储器里；或者存储在由晶体管或集成电路构成的内存储器里。不管是外存储器还是内存储器，它们都能接收和保存信息，而且能根据指令提供需要使用信息。

4. 复制符号 人可以认出“智能”二个字，并把它们复制下来，存储在记忆的某个地方。计算机则可以把符号从一组存储单元复制到另一组存储单元，从内存储器复制后送到外存储器存起来，或者相反。

5. 建立符号结构 通过找到各种符号之间的关系，在符号系统中形成各种各样的新的符号结构。人通过学习，对符号进行不同的组合，得出新的关系。例如，把“人工”和“智能”组合在一起，形成一个新的符号结构“人工智能”。同样，计算机在对符号进行操作运算的过程中，将会建立各种新的符号结构。

6. 条件性转移 一个物理符号系统，在原有的、存储在记忆中的符号结构的基础上，可以根据当前的输入进行并完成一系列活动。“条件转移”的意思是说，如果满足了某个条件，就执行（包括发动和停止）某种活动，或者说，若

条件A得到满足，就执行活动B。反过来说，若条件A未得到满足，就不执行活动B。比如说，给一个演员的指令是：

“听到锣声，先迈你的左脚，后迈右脚，每3秒钟跨一步，到戏台的边沿就停止”。在这里，“听到锣声”和“到戏台的边沿”都是条件。演员根据指令中给出的条件，锣声一响就有节奏地走步，待到戏台边沿时就停止往前走了。若没有这种条件转移的能力，那么，演员就要跌到戏台下面去。这种条件转移是计算机进行作业的重要特点。由于计算机具有这种条件转移的能力，使它能完成各种各样的极其复杂的作业。

总之，无论是人还是计算机，都具有物理符号系统的上述六种功能，都是完善的物理符号系统。纽厄尔和西蒙在这个基础上提出了一个叫做物理符号系统的假设。这个假设的意思是说，任何一个系统，如果它的行为表现出智能的话，就必然具有上述六种功能；反过来说，任何一种系统，如果能执行这六种功能，那么，它的行为就会表现出智能的特征。

他们从这一假设出发，又作出了三个附带的推论。第一，既然人具有智能，它就一定是个完善的物理符号系统。他们认为，人类的智能就是在人对符号和符号结构的操作过程中，在对信息的加工过程中体现出来的。第二，既然计算机是一个完善的物理符号系统，它具有上述六种功能，所以，计算机的行为就能够表现出智能。这一推论，可以说是人们进行人工智能研究的基本前提。第三，既然人是一个完善的物理符号系统，计算机也是一个完善的物理符号系统，它们的工作的基本原则是一致的，所以，我们完全可以用计算机来模拟人的认知过程，包括对复杂问题的解决过程。也就是说，

我们可以按照人类心理活动的操作过程来编制计算机程序，把人类记忆、理解过程的特点和解决问题的步骤、策略和经验编进计算机程序中去。这样，一方面模拟了人类的认知活动，可以进一步深入了解人类认知活动的特点、规律。因为用计算机来模拟人的认知活动，就可以利用计算机程序形式化地描述人的认知过程，建立人类各种认知过程的模型或理论。另一方面，用计算机来模拟人的认知过程，将会使机器的行为表现出人类那样的智能。这也就是上面提到的研究人工智能的心理学途径，属于人工智能研究的一个重要组成部分。

四、若干人工智能系统介绍

人工智能从诞生到现在，已经走过了32个年头。在这三十多年里，它克服了前进道路上的困难，取得了非常可观的成就，建立了不少人工智能系统。下面扼要地介绍几个例子。

（一）语言理解系统

1、SHRDLU系统 由麻省理工学院温诺格拉德（T. Winograd）研究的SHRDLU系统，是一个用来指挥模拟机械手臂进行操作的自然语言理解系统。虽然这些操作是在一个小的空间中进行的，但是，它所处理的上下文及逻辑关系是比较复杂的。使用者通过英语，指挥机械手臂操作一些积木块。这些积木块的形状、大小和颜色是不同的，如图11-2。下面是使用者与系统对话的一些例子。

、对话（1）：PICK UP A BIG RED BLOCK.

(捡起一块大的红色木块。)

OK. (是)

对话 (2) : GRASP THE PYRAMID (抓住那块棱锥体。)

I DON'T UNDERSTAND WHICH PYRAMID YOU MEAN. (我不明白你指的是哪一块棱锥体。)

对话 (3) : FIND A BLOCK WHICH IS TALLER THAN THE ONE YOU ARE HOLDING AND PUT IT IN THE BOX. (找一块比你现在拿着的更高的木块, 并把它放在箱子里。)

BY "IT" I ASSUME YOU MEAN THE BLOCK WHICH IS TALLER THAN THE ONE I AM HOLDING. OK. (我想, 你说的“它”是指那块比我正拿着的木块更高的木块。是。)

对话 (4) : WHAT DOES THE BOX CONTAIN? (箱子里装了些什么?)

THE BLUE PYRAMID AND THE BLUE BLOCK. (蓝色棱锥体和蓝色木块。)

对话 (5) : WHAT IS THE PYRAMID SUPPORTED BY? (这块棱锥体由什么东西支撑着?)

THE BOX. (箱子。)

对话 (6) : WHY DID YOU CLEAR OFF THE

RED BLOCK? (为什么你把红木块清出来?)

TO PICK IT UP. (为了把它捡起来。)

对话(7): WHY DID YOU DO THAT? (你为什么这样做?)

BECAUSE YOU ASKED ME TO.

(因为你要求我去做。)

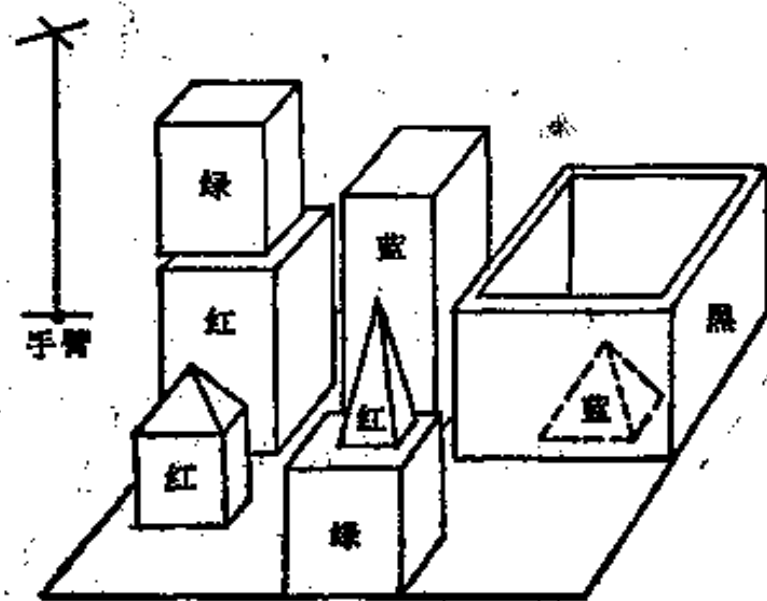


图11-2 SHRDLU指挥模拟机械手臂操作的积木世界

在对话(1)中,使用者要它拿起一块大的红色木块。因为这样的木块在这里只有一块,所以,它知道应该拿哪一块。但是,问题是在这大的红木块上,还有一块绿木块。所以,它必须把机械手臂移到这绿色木块上,并把绿木块提起来,放在空的地方。然后,机械手臂再移回来,抓住红木块,并把它提起来(如图11-3所示)。最后,用一个简单的词“OK”(是),向使用者作出反应。

在对话(2)中,使用者要求它抓住那块棱锥体。但是,这里有二块棱锥体,而且,在问题的上下文里没有提供应该抓

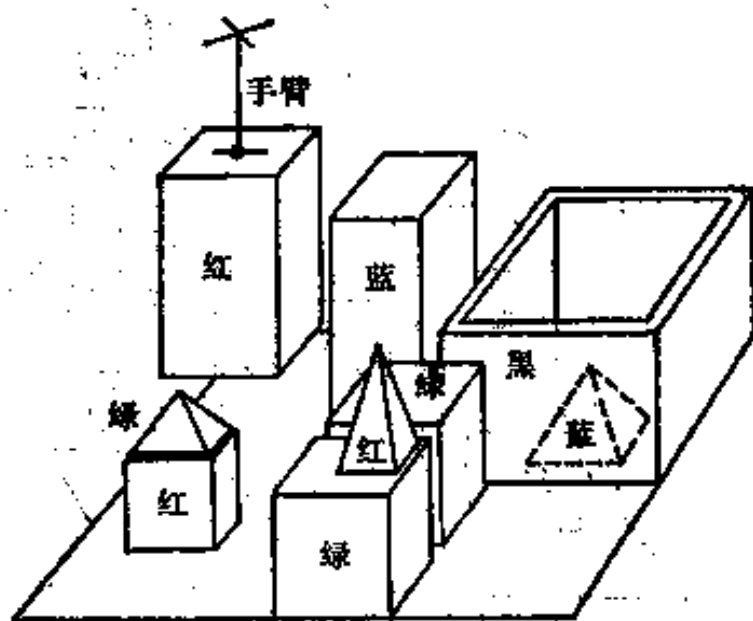


图11-3 SHRDLU系统执行“捡起一块大的红木块”指令的过程。

哪一块的线索。所以，SHRDLU明白，自己不能确定“棱锥体”这个词所指的具体对象，并把这情况告诉使用者。

对话(3)表明该系统处理较复杂的复合句的能力，也表现出处理句子的比较结构的能力，以及处理较复杂的代词的能力。它理解了使用者的意思，指挥模拟机械臂放下了正抓着的红色方木块，并抓起那块更高的蓝色木块，把它放到旁边的箱子里（如图11-4所示）。

对话(4)表示了系统的问答能力。在回答这个问题的过程中，它根据语言的上下文线索，告诉使用者，在箱子里有蓝色棱锥体和蓝色木块。在对话(5)中，使用者再次提到“棱锥体”这个词。现在，系统能够确定这个词所指的对象，因为在刚才的对话中已经提到过它了。对话(6)表明，该系统具有初步的处理动作参照的能力，能回忆以前发生的一系列动作，以及这些动作之间的因果关系。这是一种重要的技能。

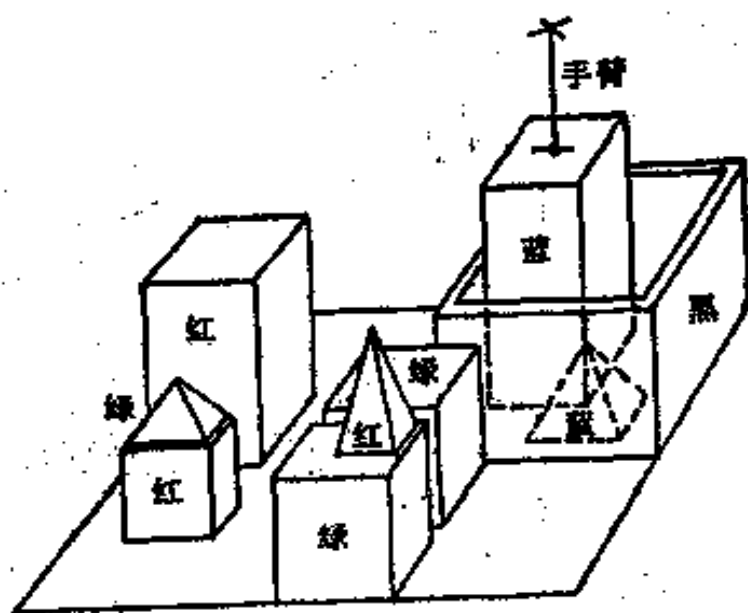


图11-4 在执行“找一块比你现在拿着的更高的木块，并把它放在箱子里”指令后的状态（出现在图11-3的动作之后）

当然，在SHRDLU的记忆里，存储着这些不同木块当时所处的位置以及它们之间关系的知识。这些知识是随着木块的移动而不断地改变着的。

2. CLUS-1系统 CLUS-1 (Chinese language understanding system 1 的缩写) 是第一个机器理解汉语的问答系统，是由心理所李家治和陈永明等研制的。该系统的作业是模拟一次动物常识课的师生对话。理解语言是离不开知识的。这个系统采用的知识表征方法，是略加改进的语义网络模型。

CLUS-1系统理解人提出的问话和组织答话主要通过以下四个步骤。

(1) 判别句型；(2) 分析句子；(3) 查询知识库；(4) 组织答话。

判别句型主要是判定问句的类型。本问答系统中使用的

问句有“一般疑问句”、“加重疑问句”等6类，共有15个不同的句型。分析句子是要确定问句中的主词、宾词，以及把主词和宾词联系起来的关系词。例如，“麻雀有翅膀是不是？”属于一般疑问句。通过对句子的分析，确定它的主词是“麻雀”，宾词是“翅膀”，要确证的是这两者之间是否存在一种“有”的关系。查询知识库就是对语义网络进行搜索。除了通过直接的搜索来回答人提出的问题以外，该系统还具有一定的推理能力。

(1) 概括的能力，可以从具体事例中概括出共同的特征，对一个较一般的概念作出一个粗略的定义。例如，问“什么是鸟？”在这里，“鸟”是一个较一般的概念，在机器中并没有关于鸟的现存定义。如果仅仅是依靠信息的提取来回答问题，结果将是列举出各种具体的鸟名。这样的答话不是问者所期望的。CLUS-1能够从麻雀、燕子等少量的实例中抽取出它们的共同特征，然后根据一定的文法型式构成句子，回答问者的问题。这样所作出的答话，可以说是对鸟下了一个粗略的定义：鸟有羽毛，有翅膀，有二条腿，会下蛋，是卵生动物。

(2) 依靠部分信息进行推理。在网络中，鸟的实例只有少数几个。按照通常的语义网络模型，如果问题中出现某个在网络中没有的鸟概念，系统就无法理解和回答这个问题。为了在一定程度上处理好这种情况，在CLUS-1的记忆网络中设置了一个待填的节点“X”。这样，如果问这个系统：“鸽子是鸟，它有羽毛吗？”这个问题，尽管在系统的记忆中并没有“鸽子”这个概念，但问话中已给出了部分信息，即它是鸟，所以，它可以填入“X”节点。这时，系统所要做的工作，就是要考察记忆中的各种鸟是否都有羽

毛。如果记忆中的各种鸟都有羽毛，系统将回答说，鸽子是有羽毛的；如果大多数鸟有羽毛，系统将回答说，鸽子很可能有羽毛，如此等等。

(3) 依靠间接信息进行推理，如果问一个成年人“马为什么不会飞？”他可以立刻回答说，因为马没有翅膀。人们从儿童时代开始，经过学习或经验积累，把翅膀和飞这两个概念联系在一起了。如果向机器提出这个问题，它怎么来处理这个问题呢？在这种情况下，CLUS-1 将间接地从会飞的动物中寻找信息，去查询那些会飞的动物是用什么东西来飞的。它发现这些动物都是用翅膀来飞的。于是，系统就把“翅膀”这个概念取出来，并转而去核查马是否有翅膀。结果，它发现马并没有翅膀，所以，它推论出马之所以不会飞，是因为它没有翅膀，并以此作为对上述问题的回答。

另外，在分析和组装句子时，遵循若干规则，例如：

(1) 分解和填补规则。对包含复数主词的输入问句。系统在分析时将把它分解为几个单数主词，并填入相应的关系词，宾词和疑问词。如果问：“麻雀、燕子和鸵鸟都是鸟吗？”它将被分解为：麻雀是鸟吗？燕子是鸟吗？等。若问“凡有翅膀的都是鸟吗？”系统将首先确定哪些动物是有翅膀的，然后把这些动物的名称填进去，结果将构成：麻雀是鸟吗？鸵鸟是鸟吗？蝙蝠是鸟吗？等等。

(2) 合并删除规则。主要用于组装答话，避免词的重复出现。例如，若 A_1 是鸟， A_2 是鸟， A_3 也是鸟，那么，其答句将是： A_1 、 A_2 和 A_3 都是鸟。

(3) 替换规则。单数主词与代词“它”，复数主词与代词“它们”，彼此可以相互替换。如果在问“麻雀、燕子和鸵鸟都有翅膀吗？”以后，接着就问“它们都会飞吗？”

这时，就需要用前面提到的那些具体的主词来替换“它们”这个代词。而在组织答话时，其过程往往相反，用代词来替换具体的主词，从而使答句显得更精炼和合乎常情。

(二) 专 家 系 统

1. 化学专家系统 美国斯坦福大学费肯鲍姆等人研制的化学专家系统DENDRAL,是用来确定有机化合物分子结构的。它能够利用化学分析家的专门知识,根据化学分析所得到的质谱数据,推断出化合物分子结构的细节。DENDRAL的作业分为三个部分:

规划:根据质谱图推导出一个约束表,其中既包括在最后的分子结构中应有的子结构,也包括不应该出现在最后的分子结构中的子结构。

生成:系统内部存储的产生式规则,根据给予的化学分子式生成各种可能的结构。这些结构中不包括那些被禁止的子结构。同时,把每一可能的结构预测出一张质谱图来。

测试:把实验所得的质谱图与上面预测的每一质谱图进行比较,如果它与某一预测的质谱图匹配得最好,那么,该预测的质谱图所代表的分子结构,就是当前实验的化合物的分子结构。

DENDRAL的作业水平达到了化学专家的水平,并且已在实际工作中得到应用。

据报道,在这个基础上发展起来的Meta-DENDRAL系统具有学习的能力,它可以通过正、反示例自动归纳出有关质谱的规则。同时,它还具有一定的创造性,曾经发现过当时化学家还不知道的新的质谱规则。

2. 医疗专家系统 由肖特利弗(E. H. Shortliffe)

设计的MYCIN系统是一个著名的医疗专家系统。这个用于医疗诊断的计算机程序，专门诊治血液感染病，并为医生提供用抗菌素来治疗疾病的方案。只要把病人的病史、症状和化验结果输入计算机，MYCIN就能够根据内部存储的医学知识（约300余条产生式规则）进行推论，作出诊断，开出药方。它还可以用英语与医生进行对话，回答有关的问题。比如说，如果医生对MYCIN作出的诊断或给出的药方有疑问，不清楚它为什么会得出这种判断或提出那种医治方案，那么，医生就可以要求MYCIN回答它自己的思路。例如，医生可以问它：“你为什么会得出这个结论的？”它就会向医生列举得出这个诊断时所经历的一步一步的推论过程。也就是说，它能记住或“理解”自己的行为，知道是怎么得出某种结论的。它成功地应用了逆向推理的控制策略。无论从系统的原理、结构以及性能上看，MYCIN都是相当成功的，但是在美国至今还没有在临床上应用它。从这里可以看到，一个专家系统是否能在实际中使用，不仅取决于它的技术性能，而且也依赖于社会效应及用户是否乐于接受等各种因素。

后来，在MYCIN系统的基础上又发展起来了EMYCIN (Empty MYCIN)。后者是一个与具体医学领域无关的系统，是把前者包含的具体医学知识抽去而留下其控制结构和产生式规则的格式的基础上构成的。EMYCIN是一个“建构性的专家系统”，也就是说，用户只要把具体领域的知识填入EMYCIN，就可以构成用于各种不同领域的专家系统。

我国自己也研制成功一些医疗专家系统，包括中医的和西医的专家系统。例如，科学院成都计算机应用研究所和成都中医学院等单位共同研制的“中医计算机诊疗系统”，能

够诊治痹症、胃病、脱发、妇科黑瘰、儿科腹泻等九种疾病，具有中医辩证施治的特点；著名中医关幼波诊治肝病的计算机程序；清华大学已研制成了许多个诊治不同疾病的西医医疗专家系统，它们已在不同程度上投入门诊使用。

其他方面也有一些专家系统。例如，斯坦福研究所杜达（R. O. Duda）等人研制的PROSPECTOR（勘探者），是一个用于地质勘探的计算机咨询系统。它吸收了有经验的地质勘探专家的知识，模拟这些专家的推理过程。它可以根据岩石标本及其他勘探数据，对可能的矿藏进行预测；必要的时候，它可以提醒用户为作出一个肯定的结论还需要什么样的数据资料。美国曾用这个系统对华盛顿州的某个钼矿进行预测，后来的实际钻探的结果完全证实了PROSPECTOR的预测。

任何一个专家系统，都是以某个专门领域的丰富的知识以及这些知识的利用为基础的。怎样获取人类专家的知识以及他们利用这些知识去解决问题的推理技巧，并把它们组织到计算机里去，这是专家系统研制的关键所在。

（三）博 奕

我们中不少人比较熟悉的下棋就是一种博弈游戏。这种博弈游戏乃是人们运用自己具有的知识去击败对方的复杂的智力活动。计算机能不能下棋，与人对奕，其回答可以说是肯定的。这方面最著名的是塞缪尔（A. Samuel）编制的跳棋程序。这个程序利用启发式搜索技术，能根据对方的走步从几种可能的走步中选出较好的一步。而且，它还有一定的学习和积累经验的能力。50年代末，塞缪尔曾在IBM704计算机上同他自己编制的程序进行了比赛。塞缪尔下跳棋是很

不错的，但是，在比赛中，尽管他绞尽脑汁地进行奋战，结果还是输给了自己编制的跳棋程序。后来，他又对该程序作了进一步的修改和提高。在60年代，这个程序与美国肯塔基州的跳棋冠军对阵，结果，把这位州冠军给打败了。这件事一下子曾成了当时报纸的重要新闻。有人认为，塞缪尔的跳棋程序，是这一阶段机器智能发展的最大成就。

人们也致力于编制下象棋的计算机程序，并希望获得同样的成功。然而，在70年代初期以前，计算机的棋艺水平仍然是很低的。后来，美国西北大学编制成一个下象棋的计算机程序，名叫CHESS4.5，它在1976年美国象棋冠军赛B组比赛中以五胜无负的成绩获胜。这件事情使人们相信，虽然计算机的象棋棋艺水平离象棋大师的棋艺水平还有相当的距离，但是，设计出一个争夺象棋世界冠军的计算机程序，原则上不是不可能的。关键的问题是要足够清楚地了解和掌握象棋大师们下棋的经验知识和技巧，并把它们变成计算机可用的规则，编进计算机程序中去。

