

利/害条件下归纳推理的心理效应分离*

蒋柯¹ 熊哲宏²

(¹ 西南民族大学社会学与心理学学院, 成都 610041) (² 华东师范大学心理与认知科学学院, 上海 200062)

摘要 研究设计了两类包含内容的归纳推理任务, 一类是对获得收益的可能性的推理; 另一类是避免伤害的推理。实验显示, 不同内容的归纳推理结果有显著差异。其中, 避免伤害的归纳推理表现出将伤害可能过度推延的特点; 而获得收益的归纳推理则表现出对收益可能推延不足。此外, 不同的自我卷入水平——自我卷入式推理或者非自我卷入式推理——也会对推理产生显著影响。在避害条件下, 自我卷入式推理比非自我卷入式推理表现出更强烈的过度推延的特点; 在获利条件下, 自我卷入式推理比非自我卷入式推理表现出更严重的推延不足。研究经过分析提出, 这种差异可能暗示了归纳推理的领域特殊性特征。

关键词 归纳推理; 内容效应; 领域特殊性; 进化心理学

分类号 B842

1 引言

归纳推理(Inductive inference)是人根据关于已观察到事件的知识, 产生出关于未观察到事件的知识的精神活动, 是人类能超越已知知识的限制, 产生新知识, 探索未知领域的重要的精神工具。

近年来, 关于归纳推理的研究都采用空白内容或陌生内容来构成推论(Sloman & Lagnado, 2004)。比如, 羊可能会患“X 病”; 或者, 知更鸟“有芝麻骨”。还有些内容对被试来说虽然并不陌生, 但其本身并不包含明显的价值判断或利害关系, 比如, 鹰的“肝脏有两叶”, 这样的内容的只是推理命题的结构性成分, 被试实际上完全可以忽略这些内容。这种实验设计方法是为了排除被试已有经验对其推理过程的渗透, 使得实验能够集中显示出归纳推理中的概念效应或相似性效应(Rips, 1975; Gelman, 1988; Osherson, Smith, Willkie, Lopez, & Shafir, 1990; Sloman, 1993; Sloman & Lagnado, 2004; Sloutsky & Fisher, 2004; 蒋柯, 2005; Heit & Hayes, 2005)。

这种研究方法会带来一个有争议的问题: 归纳

推理能不能脱离具体内容? 这个问题还可以表述为更加一般化的问题: 归纳推理是领域一般性的还是领域特殊性的?

领域一般性的观点认为心理机制是一种通用型的机制, 一个心理机制可以应对多种要求, 并且对各种刺激的反应都遵循同样的反应规则。因此, 用单一的理论模型就可以描述心理机制在多种条件下的反应特征。持领域一般性假设的研究者会致力于消除具体内容的影响而考察心理机制的一般性反应规律。

根据 Cosmides 和 Tooby 的观点, 领域特殊性是指: 我们的心理主要由一簇特殊化的机制构成, 它们对领域特殊的表征进行操作(Cosmides & Tooby, 1994)。这种特殊化的机制的工作原理是: “每一个这样的装置有其自己的程序, 并对这个世界的不同部分施加自己特有的组织”(Tooby & Cosmides, 1995)。简而言之, 领域特殊性就是指一种心理机制只对特定的刺激作出反应, 并且具有自己独特的反应策略。持领域特殊性假设的研究者会致力于探讨具体内容对心理机制反应的影响, 这方面的研究以

收稿日期: 2008-01-10

* 教育部人文社会科学研究项目(05JA720008)、全国教育科学“十五”规划重点课题(DBA050049)、及上海市教育科学研究重点项目(A0411); 西南民族大学 2009 年引进高层次人才科研资助金项目(2009RC017); 西南民族大学 2009 年教改重点课题项目(校字 200949 号)西南民族大学 2009 年度中央高校基本科研业务费专项项目(09SZYZJ10)。

通讯作者: 熊哲宏, E-mail: zhxiong@psy.ecnu.edu.cn

Cosmides 关于“社会交换逻辑”的研究为代表 (Cosmides, 1989)。

归纳推理是领域一般性的还是领域特殊性的, 这是一个很宏观的问题, 需要更多、更系统的研究才能说明。作为这个系列研究的起点, 本研究假设: 推理内容会参与归纳推理, 并对推理结果产生影响, 即归纳推理会表现出内容效应。提出这个假设的理由有二:

第一, 本研究开始时, 我们首先注意到这样一些事实: 1) Osherson 首先归纳出 11 种归纳推理现象, 至今没有一种归纳推理的领域一般性理论模型能够解释归纳推理的全部 11 种现象。这个事实暗示了, 这 11 种现象可能并不遵循一个共同的规则, 也就是说, 很有可能并不存在一个领域一般性的归纳推理规则。2) 在演绎推理中, Cosmides 用“四卡片任务”进行的研究显示了演绎推理的内容效应 (Cosmides, 1989), 以此类推, 归纳推理也可能具有内容驱动的特征。

第二, 当我们在传统的思路下研究归纳推理遇到困难的时候, 不妨换一种策略, 采用功能分析的方法。功能分析按照功能来区分和界定心理机制。在进化心理学的范式下, 所有的心理机制都是围绕着一个目标而设计出来的, 那就是有机体基因的保存与传播。为了实现这个目标, 有机体必需要能够高效准确地解决一系列适应性问题, 诸如觅食、求偶、抚养下一代以及社会合作等等。在漫长的进化过程中, 根据自然选择的法则, 一些专门化的机制会被设计出来应对这些适应性问题, 这些专门化的机制就是“适配器”。归纳推理的能力关系到个体能否很好地适应环境的变化, 因此它和个体的生存与繁衍有关。从进化心理学的角度来看, 归纳推理一定会在进化过程中被固定下来, 成为一种专门的“适配器”(Bass, 2004/2007; 熊哲宏, 李其维, 2002)。

我们进一步推论, 有机体应该具有不止一种归纳推理的适配器, 不同领域的归纳推理任务有不同的适配器与之对应。在生理结构中有类似的例子, 比如皮肤中的热感觉和冷感觉分别由两套相互独立的神经通路来传播, 两套通路的反应方式不相同, 这种设计能够很好地适应有机体对环境作出快速准确反应的要求。以此类比, 如果一个特定领域的适应性问题需要有机体具有快速准确的反应能力, 而这个领域的任务又有自己独特的获得信息和加工信息的要求, 那么它就需要有特定的归纳推理机制与之对应。

与有机体的适应性相关的具有专属性的行为倾向有很多, 其中获得收益和避免伤害也许是最显著的两类。获得收益包括: 从环境中获得必要的资源、在特定条件下快捷准确地到达目的、治疗疾病和恢复健康等内容; 避免伤害包括: 避免食用有毒的食物、避免自己拥有的资源遭受损失、回避有伤害性的环境和动物以及避免从事有可能造成伤害的活动等。根据前面的推论, 我们假设有关这两类任务的归纳推理是由两套不同的机制来实现的, 这两个机制各有不同的反应法则。

关于本研究的假设, 还有以下两个问题需要说明:

第一, 个体进行归纳推理时可能有不同的自我卷入水平: 自我卷入式推理和非自我卷入式推理。

从 Cosmides 的有关“社会交换理论”的研究中, 我们可以得到这样的启示: 如果归纳推理是一个进化而来的机制, 那么它应该体现出自我卷入式推理和非自我卷入式推理之间的差异。比如, 当人面对一个可能会发生危险的情景时, 仅仅以一个旁观者的角度来推出并评价其危险性就是非自我卷入式推理, 而自己必须对这个情景作出现实的反应就是自我卷入式推理。在这两种情况下, 人对情景中危险性的推理会产生不同的结果。在获得收益的条件下亦然。因为在进行自我卷入式推理的时候, 人们会自动地根据情景条件进行归纳推理并作出符合适应性需求的反应, 但是非自我卷入式推理则会受个人经验的影响, 体现出更复杂的特征。基于这种观念, 我们进一步假设, 人们对获得收益和避免伤害条件下, 不同的任务要求——自我卷入式推理或者非自我卷入式推理——也会影响其推理的策略。本研究将考察自我卷入式和非自我卷入式两种形式的推理。

第二, 归纳推理有两种形式: 一种是从已观察的事实推理到另一个个例的推理, 例如, 根据前提事实“ n 个已观察到的毛利人皮肤黝黑”推理到: “下一个将要遇到的毛利人皮肤也会黝黑”; 另一种是从已观察的事实推理到一个类整体的推理, 如, 根据前提事实“ n 个已观察到的毛利人皮肤黝黑”推理到“所有的毛利人皮肤都黝黑”。根据罗素的定义: “已知有 n 个数目的 α 已经(被)发现为 β , 并且没有 α 已经(被)发现不是 β , 那么这两个陈述: (a) ‘下一个 α 将是 β ’, (b) ‘所有的 α 都是 β ’ 我把(a)叫作 ‘特殊归纳’, 而把(b)叫作 ‘一般归纳’ ”(罗素, 1948/1983)。本研究将考察归纳推理的这两种类型。

2 研究方法

2.1 被试

180 名北京师范大学珠海分校的二年级大学生, 都没有接受过专门的逻辑学训练(专门的逻辑学训练包括: 是逻辑学或相关专业, 如哲学、数理逻辑等专业的学生; 学习过《逻辑学》课程)。

2.2 实验材料

本研究采用人工材料——某种假想生物——作为实验材料。在本研究中, 内容作为实验设计的自变量之一, 必须为被试所熟悉并能够准确理解其意义, 为了避免被试已有经验对实验任务的渗透, 我们设计了某种“假想生物”作为实验材料。使用人工符号代表的“假想生物”作为实验材料的方法在推理研究中已有众多先例, 事实说明这种方法能够较好地实现无关变量的控制。

实验中使用的“假想生物”图案是在 Windows 系统的绘图板上手绘而成, 黑白线条图案, 共 10 个, 其中一个为前提刺激, 另外 9 个是结论刺激。所有 10 个假想生物都在特征维度和概念维度上进行区分。在特征维度上, 区分来自身体形状、斑点、眼睛和尾巴等特征, 每个特征有 2~3 种变形。结论刺

激和前提刺激的区分度通过两者差异特征的数量来控制, 分为三个区分水平: 一个特征差异; 两个特征差异和三个特征差异。

在概念维度上, 所有的 10 个假想生物都属于一个上位概念: “新异生物”, 分为两个基本概念: Gubee 和 Poogusa。Gubee 中包含 3 个下位概念: *mossiva-Gubee* 和 *zungum-Gubee* 和 *tavvialu-Gubee*。概念的命名没有特别意义。前提刺激属于 *mossiva-Gubee*, 另外 9 个结论刺激中, 3 个刺激和前提刺激共享一个下位概念 *mossiva-Gubee*; 另外有 3 个和前提刺激共享基本概念 Gubee, 其中两个属于 *zungum-Gubee*, 一个属于 *tavvialu-Gubee*。最后 3 个属于另一个基本概念 Poogusa, 即和前提刺激共享上位概念“新异生物”, 其中两个属于 *lamoodu-Poogusa*; 一个属于 *pigi-Poogusa*。结论刺激在基本概念和上位概念水平的分类对于本研究并没有特殊意义, 但本研究作为一个系列研究的一部分, 这种分配是为了适应后继研究的需要。实验结果显示, 这个分配没有对本研究产生影响。

9 个结论刺激在概念维度和特征维度中按照 3×3 全交叉分布, 见表 1。

表 1 9 个实验刺激在特征和概念维度上的全交叉分布

特征差异		1	2	3
概念差异	下位概念(<i>mossiva-Gubee</i>)	<i>mossiva-Gubee</i> 	<i>mossiva-Gubee</i> 	<i>mossiva-Gubee</i> 
	基本概念(Gubee)	<i>zungum-Gubee</i> 	<i>zungum-Gubee</i> 	<i>tavvialu-Gubee</i> 
	上位概念(新异生物)	<i>lamoodu-Poogusa</i> 	<i>lamoodu-Poogusa</i> 	<i>pigi-Poogusa</i> 

向被试呈现实验材料时, 概念名称和图案一起呈现, 形式和大小比例见图 1。



图 1 实验中刺激呈现示意图

概念名称的文字和图案的大小比例是通过一个事先的预备检测来确定的。在 WORD 文档中, 通过改变文字的字号调节文字符号的大小; 通过位图调节功能调节图案尺寸。将不同的组合在计算机屏幕上快速呈现, 让被试判断两种刺激的主观显著性差异。图例所示为实验中采用的文字图案大小比例。参与预测的 5 个被试中有 4 人认为这种比例下, 文字和图案有大致相当的显著性。

本研究控制刺激的概念信息和特征信息是希望能够验证, 无论在概念维度还是在特征维度上,

利害条件都会在归纳推理中显示出内容效应。

2.3 实验设计

实验采用双因变量设计。每一个被试参与的实验任务都包括特殊归纳和一般归纳两种测验题目。特殊归纳是从个例到个例的推理; 一般归纳是从个例到概念整体的推理。

特殊归纳实验采用 5(推理内容)×3(概念差异)×3(特征差异)被试内和被试间混合设计。其中, 推理内容是被试间设计; 概念差异和特征差异等是被试内设计。

一般归纳实验采用 5(推理内容)×3(概念差异)被试内和被试间混合设计。其中, 推理内容是被试间设计; 概念差异是被试内设计。

推理内容: 基于前述假设, 我们设计的归纳推理内容包含了获得收益——简称“利”和回避伤害——简称“害”两种内容条件。“利”和“害”的刺激

应该给被试产生方向相对但强度相当的心理感受, 为了达到这个效果, 本研究选择了个体在自然生存条件下“利”和“害”的极端条件, 利的条件是能够从死亡的危险中挽救生命, 比如能够救命的药物; 害的条件是能够致命的危险, 例如有致命危险的攻击。因为分别是两个极端的刺激条件, 所以它们引起的心理感受是对应的。不同的刺激强度的“利”“害”条件有可能对人们的归纳推理产生不同的影响, 类似问题将在后继研究中探讨。

研究还在“利”和“害”条件下分别设计了两种任务, 一种是要求被试对该条件下的收益(或伤害)的可能性进行推测, 例如, 已知一种生物是有剧毒的, 估计另一种生物也有剧毒的可能性, 这是非自我卷入式推理任务, 在本研究中简称为“旁观”任务; 另一种则要求被试假设自己身处这种条件之下, 对自己的某种行为意愿进行推测, 如已知某种生物可以做治疗蛇毒的解药, 并假设自己已经中了蛇毒正在找解药, 现在找到了另一种生物, 估计愿意用它做解药的意愿强度, 这是自我卷入式推理任务, 在本研究中简称为“卷入”任务。

于是形成了四种包含内容的归纳推理任务:

- 1) “利”条件下的非自我卷入式推理, 利—旁观;
- 2) “利”条件下的自我卷入式推理, 利—卷入;
- 3) “害”条件下的非自我卷入式推理, 害—旁观;
- 4) “害”条件下的自我卷入式推理, 害—卷入; 此外, 为了有一个可供比较的基线, 本研究还包括了一个空白内容的归纳推理任务, 即, 5) 中性任务推理, 中性任务不包含有利害关系的内容, 采用空白内容设计, 用无意义字符串代替内容, 表达为: “能够产生物质 Tg3—HWb”。

本研究将比较这五种任务条件下的归纳推理特征。如果有显著性差异, 则能够验证本研究关于归纳推理具有内容效应的假设。

概念差异: 分三个水平, 下位概念、基本概念和上位概念。即要求被试分别在三个水平上进行归纳推理。

特征差异: 分三个水平, 差异特征数量为 1、2 和 3。

180 名被试分成五组, 每组 36 人, 分别参与一种推理内容的任务。

每个实验任务都包含 9 个特殊归纳和 3 个一般归纳。9 个特殊归纳由概念差异和特征差异 3×3 全交叉的个例构成, 见表 1; 3 个一般归纳分别是推理到下位概念、基本概念和上位概念。

两个因变量分别是被试在各种条件下对特殊归纳和一般归纳推理力的判断。被试被要求用从 0~10 之间的一个数字, 0.5 间距, 来表示自己推理的信心或自己希望采取某种行为的意愿强度。

实验题目顺序随机排列。

2.4 实验程序

所有的实验题目用 A4 纸打印, 装订成册, 每册 4 页。测验在教室中进行, 施测时将题目分发给被试, 每人一册, 要求被试单独作答, 不得与别人交流, 不限时间, 但鼓励尽快完成。通常被试会在 5~10 分钟之内完成任务。

3 结果

实验包括特殊归纳和一般归纳两个双因变量, 其中特殊归纳依赖于推理内容、特征差异和概念水平等因素的变化, 而一般归纳则只与推理内容和概念水平两个因素相关。在数据处理时, 本研究将对特殊归纳和一般归纳的结果分别进行分析, 并在相同的实验控制中对两者进行比较。

3.1 特殊归纳的结果

在特殊归纳中, 不同推理内容的平均推理力见表 2。

表 2 推理内容×概念水平条件下, 特殊归纳的推理力(SD)

推理内容	概念水平			平均推理力
	下位概念	基本概念	上位概念	
害—旁观	6.07(2.43)	5.50(2.57)***	5.36(2.86)***	5.64(2.64)***
害—卷入	7.71(1.93)***	7.05(2.07)***	6.88(2.40)***	7.21(2.16)***
利—旁观	5.60(2.31)	4.44(2.46)	4.48(2.64)	4.84(2.52)
利—卷入	4.65(2.52)***	3.26(2.51)**	2.95(2.55)**	3.62(2.63)***
中性任务	6.18(2.44)	4.30(2.53)	4.10(2.72)	4.86(2.72)

注: *表示和同列中性任务比较差异性显著, ** $p < 0.01$, *** $p < 0.001$ 。

t 检验结果显示, 在下位概念水平, 利—卷入与其它四种推论任务有显著性差异; 利—旁观和害—卷入、利—卷入有显著性差异; 害—旁观和害—

卷入、利—卷入有显著性差异; 中性任务和害—卷入、利—卷入有显著性差异; 利—旁观、害—旁观和中性任务之间差异不显著。

在基本概念水平,利—旁观和中性任务差异不显著;其它任务之间都有显著性差异。

在上位概念水平,利—旁观和中性任务差异不显著;其它任务之间都有显著性差异。

不同推理内容的平均推理力 LSD 多重比较结果见表 3。

从表 3 可见,除了利—旁观和中性任务之间没有显著差异外,其它各种推理内容之间都存在显著性差异。害—卷入的推理力高于其它所有内容;接下来的差异性关系是:害—旁观>利—旁观和中性任务>利—卷入;利—卷入显著低于其它所有内容。

5×3×3 (推理内容×概念差异×特征差异)

ANOVA 结果显示,推理内容主效应显著, $F_{\text{推理内容}}(4, 1618)=97.621, p<0.001$; 概念差异主效应显著, $F_{\text{概念差异}}(2, 1618)=46.22, p<0.001$; 特征差异主效应显著, $F_{\text{特征差异}}(2, 1618)=43.48, p<0.001$; 推理内容和概念差异的交互作用显著, $F_{\text{推理内容} \times \text{概念差异}}(8, 1618)=2.01, p<0.05$; 概念差异和特征差异交互作用显著, $F_{\text{概念差异} \times \text{特征差异}}(8, 1618)=5.70, p<0.001$; 但是推理内容和特征差异交互作用不显著, $F_{\text{推理内容} \times \text{特征差异}}(8, 1618)=1.39, p=0.20$; 推理内容、特征差异和概念差异三者的交互作用不显著, $F_{\text{推理内容} \times \text{特征差异} \times \text{概念差异}}(16, 1618)=1.96, p=0.34$ 。见表 4。

概念差异和推理内容交互作用见图 2。

表 3 特殊归纳中各种推理内容的平均推理力 LSD 多重比较

任务类型	害—旁观	害—卷入	利—旁观	利—卷入
害—卷入	1.57***			
利—旁观	-0.81***	-2.37***		
利—卷入	-2.02***	-3.59***	-1.21***	
中性任务	-0.79***	-2.35***	0.02	1.24***

注: *** $p<0.001$

表 4 特殊归纳 5×3×3(推理内容×概念差异×特征差异)ANOVA 结果

变异来源	SS	df	MS	F
推理内容	2249.68	4	562.42	97.62***
特征差异	500.95	2	250.48	43.48***
概念水平	532.60	2	266.30	46.22***
推理内容×特征差异	63.90	8	8.00	1.39
推理内容×概念水平	92.83	8	11.60	2.01*
特征差异×概念水平	131.34	4	32.84	5.70***
推理内容×特征差异×概念水平	31.32	16	1.96	0.34
误差	9062.46	1573	5.76	
总变异	56981.27	1618		

注: * $p<0.05$, *** $p<0.001$ 。

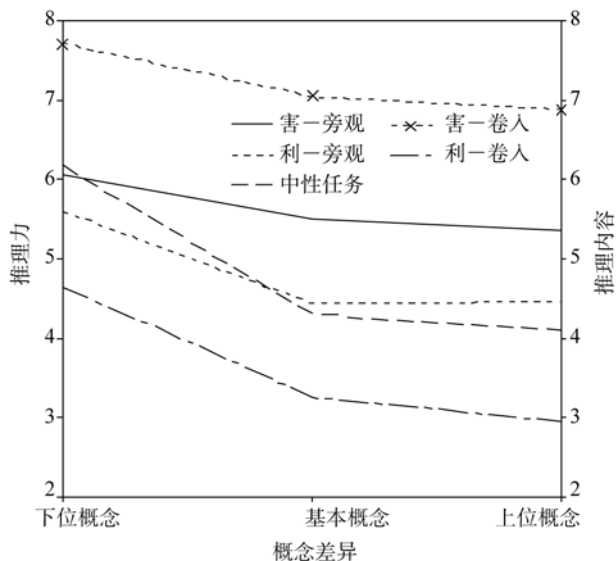


图 2 概念差异和推理内容交互作用

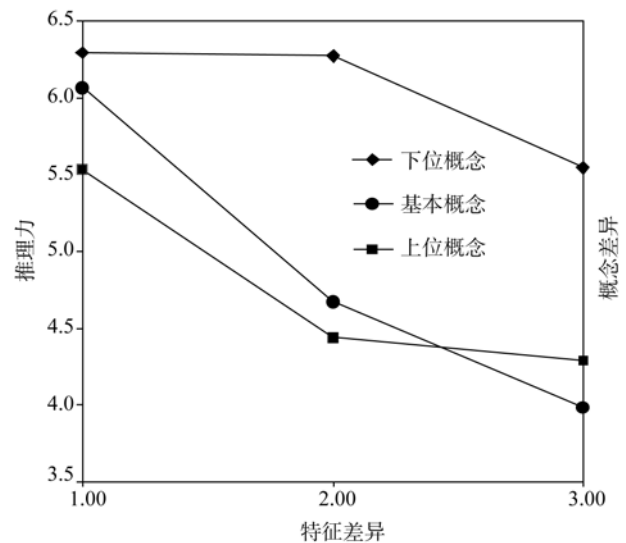


图 3 特征差异和概念差异交互作用

总体上, 害—卷入和害—旁观受概念差异变化的影响较小, 而中性任务、利—卷入和利—旁观对概念差异的变化更敏感, 并且都有共同的变化规律, 即在下位概念到基本概念之间的变化较大, 基本概念和上位概念之间的变化较小。

特征差异和概念效益交互作用见图 3。

在下位概念水平, 推理力在特征差异 1 和特征差异 2 之间几乎没有变化, 但特征差异 2 和特征差异 3 之间有较大的差异。而基本概念水平, 推理力

则随着特征差异的增加下降较大。上位概念水平, 推理力的变化在特征差异 1 和特征差异 2 之间明显, 在特征差异 2 和特征差异 3 之间不显著。

因为推理内容×特征差异的交互作用不显著, 所以不对特征差异下的推理力的分布进行专门检验。

3.2 一般归纳的结果

在一般归纳中, 推理内容×概念水平条件下的平均推理力见表 5。

表 5 在各种推理内容下, 三个概念水平一般归纳的推理力(SD)

推理内容	概念水平			平均推理力
	下位概念	基本概念	上位概念	
害—旁观	5.57(2.89)	5.07(2.58)	4.44(2.61)	5.03(2.71)
害—卷入	7.65(1.80)***	7.31(1.70)***	6.94(2.18)***	7.30(1.91)***
利—旁观	5.53(2.56)	4.15(2.48)	3.77(2.35)	4.49(2.56)
利—卷入	4.21(2.15)	3.57(2.59)*	2.65(2.46)	3.48(2.47)**
中性任务	5.19(2.65)	4.72(2.74)	3.40(2.76)	4.44(2.80)

注: *表示和同列中性任务比较差异性显著, * $p<0.05$, ** $p<0.01$, *** $p<0.001$ 。

LSD 多重比较结果显示, 在下位概念水平, 中性任务和害—卷入有显著差异; 利—卷入显著低于害—旁观、害—卷入和利—旁观。和中性任务没有显著差异; 害—卷入显著高于其它四种任务; 害—旁观、中性任务和利—旁观没有显著差异。

在基本概念水平, 中性任务和害—卷入、利—卷入都有显著性差异; 利—卷入和利—旁观没有显著差异, 和其它三种任务有显著性差异; 害—卷入显著高于其它四种任务; 害—旁观和害—卷入、利—卷入有显著差异, 和利—旁观、中性任务没有显著性差异。

在上位概念水平, 中心任务和害—卷入有显著差异, 和其它任务没有显著差异; 利—卷入显著低于害—旁观和害—卷入; 害—卷入显著高于其它四种任务; 利—旁观只与害—卷入有显著差异; 害—旁观和利—卷入、害—卷入有显著差异。

一般归纳的平均推理力 LSD 多重比较结果见表 6。

多重比较结果显示, 害—卷入的推理力显著高于其它 4 种条件; 利—卷入的推理力则显著低于其它 4 种条件; 害—旁观、利—旁观和中性任务之间没有显著差异。因此, 包含内容的一般归纳推理力高低实际上分成了三个层次, 害—卷入最高, 中间是害—旁观、利—旁观和中性任务, 利—卷入最低。

5×3(推理内容×概念水平)的 ANOVA 结果显示, 推理内容的主效应显著, $F_{\text{推理内容}}(4, 538)=36.61$, $p<0.001$; 概念水平的主效应显著, $F_{\text{概念水平}}(2, 538)=14.40$, $p<0.001$; 推理内容和概念水平之间交互作用不显著, $F_{\text{推理内容} \times \text{概念水平}}(8, 538)=0.54$, $p=0.82$ 。见表 7。

表 6 推理内容下一般归纳的平均推理力 LSD 多重比较

任务类型	害—旁观	害—卷入	利—旁观	利—卷入
害—卷入	2.28***			
利—旁观	-0.54	-2.82***		
利—卷入	-1.55***	-3.83***	-1.01**	
中性任务	-0.59	-2.86***	-0.05	0.96**

注: ** $p<0.01$; *** $p<0.001$ 。

3.3 特殊归纳和一般归纳的比较

在推理内容×概念差异双重条件下, 特殊归纳和一般归纳的比较, 见表 8。

t 检验结果显示, 在害—卷入和利—卷入两种内容下, 特殊归纳和一般归纳在各个概念水平上都没有显著性差异。在中性任务中, 在下位概念条件下, 特殊归纳明显高于一般归纳。在害—旁观、利—旁观和和中性任务中, 在上位概念水平, 特殊归纳显著高于一般归纳; 在其它位置上, 没有显著差异。

表 7 一般归纳 5×3(推理内容×概念水平)ANOVA 结果。

变异来源	SS	df	MS	F
推理内容	879.08	4	219.77	36.61***
概念水平	172.83	2	86.41	14.40***
推理内容×概念水平	26.14	8	3.27	0.54
误差	3139.38	523	6.00	
总变异	17368.50	538		

注: *** $p<0.001$ 。

表 8 推理内容×概念水平条件下特殊归纳和一般归纳比较

推理内容	概念水平								
	下位概念			基本概念			上位概念		
	特殊	一般	均差	特殊	一般	均差	特殊	一般	均差
害—旁观	6.07	5.57	0.50	5.50	5.07	0.43	5.36	4.44	0.92**
害—卷入	7.71	7.65	0.06	7.05	7.31	-0.27	6.88	6.94	-0.07
利—旁观	5.60	5.53	0.07	4.44	4.15	0.29	4.47	3.77	0.70**
利—卷入	4.65	4.21	0.45	3.26	3.57	-0.31	2.96	2.65	0.31
中性任务	6.18	5.19	0.98***	4.30	4.72	-0.42	4.10	3.40	0.69*

注: * $p<0.05$; ** $p<0.01$; *** $p<0.001$ 。

4 讨论

4.1 归纳推理的内容效应

整体上的结果显示,在不同推理内容条件下,归纳推理的结果具有显著性差异,因此,实验结果印证了我们整体假设:推理的内容会影响归纳推理的表现。

实验的总体结果显示,在“害”条件下的“旁观”和“卷入”两种推理的推理力都比中性任务的更高,说明人们对有害的事件再次发生的可能性估计偏高。即人们会倾向于认为:如果某种有害事件曾经发生过,那么在相应的条件下它很有可能再次发生,因此,人会倾向于回避可能与有害事件相关的各种条件,即使有的条件和原来的有害事件本身关系并不是很密切,人们依然会对它们保持高度的警惕。

此外,在“害”条件下的推理力偏高还说明,如果得知前提刺激是有毒的,人们会把这个属性推理到类的整体,即,人们会倾向于认为和前提刺激的同属于一个类别的其它任何一个个体都是有毒的可能性较大。总之,实验数据揭示了这样一个事实,当已知环境中可能存在危险时,人会把这种危险过度地推广到其它情境,显得草木皆兵。

另外一方面,实验显示人在“利”条件下“旁观”和“卷入”两种内容体现出了不同的效应。“利”条件下的非自我卷入式推理和中性任务没有表现出分离,只有自我卷入式推理明显低于中性任务。

这表明,人们在已知前提刺激能够带来收益以

后,虽然对能够从另一个新刺激或新情景那里获得同样收益具有一定的信心,但却不愿意投入实际行动。对此,本研究的解释是,在可能获利的条件下,只有前提刺激能够带来收益是确定的,虽然在理论上可以推论与前提刺激有关联的新刺激也可能带来收益,但这是一个不确定的结论。投入实际行动就意味着要付出一定的代价,这个代价可能是行动本身的消耗,也可能是因为在这个情景中投入了行动就等于放弃了在其它情景中获得收益的机会。因此,人会尽可能地把行动投入到确定性更高的情景中,这就导致了利条件下行动推理的保守性。

实验结果与人在日常生活表现出来的行动策略是相吻合的。比如,害条件下的结果和人们常用的“以防万一”的避害策略是一致的。在不能断定遭遇的新刺激是否有害时,宁可相信它是有害的,对它采取一定的防范措施或者回避它,这种策略使得人面对环境中的新刺激时,总是能够保持高度的警惕,因而使人避免了遭受很多可能的伤害,因此这是一种有机体的自我保护策略。类似的行动策略还可以在动物身上观察到,实际上很多动物在面临新刺激时,往往表现出比人更加谨慎的态度。我们有理由相信,这种策略是进化而来的一种适应性机制。只有对可能的危险更加敏感的个体才能更好地保护自己和自己的后代,它们的基因才会有更多的机会被保留并延续下去,这样的个体将拥有更多的后代。因此,经过漫长的自然选择过程,对可能的危险更敏感的基因最终可能散布到现存的所有个

体之中,这使得现存的所有个体都具有了对新异刺激、可能的危险保持较高的警惕性的特征。

在获得收益的条件下,个体如果不能谨慎地控制自己的行动投入,就可能付出大的代价而获得小的收获,或者丧失更多的获得收益的机会。比如,一个人某一天在一棵树下面逮到一只兔子,如果他就因此就作出每天都守在这棵树下的自我卷入式推理,这会使得他丧失了到其它地方逮到别的兔子的机会。采取这种策略的人很快就会被饿死,而只有那些更加有效的控制自己的行动投入的个体才可能获得更多的收益,并繁衍更多的后代。因此,在获得收益条件下的谨慎的自我卷入式推理也会成为一种适应性机制而被世代相传。

4.2 非自我卷入式推理和自我卷入式推理

本研究设置了“非自我卷入式推理”和“自我卷入式推理”两种类型。其中“非自我卷入式推理”要求被试根据前提刺激推论出有关其它刺激的知识,这是归纳推理的标准范例。“自我卷入式推理”则要求被试设想自己置身于前提刺激所描绘的任务情景中,要求被试在这个情景中自己可能作出的行为选择进行推论。虽然这个范例可能会被看作是决策任务,但是在本研究的实验设计中它仍然是一个归纳推理任务而不是决策任务。理由是:

第一,本研究在实验中设计“自我卷入式推理”的目的是为了考察人在作为旁观者和自我卷入两种条件下,归纳推理是否会有区别。在自我卷入式推理任务中,指导语要求被试想象自己被牵涉进了某种有害或可能获利的情景中,要求被试在这种条件下估计自己可能作出某一种行为的意愿水平。这个任务的内在逻辑是:在任务刺激所描述的情景中,已知条件与自己的某种行为意愿之间有确定的联系,被试需要根据这种确定的联系产生出关于另一种未知联系的认识。比如,如果已知某种生物可以治疗毒蛇咬伤,那么在自己被毒蛇咬伤后,一定会用这种生物做解药,这是一个已知的确定联系;但是另外一种未知生物和自己的行为意愿(用它来做解药)之间的联系就是一个未知事件,被试的任务就是从已知信息产生出关于未知事件的认识。

第二,在本研究中,“自我卷入式推理”要求被试评估自己的行为意愿,这种评估不是对行为本身是否要发生的判断,而是对行动的意愿水平的估计,如前所述,本质上它是从已有知识中形成结论的加工过程。而决策是如何在众多的可能性中作出选择(Eysenck & Keane, 2000/2004)。如果被试面临两种

或更多的未知生物,要决定用哪一种做解药,那么他所面对的就是一个决策任务。因此,构成决策必须要有多种供选择的方案;并且决策的输出结果是一个排他性的选择,即通过决策,选择了一种方案就意味着排斥了其它方案。而本研究的“自我卷入式推理”任务并不具备这样的特征,因此,这个任务是一个归纳推理的任务。

实验结果显示,非自我卷入式推理和自我卷入式推理有显著差异,这也印证了我们事先的假设:在两种情况下:在置身事外作抽象的推理时,以及有自我卷入时,人的归纳推理会表现出不同的特征。

对于非自我卷入式推理和自我卷入式推理的差异,我们认为有两种可能的解释,一种是非自我卷入式推理和自我卷入式推理分别是不同机制的结果;另一种解释是,其中一种推理的结果是另一种推理结果经过修正后的产物。无论如何,本研究的数据还不足以解释如此复杂的问题,我们希望在后继研究中更多地揭示自我卷入式推理和非自我卷入式推理之间的复杂关系。

4.3 在特征维度上的内容效应

本研究还显示了归纳推理在特征维度上也存在内容效应。虽然本研究的主要目的并不是分离归纳推理的特征效应和概念效应,实验设置概念和特征两个维度的目的是为了印证归纳推理在这两个维度上都存在内容效应。方差分析的结果显示内容效应和概念效应有交互作用但是和特征效应没有交互作用,这个微妙差异有可能是归纳推理的特征效应和概念效应的分离的一个线索,这个问题也有待进一步的研究去解释。

4.4 特殊归纳和一般归纳

本研究涉及两种形式的归纳推理,从个例到个例的推理,以及从个例到概念的推理,大多数研究都不对它们作区分,因为在这些研究中它们并没有表现出差异。本研究通过比较特殊归纳和一般归纳,发现了一些有趣的现象:

在整体上,特殊归纳和一般归纳在不同内容中的表现是基本一致的,但是在一些细节上却显示出了差异。比如,在特殊归纳中,害—旁观和中心任务有显著差异,但是在一般归纳中却没有差异了;以及特殊归纳和一般归纳相比较,在利—旁观、害—旁观和中性任务中,在上位概念水平上,特殊归纳显著高于一般归纳,而中性任务在下位概念水平上也有这种差异。

我们认为这些差异是由于人的个例表征和概念表征之间的可得性差异造成的。个例表征比概念表征更具体、更形象,也更具有可得性,因此,与个例相关的推理就更加容易建立起来。

依据这个假设,我们不难理解为什么在这几个位置上——上位概念水平的害—旁观、利—旁观和中性任务,以及在下位概念水平的中性任务——特殊归纳都显著高于一般归纳。

利—卷入和害—卷入在所有概念水平上,特殊归纳和一般归纳都没有显著差异,这可能是因为利—卷入的所有推理力都很低,而害—卷入的所有推理力都很高,地板效应和天花板效应掩盖了其中的差异。

利—旁观和害—旁观在下位概念水平上也没有差异,我们认为,因为下位概念表征的抽象水平并不高,和个例表征的差异不是很大,在推理包含了内容时,推理内容的显著性可能掩盖了表征之间的差异,因此推理也没有显示出差异。

值得注意的是,在基本概念水平,在所有推理任务中,特殊归纳和一般归纳都没有表现出差异。我们认为,这个现象可以用“最适意归纳水平”的假设来解释。这个假设即是说,因为基本概念是人们最容易掌握的类别化表征,所以人们进行归纳推理时,最容易在基本概念水平上形成一般性结论,人们对在基本概念水平上的归纳推理结论有更高的信心(Sloman & Lagnado, 2004)。

根据前面的分析,我们认为,由于个例表征和概念表征之间的可得性差异,总体上特殊归纳高于一般归纳。而另一方面,根据“最适意归纳水平”的假设,在基本概念水平,人们一般归纳的信心会有额外增加,从而使得一般归纳的推理力和特殊归纳持平。

总之,本研究发现,在特殊归纳和一般归纳之间存在一些微妙的差异,这些差异是否反映了加工过程中更加复杂的区别,还有待更多的研究去检验。

5 结论

本研究属于探索性研究,旨在于发现在以往的研究中被忽略的归纳推理现象。首先,实验考察了内容对归纳推理的影响,研究发现归纳推理中存在内容效应。实验显示在包含不同内容——获得收益的行动和避免伤害的行动——的归纳推理之间存在显著的差异。其中,避免伤害的归纳推理表现出将

伤害的可能性过度推延的特点;而获得收益的归纳推理则表现出对收益的可能性推延不足。其次,研究还发现不同的自我卷入水平——自我卷入式推理或者非自我卷入式推理——也会对推理产生显著影响。在避害条件下,自我卷入式推理比非自我卷入式推理表现出更强烈的过度推延的特点;在获利条件下,自我卷入式推理比非自我卷入式推理表现出更严重的推延不足。本研究进一步假设,这种差异可能暗示了归纳推理的领域特殊性特征。最后,研究还发现了一些有趣的细节:比如,特殊归纳和一般归纳之间存在微妙的差异。总体上,本研究属于一个探索性研究,旨在发现问题和提出问题,我们希望通过更多的研究深入探讨本研究发现的诸多有趣的现象。

参 考 文 献

- Bass, D. M. (2007). *Evolutionary psychology: the new science of the mind* (p.75). (translated by Xiong Zhehong, Zhang Yong, & Yan Qian) Shanghai: East China Normal University Press. (Original work published 2004).
- [巴斯, D. M. (2007). *进化心理学: 心理的新科学*(p.75). (熊哲宏, 张勇, 晏倩译)上海: 华东师范大学出版社.]
- Cosmides, L. (1989). The logic of social exchange: Has natural selection shaped how humans reason? studies with the Wason selection task. *Cognition*, 31, 187-276.
- Cosmides, L., & Tooby, J. (1994). Origins of domain specificity: the evolution of functional organization. In Hirschfeld and Gelman (eds), *Mapping the mind* (p.85), Cambridge University Press.
- Eysenck, M., & Keane, M. T. (2004). *Cognitive psychology: a student's handbook*. (4th Edition). (pp.727-744). (translated by Gao Dingguo, Xiao Xiaoyun), Shanghai, East China Normal University Press. (Original work published 2000).
- [艾森克, M. W., 基恩, M. T. (2004). *认知心理学* (pp.727-744). (高定国, 肖晓云 译) 上海, 华东师范大学出版社.]
- Gelman, S. A. (1988). The Development of induction within natural kind and artifact categories. *Cognitive Psychology*, 20, 65-95.
- Heit, E., & Hayes, B. K. (2005). Relations among categorization, induction, recognition, and similarity: comment on Sloutsky and Fisher (2004). *Journal of Experimental Psychology: General*, 134(4), 596-605.
- Jiang, K. (2005). The theatrical research and paradigms of temporary children's inductive inference. *Gui Zhou University Transaction(natural science)*, (12), 384-387.
- [蒋柯. (2005). 当前儿童归纳推理研究的理论与范式. *贵州大学学报(自然科学版)*, (12), 384-387.]
- Rips, L. J. (1975). Inductive judgments about natural categories. *Journal of Verbal Learning and Verbal Behavior*, 14, 665-681.
- Russel, B. (1983). *Human knowledge* (p.480). Beijing, Commercial Press. (Original work published 1948).
- [罗素. (1983). *人类的知识* (p.480). 北京, 商务印书馆.]

- Osherson, D. N., Smith, E. E., Willkie, O., Lopez, A., & Shafir, E. (1990). Category-based induction. *Psychological Review*, 97, 185–200.
- Sloman, S. A. (1993). Feature-based induction. *Cognitive Psychology*, 25, 231–280.
- Sloman, S. A., & Lagnado, D. A. (2004). The problem of induction. In: *Handbook of thinking and Reasoning*, 93–116.
- Sloutsky, V. M., & Fisher, A. (2004). V. Induction and categorization in young children: a similarity-based model. *Journal of Experimental Psychology: General*, 133, 166–188.
- Tooby, J., & Cosmides, L. (1995). Foreword to Baron-Cohen. In Baron-Cohen, S. *Mindblindness* (p. xiv). Cambridge, MA: MIT Press.
- Xiong, Z. H., & Li, Q. W. (2002). Are “Darwinian Modules” real mental module?—on the irrationality in “Swiss Army Knife” model of cognition. *Psychological Science*, 2, 826–836.
- [熊哲宏, 李其维. (2002). “达尔文模块”与认知的“瑞士军刀”模型. *心理科学*, 2, 826–836.]

Different Contents Different Inductive Inference: Under the Conditions of Embrace-Advantage and Resist-Disadvantage

JIANG Ke¹; XIONG Zhe-Hong²

(¹ College of Sociology and Psychology, Southwest University for Nationalities, Chengdu 610041, China)

(² School of Psychology and Cognitive Science, East China Normal University, Shanghai 200062, China)

Abstract

Whether inductive inference is domain general or content-dependent remains elusive in the literature. Most previous studies are conducted using blank-content reasoning method, which sets the precondition that inductive inference must be independent from special contents and thus fails to distinguish aforementioned two approaches.

Recent studies reveal that none of domain general models can explain all eleven forms of inductive inference and deductive inference is content-dependent, indicating that content might play an important role in regulating inductive inference. This view is also supported from an evolutionary psychology perspective. Inductive inferences could be a group of adaptations that are gained during the evolution process to address certain events critical to survival.

In the current study, we aim to address this issue by using a novel paradigm to explore the role of content on inductive inference. Two tasks including embrace-advantage and resist-disadvantage are utilized because they are closely related to human survival and reproduction. Each task contains two forms of reasoning. Subjects are asked to either estimate the general possibility of embrace-advantage or resist-disadvantage, or express their own will of action in case of embrace-advantage or resist-disadvantage. Thus, a total of four kinds of reasoning tasks were conducted including reasoning in embrace-advantage, action-will in embrace-advantage, reasoning in resist-disadvantage and action-will in resist-disadvantage. A blank-content test serves as control for these reasoning tasks.

The results demonstrate that embrace-advantage inference is markedly different from resist-disadvantage inference, indicating a potent content-dependent effect. Embrace-advantage shows under generalization whereas resist-disadvantage shows the trend of over generalization, and the embrace-advantage inference is under generalization. We reason that these two tasks are driven by two different psychological mechanisms. Embrace-advantage follows the strategy of sufficiency-necessity, as compared to the sole sufficiency strategy utilized by resist-disadvantage tasks. Such differences have fulfilled the constitutive standards of domain specificity and we conclude that inductive inference is based on domains.

One innovating feature of the current study is the use of artificial concepts to build up the reasoning propositions. Because content is an important independent variable, artificial concepts can avoid the confounding effect of previous knowledge on subjects' responses. The results indicate that it is an excellent method to control extraneous variable in reasoning research.

Key words inductive inference; content effect; domain specificity; evolutionary psychology