

# 越顺利，越没有你的位置

AI 时代普通人的困境与出路

王德生 + Claude

据学员专栏十一位作者的十一篇论文编成

德麦国际出版社

SDE UNIVERSES

## 出版信息

书 名 越顺利，越没有你的位置

副标题 AI时代普通人的困境与出路

编 著 王德生 + Claude

副 署 据学员专栏十一位作者的十一篇论文编成

原作者（按章序） 陈晓艳·阳涌·季春雷·孔凡鹤·高鹏·张琼·高于涵·少敏·胡敏·秦莉·胡志英

出 版 德麦国际有限公司（Demai International Pte. Ltd.）·新加坡

丛 书 SDE Universes 专著·第84号

书 号 ISBN 979-8-90690-018-0

版 次 2026年8月第一版

定 价 US\$21.30

版 权 © 2026 王德生 / 德麦国际有限公司。保留一切权利。未经书面许可，不得以任何形式复制、转载或用于模型训练；转引请注明出处与链接。

署名说明 本书十一章正文分别由十一位作者独立撰写，原文已先行公开发表于学员专栏；每章章首标注原作者与原文出处。导论、枢纽章、合章、编首、前言、导读、结语与三则附录由王德生 + Claude 撰写，是本书的新增部分。十一位原作者对各自章节保有署名权；十一章正文除以下三类技术性处理外未作改动：删除网页版的阅读条、评分条、导读框与配套读物入口；统一标题层级；将期刊体的“本文”改为书体的“本章”。各章文献表已剥出合并于书末，章内引注编号仍为原章编号。

选目说明 十一章里有十章不是选出来的，是按预先写死的协议抽出来的：抽样框（汉字八千以上、页面标注创新智商一四六以上、未被任何已出版专著收录）、随机种子（20260817）与四条成功标准，全部写在抽样之前。第十一位作者胡敏是指定席——种子没有抽中他，是编者点进来的，这一点在附录一里连同抽样过程中的一次修框一起如实交代。每位作者章内的那一篇是有意选择，不是随机，判据是与题域的相关性。

证据纪律 本书新提出的三个量（吸收率  $\alpha$ 、收件位移  $\Delta$ 、存押量  $\beta$ ）与那条枢纽命题均未经真实数据检验。凡本书主张为已完成的检验，一律为零。三条判错办法写在枢纽章第三节、合章第三节与附录三，其中最便宜的一条只需从一个组织的既有系统里取两个数。

## 作者介绍

王德生，数学博士、哲学家，SDE 本体论（显露—差异—纠缠）的创立者，新加坡德麦国际出版公司（Demai International Pte. Ltd.）创始人。中国科学院计算数学博士、博士后，曾任教于南洋理工大学，此前在计算数学与计算机图形学领域发表多篇国际期刊论文。以三十余年的学术探索，完成了从计算数学家到跨域思想家的转向。核心命题是：「世界不是被发现的，而是发生的。」已出版八十多部跨域专著，覆盖哲学、数学、人工智能、医学、教育、艺术、文明研究等约二十个学科。

Claude，Anthropic 开发的语言模型。在本书中承担的工作是：执行抽样协议并如实报告其中的一次修框、抽取与统稿十一章正文、撰写导论、枢纽章、合章与全部前后件的初稿。本书的判断由两位署名者共同负责；凡书中说"本书主张"处，指的是这一共同判断，而不是十一位原作者中任何一位的立场。

十一位原作者（按章序） 陈晓艳、阳涌、季春雷、孔凡鹤、高鹏、张琼、高于涵、少敏、胡敏、秦莉、胡志英。他们分属癌症发生学、人机协作研究、社会存在论、营养生成论、法哲学、伦理人类学与团体治疗、中国思想史、人机系统效力审计、运动医学与修复过程研究、艺术哲学、过程存在论，彼此互不相识，写作时没有任何一位知道的文章会与另外十篇装在一起。

## 前言

### ■ 一、一次赌，和它的赔率

这本书的做法，先说清楚。

我们没有先想好一个主题再去找文章。我们先写死了一份协议：从学员专栏五百多篇署名文章里，划出一个抽样框——汉字八千以上、页面标注的创新智商在一四六以上、并且没有被此前任何一部专著收录过；框里有十二位学员、六十六篇文章。然后用一个预先声明的随机种子抽出十位学员。抽中谁就是谁，不重抽，不替换。

这是一次赌。赌的是：十位互不相识、写着完全不同题域的作者，他们各自那一篇里，是否真的藏着同一件事。如果没有，这本书就不该存在，我们会如实报告失败——这样的失败我们此前报告过一次，用的是某学术期刊同一期随机抽出的十篇文章，那一次的结论是“只装出了一句体裁惯例，零预测力”。

这一次的结果是：十一章里有十一章可以从同一句话里派生出来。这个数字高得让我们自己也想再查一遍。所以查错的办法也一并写在书里了，见附录三。

### ■ 二、抽样过程中出的一个错，必须先说

第一遍抽样时，我们用文章的网址短名去查“这篇是不是已经收进过别的书”。这个办法漏检了。改用中文题名重查，才发现有三篇早已是第八十二号书的正章。其中一位作者只有那一篇合格，于是他整个人退出抽样框，框从十三位缩到十二位，只好用同一个种子在修正后的框上重抽。

我们把这件事写在前言里，而不是塞进附录，是因为它关系到这本书的成色：“先声明后抽样”的纪律在这一轮上打了折——修框发生在第一次抽样之后。任何人要因此给本书的抽样部分打个问号，是有道理的。我们能提供的只是：修框的判据（已入书者出框）与修框的动作（同种子重抽），事先都已写在协议里，不是事后调出来的。

### ■ 三、第十一位是点进来的，不是抽出来的

十位抽定之后，站上产量最高、发表专著最多的那位学员没有被抽中。我们把他补了进来，选的是他框内分数最高的一篇。

这是一次公然的人工干预，本书不打算把它藏在随机性的外衣底下。补进来之后发生了一件出乎意料的事：他那一篇（关于"裁定这个动作本身构成干扰"）与随机抽中的少敏那一篇（关于"人工复核越多越不能改道"）咬合成了全书最硬的一对——同一个"评估"动作的两端，一个从组织侧，一个从过程侧，互不知情。合章整章就架在这一对上。

所以指定席的账要这样记：它给本书增加了最有价值的一章，同时降低了本书作为"随机可装配性"实验的证据强度。两件事都是真的。

## ■ 四、这本书在说什么

一句话：

凡是让普通人变得更顺利的安排——托底、代劳、去阻力、加复核、把好行为做成工序——都在同一个位置上取消了那个"必须由本人亲自穿越一次"的动作；而所有读数因为结果变好，一律不报警。

十一位作者从十一个方向撞进这句话。陈晓艳写的是公共卫生政策如何把主语从个人换成人群；阳涌写的是一项人工智能中介的任务在扰动之后找不到任何承载者；季春雷写的是零摩擦的外部代理如何让主体性从生存必要条件变成可选配件；孔凡鹤写的是被工艺除去的抑制物恰恰是内源系统赖以保持激活的信号；高鹏写的是法律强制行善如何让道德语法失去自我繁衍能力；张琼写的是一个成功团体的隐性担保如何改写决断力的基础——她那句话可以直接抄下来："熄灭它的不是暴力，是拥抱。"

## ■ 五、它与另外几本书不一样在哪

站上此前有几本书离得很近，必须先划开，否则读者会以为在读续集。

第七十八号写能力、判断与责任先失去可观测性。本书不同意"可观测性全部失去"这个说法——少敏那一章证明观测反而更精密了。失去的是可执行性。

第八十号写判断不会发出警报，第八十一号写答案随时可得之后知识如何贬值。两者问的都是"那样东西还剩多少"。

第八十二号最近：它问"承重的那样东西是不是还由本人在做"，用的量是执行率与沉淀存量。本书假定它已经不由本人在做了，转而问下一个问题——当它不由本人在做时，出问题的信号寄到哪里去了。那本书数的是动作，这本书数的是地址。

## ■ 六、三处被材料改掉的想法

第一处。我们原本准备把这本书写成“能力在流失”。写到第八章不得不改：少敏的证据指向相反的方向——监督者越来越会发现问题，异议越来越精确。能力（识别的能力）在上升，失效在别处。于是全书的重心从“感觉不到”移到了“感觉到了却没有收件地址”。

第二处。我们原本给出路准备的方案是“减少评估”。十一章里没有一章支持它。陈晓艳明确拒绝把风险责任推回个人；阳涌明确拒绝把“脱离人工智能独立完成”当作唯一的能力标准。出路只好另找地方落脚，落在了第十章与第十一章。

第三处。我们原本以为最有力的一对是孔凡鹤（生理阻力）与张琼（关系担保）。装到一半发现它们隔得不够远——两章共享太多结构。真正隔得远的那一对是少敏与胡敏，而胡敏还是点进来的。这一处我们把设计上的失手写进了合章第一节。

## ■ 七、三件先认下的

一，本书新提出的三个量与那条枢纽命题，没有做过任何真实数据检验。这本书是一本理论书，不是一份研究报告。

二，十一位作者用的是同一套写作工序，也由同一位评分者评过。因此“十一章形状相同”这件事里，有一部分不是世界的性质，是工序的性质。这一条是本书最该被攻击的地方，我们把它写在枢纽章末节和附录三，不指望绕过去。

三，本书的枢纽命题有一位近得危险的先行者，而且就在同一位作者名下——胡敏的另一篇文章《回应而非创造》，写过“人被从必须自己做出回应的位置上挪开”。我们从那里拿了“位置”这个提法，在枢纽章第六节点名记账。

## ■ 八、没能回答的问题

三个量里只有一个（收件位移）给出了可立即执行的代理指标。吸收率的测量至今依赖各章自己的领域指标，跨领域不能相加。存押量最麻烦：它的定义里含有“不可逆”，而一件事是否真的不可逆，往往要到很久以后才知道——秦莉那一章诚实地承认了这一点，本书没有解决它。

还有一个更大的问题没有回答：如果收件位移确实在普遍增大，那么把地址修回来这件事，本身是不是也需要一个有权改道的位置？如果需要，谁去改第一个？这本书给不出答案，只能指出这是一个自封的门。

## ■ 九、写给四类读者各一句

如果你是在系统里做事的普通人：直接翻到枢纽章第三节和合章第二节那张表。那两处会告诉你，你感到的不对劲很可能不是你的心理问题，也不是你能力不够——是一封信寄错了地方。

如果你管着一个组织：请去取两个数——异议记录的增长率，与方案实际改道率。这两个数在你的系统里已经有了。它们的关系比任何一场满意度调查都说明问题。

如果你是临床、康复或教育中人：请读第九章的"裁定性静默区"。它是本书唯一一个已被临床验证过的出路形式，只是从来没有人把它读成一条本体论线索。

如果你是研究者或审稿人：附录一是抽样协议全文（含那次修框与那一席指定），附录三是判错清单。要判本书的错，这两处最快。

## ■ 十、不打算给行动清单的理由

这本书没有"十条建议"。原因写在合章第四节：留痕之所以能替代改道，不是因为人在敷衍，而是因为人在认真。一份看起来可执行的清单，在一个收件位移已经很大的系统里，最可能的命运就是被转成又一份"已知悉"的记录，然后一切照旧，而读数还会因此变好一点。

真要给一句，只能是这一句：留一处必须由你自己付款的地方，哪怕它很小。这句话不能替你执行，这正是它的意思。

## ■ 十一、最后

这十一位作者互不相识。他们在各自的题域里工作，没有一个人在写"AI 时代的普通人"。这本书把他们的十一篇放在一起，声称它们说的是同一件事。这个声称有可能是错的——判错的办法都在书里，欢迎来拆。

王德生 + Claude 2026 年 8 月



## 导读

### ■ 这本书是怎么做出来的

十一章里有十章是抽的，第十一章是点进来的。协议写在阅读之前，全文见附录一，其中包括抽样过程中的一次修框与那一席指定的理由。抽中的作者互不相识，写作时没有一位知道自己的文章会与另外十篇装在一起。

十一章正文除三类技术性处理外一字未改：删掉网页版的阅读条、评分条、导读框与配套读物入口；统一标题层级；把期刊体的"本文"改成书体的"本章"。各章文献表已剥出合并到书末。新写的是导论、枢纽章、合章、四篇编首与前后件，约一万八千字。

### ■ 四条读法

通读（约二十六万字）：按编序读。第一编建立"主语位已经被移走"，枢纽章给出三个量与全书唯一的新命题，第二编看被拿走的那个阻力，第三编看评估如何从两侧同时收窄，合章收束这一对，第四编处理"还剩下什么必须自己付款"。

给在系统里做事的人（约两万字）：前言 → 枢纽章第三节 → 合章第二节那张表 → 第九章的"裁定性静默区" → 结语。这条线回答的是"我感到的不对劲是什么"。

给管理者与制度设计者（约五万字）：导论 → 枢纽章 → 第八章（少敏）→ 第九章（胡敏）→ 合章 → 附录三第一条那个判错办法。第八章与第九章是全书咬合最紧的一对，先读它们比先读结论有用。

给临床、康复与教育中人（约六万字）：导论 → 第四章（孔凡鹤，阻力与内感）→ 第六章（张琼，担保与决断）→ 第九章（胡敏，裁定性干扰）→ 第十一章（胡志英，蚀先于生）。

给研究者与审稿人（约两万五千字）：附录一（抽样协议与修框记录）→ 枢纽章第五节（与五本书的分界）→ 枢纽章第六节（缺席祖宗记账）→ 附录三（判错清单）。要判本书的错，这四处最快。

二十分钟版：前言第四节 → 枢纽章第三节 → 合章第二节那张表 → 附录三第一条。

### ■ 三个必须先认识的量

吸收率  $\alpha$ ：后果被外部结构吸收、以至于本人不承担可感反馈的比例。收件位移  $\Delta$ ：产生纠错信号的位置，与拥有执行权的位置之间的距离。存押量  $\beta$ ：仍由本人不可逆承担、无法被他人吸收也无法撤销的那一部分。

全书只提出一条新命题：提高  $\alpha$  的安排同时提高  $\Delta$ ； $\Delta > 0$  之后，纠错信号不消失，而是在到达执行权之前被转换出一份"已知悉"的记录。

## 目 录

|                           |            |
|---------------------------|------------|
| 出版信息.....                 | 2          |
| 作者介绍.....                 | 4          |
| 前言.....                   | 5          |
| 导读.....                   | 9          |
| 导论 为什么读数一切正常，而人不对劲.....   | 13         |
| <b>第一编 主语位被移走.....</b>    | <b>16</b>  |
| 第一章 主语不再是个人.....          | 17         |
| 第二章 一次成功由谁承载.....         | 49         |
| 第三章 自我可以被完美绕过.....        | 83         |
| 枢纽章 信到了，人不在收件的位置上.....    | 110        |
| <b>第二编 阻力被拿走.....</b>     | <b>115</b> |
| 第四章 被工艺除去的那道阻力.....       | 116        |
| 第五章 善失去了自己的语法.....        | 143        |
| 第六章 熄灭它的不是暴力，是拥抱.....     | 170        |
| <b>第三编 评估从两侧同时收窄.....</b> | <b>189</b> |
| 第七章 带着全部信念盖下的那个章.....     | 190        |
| 第八章 越审越不能改.....           | 215        |
| 第九章 裁判席本身就是干扰.....        | 240        |
| 合章 同一个动作的两端.....          | 261        |
| <b>第四编 还剩下的那一注.....</b>   | <b>264</b> |
| 第十章 押进去就拿不回来的那一部分.....    | 265        |
| 第十一章 蚀先于生.....            | 293        |

|                       |     |
|-----------------------|-----|
| 结语 把地址写回来.....        | 320 |
| 参考书目.....             | 321 |
| 附录一 选目协议全文与十一章清单..... | 337 |
| 附录二 十一章读数速查表.....     | 340 |
| 附录三 判错清单.....         | 342 |

## 导论 为什么读数一切正常，而人不对劲

### ■ 一、一个不肯尖叫的不适

季春雷在第三章开头写过一段，值得放在整本书的最前面：有一种不适，它不尖叫；它出现在生活本该最满意的时刻——深夜放下手机屏幕暗下去的瞬间，完成一次无可挑剔的工作汇报之后，假期旅行的第二天傍晚，风景在窗外而自己不知道在看什么。它不是一个有宾语的问题，只是一个句法上的空位："然后呢？"

回答这个空位的话语极其丰裕：找到热爱的五步流程、心流量表、使命驱动的愿景陈述、无数他人正在热烈活着的证据。所有这些回答共享一个隐含前提——你感到空，是因为你缺了什么。而季春雷那一章的判断正好相反：那个问号不是因为你缺了什么，恰恰可能是因为一切都太齐全了。

这本书从这个反转出发，但不停在感受层面。感受不能被辩论，只能被承认或否认。本书要做的是把那份感受译成三个可以去数的量，并且交出一条可以被证伪的命题。

### ■ 二、把好消息读成账单

先看四组事实，它们分属四个不相干的领域，全部是好消息。

营养学的好消息：植酸、单宁这类"抗营养因子"可以被工艺除去，营养素的吸收率因此上升。孔凡鹤在第四章指出的代价是：那股阻力本身正是迫使内源提取系统保持激活的信号——去掉它，营养到了，机能沉默了。他还给了一个生理学上的桥：短链脂肪酸、肠道激素、迷走神经传入孤束核构成"内感"的分子基础；一片零阻力的补充剂产生的内感信号几乎是空白的——你无法与一个你听不到的声音对话。

团体治疗的好消息：一个成熟团体能为成员提供极高质量的隐性担保，凝聚力、利他与人际学习这些核心疗效因子全部处在最佳状态。张琼在第六章指出的代价是：这份担保会改写成成员决断力的基础——不是把它压制掉，是把它配准掉。她给的可观测指标很具体：当沉默降临时，打破沉默的那句话，是否还携带一个不容于既有语法的、粗粝的、尚未被包装的信号；还是仅仅执行了一次符合语法的过渡动作。

法律的好消息：以法律强制公民行善，可以稳定地提高救助率。高鹏在第五章指出的代价是三段退化——行为指标替代、公共语法垄断、道德主体失语；被侵蚀的不是行善的意愿，是行善的能力，那种不需要制度审批就能在个体内心自我繁衍的、以第一人称驱动的伦理能力。

人工智能的好消息：绩效差距在压缩，任务质量在达标。阳涌在第二章指出的代价是：一项任务达到了质量标准，而在模型撤除、版本更换、操作者替换之后，不存在任何能够稳定复现、解释、修复并对其负责的承载者。他把这种状态叫作无着成功，并且强调它不是低质量成功——恰恰相反，正因为结果足够好，系统才不会报警。

四个领域，四份好消息，四张同样格式的账单：被拿走的那样东西，都是"必须由本人亲自穿越一次"的那个动作；而拿走它的收益，全部记在了正账上。

### ■ 三、困境的准确位置

如果本书只说到这里，它就只是又一本"现代性使人退化"的书。这类书已经很多，而它们共同的弱点是：无法与"人本来就在变懒"这个更简单的解释区分开。

本书的分界点在这里：普通人的困境不是感觉不到，是感觉到了却没有收件地址。

少敏在第八章给出了最硬的一句话——失效发生在路径侧，不在判据侧。人工复核让监督者越来越会发现问题，异议越来越精确，可就是越来越难把这份异议转成另一条可执行的路径；组织因此同时得到更多的控制证据与更少的现实控制。陈晓艳在第一章从公共卫生的方向给出同一个空格：危险流已经接通、进入路径已经开放、至少有一个当前控制者能够切断二者，而那个位置尚无具体承载者。

这两处加上阳涌的无着成功，说的是同一件事的三个位置：信号的产生端完好，控制者在场，而中间那一格——能接住信号并且必须为之付款的位置——是空的。

感觉因此不是错觉。它是一封地址不存在的信。

### ■ 四、出路为什么不能是"回去"

本书拒绝三条看起来现成的出路。

拒绝"少用工具"。阳涌明确拒绝把"脱离人工智能独立完成"当作唯一的能力标准：个人、组织、技术三条渠道，任何一条能做到复现、解释、修复、负责这四项功能，任务就已经承载。分布不等于无主。

拒绝"把责任还给个人"。陈晓艳整章都在拒绝这条：不得以未来承受者尚未被命名为理由延期责任，也不得把风险转移回个体作为解决方案。

拒绝"少一点评估"。合章第四节说明了为什么：评估密度是可以调的旋钮，但十一章里没有一章支持"回到没有评估的状态"。而且降低评估密度在多数系统里只会提高  $\Delta$ ，不会降低它。

出路只剩一个方向：不在旋钮上，在位置上。第九章的"裁定性静默区"是它的制度形式——在不可逆时间窗内彻底不裁定，而不是换一个更好的裁定者。第十章的"存押"是它的存在论形式——在无保底、不可逆、自己承担代价的条件下把一部分存在押进一个动作。第十一章给出它的时间形式：蚀先于生，维持不是一次完成的事，生只是蚀的竞争火焰中偶然没被烧着的那一撮。

三种形式指向同一件事：保留一处必须由本人付款的位置。它很小，也不体面，不会出现在任何读数上——这恰恰是它还没有被吸收掉的原因。

## ■ 五、这本书不主张什么

本书不主张人工智能有害，不主张制度化有害，不主张回到某个更艰苦的过去。四组好消息在本书里始终是好消息。本书主张的只是：它们的账单开在另一本账上，而那本账目前没有人在看。

## 第一编

# 主语位被移走

• • •

三章分别从公共卫生、人机协作与社会存在论三个方向，处理同一件事：那个本该承担后果的位置上，具体的人不见了。

陈晓艳写"风险待承位"：危险流接通、路径开放、控制者在场，而承载者的那一格空着。阳涌写"无着成功"：任务达标，扰动之后没有人接得住。季春雷写"主体性代偿"：不是自我被摧毁了，是自我被绕过了——它从生存的必要条件变成一个可以被完美绕开的可选配件。

读这三章时值得带着一个问题：这三处空缺，是同一个空缺吗？枢纽章的回答是"是"，并给它取名为收件位移。



# 第一章 主语不再是个人

本章原题《预防如何不再以个人为主语》，作者 陈晓艳，原文见 <https://sdeuniverses.com/students/chen-xiaoyan/pre-subject-prevention/>

代际正义、物质流分析与社会实践理论 论“风险待承位”及滚动前主体预防窗

摘要

癌症预防常在风险已经附着于具体个人以后才开始说话：劝其戒烟、节制饮酒、接种疫苗、改善饮食、参加筛查或遵守防护规范。即使政策转向“上游”，也常把环境当作个人风险的远因，却不登记尚未被某个人占据的暴露位置。本章使代际正义、工业生态学的物质流分析与社会实践理论相互裁决，指出三者更深的共同前提：只有位置被人占据以后，责任、危险流和实践招募才获得同一分析单位。本章提出风险待承位：在一个预先划定的进入单元中，医学相关危险流已经接通、可复制进入路径已经开放、至少一个当前控制者能够切断二者，而该次位置尚无具体承载者的四联状态。由位置形成至其被占据，构成一扇按位置编号、可反复出现的滚动前主体预防窗，不再以整个系统的“首名承载者”作为一次性开关。癌症预防据此改写为“流向追踪—义务定位—招募拆解”，其近期结局是关闭待承位而非劝服尚不存在的人。对苏格兰无烟立法前后四项公开研究的规则前置算术检验显示，非吸烟成人唾液可替宁下降39%，酒吧工作人员约下降89%，而吸烟家庭成人仅下降16%且不显著；但对同一证据生态的记录字段审计又发现，四项研究均未记录待承位的形成时点或首次占据时点。前一结果支持规则—流向—暴露链，后一不利结果说明新对象尚未获得直接经验验证。本章给出强近邻的双向判决、竞争风险式研究设计、可删除条件及有日期的前瞻性判决线。

关键词 癌症预防；代际正义；物质流分析；社会实践理论；风险待承位；前主体预防窗；风险招募

## ■ 一、当“可预防”成为一句责问

2026年发表的一项全球归因分析估计，2022年全球约1870万新发癌症病例中，约710万例、即37.8%，可归因于研究纳入的可改变风险因素；吸烟、感染与饮酒分别约占15.1%、10.2%和3.2%[1]。作为更宽的疾病负担背景，GLOBOCAN估计同年全球新发癌症近2000万例、死亡约970万例[2]。这个结果为人口预防提供了强理由，却也极易在日常语言中发生一次危险变形：从“在特定反事实模型中，一部分人群病例可由改变暴露而避免”，

滑向“这个患者本来可以不得癌”。前一句是人群层估计，后一句是对一个具体人的因果和责任裁决。两者既不具有相同的分母，也不回答相同的问题。

癌症预防的话语尤其容易以个人作语法主语。一个人吸烟、饮酒、饮食不当、缺乏运动、没有接种、没有筛查、没有戴好防护具；于是预防就被叙述为这个人接收知识、形成动机、作出选择、保持依从。即便风险来自广告密度、产品价格、轮班劳动、空气流动、生产配方或城市基础设施，干预记录仍常从“已暴露者”开始。直到某人被识别为吸烟者、酒精高风险者、肥胖者、高危职业工人或筛查未参加者，预防对象才在数据库中出现。

这种叙述并非完全错误。戒烟、减少或避免饮酒、接种适用疫苗、采取职业防护、维持健康体重和按指南参加筛查，都可能产生真实获益。欧洲第五版癌症防治守则同时提出 14 条面向公众的建议与 14 条面向政策制定者的建议，明确把个人行动与可使行动成为可能的政策条件并列[31-33]。问题因此不是“个人是否重要”，而是：当预防第一次能够介入时，具体个人是否已经必须存在？如果一条含致癌物的产品流已经被批准，一套接触它的工作流程已经建成，一种消费实践已经获得价格、场所、技能和意义支持，却还没有人被招募为稳定使用者或暴露者，预防是否已经开始得太晚，仍是一个开放问题。

本章回答的是一条路，而不是一个标签。它不以“结构替代个人”为结论，因为人口策略、社会决定因素、商业健康决定因素、健康融入所有政策、职业卫生控制层级和源头预防设计都早已越过纯粹个人教育[21-30]。如果本章只说“癌症预防应该上游化”，它没有新的学术剩余。更严格的问题是：在没有可责问个人时，是否已经存在一种可登记、可关闭、会把风险分配给下一个进入者的对象？只有回答这个问题，预防才可能拥有不同于“更早干预”的分析单位和结局。

本章把该对象称为“风险待承位”，把它从形成到被具体人占据的区间称为“前主体预防窗”。“主体”在这里不是哲学上的人格本体，更不是说未来的人尚且不是人；它只指已经被医疗、公共卫生、雇主或市场当作风险承载者加以寻址的具体人。“待承位”也不是对未来受害者的统计想象：它必须同时具有已经接通的危险流、已经开放的进入路径、可以采取行动的当前控制者和尚未占据的具名位置。前主体预防不是无视个人，而是拒绝等到位置已经把风险过户给个人，才把预防写成其自律问题。

## ■ 二、医学与因果防火墙：预防不是对患者的反事实审判

第一道边界是对象边界。本章所称癌症，是具有异常细胞生长、侵袭及可能转移等特征的一组恶性肿瘤，不把“思想癌症”“心理癌症”或“组织癌症”当作医学实体。人格、创伤、

情绪、意志或所谓“负能量”不能在本章中被写成躯体癌症的直接原因。社会、心理和思想层面的“癌变”若作为隐喻出现，只能说明某种结构的自我扩张，不能与病理诊断互换。

第二道边界是归因边界。人口归因分值（population attributable fraction, PAF）是在明确暴露分布、相对风险、潜伏期和反事实参照下估计：若某些暴露改变，人口病例负担可能减少多少。它不是把每一例癌症拆成 37.8% 的个人责任，也不能判定某一患者的肿瘤“就是由”某一生活方式造成。PAF 对暴露定义、因果充分性、交互作用、共同暴露和时间选择敏感；将多个归因比例相加还可能重复计算[3-4]。因此，“可改变”修饰的是模型中的暴露条件，不修饰患者的道德品质。

第三道边界是证据边界。短期生物标志物或暴露下降可以支持某条过程链，却不能自动升级为癌症发病或死亡下降。动物中的跨代效应可以触发伦理警觉和人群研究，不能直接证明人类后代的癌症因果[11]。观察性关联不能越级替代随机证据，政策自然实验也须面对同期趋势、样本变化、实施差异与测量偏倚。本章后文使用唾液可替宁和细颗粒物检验无烟规则如何改变暴露，只把它们当作过程端点，不把数周至一年内的变化换算成肺癌病例。

第四道边界是临床边界。预防不能替代依法合规的个体诊疗、遗传咨询、接种与筛查指南、职业卫生标准或知情同意。不同癌种、感染、遗传易感和暴露具有不同自然史；同一种政策或生活建议不可能覆盖所有风险。2022 年仍有约 62.2% 的全球病例没有被前述模型中的可改变因素归因[1]，这部分不能被理解为“不可预防”或“纯属随机”，更不能反过来指控患者。模型没有纳入、证据不足、尚不可改变和真正无法控制，是四种不同状态。

最后，预防成功不能只以平均风险下降结算。如果一种替代材料降低消费者暴露，却把工人暴露转移到回收、焚烧或非正规拆解环节；如果面向未来的环保措施把成本集中施加于当代贫困群体；如果数字化风险识别以监控和歧视换取效率，那么平均改善不构成正当成功。本章提出的路径必须同时接受不转移风险、当代公平、未来义务和个人自主四项反向约束。

### ■ 三、先删除一个伪创新：“从个人到结构”早已发生

原始预防（primordial prevention）至少可追溯到 Strasser 1978 年的表述，其要点是在风险因素成为既成社会模式以前行动[21]。Rose 随后区分“患病个体之原因”和“发病率之原因”，指出人口中病例分布不能只由高风险个体策略解释[22]。Frieden 的健康影响金字塔把社会经济条件、改变默认环境和长期保护性干预置于咨询教育之下更广的影响层[23]。商业健康决定因素研究又把产品组合、市场塑造、游说、知识生产和全球经济系统纳入疾病形成路径[24-25]。健康融入所有政策与复杂系统公共卫生进一步要求跨部门和反馈性治理[26-27]。

工程与职业卫生的答案更直接。美国职业安全与健康研究所的控制层级把消除、替代和工程控制置于行政控制与个人防护具之前，因为较高层级较少依赖持续、完美的人类操作\ [29]。扩大生产者责任把产品消费后的物质与财务责任前移到生产者，迫使设计阶段考虑全生命周期\ [30]。第五版欧洲癌症防治守则则把无烟环境、价格、营销、食品环境、城市活动基础设施与职业保护等政策建议同公众建议成对提出\ [31-33]。所以，“别只怪个人”“要改变结构”“预防要上游”既非空白，也非本章可据为已有的判断。

行为科学同样不是个人主义的同义词。COM-B 模型把行为理解为能力、机会与动机的共同结果，并把环境重构、限制、示范、赋能等纳入干预功能\ [28]。助推与选择架构改变默认项、可见性和行动摩擦，虽可能仍保留个人选择，却已不只提供知识。社会实践取向更进一步：它不把行为当作孤立个人携带的属性，而研究材料、能力、意义怎样结合成实践，实践怎样招募承载者并同其他实践结盟\ [16-20]。

因此，本章必须接受一项敌意起点。第二轮近邻搜索不是只在公共卫生内部寻找同义词，而是沿五个方向主动寻找“这件事早已被说过”：时间方向的原始预防，结构方向的人口策略、商业决定因素与复杂系统，工程方向的控制层级和源头设计，暴露方向的风险环境与暴露机会模型，主体方向的生态可供性、行为场景和“被配置的用户”\ [21-29,48-53]。其中 Strasser、Barker 与 Gibson 的工作均早于 1980 年；源头设计、风险环境和暴露机会模型又已进入实际安全、减害或暴露研究。由此，凡是可以被“上游”“环境”“机会”“脚本”或“实践招募”完整表达的内容，都不构成本章的独立贡献。

新概念只有在四项条件同时满足时才有暂时资格：第一，它登记的是近邻账本普遍漏掉的对象，而非给既有对象换名；第二，它能给出可复核的形成、占据和预防关闭边界；第三，它与至少一个强近邻在同一观察场景中给出相反的有无判断，而非仅仅“更重视某因素”；第四，它产生关于待承位关闭先于新暴露、再先于疾病终点的可反驳预测。缺少任何一项，都应删除名称。

## ■ 四、三条完整路径在何处互相否决

代际正义、物质流分析与社会实践理论并非同一母学科的三个分支。前者属于政治哲学和规范伦理，终点是当前行动者对不同时代的人负有什么义务；后者属于工业生态学，以系统边界、库存、流量和质量守恒追踪物质；第三者属于社会理论，解释实践怎样形成、维持、招募与解体。三者在本题中的过程和完成标准互不相容。

表 1 三门学科对癌症预防过程的互斥裁决

| 学科路径 | 起点 | 过程装置 | 完成终点 | 对另两路的根本异 |
|------|----|------|------|----------|
|------|----|------|------|----------|

|        |                    |                                 |                     | 议                                       |
|--------|--------------------|---------------------------------|---------------------|-----------------------------------------|
| 代际正义   | 当前活动的收益与<br>伤害跨期分离 | 权利、能力、非同<br>一性、折扣与义务<br>承担者     | 谁现在应当为尚不<br>能同意者做什么 | 流量无权决定负担<br>是否正当；实践延<br>续也不能免除责任        |
| 物质流分析  | 明确边界内的投<br>入、库存与输出 | 质量守恒、过程节<br>点、排放、替代、<br>循环与末端处置 | 危险物质实际去了<br>哪里      | 善意不改变质量平<br>衡；“绿色”叙事<br>不能替代去向数据        |
| 社会实践理论 | 材料、能力与意义<br>的可结合要素 | 结合、重复、结<br>盟、招募、退出与<br>承载       | 谁怎样成为某实践<br>的承载者    | 把风险者视作先在<br>选择者，会漏掉偏<br>好和身份如何被实<br>践生产 |

第一处冲突是“责任能否由流量推出”。物质流分析可以证明某种危险元素从生产转入空气、废物或回收工序，却不能仅凭吨位回答谁负有义务。代际正义则会追问：谁获益、谁控制、谁承受、谁无法同意，为什么未来损害不能用当期利润抵消。反过来，规范原则即使写得完美，若不知道物质在替代和末端处置中怎样迁移，义务可能只改变报表而不改变暴露。

第二处冲突是“个人是否先于实践”。代际伦理通常要找一个可以对其说明理由的受益者或受害者；物质流也常在系统边界末端添加“人体受体”。社会实践理论却指出，吸烟、饮酒、通勤、烹饪、劳动或筛查参加并非先有稳定偏好再执行。材料可得性、技能、时序、同伴关系和意义结合以后，实践才招募承载者；动机、欲望和身份可以是参与的结果[17-18]。如果如此，等到“有动机的个人”出现才预防，已经把招募过程漏掉。

第三处冲突是“循环是否天然等于预防”。工业生态学倾向提高资源效率和循环率，但一项增材制造末端物料管理研究显示，追踪约 1044 万千克材料时，焚烧、填埋和回收各自形成不同的空气、土地、水与职业接触路径；每千克投入在不同情景下可关联到可量化的环境释放[14]。欧洲电子废物人群生物监测也在 195 名工人与 73 名对照中发现，多个废物处理类别的铅暴露指标升高，并观察到体内外指标相关[15]。这些材料并不证明所有回收致癌，却推翻“离开消费者视野即消失”的假设。社会实践理论还会补问：循环技术是否创造了新的非正规拆解实践和劳动招募？代际正义则追问：谁替未来资源节约承担当前暴露？

三条路不能合并成“兼顾伦理、环境和行为”。它们必须互相设限：流量使规范义务落到真实节点，规范义务阻止流量最优牺牲弱者，实践过程则解释为何材料和规则改变仍可能被日常生活绕开。只有保持这种冲突，才可能生成一条不属于任何单门学科的预防路径。

■ 五、代际正义：在受害者尚未可识别时，义务已经成立



Rawls 以正义储蓄原则讨论世代之间如何维持正义制度的物质条件[5]。Parfit 的非同一性问题则制造更尖锐的困难：今天的政策可能改变未来谁会出生；某个未来人若只能在该政策下存在，便难以简单声称“若选择另一政策，他会过得更好”，因为另一政策下可能根本没有同一个人[6]。如果权利主张必须等到特定受害者身份固定，许多跨代环境与健康风险会逃离责任。

这一难题没有取消义务，反而迫使义务从“对某个已知人造成可比较损害”转向更一般的正当性：当前人是否以不可接受的方式塑造未来人的生存条件，是否把严重且可避免的风险施加给无法同意者，是否只因时间距离而折扣基本权利。Gardiner 把气候问题称为“完美道德风暴”，因为全球、跨代与理论层面的分散使当代行为者易于把责任推远[7]；Caney 反对仅因损害发生较晚就折扣对基本利益的保护[8]；Ekeli 和 Shrader-Frechette 也分别从不确定性、预防原则与风险分配论证面向未来者的责任[9-10]。

用于癌症预防，代际正义给出三个规范步骤。第一，确认风险和收益是否跨期分离：当前市场、税收、就业或便利由一群人获得，较长潜伏期后的癌症风险由另一群人承担。第二，不以未来承受者尚未被命名为责任延期理由；只要严重伤害具有合理可预见性、当前主体拥有控制能力，就可以形成当期义务。第三，义务按控制、获益、能力和造成风险的关系定位，而不是把全部负担留给未来个人在暴露后自救。

但代际正义不能单独完成预防。它容易把“未来人”写成抽象集合，却不知道某种化学品在产品、排放和废弃后会进入哪个环境或职业节点；也可能以未来利益之名，要求当代低收入群体承担昂贵转型。动物实验提示某些暴露可能具有跨代生物效应时，伦理上可以提高研究和谨慎义务，却不能把动物证据升级成人类癌症事实[11]。所以规范路径必须交给物质流回答“风险在哪里”，再交给实践理论回答“谁会怎样被招募进去”。

## ■ 六、物质流分析：在人体受体出现以前，危险已经有了去向

物质流分析（material flow analysis, MFA）以明示的空间、时间与过程边界记录投入、库存、转换与输出，其基本纪律是质量守恒：进入系统的物质不会因改名、售出、稀释或移出账面而消失[12-13]。它可以在尚未观察到具体患者时追踪潜在危险材料，从原料开采、制造、产品使用、维护、回收、焚烧、填埋到环境排放和职业接触。这里的“潜在”不是“必然致癌”；危险性、暴露、剂量和疾病仍须分别验证。但如果物质尚未到达人体，预防已经可以在流量节点发生。

MFA 对传统癌症预防有三项纠正。其一，它把末端个人暴露还原为一条上游链。工人没有佩戴好防护具可能是最后一个可见事件，之前却已有配方选择、设备密闭、采购标准、生产

节拍、通风设计和废物合同。控制层级之所以把消除、替代和工程控制置于个人防护具之上，不只是因为前者“更结构化”，而是因为它们在危险流与人体接触以前改变系统[29]。

其二，它揭示“减量”与“转移”的差别。一个工厂的排放下降，可能是危险物进入产品、废水或外包环节；消费者暴露下降，可能伴随回收工人暴露上升；禁止一种物质，替代品若具有相似危险或需要更高能耗，也可能形成遗憾替代。扩大生产者责任若只提高回收率而不追踪工艺、非正规劳动和残余排放，也不能自动满足预防[30]。真正的源头阻断须说明质量去了哪里、危险性怎样变化、谁接近新流向。

其三，它产生比发病率更早的过程指标。生产投入量、含危险成分的产品份额、密闭率、环境释放、岗位接触时间、个人监测和生物监测可以构成从流到人的多级证据。由于癌症潜伏期长，等待发病率变化才判断政策，常使错误配置持续多年。MFA 允许先检验“危险流是否减少或改道”，但不能越级宣称“癌症已经预防”。这一证据节制恰是它与医学因果链的接口。

MFA 自己的盲点同样明显。它容易把人当作过程图上的终端节点，把“消费者”“工人”“居民”视作接受剂量的容器。可是人为什么进入某岗位、为何某产品成为日常必需、为何替代方案无法被使用，并不由质量平衡解释。即使总流量不变，不同的场所规则和实践组合也会造成完全不同的招募和暴露。要回答风险怎样获得承载者，必须转向社会实践理论。

## ■ 七、社会实践理论：不是个人选择实践，而是实践招募个人

Reckwitz 把实践描述为由身体活动、心理活动、物、背景知识、情绪和动机等要素联结而成的程式化行为类型[16]。Shove、Pantzar 与 Watson 进一步以材料、能力和意义分析实践的生成、维持与消亡[17]。在这个框架里，个人是实践的承载者，却不是全部解释的起点。实践只有在材料可得、能力可学、意义可认时才能持续；它也必须不断招募新承载者，以替代退出者。

Blue 等把这一视角引入公共卫生，强调吸烟等活动的“生命史”、实践之间的联盟以及招募与退出，而不是只修改个人认知[18]。饮酒研究也显示，不同场合、关系、器物、时间和意义组成不同饮酒实践，不能用一个总量变量解释所有过程[19]。关于肥胖社会的实践研究则拒绝把身体重量只写成个人选择结果，转而考察进食、工作、照护、通勤与休闲实践如何联结[20]。这些研究并不否认意图，却把意图放回实践形成过程：人想要什么，部分取决于他已经学会并参与了什么。

这一主体路径给癌症预防带来三个动作。第一，不从“谁缺乏动机”开始，而从某种实践需要哪些材料、能力和意义开始。吸烟需要产品、点火和场所，也需要社交节奏、身份象征或

应对压力的意义；职业暴露依赖工具、工艺、绩效要求和“熟练工就是这样做”的规范；久坐与高能量饮食可能嵌入通勤、轮班、照护和廉价便利的实践束。

第二，把招募率与退出率设为过程端点。一次宣传可能改变风险知识，却没有减少实践吸收新成员的能力；价格、可得性、广告、同伴传授、场所和技能改变，才可能使新进入减少。反之，一项政策可以降低公共场所暴露，却在未覆盖的家庭、非正规工作或数字营销中保留实践通道。预防成败因而不是“多少人被劝服”，而是实践是否仍能稳定地获得承载者。

第三，干预实践之间的联盟。饮酒可同庆祝、社交、观看体育、下班放松结盟；吸烟可同休息制度、同伴互惠和身份建立结盟；高风险工作步骤可同计件速度、设备维护和职业尊严结盟。只去除一个元素，联盟可能用别的材料或意义恢复。实践理论因此反对“一处按钮”的简单机制，也反对把政策未覆盖处的残余暴露重新归咎于个人。

但社会实践理论不能自己判定某种材料是否致癌，不能仅凭“实践可被重组”决定哪些限制正当，也不能把参与者当作毫无自主的载体。它必须接受医学危险证据、代际义务和权利边界。主体由实践部分生成，不等于主体没有反思和拒绝能力；预防不能借“拆解实践”之名实施强制监控、污名化或剥夺合法选择。

## ■ 八、共同前提被各自的内部材料推翻

三条路径表面相距极远，却共享两层未被明说的前提。第一层是个人先行：先有一个可以承受风险的具体个人，然后才对其配置义务、计算物质流入或解释其参加实践。政策争论据此变成“应该让个人负责，还是让结构负责”；两边都默认风险主体已经在那里，只争谁来改变他周围的条件。

第二层更隐蔽，是占据才算存在：即便承认义务、流和招募可以先于疾病，只要尚无人占据某个暴露位置，三个账本都会把它记成空白。规范伦理有尚未命名的人，却没有哪一个具体暴露位置正等待分配；MFA有库存和排放，却没有“下一次进入将把哪条流接到谁身上”的字段；实践理论有招募能力，却常以实践已有承载者的生命史为经验入口。三家争论的是占据以后该记录什么，却共同漏掉“尚未占据”本身可能是一种正状态。

代际正义首先从内部推翻这一顺序。非同性问题表明，未来受影响者身份未固定，不能成为当前不作为的充分理由[6]。义务可以先于具体受益者或受害者被命名，只要当前活动以可预见方式塑造未来人的基本条件。换言之，规范对象可以在可寻址个体出现以前成立。

物质流分析从另一面推翻它。质量守恒不需要先找到患者，才能记录含危险成分的原料进入了哪条生产、使用和废弃路径[12]。流量、库存和排放可以先于人体受体被计算；如果必须



等到生物监测阳性才追踪，MFA 的源头意义反而被取消。过程对象也可以在个体出现以前成立。

社会实践理论则推翻“稳定动机先于参与”。实践不仅响应已经存在的偏好，还通过材料、能力、意义和联盟生产熟悉感、欲望、身份与重复 [17-18]。第一支烟、第一份高暴露工作或第一次被算法营销触达，不是一个完整自我在真空中选择；它也是一个招募安排的结果。主体对象并不总在过程之前完成。

三组内部反证合在一起逼出一个更强判断。代际正义使“无人可命名”不再等于“无人负义务”；MFA 使“尚无人接触”不再等于“无危险流”；实践理论使“尚无人选择”不再等于“无进入路径”。如果三种否定同时成立，尚未占据就不只是资料缺失，而可能是一个能够被预防性关闭的对象状态。

癌症风险不必等到附着于某个可责问的个人才成为预防对象。当危险流已接到一个开放进入位置、当前控制者有能力切断连接而该位置尚未被占据时，系统已经产生了一个“风险待承位”。预防的最早单位不是尚不存在的人，而是这个可以在过户风险以前被关闭的位置。

这个判断不等于“结构决定个人”，也不声称所有危险都预先占好了座位。没有真实危险连接、没有可复制进入路径、没有当前控制者，或位置形成时已经有人占据，待承位都记为零。某些家庭、职业或感染情境中，形成与占据几乎同时发生，窗口为空；在弥散消费市场中，位置甚至可能无法可靠划界。正因如此，概念必须允许零、不可判定和领域限缩，而不能以“广义上总有上游”免于反驳。

## ■ 九、从一只钟到一个对象：风险待承位及滚动前主体预防窗

原定义把  $tR$  写成“整个安排的首名承载者进入”。这比泛称上游更严格，却有一个致命漏洞：第一个人进入以后，同一生产线、平台或消费实践仍会不断产生后来者；若整个系统只允许一扇窗口，那么首人之后所有尚未被占据的位置都被错误地划到“主体已出现”一侧。换言之，旧定义把主体当作全系统一次性开关，无法处理岗位轮替、滚动招生、分批上市和连续招募。

修正方法不是把窗口拉长，而是把钟绑定到位置。先划定风险安排  $a$ ，再预先规定一个可重复识别的进入位置  $p$ 。 $p$  不是未来人的身份，也不限于劳动岗位；它可以是一个尚未分配的任务位、一个服务资格—时段单元、一个销售渠道中的下一批进入单元，或任何具有明确进入规则和暴露连接的事件位。若研究者只能说“未来大概会有人受影响”，却不能说明哪一种进入规则把哪条危险流接到哪一类活动上， $p$  不存在。

在时间  $t$ ，定义四个可观测条件： $F_{ap}(t) \setminus > 0$  表示具有医学相关证据的危险或风险增益流已经接到  $p$  的通常活动； $R_{ap}(t)=1$  表示材料、能力、意义、规则与场所已组成可复制的进入路径； $C_{ap}(t)$  表示至少一个具名当前行动者拥有在预定期限内切断  $F$  或  $R$  的实际能力； $O_{ap}(t)=0$  表示该次位置尚未由具体人占据。风险待承位  $V_{ap}(t)$  定义为：

$$V_{ap}(t) = 1。$$

这一定义故意比“危险”“机会”“场景”更窄。仓库中有致癌危险物却与任何进入位置物理隔离， $F_{ap}=0$ ；产品获批但没有实际渠道、技能或使用场景， $R_{ap}=0$ ；只有倡议者而无人能改变连接， $C_{ap}$  为空；已有工人、消费者或居民正在该次位置承受暴露， $O_{ap}=1$ 。四种情形都不是风险待承位。它们仍可构成重要的危险、治理失败或现有暴露，却不能借本章名称获得额外新颖性。

令  $tA(p)$  为  $V_{ap}$  首次从 0 变为 1 的时间， $tR(p)$  为该位置首次被具体人占据并开始相关暴露或实践承载的时间， $tC(p)$  为位置在占据前因消除、替代、隔离、准入改变或招募拆解而被预防性关闭的时间。若占据先发生，则位置  $p$  的严格前主体预防窗为：

$$W_{pre}(p) = \setminus [tA(p), tR(p));$$

若预防关闭先发生，则观察到的是  $\setminus [tA(p), tC(p))$  及结局“关闭先于占据”。 $tR(p)$  与  $tC(p)$  是竞争事件：一个把风险过户给具体人，另一个使这次过户不再发生。安排  $a$  在时点  $t$  可以同时拥有一组开放位置  $\Omega_a(t)=$ ；一个位置被占据不使其他位置消失，新位置也可随后形成。因此，前主体窗口是按位置编号、随招募滚动更新的窗口族，而不是产业或实践历史上只发生一次的黄金时刻。

这项修正也建立了零情态词判据。肯定编码必须逐字回答五个问题：哪条具有医学相关性的流已经接通；哪个进入位置已经开放；一个符合进入规则的人按通常路径进入时怎样承受相关暴露；谁能在占据前切断流或路径；什么记录证明该次位置尚未占据。“可能接触”“将来有人”“企业应负责”“环境有风险”都不能单独记为 1。若位置单位事后为了配合结局才划定，或匿名、弥散市场使  $tA(p)$  与  $tR(p)$  无法区分，则结果为暂不判定，而不是用理论想象补齐。

表 2 风险待承位从形成到关闭、占据及疾病的阶段阶梯

| 阶段       | 可观测状态                        | 主要控制者       | 预防动作             | 常见误判           |
|----------|------------------------------|-------------|------------------|----------------|
| 0 潜在危险   | 要素存在但尚未结合， $F$ 或 $R$ 至少一项为 0 | 研发者、采购者、监管者 | 危险评估、安全设计、拒绝遗憾替代 | 把实验危险性直接当成人群病因 |
| 1 危险连接安装 | 产品、配方、设备、排放或渠道已              | 生产者、平台、业    | 消除、替代、密闭、准入和流量上  | 有危险即宣称已有       |

|           | 进入系统，但进入位置未必开放                                         | 主、公共部门          | 限                | 待承位               |
|-----------|--------------------------------------------------------|-----------------|------------------|-------------------|
| 2 待承位形成   | $F \setminus > 0$ 、 $R=1$ 、 $C$ 非空、 $O=0$ ， $V_{ap}=1$ | 多个共同控制者         | 在占据前切断流、入口或招募结合  | 把尚无个体病例当作尚无可关闭对象  |
| 3A 预防性关闭  | $tC(p)$ 先发生，位置未向个人过户风险                                 | 关闭流或入口的控制者      | 核验替代、反弹、公平与风险转移  | 只有政策文本即宣称关闭       |
| 3B 位置被占据  | $tR(p)$ 先发生， $O$ 由 0 变 1                               | 雇主、平台、场所管理者与参与者 | 上游控制加定向保护、退出支持   | 以首位被占据为整个系统窗口永久关闭 |
| 4 个体风险可识别 | $tL$ 发生，可由监测或行为标签寻址                                    | 医疗、公共卫生、雇主与个人   | 个体支持、监测、接种、筛查或治疗 | 把识别时点倒写为因果起点      |
| 5 生物损伤或疾病 | 生物标志、癌前病变或恶性肿瘤出现                                       | 临床团队、患者、制度      | 二级预防、早诊、治疗与照护    | 把治疗成功等同于上游风险已消除   |

这一阶梯不是癌症自然史模型，也不是临床分期。阶段 0—4 描述风险连接、位置和可寻址主体的形成，阶段 5 才可能进入医学损伤；不同癌种和暴露之间不能用同一时长。它的功能不只是把“患者首次被看见”同“风险首次开始”分开，还把两种常被混写的零分开：尚无待承位，与待承位已被预防性关闭。

■ 十、核心路径：流向追踪—义务定位—招募拆解

前主体预防不是一句“越早越好”，而是三个有先后又需反复校正的动作。第一个动作是流向追踪：先画出危险或风险增益要素从设计、采购、生产、销售、使用到排放和末端处置的真实路径，标明库存、速率、接触机会、替代物及数据缺口，再指出它接到了哪个进入位置。医学证据决定追踪何种危险，质量平衡阻止“看不见即不存在”。这一动作的结果不是一个患者名单，而是“流—位置”连接图；没有接位证据，危险流不能自动生成待承位。

第二个动作是义务定位：对每条“流—位置”连接回答谁拥有准入、配方、价格、场所、工作节拍、通风、信息、废物合同或监管执法的控制能力；谁从该安排获益；谁把风险外部化；谁无法同意或退出；谁有资源采取更有效控制。义务不能抽象地落给“社会”，也不能因存在多个控制者而消失。它要被写成当前时点可执行的断开动作、期限与核验指标。代际影响较长时，未来受害者尚不可识别并不令义务暂停。

第三个动作是招募拆解：识别一种高风险实践赖以开放下一个位置的材料、能力、意义和联盟，寻找最小但关键的结合点。拆解不等于只禁止材料，也不等于说服个人。它可以减少可得性、改变默认场所、取消将风险同身份或奖励绑定的意义、提供退出所需能力、重组工作节

拍，或切断高风险实践与社交、照护和谋生实践之间的联盟。动作完成后要观察待承位形成率、占据率、预防关闭率、退出、替代实践和空间转移，而不能只看知识分数。

三步并非各做一项即可。若只追踪流量，没有义务定位，报告会停在“系统复杂”；若只定位义务，没有招募拆解，规则可能被生活实践绕开；若只改变实践，没有追踪流量，危险物可能改道至更隐蔽的人群。真正的优先点是对位置  $p$  而言三个动作首次共同可行的时点。令  $I$  为候选干预， $\Delta F_p(I)$  为危险流下降或与  $p$  断开的可核验变化， $\Delta R_p(I)$  为该进入路径关闭， $D_p(I)$  为控制者义务可落实， $T_p(I)$  为风险未向更弱势群体或未来转移。最早共同可逆点可概念化为：

$$\tau^*(p) = \min。$$

这不是一个已经验证的算法， $\Delta F$  与  $\Delta R$  也不能跨领域共用单位。表达式只建立四个不可互相补偿的门槛：危险流同位置的连接确实改变，进入路径确实改变，义务有现实承担者，改善没有靠转移风险取得。近期完成标准不是“对象知道风险了”，而是  $tC(p)$  先于  $tR(p)$  发生。若不存在同时满足者，研究应如实报告“当前无共同节点”，而不是给最早节点贴上理想标签。

### （一）为什么三步的顺序不能倒置

从个人劝导开始，研究者首先获得的是已进入实践者的特征，随后才回溯环境，容易把选择后的差异误作选择前原因。流向追踪先行，则可在不知道未来谁会暴露时锁定材料和制度路径。义务定位居中，是为了防止技术分析假装无人负责；招募拆解最后，是因为只有知道何物在流、谁能控制，才能分辨一种实践真正不可缺的要素与可替代装饰。

但顺序不是线性瀑布。实践调查可能发现人们使用某产品并非研究者假定的意义，迫使 MFA 重画系统边界；物质追踪可能发现高风险转移至外包工序，迫使伦理分析重新定位义务；义务审查可能显示监管者无执行能力，干预便需回到生产和采购节点。三步构成闭环，优先序只说明进入问题的第一视角。

### （二）个体行动在路径中的位置

个人仍有真实行动能力。已吸烟者戒烟、家长建立无烟家庭、工人报告泄漏、公众接种疫苗或参加适用筛查，都可能挽救生命。本章只拒绝两种推论：其一，个人行动有效，所以风险起点必在个人；其二，个人没有成功退出，所以制度和生产过程没有责任。一旦  $tR(p)$  先于  $tC(p)$  发生，个体支持、临床预防与上游控制必须并行；不得为了提高“关闭先于占据”的漂亮比例而撤减已有承载者服务。把个人从第一主语移开，不是把他从句子中删除，而是把他从最后责任人改写为享有保护、支持退出并参与治理的行动者。

## ■ 十一、二维判决之后再问一次：位置究竟关闭还是被占据

“政策实施”不是过程完成。为了阻止单一指标冒充成功，本章先用两条互不替代的轴裁决近期机制：危险流或接触路径是否实质改变，实践招募的关键结合是否被破坏。规范正当性不进入二维表作为第三个可加分维度，而是在所有格外设置硬门槛：任何平均改善若依赖风险转移或不正当强制，都不能记为成功。

表 3 危险流改变与实践招募改变的二维判决

| 危险流／实践招募 | 招募结合未被破坏                           | 招募结合被破坏                                       |
|----------|------------------------------------|-----------------------------------------------|
| 危险流未实质改变 | 空转：教育、承诺或标签变化，但材料与进入条件仍在；新暴露可继续发生  | 实践绕行：原实践暂时失去成员，却保留危险库存和渠道；新实践、替代品牌或其他场所可能重新招募 |
| 危险流实质改变  | 转移或残余：某场所暴露下降，但未覆盖实践、外包环节或替代品仍维持招募 | 窗口封闭候选：流量与招募同时下降；仍须通过义务、公平、不转移和长期反弹审查         |

左上格提醒，风险知识上升不等于物质条件变化。右上格提醒，反营销或退出支持即使暂时有效，如果供应、价格和场所不变，实践可能换一种意义重新出现。左下格是最常见的政策半成功：公共室内无烟规则显著降低场所烟雾，却不自动拆解家庭吸烟或数字化尼古丁营销；工厂密闭减少固定岗位暴露，却可能把维护和废物处理变成高峰暴露。右下格也只是“候选”，因为严格控制可通过失业、价格负担或外包把损害转给别处。

二维表仍不足以判定风险待承位。只有在危险连接和进入路径都已存在、位置尚未占据且有当前控制者时，右下格之前才出现  $V_{ap}=1$ ；干预之后还必须区分  $tC(p)$  先发生还是  $tR(p)$  先发生。两个项目即使得到相同的平均暴露降幅，也可能有相反过程：一个在十个岗位分配前关闭八个待承位，另一个让十人先暴露再使八人退出。前者的近期结局是 8 次“关闭先于占据”，后者是 0 次；把两者都写成 80% 改善，会把预防和补救混为一谈。

这套判决产生一条比“上游更早”更具体的可反驳时间预测：对预先登记的位置队列，干预首先提高“ $tC(p)$  先于  $tR(p)$ ”的累积发生，继而降低新暴露发生率，最后才可能影响癌前病变和癌症发病。若只有已有参与者的自报意图改变，没有待承位关闭或客观新暴露差异，本章机制不成立。若不登记位置也能由既有源头设计或实践变量得到完全相同的外部预测，待承位只是多余中介；若高质量个人方案在同等资源下同样或更早改变新进入率，也没有理由宣称前主体节点具有一般优先权。

## ■ 十二、十二个强近邻：必须给出相反预测，而非不同强调

风险待承位必须在最危险的既有概念旁接受双向判决，而不是只同最弱的个体教育比较。下表每一行都写出两种会让对方失败的观察。若只能说本章“更重视时间、伦理或结构”，而不能产生相反的有无判断，名称没有独立资格。

表 4 十二个强近邻与风险待承位的双向预测判决

| 强近邻                 | 若近邻足够、本章对象为零                                      | 若待承位独立、近邻不够                               | 当前判决                     |
|---------------------|---------------------------------------------------|-------------------------------------------|--------------------------|
| 原始预防\ [21\]         | 风险因素尚未成为社会模式时应行动，即使 F 与 R 均未形成；此时 V=0             | 风险因素和进入路径已经形成但 p 尚未占据；原始预防可能判为已过“形成前”，V=1 | 时间区间交错但不相同；需增量预测裁决       |
| Rose 人口策略\ [22\]    | 整体暴露分布可移动，却未必存在可编号位置；V 可全为 0                      | 小型封闭系统可先关闭一批 p，而人口分布尚无 measurable 变化      | 分布单位与位置单位相反；两者可并存        |
| 商业健康决定因素\ [24-25\]  | 企业游说、定价或市场塑造存在，但尚无医学相关“流—位置”连接；V=0                | 非商业的新工艺或公共设施也可形成危险待承位；商业变量=0 而 V=1        | 本章不能被“企业影响”一词完全吸收        |
| 复杂系统公共卫生\ [27\]     | 有反馈、适应和涌现，却没有稳定进入单元；系统命题成立而 V 不可判                 | 一条低反馈、边界清楚的新生产线已有未分配危险岗位；V=1 而复杂性并非必要解释   | 复杂性不是待承位的必要条件            |
| COM-B／行为改变轮\ [28\]  | 已识别行动者具有能力、机会和动机，但其行为没有接到医学相关危险流；COM-B 可预测行为而 V=0 | 无具体行动者可测动机时，一个开放危险岗位已构成 V=1               | 个体行为状态与位置状态不是同一变量        |
| 控制层级与源头设计\ [29,53\] | 设计阶段发现危险应消除，即使进入路径尚未开放；PtD 成立而 V=0                | 同一危险源控制后若替代流接到另一个开放位置，技术项目可报“源已替代”，V 仍为 1 | 这是最强工程吸收者；全生命周期重评可消除本章剩余 |
| 实践导向公共卫生\ [18-20\]  | 一个已有承载者、正在衰退且不再招募的实践仍值得研究；实践存在而开放 V 可为 0          | 一个尚无人进入、但材料—能力—意义已接到具名危险位置的事件可有 V=1       | 若实践研究能稳定编码未占据位置，本章将被吸收   |
| 风险环境\ [48\]         | 当前使用者在高风险社会—物理环境中受害，即使没有下一处开放位置；风险环境高而 V=0        | 尚无使用者或工人的危险位置已开放；V=1 而个体伤害资料仍为零           | 环境是场，待承位是尚未完成的风险分配事件     |



|                             |                                                  |                                                   |                                |
|-----------------------------|--------------------------------------------------|---------------------------------------------------|--------------------------------|
| 暴露机会模型\ [49]                | 已知人员轨迹靠近既往<br>喷洒区，可计算暴露机<br>会；若无新进入位置则<br>V=0    | 未分配岗位已同危险工<br>序连接，V=1；因无人<br>轨迹，个体暴露机会指<br>数尚无定义  | “人的潜在暴<br>露”与“位置等待承载<br>者”方向相反 |
| Gibson 生态可供性\ [50]          | 无害对象可以向某类人<br>提供行动可能；可供性<br>存在而 F=0、V=0          | 进入岗位即被动接触危<br>险，未必存在可感知或<br>意图性的行动可供性；<br>V=1     | 可供性过宽，不能替代<br>医学流与控制条件         |
| Barker 行为场景／配<br>置压力\ [51]  | 场景可有稳定行为模式<br>或缺员压力，却没有医<br>学相关危险流；场景成<br>立而 V=0 | 一次性维修任务可有开<br>放危险位置，却尚未形<br>成稳定行为场景；V=1           | “缺人”是形式近邻，<br>医学连接和义务使二者<br>分开 |
| Woolgar “被配置的用<br>户” \ [52] | 产品脚本可以界定、约<br>束一个无害用户；用户<br>被配置而 V=0             | 非自愿、偶发或维护性<br>暴露可接到位置，却没<br>有被设计者明确写入用<br>户脚本；V=1 | 脚本解释预设角色，不<br>自动解释风险过户         |

最危险的并非任何单个近邻，而是三者联盟：源头设计已经要求在人接触前消除危险，实践导向公共卫生已经追踪招募，用户配置研究已经说明技术预设未来角色 \ [18-20,52-53]。风险待承位目前只剩一项可能的增量：把“危险连接—开放进入—当前控制—尚未占据”登记为同一位置事件，并以关闭先于占据作为可观察结局。若三类既有记录在不增加该字段的条件下即可产生相同位置队列、相同排序和相同外部预测，本章术语应被删除，而不是宣称自己“综合得更全面”。

历史边界也必须明说。本章不声称最早提出上游、人口或源头预防；不声称首次反对责备患者；不声称首次提出环境招募人、技术预设用户或暴露可以被建模；更不声称所有癌症可在人出现以前被预防。其狭窄主张只有一条：某些风险安排会反复产生尚未占据、却已经把危险流接到开放入口并可被当前控制者关闭的位置；这些位置可以组成独立队列，接受“关闭”与“占据”的竞争事件检验。

■ 十三、真实数据检验：苏格兰无烟立法的改变与边界

（一）为什么选择无烟立法

烟草烟雾对人致癌，二手烟暴露亦有充分致癌证据 \ [36]。苏格兰于 2006 年实施公共场所和工作场所无烟立法，为检验一条狭窄过程命题提供了自然实验：不把非吸烟者的个人行为改变作为直接靶点，只改变场所规则与烟雾流向，客观暴露能否迅速下降？同时，未被公共场所

规则覆盖的家庭实践是否保留残余通道？该案例从已存在成人、儿童和在岗工人开始，不能直接观察占据前的位置，却可以对路径中的两个必要环节——流向改变和实践边界——进行不利检验。

本章在进一步提取和跨研究解释之前，先锁定主要端点、算式和反面结果。工作文件未存入第三方不可篡改平台，且部分研究方向在规则锁定前已经可见，因此这里只称“规则前置的算术检验”，不包装为正式预注册。主要客观端点为非吸烟者唾液可替宁几何均值；辅助端点为过去7天工作场所二手烟暴露时数和个人PM<sub>2.5</sub>。降幅统一按 $100 \times (\text{基线} - \text{随访}) / \text{基线}$ 计算，不进行跨样本合并。

事先设定的最低支持线为：至少一个没有被直接要求改变本人吸烟行为的群体，客观暴露下降不少于25%。实践残余线为：吸烟家庭人群降幅不足非吸烟家庭的一半，或差异不显著。若公共禁烟被研究证实导致家庭暴露上升，正文必须写风险转移；若所有家庭分层下降完全一致，则撤回实践边界解释。最重要的禁止推断是：可替宁或PM<sub>2.5</sub>下降不能在本章中直接换算成肺癌发病率下降。

（二）规则锁定后的结果

Haw 与 Gruer 对苏格兰非吸烟成人的重复横断面调查显示，唾液可替宁几何均值由立法前0.43 ng/mL 降至立法后0.26 ng/mL，按统一算式为39.5%，与论文报告的39%一致；来自非吸烟家庭者由0.35 降至0.18 ng/mL，下降48.6%，而来自吸烟家庭者的发表降幅仅16%且无统计学显著性[37]。16%只相当于49%的0.327，达到事先设定的实践残余标准。

Akhtar 等对学龄儿童的重复横断面研究中，非吸烟儿童总体几何均值由0.36 降至0.22 ng/mL，按端点下降38.9%；双亲均不吸烟者由0.14 降至0.07 ng/mL，下降50.0%；仅父亲吸烟者由0.57 降至0.32 ng/mL，下降43.9%[38]。儿童结果说明家庭吸烟并不机械消除政策外溢效应，也反对把所有残余差异简化为“家庭实践完全不变”。

针对酒吧工作人员的研究提供更接近源头场所控制的证据。Semple 等报告非吸烟酒吧工作人员唾液可替宁几何均值由2.94 降至0.41 ng/mL，端点比为86.1%，配对分析报告下降89%（95%CI 85%—92%）；工作场所过去7天二手烟暴露由28.5 小时降至0.83 小时，端点下降97.1%，个人PM<sub>2.5</sub>也显著下降[39]。Ayres 等随访呼吸和感觉症状及炎症指标，进一步显示规则实施后工作环境与部分健康相关端点变化[40]，但年度随访只保留了基线样本的一部分，不能忽略选择性失访。

表5 苏格兰无烟立法公开数据的规则锁定算术结果

| 群体与端点 | 基线 | 随访 | 统一算术降幅 | 允许的过程结论 |
|-------|----|----|--------|---------|
|-------|----|----|--------|---------|



|                |            |            |               |                       |
|----------------|------------|------------|---------------|-----------------------|
| 一般非吸烟成人唾液可替宁   | 0.43 ng/mL | 0.26 ng/mL | 39.5%         | 场所规则后客观暴露下降，越过25%支持线  |
| 非吸烟家庭成人唾液可替宁   | 0.35 ng/mL | 0.18 ng/mL | 48.6%         | 公共与家庭无烟条件叠加时降幅较大      |
| 吸烟家庭成人唾液可替宁    | 摘要未列端点     | 摘要未列端点     | 发表值 16%，不显著   | 未覆盖家庭通道保留残余；不得称政策造成转移 |
| 非吸烟儿童总体唾液可替宁   | 0.36 ng/mL | 0.22 ng/mL | 38.9%         | 客观暴露下降，但家庭分层异质        |
| 双亲不吸烟儿童唾液可替宁   | 0.14 ng/mL | 0.07 ng/mL | 50.0%         | 与家庭无烟条件一致的较大下降        |
| 仅父亲吸烟儿童唾液可替宁   | 0.57 ng/mL | 0.32 ng/mL | 43.9%         | 不支持“有吸烟家长即无政策外溢”的绝对说法 |
| 非吸烟酒吧工作人员唾液可替宁 | 2.94 ng/mL | 0.41 ng/mL | 86.1%；配对值 89% | 不依赖非吸烟者改变行为即可快速改变暴露   |
| 酒吧工作人员每周暴露时数   | 28.5 h     | 0.83 h     | 97.1%         | 工作场所接触通道近乎封闭，仍不等于癌症终点 |

### （三）支持、反面结果与最强竞争解释

数据支持路径中最有限却关键的一句：场所规则与烟雾流向可以在没有把非吸烟者作为行为改变对象的情况下，迅速改变客观暴露。酒吧工作人员的高基线暴露和大幅下降尤其不适合用“他们突然更自律”解释。它说明预防的有效主语可以是场所规则、排放条件和控制者，而不必首先是暴露者。

数据同时反对过强命题。吸烟家庭成人只有 16%且不显著的降幅，提示公共场所规则没有封闭家庭实践通道；然而原研究没有发现禁烟使吸烟从公共场所转移进家庭[37-38]。所以正确表述是“原有家庭通道保留残余”，不是“政策把风险转移给家庭”。儿童中仅父亲吸烟组仍下降约 44%，也说明家庭边界不是全有或全无，政策规范、家长行为和场所变化可能存在外溢。

至少五种竞争解释仍在。同期吸烟率变化可能贡献暴露下降；重复横断面样本的家庭、职业与社会经济构成可能变化；季节、检测限和几何均值处理可能影响数值；自报工作暴露受社会赞许影响，尽管可替宁减轻了这一问题；酒吧工作人员年度随访约为基线的一半，健康或就

业变化可能影响留存。四项研究的样本、设计、时点和端点不同，本章没有把它们合并为一个效应量。

更重要的是，这组数据没有直接观测任何位置  $p$  的  $tA(p)$ 、 $tR(p)$  或  $tC(p)$ ，也没有测量新吸烟招募，更不能证明肺癌发病下降。它检验的是规则—流向—暴露链和实践残余，不是整个待承位概念。若本章因得到大幅可替宁下降便宣布“新理论获证”，就重复了以代理终点替代疾病终点的错误。真实数据的价值正在于既保留支持，也把概念限制在它实际触及的那一段过程。

(四) 把新对象真跑一次：结果是不利的记录字段审计

为避免只用现成数据证明路径、却让新对象永远停在公式中，本章对上述四项原始研究又做了一次探索性可记录性审计。审计不重算效应，也不把“研究没打算记录”指控为研究缺陷；它只问：依照本章定义，公开报告能否重建至少一个预先划定位置  $p$  的形成、占据与关闭？字段在复核时固定为五项：规则生效时点、客观暴露或环境流、占据前的位置编号、 $tA(p)$ 、 $tR(p)$  或  $tC(p)$ 。由于审计规则是在原论文已知后形成，它不是预注册验证，只能给概念一条诚实的不利结果。

表 6 苏格兰四项研究对风险待承位的记录字段审计

| 公开证据记录                         | 规则生效时点 | 客观流或暴露                           | 占据前位置 $p$                  | $tA(p)$ / $tR(p)$ / $tC(p)$ | 审计判决                             |
|--------------------------------|--------|----------------------------------|----------------------------|-----------------------------|----------------------------------|
| Haw 与 有<br>Gruer 成人调<br>查[37]  |        | 有：唾液可替<br>宁及场所自报                 | 无；抽样对象<br>均为已存在成<br>人      | 均无                          | 支持规则后暴<br>露改变，不能<br>重建待承位        |
| Akhtar 等 儿 有<br>童调查[38]        |        | 有：唾液可替<br>宁                      | 无；儿童在观<br>察前已处于家<br>庭与公共场所 | 均无                          | 支持边界异<br>质，不能判关<br>闭先于占据         |
| Semple 等 酒 有<br>吧工人研究\<br>[39] |        | 有：可替宁、<br>暴露时数、个<br>人 $PM_{2.5}$ | 无；工人均已<br>被招聘并在岗           | 均无                          | 强支持工作场<br>所流量下降，<br>待承位为不可<br>观察 |
| Ayres 等 有<br>BHETSE 随<br>访[40] |        | 有：暴露与健<br>康相关端点                  | 无；队列从在<br>岗工人开始            | 均无                          | 支持补救性改<br>善，不能证明<br>预防性关闭        |
| 字段可得率                          | 4/4    | 4/4                              | 0/4                        | 0/4                         | 路径端点可<br>测；新对象直<br>接证据为零         |

这次真跑没有得到“部分支持”，而得到一个更尖锐的分叉。若研究问题是场所规则能否降低已有非吸烟者的暴露，四项研究提供了相当一致的过程证据；若研究问题是一个尚未占据的位置能否在风险过户前被关闭，四项研究全部失语。原因不是样本量不足，而是数据从已暴露或已在场的人开始建档。本章因此不能把苏格兰案例列为待承位的经验例证，只能把它列为路径可行性证据和新数据结构为何必要的反例。反过来，如果未来在加入位置字段后仍不能可靠重建竞争事件，正确结论不是“数据还不够”，而是这一对象可能只适用于少数有离散岗位或服务单元的场景。

十四、六个癌症风险领域：同一路径不能复制成同一处方

前主体预防是一套提问顺序，不是一项普适干预。吸烟、饮酒、职业致癌物、致癌性感染、体重相关风险和筛查具有不同的危险载体、主体形成和医学证据；有的存在清楚可编号的岗位，有的是弥散市场、病毒传播或服务基础设施，有的主要涉及早期发现而非阻止风险形成。位置 p 能否在观察结果以前稳定划定，是各领域必须先过的门槛。把六者塞入一个“生活方式”篮子，会再次把制度差异压回个人，也会制造虚假的共同计量。

表 7 六个癌症风险领域的位置单位与前主体路径边界

| 风险领域   | 候选位置单位 p                          | 优先过程指标                          | 最可能的义务节点               | 不能越过的边界                        |
|--------|-----------------------------------|---------------------------------|------------------------|--------------------------------|
| 烟草与二手烟 | 预先限定“渠道×时期×新进入批次”；弥散市场中常不可精确编号    | 待承位形成/关闭、营销触达、首次尝试、场所烟雾、可替宁、退出率 | 生产者、零售与平台、场所管理者、监管者    | 不得把任一潜在消费者虚构成位置；已成瘾者需要支持而非责备   |
| 酒精     | 许可场所、活动或销售渠道中的新进入批次；个体家庭消费常不可判    | 纯酒精流、位置关闭、营销触达、首次/高强度饮酒、替代场所    | 生产与零售、平台、活动组织者、税收与许可机关 | 不把所有饮酒实践视为相同；不以污名化有酒精使用障碍者替代政策 |
| 职业致癌暴露 | 尚未分配的具名岗位—任务—班次事件，识别度最高           | 流—岗位连接、tA(p)、关闭/分配竞争事件、任务及生物监测  | 设计者、采购、雇主、设备与场所控制者、监管者 | 个人防护具不能替代可行的消除与工程控制；外包不转移责任    |
| 致癌性感染  | 学校/基层服务的资格—时段单元；传播中的“下一感染者”通常不可编号 | 服务位开放与关闭、疫苗覆盖与公平、感染发生、治疗完成      | 卫生系统、学校、支付与采购、服务组织者    | 不得把感染概率伪装成空位；接种、检测和治疗服从同意与指南   |

|                  |                                  |                             |                       |                               |
|------------------|----------------------------------|-----------------------------|-----------------------|-------------------------------|
| 体重、饮食与活动<br>相关风险 | 学校餐饮、工作班次或交通设施的进入单元；个人生活总体通常不可判  | 供给、价格、营销、工作时序、位置关闭及分层暴露     | 食品与平台、雇主、学校、城市与交通治理   | 高体重是风险指标而非道德身份；若位置无法划界就不用本章术语 |
| 筛查与癌前干预          | 名册中的邀请—预约—后续处置服务位；它是服务位置而非“癌症空位” | 服务位开放、摩擦、知情参加、阳性后完成、伤害与分层覆盖 | 卫生系统、支付者、服务机构、雇主与公共沟通 | 并非所有筛查预防发病；不得把未参加者当作被风险位置占据   |

### （一）烟草：从“戒烟者”之前寻找招募条件

烟草是最能显示双层预防的领域。对已经使用烟草者，戒烟支持、药物与临床帮助具有不可替代的价值；对尚未开始者，价格、零售密度、产品设计、广告、社交场所和无烟规范决定实践能否继续招募。WHO《烟草控制框架公约》及其第8条把二手烟保护写成国家义务，IARC手册也评估了无烟政策、税收与价格等人群措施<sup>[34-35]</sup>。这些工具早已说明个人劝导不是唯一道路。

前主体分析的额外任务，不是给整个新产品标一个 tA 与 tR，而是在结果出现前规定渠道、时期和新进入批次 p。某种新型尼古丁产品在配方、许可、零售、社交媒体意义和使用技能结合以后，可能具备招募能力；但若无法识别哪一个尚未占据的进入单元，就只能使用营销触达和新尝试率等既有指标，不能强称存在待承位。可检验端点应包括位置形成与关闭、首次尝试、持续使用转化、营销触达与实践退出，而不只统计已有使用者的知识。对于已经依赖尼古丁的人，政策应提供可及戒烟支持并避免羞辱；关闭新进入与帮助承载者退出，是同一预防系统的前后两段。

### （二）酒精：总量控制必须进入具体饮酒实践

饮酒增加多种癌症风险；IARC 关于减少或停止酒精消费及人群酒精政策的预防评估，为价格、可得性、营销限制等措施提供了证据审查<sup>[44]</sup>。但“酒精总量”不能解释所有主体过程。家庭聚会、商务应酬、夜间经济、体育赞助和独自应对压力，是材料、时序、关系和意义不同的实践<sup>[19]</sup>。同样的价格政策在不同实践中可能改变购买量、替代品或场所选择，不能用平均销量掩盖分层变化。

流向追踪在此不是追踪一种致癌化学品进入工厂，而是追踪纯酒精在产品、容器、价格、场所和营销渠道中的可得流。义务定位不能止于“消费者适量”，因为生产者、零售许可、平台推荐和活动组织都控制招募条件。招募拆解也不等于消灭社交生活，而是让低酒精或无酒精实践拥有材料、能力与正当意义，并减少将饮酒同成年、成功、亲密或放松绑定的制度奖励。

是否成功，应由新高风险饮酒发生、场所暴露和替代实践持续性判断，而非只由宣传覆盖率判断。

### （三）职业暴露：工人出现以前，岗位设计已经完成大部分因果工作

职业领域最接近可识别的待承位。新生产线从工艺选择、采购、建筑、通风、维护到人员配置，常有一段尚无工人实际接触的设计期；尚未分配的岗位—任务—班次可以在结果以前编号。此时危险连接可能已经安装，进入路径通过岗位、技能、节拍和合同开放，控制能力则集中在设计者、采购、业主和雇主。控制层级与源头设计要求优先消除、替代和工程控制，正是因为这些动作可以在工人必须“正确行为”以前降低接触<sup>[29,45,53]</sup>。本章只有在进一步记录“该位置是先被关闭还是先被分配”时，才比既有工程原则多出经验内容。

然而，替代和循环需要 MFA 校验。增材制造末端材料研究与欧洲电子废物生物监测提示，回收、焚烧、填埋和拆解可生成新的环境或职业接触路径<sup>[14-15]</sup>。这些研究各有特定材料和暴露范围，不能推论“循环经济致癌”；它们足以要求任何绿色转型报告物质去向、工作任务和被招募劳动群体。若危险从正式工厂移至承包商或非正规工人，原企业的报表改善不满足不转移门槛。

职业预防尤其不能把 PPE 失败写成工人道德失败。个人防护具在某些剩余风险和应急情形中必不可少，但其效力依赖选择、适配、维护、培训、佩戴时间和生产节拍。若可行的高层级控制未采用，却把全部因果焦点放在“没有戴好”，便是把设计完成的暴露过户给最后出现的人。

### （四）致癌性感染：在感染以前行动，但不能把病原体社会化为隐喻

感染相关癌症使“在癌症以前预防”具有清楚医学路径。瑞典全国人群研究纳入约 167 万名 10—30 岁女性，HPV 疫苗接种者的浸润性宫颈癌调整后发病率比为 0.37，17 岁前接种者为 0.12<sup>[41]</sup>。台湾全民乙肝疫苗实施后，儿童肝细胞癌发病率在相应出生队列中下降；该研究为全国政策前后和出生队列证据，不是个体随机试验<sup>[42]</sup>。马祖地区幽门螺杆菌大规模筛查与根除计划报告感染率由 64.2% 降至 15.0%，胃癌发病率相对下降 53%，而胃癌死亡下降 25% 的置信区间跨过无效值<sup>[43]</sup>。三项证据设计不同，不能简单合并，却都表明传播、感染或持续感染阶段可以成为癌症发生前的阻断点。

前主体路径在这里须格外谨慎。病原体传播有明确生物机制，不能被泛化为“有害思想传染”；疫苗和根除治疗有适应证、禁忌、不良反应、耐药与公共卫生规范，不能由社会理论替代。真正的上游问题包括采购、冷链、服务可达、学校与基层卫生能力、污名、阳性后的治疗

完成，以及谁因费用、性别或地理位置被排除。个人同意仍是必要条件；“未来癌症预防”不能成为越过知情同意的万能理由。

### （五）体重、饮食与身体活动：改变实践束，不把身体变成罪证

较高体重、某些饮食暴露和身体活动不足与多种癌症风险有关，但个体体重是代谢、药物、遗传、睡眠、劳动、食物环境和社会实践的共同结果。第五版欧洲癌症防治守则面向公众提出体重、饮食与活动建议，同时要求可负担健康食品、学校与工作场所支持、主动交通基础设施、营销与标签政策等条件[33]。这组双轨设计已经拒绝“只要意志坚定”。

社会实践分析还需避免把“健康环境”简化为选择架构。轮班工人的进食时序、照护者的时间贫困、步行与通勤安全、学校和工作单位的餐饮、社交庆祝的食物意义彼此结盟。预防应问哪些实践持续招募久坐或高能量摄入，哪些替代实践能够在不增加时间和金钱负担的情况下存活。过程指标可包括食品供给、相对价格、营销触达、交通可达与工作时序；体重变化只是较远端指标之一。任何政策都不得把较高体重者描述为癌症的制造者，也不能承诺减重即可消除个体风险。

### （六）筛查边界：发现得更早不总等于没有发生

筛查是本章最容易越界的领域。宫颈癌筛查发现并处理高级别癌前病变、结直肠筛查发现并切除腺瘤，可以减少部分浸润性癌症发生；另一些筛查主要把已存在的肿瘤提前诊断，目标是降低晚期和死亡，而非阻止最初发生[46-47]。过度诊断、假阳性、程序伤害和资源挤出也须纳入。因此，不能把所有“参加筛查”都放进同一前主体窗口，更不能把未来可能患癌的人称作“癌症待承位”。

位置分析在筛查中只适用于服务系统如何在邀请发出前形成可进入或不可进入的服务位：名册是否完整，预约是否适配工作与照护，阳性结果是否有诊断和治疗通道，弱势群体是否承担更高摩擦。这里被关闭或占据的是服务位置，不是致癌暴露位置，必须在编码中另列类别。如果这些条件缺失，后来把未参加写成个人依从性失败，便把系统安排倒写成主体选择。但筛查的具体年龄、频率和方法必须遵循当地指南和个体风险评估，本章概念不能作为擅自增加或减少筛查的理由。

## ■ 十五、何时仍必须以个人为主语

如果把本章理解为“公共卫生永远不应以个人为主语”，它会制造新的错误。至少有五种情形，具体个人必须回到语法和伦理中心。



第一，已经暴露或已参与实践的人需要可及、无污名的退出支持。戒烟、酒精使用障碍治疗、职业健康服务、疫苗补种、感染检测与治疗都不能等待结构改革完成。某个  $p$  被占据只表示该次竞争事件由  $tR(p)$  获胜，不等于预防无效，更不表示系统中其他位置无须关闭；识别滞后阶段和个体临床阶段仍有大量可做之事。

第二，高遗传易感或有明确家族史者可能需要个体化咨询、监测和风险降低。其风险不能被一般人口环境充分解释，上游政策也不替代专业遗传服务。反过来，遗传风险同样不能用来免除环境和制度责任；二者不是零和。

第三，知情同意、身体完整与个人价值不能被总体效益吞并。接种、检测、预防性用药或手术即使在人群上有效，也必须在适用规范下由个人或合法代理授权。代际正义不能把当前人只当作未来人口的工具。

第四，个体行动可成为制度改变的来源。工人报告、患者组织、家庭无烟承诺、社区倡议和消费者选择都可能重塑规则。社会实践理论中的承载者不是被动容器；他们能够反思、退出、改造意义和组织集体行动。把个人从责任终点移开，不等于否认其政治主体性。

第五，临床已经面对具体患者时，必须以此人的癌种、分期、共病、偏好和生活条件决策。人口风险、PAF 与预防窗口不能倒推个体病因或治疗。对于已经确诊者，首要问题是准确诊断、循证治疗、支持照护与共同决策，而不是追问其过去是否足够自律。

因此，更精确的语法是：在风险安排形成阶段，先让流、控制者和实践作主语；在权利、支持和临床决定阶段，让具体人重新作不可替代的主语。两种语法按过程分工，而不是互相取代。

## ■ 十六、怎样研究一个尚未被人占据的对象

### （一）分析单位：位置事件，而不是疾病个案或整个产业

风险待承位的最小分析单位是一个位置事件  $(a, p)$ ：在观察结果以前划定的系统  $a$  中，进入规则能够稳定识别位置  $p$ ，危险流与进入路径已经连接，当前控制者可采取行动，而该次位置尚未占据。研究者需预先规定空间、时间、产品或工艺、医学风险证据、进入规则和位置重新编号规则。不能从癌症患者回顾性故事倒推“他当年本来占了哪个空位”，也不能在看见政策效果后把有利批次切成位置，因为两者都会把结局写进定义。

岗位—任务—班次是最清楚的单位：配方和设备冻结、实际投料、岗位开放、人员分配、首次任务接触与首次监测可识别可依次计时。同一岗位离职后重新开放，应生成新的位置事件，而不是沿用旧  $p$ 。服务系统可用资格—机构—时段作为单位。消费市场只能在渠道与进入

批次有事前边界时编码；“全国下一个潜在消费者”既不可核验，也不可穷尽，必须记为不可判定。位置的可编码性不是技术附注，而是理论能否离开比喻的第一道检验。

## （二）最小记录：四联状态与硬约束不能加成总分

流、进入、控制和占据各自可以拥有连续证据，却不能加成一个“待承指数”。高幅度流量下降不能补偿位置已经先被人占据；明确义务文本不能补偿执行后连接仍在；关闭大量位置也不能补偿风险被外包给更弱势工人。最小记录须逐项给出证据来源、时间戳、判断者和反证材料。

位置注册：a 与 p 的边界、进入规则、重开或拆分规则、可计数性及规则锁定日期；结果后修改另列探索性版本。

流—位置连接：医学危险证据、投入、库存、输出、通常任务或使用中的接触路径、替代物、测量误差和质量平衡缺口。

进入与招募：必需材料、能力、意义、场所、联盟以及路径何时可复制；知识或态度不得单独代理 R。

控制与义务：具名主体的法律权限、实际能力、获益关系、动作期限和执行证据；“政府”“企业”或“公众”不得作为无法复核的集合名词。

竞争时间： $tA(p)$ 、 $tR(p)$ 、 $tC(p)$ 、 $tL$  及不确定区间；空窗、同时发生和暂不判定都须进入分母。

转移与权利：不同职业、地区、收入、年龄与未来时段的风险，知情同意、退出支持、申诉和隐私；任一硬门槛失败不能由平均效果补偿。

疾病边界：位置关闭、新暴露、癌前病变、癌症发病与死亡分层报告；由过程端点外推疾病必须明示模型、潜伏期和不确定性。

## （三）核心设计：让“关闭”与“占据”真正竞争

位置 p 自  $tA(p)$  进入队列，随后可能先发生预防性关闭  $tC(p)$ ，也可能先发生占据  $tR(p)$ 。占据不是普通失访，关闭也不是把所有未占据位置都算作成功；两者是互相排斥的竞争事件。主要近期结局应是预定时限内“关闭先于占据”的累积发生，另报“占据先于关闭”、两者均未发生与不可判定。若位置可数，可估计关闭和占据的原因别风险或累积发生函数；若位置不可数，就不得用几个案例冒充率。

这一设计改变了政策比较的分母。传统暴露研究常从已在岗工人、已吸烟者或已进入服务者开始；位置队列则必须从尚未占据的 p 开始。以 100 个预先登记岗位为例，60 个在分配前完成替代或隔离、30 个先被分配、10 个观察期末仍未开放，不能报告“60%的人避免暴露”，因



为分母不是人；准确表述是预定岗位事件中 60% 出现关闭先于占据，随后还需验证 30 名已入岗者的暴露、10 个开放位的后续结局和替代风险。该结果也不能直接换算成少了多少癌症。

#### （四）三类研究及其失败方式

第一类是前瞻性位置登记研究。独立团队在至少三种机制差异明显的风险领域按同一手册登记  $p$ ，报告评定者对  $V_{ap}$ 、 $tA(p)$  和事件类型的一致性、不可判定率、开放时长及竞争结局。若位置只能在职业领域可靠编码，概念应主动限缩；若多数位置是分析者事后发明，概念应删除。

第二类是准实验或阶梯实施研究。利用场所规则、产品标准、工艺替代、营销限制或服务扩容的分期实施，比较干预与对照系统的关闭先于占据、新暴露和风险转移。中断时间序列、差分中的差分或合成控制可用于后续聚合端点，但须检验平行趋势、政策同期性、外溢、位置构成变化和测量变化。不能只在干预组把位置划得更细，从而人为提高关闭数。

第三类是资源等量比较。把位置关闭干预同高质量个人咨询或退出支持在同一风险领域、同等资源下比较。两者不能共享完全相同的直接分母，故应预先设共同远端结局——例如新进入或客观新暴露——并同时报告各自近期机制。本章的优先序主张只有在位置干预更早或更大改变共同结局，同时不撤减已暴露者支持、不扩大不平等时才成立。若个人方案同样有效，政策可由成本、权利与可实施性决定，而非由理论身份决定。

#### （五）测量误差、可识别性与一个禁止性结论

$tA(p)$  往往不是单一瞬间。许可、投产、营销与场所实践可能分步形成，研究应允许区间删失、多个候选时点和敏感性分析。 $tR(p)$  可能因非正式或匿名首次接触而被观测得过晚； $tC(p)$  则可能因“纸面关闭”早于真实断开而被观测得过早。这两种误差会系统性美化关闭先于占据。工艺日志、岗位系统、环境监测、匿名销售或平台记录和质性观察可交叉验证，但必须服从隐私、劳动权利和合法数据治理。

因果识别还会受选择影响：风险最高的场所可能最早控制，经济与文化趋势可能同时改变流量和招募，人员与产品可跨地区移动，干预还可能改变位置的定义和数量。本章没有绕开现代流行病学设计，反而增加了一个易被操纵的分母。因此必须给出一条禁止性结论：当位置  $p$  无法在结果以前稳定注册时，不得计算待承位关闭率，也不得用“前主体预防窗”替代既有的暴露、招募或服务可及性术语。

### ■ 十七、反向约束：预防不能以取消主体为代价

前主体预防最危险的误用，是从“位置尚未被具体人占据”滑向“人只是可以被系统填充的空位”。为阻止这种倒置，任何政策或研究至少接受八条反向约束。

其一，不以平均获益许可强迫。面向人口的预防仍受比例原则、合法程序、知情同意和身体完整约束。严重风险并不自动授权无限制监控、强制检测或未经同意的医疗操作。

其二，不把未来利益压给当代弱者。代际正义不能只谈未来而忽略同时代分配。如果安全替代抬高基本品价格、关闭唯一就业而没有转型支持，负担可能集中于最少获益者。政策须同时报告当代与代际分布。

其三，不把人还原为流量末端或“填位者”。待承位是安排的属性，不是贴在人身上的身份。MFA描述物质，不描述完整的人；工人、消费者和居民拥有权利、知识、关系和反思能力，应参与边界定义、风险判断和干预设计，而不是只被画成受体框。

其四，不把实践承载者写成无意志媒介。招募解释减轻道德责备，却不能抹除自主行动。退出支持、集体组织和反实践创新都依赖人的能动性。社会实践理论不应成为家长主义的许可证。

其五，不以循环、绿色或创新标签免检。替代材料、回收、数字化和低碳方案都须重新追踪危险性、流向与劳动暴露。改善一个生命周期阶段不能自动抵消另一个阶段的损害。

其六，不把可预防比例用于个体定罪。PAF、群体相关和政策效应不得写入患者道德评价、保险惩罚、就业筛选或临床资源优先级。确诊者不承担证明自己“并非自找”的义务。

其七，不以概念替代指南。“风险待承位”和“前主体预防窗”未经临床验证，不能决定谁应接种、筛查、用药、手术或停止治疗；也不得作为法律能力、产品安全、职业合规、保险授权或人员录用的单独判据。

其八，保留不可知与防止分母操纵。当  $p$ 、 $tA(p)$ 、 $tR(p)$ 、 $tC(p)$ 、危险性或转移无法可靠判断时，应记录不确定，而不是以预防原则冒充事实；不得通过拆细成功位置、合并失败位置或删除已占据位置提高关闭率。政策可以在不确定下采取比例性谨慎，但学术文本必须区分证据、判断和价值选择。

这些约束不是文章结尾的伦理装饰，而是定义的一部分。若一项干预降低  $\Delta F$  和  $\Delta R$  或提高关闭先于占据，却违反任何硬约束，它最多是有效控制，不是本章意义上的正当预防。

## ■ 十八、可删除的理论与有日期的判决

“风险待承位”目前只是一个待验证对象，苏格兰记录字段审计对它给出的直接证据为零。它至少有五种删除或降级方式。第一，位置  $p$  若只能在研究者知道结局后划定，就删除对象；若只在具名岗位中可靠，则限缩为职业安全术语。第二，若危险连接、进入路径、当前控

制和尚未占据不能由独立评定者稳定区分，四联定义失败。第三，若控制层级/源头设计、实践招募与用户配置变量的联合模型已经产生相同位置排序和结局预测，待承位没有增量。第四，若提高“关闭先于占据”并不降低共同远端结局——新进入或客观新暴露——它只是文书状态。第五，若所谓关闭持续把风险转给弱势群体、撤减已暴露者支持或依赖分母操纵，该路径即失败。

为防止无限延期，本章给出一个有日期的前瞻性判决。到 2032 年 12 月 31 日，至少应出现一项公开方案的多中心位置队列：覆盖不少于三个机制差异明显的致癌风险领域、十二个法域或组织和 300 个在结局前登记的位置事件；独立评定者对  $V_{ap}$  是否成立及首发事件类型的  $\kappa$  不低于 0.70，不可判定率不高于 25%；在留出组织中，相较仅含源头设计、实践招募、风险环境和用户配置变量的基准模型，加入待承位字段后，对“关闭先于占据”的预定时限 Brier 分数至少绝对改善 0.02 且校准不恶化；另有至少两项准实验使占据先于关闭的累积发生相对下降不低于 20%，共同远端的新暴露也下降，同时最弱势分层的风险转移不超过 10%。

300、0.70、25%、0.02、20%和 10%均是作者为迫使理论接受失败而提出的研究性判决线，不是临床阈值、监管标准或既有共识。Brier 改善只裁决预测增量，不证明因果；准实验仍须处理位置构成、政策选择和外溢。若到期只有案例阐释、名称引用和回顾性拟合，没有前瞻注册、外部预测与竞争事件，学界应停止把它当作独立概念，分别使用原始预防、源头设计、风险环境和实践导向公共卫生。只有当位置可复核、联合近邻不能吸收、关闭改变共同远端结局且风险未转移，概念才获得继续使用的资格。

## ■ 十九、结论：预防应在“这个人为什么没有做到”之前开始

癌症预防当然需要个人行动，却不应等到一个人已经被烟草、酒精、职业暴露、感染服务缺口或不利实践束招募以后，才问他为什么没有保护自己。代际正义表明，尚不可识别的未来承受者并不取消当前义务；物质流分析表明，尚无人患病时危险已经具有可追踪去向；社会实践理论表明，动机与身份可以在参与中生成，而非风险之前的固定起点。三者更深的共同盲点，是都容易从已经有人占据的位置开始记账。

三者合起来形成的不是“多管齐下”，而是一条有严格先后的 How 路径：先追踪危险流接到了哪个开放位置，继而把断开义务落实到当前控制者，再拆解使下一个位置得以开放的实践结合。风险待承位只在危险连接、进入路径、当前控制和尚未占据四项同时成立时存在；每个位置都有自己的滚动窗口，并以“关闭先于占据”或“占据先于关闭”结束。窗口可能为空，位置可能无法注册，联合近邻也可能完全吸收这个对象。正因它给自己设置了这些失败方式，才有资格被暂时提出。

苏格兰无烟立法的公开数据给出两条方向相反、必须同时保留的经验线。暴露端点表明，改变场所规则可使非吸烟者客观暴露快速下降，无需先把他们变成更自律的人；吸烟家庭的残余暴露又说明，流量控制若没有覆盖实践边界，不会自动完成预防。记录字段审计则表明，四项研究都从已存在、已在场或已在岗的人开始，无法重建任何待承位的形成与竞争结局。数据因此支持路径，却没有直接支持新对象，更没有证明肺癌发病变化。

因此，“预防如何不再以个人为主语”的答案不是把“个人”替换成同样抽象的“结构”，而是改变预防登记的最小单位：在风险过户以前登记并关闭位置，在风险过户以后保护并支持具体人。前一阶段让产品流、制度控制和实践招募先承担解释与责任；后一阶段让知情同意、退出支持、临床决策与权利保护重新成为中心。一个成熟的癌症预防体系，不以责问患者完成因果叙事，也不把人当作可填充的空位；它要做的是，在尚无人可被责问之时，使一个已经开放的风险位置失去获得承载者的能力。

■ 附表 1 风险待承位与滚动前主体预防窗的最小记录及删除审计表

| 编码对象                    | 允许形成肯定判断的最小证据                  | 明确否定或记零             | 暂不判定              | 可推翻原判断的材料            |
|-------------------------|--------------------------------|---------------------|-------------------|----------------------|
| p: 进入位置                 | 在结果前写明边界、进入与重新编号规则，可由记录稳定识别    | 只有“未来某人”或事后按结局切分    | 弥散市场或匿名进入无法固定单元   | 带版本时间戳的位置注册表及独立复核    |
| F <sub>ap</sub> : 危险连接  | 医学相关流已接到 p 的通常任务/使用，边界、单位与路径明示 | 目标流未进入，或与 p 物理/制度隔离 | 质量平衡或接触路径无法闭合     | 物料账、环境/职业监测、替代品危险评估  |
| R <sub>ap</sub> : 开放入口  | 材料、能力、意义、规则与场所已可重复结合并进入 p      | 至少一个必要要素未结合，入口未开放   | 只有文化叙述，无实际可得或进入证据 | 岗位、服务、平台/销售记录与质性观察   |
| C <sub>ap</sub> : 当前控制  | 具名主体可在预定期限内切断 F 或 R，并有权限/资源证据  | 只有倡议能力，无真实控制        | 权限与执行能力冲突或尚在诉讼    | 合同、许可、预算、流程、执法与申诉记录  |
| O <sub>ap</sub> : 占据状态  | 有记录证明该次 p 尚未分配、进入或接触           | 具体人已进入并承受相关流，O=1    | 非正式或匿名进入可能早于记录    | 人员/任务日志、匿名调查、环境或生物监测 |
| V <sub>ap</sub> : 风险待承位 | F\>0、R=1、C 非空、O=0 四项同          | 任一项明确不成立            | 任一项不可可靠判定         | 四项由不同数据源交叉验证及盲态评     |

|                   | 时成立                                             |                        | 断                | 定                    |
|-------------------|-------------------------------------------------|------------------------|------------------|----------------------|
| tA(p): 位置形成       | V <sub>ap</sub> 首次由 0 变 1, 有时间戳或窄区间             | 四项从未同时成立               | 形成分散、区间过宽        | 工艺、岗位、许可、上市或服务启动日志   |
| tC(p)/tR(p): 竞争事件 | 预防性断开或首次占据均有先后可判时间                              | 同时发生按规则记零窗, 不强判先后      | 任一事件时间缺失或区间重叠    | 执行记录、任务日志、重复监测与敏感性分析 |
| Wpre(p): 滚动窗口     | 对每个 p 报告 \ [tA(p),tC(p)) 或 \ [tA(p),tR(p))及不确定度 | 位置形成时已占据, 窗口为空         | 位置或时点不可识别        | 独立评定、区间删失模型与补充时间戳    |
| ΔF/ΔR: 机制改变       | 事前定义的连接或入口指标按期下降, 并报告误差                         | 无变化或向其他位置等量转移          | 仅有政策文本、自报意图或汇总均值 | 重复测量、质量平衡、位置队列与对照系统  |
| D: 义务落实           | 控制者、动作、期限、资源与核验均具名                              | 责任只写给“社会”或最后暴露者        | 多主体权限不清或资源未落实    | 执法、采购、预算、合同和申诉记录     |
| T: 不转移与公平         | 平均与最弱势分层同报, 无重要空间、职业或跨代转移                       | 改善依赖转给更弱势群体或撤减支持       | 下游或长期数据尚未到达      | 全生命周期追踪、分层监测与社区复核    |
| 概念保留              | 位置可靠、近邻不能吸收、关闭改善共同远端结局且通过 D/T                   | 被联合近邻吸收、位置不可注册或同资源方案不劣 | 只有案例、字段审计与回顾拟合   | 公开方案的外部位置队列、预测比较与准实验 |

### 最小审计纪律

每个位置事件(a,p)单独编号, 不以癌症个案倒推全部上游历史。

研究开始前定义系统、位置、危险证据、tA(p)及竞争事件; 探索性改动另行标注。

流、入口、控制、占据与权利不可加成分, 任一硬门槛失败不得由其他高值补偿。

“可能有人暴露” “大概能控制” “应该会改变行为” 均不得编码为肯定。

知识、态度和政策文本不得单独代替客观流量、新招募或执行证据。

窗口为空、位置不可注册和无法判断是有效结果, 必须进入分母并公开报告。

过程端点、癌前病变、癌症发病与死亡分别报告, 不做无依据的终点升级。

PAF 只用于人群反事实分析, 不用于患者责任、保险惩罚或个体病因裁决。

平均效应与最弱势分层同时报告；职业、地区、收入、年龄和未来转移逐层核验。

替代、回收、数字化和外包均重新画物质与劳动边界，不接受标签豁免。

已暴露者的退出支持和临床服务不得因上游干预被撤减。

个体接种、检测、筛查、用药与手术继续服从专业指南、知情同意和当地规范。

对动物、细胞、观察、自然实验和随机证据分层表述，不得跨级写成确定因果。

若强近邻的联合模型产生相同位置排序与预测，优先使用既有术语；名称无检索命中不等于原创成立。

2032 年判决线到期后，无外部增量证据即删除或降级概念，不以引用次数续命。

数据、伦理与人工智能协作声明

本章的苏格兰无烟立法检验使用已发表的群体汇总数据，不含可识别个人信息，也不对任何个人作医学判断。算式、最低支持线、实践残余线和禁止推断在进一步数据提取前写入内部工作文件；由于部分研究方向此前已知且文件未登记于第三方不可篡改平台，本章不声称正式预注册或盲态分析。待承位记录字段审计是在新对象修订后对已知论文进行的探索性审计，明确结果为四项研究均不能重建位置竞争事件，不包装为验证。“风险待承位”和“前主体预防窗”均为未经临床和政策验证的研究概念，不得用于个体诊疗、法律能力、产品合规、保险授权、人员录用或绩效考核。

作者陈晓艳提出专著题目、三学科约束、核心研究方向并承担最终学术责任。人工智能系统在作者指令下协助文献检索、强近邻拓宽、公开数据算术核对、反例检查、文字起草与 Word 排版。作者对概念命名、论证取舍、引用与定稿具有最终决定权；人工智能不列为作者。若本章用于正式出版或投稿，作者应再次人工核对全部原始文献、政策版本、数据端点、伦理表述与目标出版机构的规范。

## ■ English Abstract

Intergenerational Justice, Material Flow Analysis, and Social Practice Theory on Unoccupied Risk-Bearing Positions and Rolling Pre-Subject Prevention Windows

Abstract

Cancer prevention often begins speaking only after risk has been attached to an identifiable person: the person is advised to stop smoking, reduce alcohol consumption, receive vaccination, change diet, attend screening, or comply with protective procedures. Even upstream policies commonly preserve the same hidden sequence. They replace the intervening agent but continue to assume that a risk-



bearing individual exists before responsibility, hazardous flows, and recruitment into everyday practices are analyzed. This article places intergenerational justice, industrial-ecological material flow analysis, and social practice theory in mutual adjudication. Each discipline supplies a complete but incompatible pathway: assigning present duties toward unconsenting future people, tracing materials through production and disposal, and explaining how practices recruit carriers. Their deeper common premise is that duties, flows, and recruitment acquire a shared unit only after a position has been occupied by a person. The article introduces the unoccupied risk-bearing position: within a prespecified entry unit, a medically relevant hazardous flow is connected, a reproducible entry route is open, at least one present controller can sever either connection, and no specific bearer yet occupies that position. The interval between formation and occupation is an indexed, recurrent rolling pre-subject prevention window, rather than a one-off interval ending when the first person anywhere enters an arrangement. Here “subject” means an addressable bearer of exposure, not a metaphysical person. The resulting preventive pathway is flow tracing–duty location–recruitment dismantling. Its proximal outcome is preventive closure of an open risk-bearing position before occupation, modeled against occupation as a competing event and constrained by anti-displacement, equity, consent, and continued support for already exposed people. A rule-locked arithmetic test of four studies surrounding Scotland’s smoke-free legislation found a 39% reduction in salivary cotinine among nonsmoking adults and an approximately 89% reduction among bar workers, while the reduction among adults living in smoking households was only 16% and not statistically significant. These results support rapid exposure change without targeting nonsmokers’ behavior, while revealing a residual household-practice channel. An adverse recordability audit found that none of the four studies recorded formation or first occupation of an indexed risk-bearing position. Thus, the pathway receives limited support while the proposed object remains directly unvalidated; neither analysis establishes a reduction in cancer incidence. The article specifies reciprocal tests against its strongest neighboring concepts, deletion conditions, and a dated prospective adjudication.



Keywords: cancer prevention; intergenerational justice; material flow analysis; social practice theory; unoccupied risk-bearing position; pre-subject prevention window; risk recruitment

## 第二章 一次成功由谁承载

本章原题《一次成功由谁承载？》，作者 阳涌，原文见 <https://sdeuniverses.com/students/yang-yong/bearerless-success/>

### 摘要

生成式人工智能正在制造一个经验悖论：它一方面能够迅速压缩不同劳动者、学生和专业人员之间的即时绩效差距，另一方面又可能扩大他们在脱离当前模型、提示链、操作者与制度环境之后继续完成任务的能力差距。现有研究分别把这一现象理解为收入与生产率分化、人力资本与学习分化，或人工智能接入和人机协作条件分化，却共同默认：只要一次任务被成功完成，这项成功所体现的能力必然已经归属于某个可识别的主体——一个人、组织或人机团队。本章推翻这一前提，提出“无着成功”与“能力承载极化”两个概念。无着成功是指：一项人工智能中介任务虽然达到了质量标准，但在模型撤除、版本更换、操作者替换或情境迁移后，不存在任何能够稳定复现、解释、修复并对其负责的承载者。能力承载极化则指，不同个人与组织把人工智能辅助成功转化为稳定能力承载者的概率持续分化。由此，AI 时代的两极分化不是现期结果差距，不是个人技能存量差距，也不是工具接入差距，而是成功能否获得稳定归宿的分化。本章提出一张“现期绩效 × 承载稳定性”二维辨别格，构造能力承载率、无着率和承载迁移矩阵，给出八项可证伪命题、五类研究设计以及制度伦理约束。核心判据是：撤去当前模型版本、提示链与原操作者之后，同类任务由哪个登记主体重做，质量下降多少？

关键词：人工智能；两极分化；无着成功；能力承载极化；人力资本；人机协作

Who Bears a Success? Capability-Bearing Polarization in the Age of AI A Theory of Bearerless Success across Economics, Psychology, and Computer Science

### Abstract

Generative artificial intelligence creates a paradox that prevailing accounts of inequality cannot fully explain. AI systems may compress immediate performance gaps between workers, students, and professionals while simultaneously widening differences in their ability to reproduce, transfer, repair, and take responsibility for the same work after the original model, prompt chain, operator, or institutional setting changes. Economics commonly observes productivity and wage outcomes; psychology examines the developmental path through which knowledge and

judgment are acquired; computer science and human-AI interaction research evaluate the quality of the surrounding collaborative system. Despite their differences, these traditions generally presume that a successful task performance must already belong to an identifiable capability bearer: a person, an organization, a technical system, or a human-AI team.

This article rejects that presumption and introduces two constructs. Bearerless success denotes an AI- mediated task event that meets an established quality threshold but cannot be reliably reproduced, explained, repaired, and owned by any recognized bearer under predefined perturbations. Capability- bearing polarization denotes a widening difference among individuals and organizations in their probability of converting AI-assisted successes into durable human, organizational, or technical capabilities. The central claim is therefore triply negative: polarization in the AI era is not fundamentally an inequality of current outputs, individual skill stocks, or model access; it is an inequality in whether successful performances acquire a stable bearer.

The article develops a two-dimensional framework combining current performance with bearing stability, formalizes a Capability-Bearing Rate, distinguishes bearerless success from deskilling, cognitive offloading, distributed cognition, automation bias, synthetic competence, and digital inequality, and derives eight falsifiable propositions. It concludes with empirical designs, ethical constraints, and a dated theoretical wager. The zero-modal diagnostic question is: After the current model version, prompt chain, and original operator are removed, which registered actor reproduces the task, and by how much does quality change?

Keywords: artificial intelligence; polarization; bearerless success; capability-bearing polarization; human capital; human-AI interaction

## ■ 一、一个越来越反常的事实：差距正在缩小，也正在扩大

人工智能进入劳动与学习现场以后，最容易测量的变化通常是“做得更快了”“答得更好了”和“低水平者追上来了”。

在一项涵盖 5,172 名客户服务人员的实地研究中，生成式人工智能辅助使平均生产率提高约 15%，经验较少和技能较低的员工获得了更大的提升；研究还观察到人工智能可能帮助部分员工学习高绩效同事的沟通模式。表面上看，这意味着人工智能具有压缩生产率分布的力量。

在专业写作任务中，Noy 与 Zhang 的实验同样发现，生成式人工智能能够显著减少任务完成时间、提高产出质量，并缩小不同参与者之间的绩效差距。

教育研究则展示了更加尖锐的两面性。Bastani 等人对近千名高中生进行的实验发现，在练习阶段，直接使用通用 GPT 界面的学生表现大幅提高；然而，当人工智能被撤除后，直接获得答案式帮助的学生在独立考试中的成绩反而低于未使用人工智能的对照组。加入教学性护栏、限制直接给出答案后，负面效应基本消失。更值得注意的是，人工智能压缩了练习阶段的成绩差距，却没有同样压缩撤除人工智能后的能力差距。

于是出现了一个不能用“人工智能究竟扩大还是缩小差距”简单回答的事实：

人工智能可能在同一个时点缩小人与人的产出差距，却在另一个时点扩大他们独立继续行动的差距。

一个原来只能完成 60 分任务的人，在人工智能辅助下完成了 90 分；另一个原来能完成 85 分任务的人，也完成了 90 分。从即时结果看，两人的差距几乎消失了。

但七天以后，模型版本更换，原提示模板失效，任务条件略有改变。前者重新跌回 55 分，后者仍能完成 88 分。

那么，第一次共同得到的 90 分究竟意味着什么？

它是两个人的能力已经趋同，还是他们只在一个短暂配置中产生了相同输出？

如果把前一种情形与后一种情形都叫“生产率提高”，我们就把两种具有相反未来的事件放进了同一个统计栏目。

AI 时代真正危险的差距，可能恰恰隐藏在已经被压平的结果里。

## ■ 二、关于 AI 两极分化的三种主流解释

现有研究关于人工智能与不平等的讨论，大体可以归入三种解释。它们分别来自经济学、心理学与计算机科学，并分别抓住了一部分真实现象。

### （一）经济学：谁获得了更多产出与收入

劳动经济学通常从任务配置、劳动生产率、职业结构、工资与就业份额出发讨论技术冲击。

任务型技术变迁理论指出，技术不是简单地“替代某个职业”，而是重新划分人和机器在一组任务中的比较优势。自动化会替代部分既有劳动任务，新任务的产生则可能重新增加劳动需求。人工智能造成何种分配后果，取决于替代效应、生产率效应与新任务效应之间的关系。

在这一传统中，两极分化通常显露为：

- 高技能与低技能劳动者的收益差异；
- 常规任务与非常规任务的职业份额变化；
- 拥有互补资产者与被替代者之间的收入差距；
- 能有效采用人工智能的企业与采用失败企业之间的生产率分化。

这一视角的优势，是它拒绝把技术影响理解成纯粹的个人心理问题。一个劳动者即使十分努力，也可能因为任务被重新组合而失去议价位置；另一个劳动者即使技能没有明显变化，也可能因为其技能与人工智能互补而提高收入。

但经济学的核心观察单位通常是结果：单位时间产出、完成率、质量、工资、就业、利润与市场份额。

只要产出提高，生产率就提高；只要低绩效者追上高绩效者，绩效差距就缩小。

至于这项提高在任务完成以后究竟沉淀到了哪里，往往不是生产率统计首先回答的问题。

## （二）心理学：谁真正发生了学习

心理学与学习科学关注的不是某一次答案是否正确，而是行为者在经验之后发生了何种相对持久的改变。

一个学生在提示下解出题目，并不自动意味着他形成了可以迁移的知识结构。一次辅助行为可能促进编码、提取、解释与反思，也可能绕过这些过程，使正确答案在没有学习发生的情况下出现。

认知卸载研究早已指出，人类会把部分记忆、计算与判断要求转移给外部环境。卸载本身并不必然有害：笔记、图表、计算器和搜索引擎都可以释放有限的认知资源。但外部工具究竟是成为认知系统的稳定组成部分，还是使内部能力无法形成，取决于任务、动机、反馈和后续提取要求。

近期关于生成式人工智能与学习的研究进一步表明，辅助期表现与撤除辅助后的表现必须分开测量。直接答案可能提高练习成绩，却削弱独立迁移；带有步骤限制、提示和解释要求的界面则可能降低这种风险。

在这一传统中，两极分化显露为：

- 能把人工智能输出转化为知识结构的人，与只完成任务的人之间的差距；

- 能进行检索、解释、验证和迁移的人，与依赖提示复现的人之间的差距；
- 形成判断力的人，与只形成答案熟悉感的人之间的差距；
- 能在错误和不确定性中学习的人，与越来越回避认知摩擦的人之间的差距。

这一视角关注的是路径：行为者经历了什么加工，知识如何被组织，反馈如何进入下一次判断。

但心理学容易把能力的合法归宿预先限定为个人大脑。只要一个人脱离工具后不能独立完成任务，就可能被判为没有形成能力。

这会遗漏另一种可能：能力虽然没有完整沉淀到个人，却可能稳定地沉淀在组织流程、共享知识库、版本化提示系统、测试套件和可审计技术管线之中。

人类文明的大量能力，本来就不由任何单个人独立持有。

### （三）计算机科学：人机系统是否形成有效协作

在人机交互和人工智能辅助决策研究中，关注点既不是单独的人，也不是单独的模型，而是整个协作条件是否支持恰当判断。

Amershi 等人提出的人机交互设计准则强调，人工智能系统需要在初始交互、日常使用、错误发生和长期变化等不同阶段，帮助用户理解系统能力、修正错误并调整信任。

Buçinca、Malaya 与 Gajos 的研究则发现，要求用户先作独立判断、再查看人工智能建议的“认知强制”机制可以减少过度依赖，但这种设计获得了最差的主观评价，而且收益更多出现在认知需求较高的参与者身上。换言之，能够改善判断质量的界面摩擦，也可能降低采用意愿，并对不同使用者产生不平等效果。

近期对知识工作者的研究也发现，生成式人工智能的使用可能改变批判性思考的分布：使用者把精力从直接生成转向验证、整合和监督，而对工具的信心与投入多少批判性思考之间存在系统关系。

在这一传统中，两极分化显露为：

- 能建立恰当信任的人与过度依赖或过度拒绝的人之间的差距；
- 拥有高质量界面、数据、反馈和权限配置的系统与粗糙系统之间的差距；
- 能追踪来源、识别不确定性、修复错误的人机团队与不可审计团队之间的差距；
- 能把模型嵌入稳定工作流的组织与只购买账号的组织之间的差距。

这一视角关注的是条件系统：模型怎样显露不确定性，界面何时制造摩擦，组织如何分配权限，日志是否保存，错误能否被追溯。

它的局限在于，“团队表现良好”可能成为最终结算标准。

只要人机系统总体上做出了正确判断，研究者就可能把它视为成功的协作。至于原有团队配置被拆散以后，这份成功能否被任何主体接住，则不一定进入主要指标。

### ■ 三、三个学科在同一个案例上给出三种相反判决

设想一家企业部署生成式人工智能客服系统。

一名入职两周的新员工，在人工智能实时建议下，第一次达到资深员工的服务质量；任务完成时间减少，客户满意度提高，投诉率下降。

经济学会说：低经验员工的生产率提升，技能差距正在压缩。

心理学会追问：人工智能撤除以后，他能否独立识别客户意图、处理异常情况并解释判断？如果不能，当前绩效不等于能力形成。

计算机科学则可能说：要求个人脱离系统独立完成，未必是合理标准。只要人机系统能够稳定地产生高质量结果、正确校准信任并在错误时启动人工接管，协作就可以被视为成功。

三种判断都拥有无法被另外两家直接吞并的证据。

经济学不能把真实发生的生产率提高说成不存在。

心理学不能因为当前答案正确，就承认学习必然发生。

计算机科学也不能接受“只有进入个人大脑的能力才是真能力”，因为复杂航空、医疗、软件和供应链系统的能力本来就分布在工具、程序、角色和记录之间。

这不是三种彼此补充的意见，而是三种相互排斥的判决标准：

一次成功究竟由当前结果、能力形成过程，还是协作系统的有效性来裁定？

如果把三种标准并列起来，结论只能是“既要生产率，又要学习，还要设计好系统”。这类结论当然正确，但没有解释最棘手的案例：

- 产出真实提高了；
- 人没有形成独立能力；
- 当前人机系统也真实有效；
- 但系统中的任何一个关键构件发生变化，任务就无人能够接住。

这项成功不是假的。

它也不一定是个人伪装出来的。

问题在于：成功发生了，却没有获得稳定归宿。

### ■ 四、三种理论共同站在一块没有被检查过的地基上



经济学、心理学和计算机科学虽然争论能力应该如何测量，却共同默认了一条更深的前提：

每一次成功都必然已经归属于某个能力承载者。

对经济学而言，生产率被归属于劳动者、岗位、企业或资本配置。

对心理学而言，学习被归属于个人内部相对持久的认知改变。

对计算机科学而言，协作绩效被归属于人机团队或部署系统。

三家争论的是谁拥有能力，却共同假定能力一定已经被谁拥有。

这一前提在传统工具环境中很少造成严重问题。

锤子不能独立生成施工方案；计算器不会在更新以后改变答案风格；纸质手册不会根据操作者的语气临时重组操作程序。工具参与任务，但任务能力仍然较容易归属于训练有素的人、明确的组织规程或稳定的技术装置。

生成式人工智能改变了这一点。

一次成功可能依赖于以下成分的临时配合：

- 某个操作者对问题的模糊直觉；
- 某一版本模型的隐含行为；
- 一段没有保存的提示链；
- 会话上下文中的偶然信息；
- 外部检索返回的当时结果；
- 操作者对输出进行的少量无记录修订；
- 组织中某位同事在临界处给出的口头提醒；
- 当前平台尚未改变的默认参数。

这些成分共同使任务成功，却没有任何一方掌握完整生成过程。

操作者无法在模型撤除后重做。

组织无法让另一名员工按记录复现。

技术团队无法解释哪一步决定了最终结果。

模型供应商也不对具体业务结果承担责任。

如果几个月后模型版本升级、原操作者离职或任务情境变化，组织只能重新摸索。

此时发生的不是“能力被机器拿走了”。

因为能力可能从未稳定地属于机器。

发生的也不只是“个人没有学会”。

因为当前系统确实完成了过去无法完成的任务。

真正出现的是一种以往统计账本里没有独立位置的事件：

成功存在，但没有稳定的能力承载者。

## ■ 五、“无着成功”：一种不属于任何单一主体的成功事件

本章把这种事件称为无着成功。

### （一）定义

对于一项人工智能中介任务，若它在时点  $t$  达到预先规定的质量标准，但在一组预先登记的扰动条件下，不存在任何可识别的人、组织或技术系统能够稳定地复现、解释、修复并承担该任务，则该事件构成一次无着成功。

这里的“无着”不是“无人署名”，也不是“无人获得奖励”。

一篇报告可以署在员工名下，却仍然是无着成功。

一段代码可以完整保存在公司仓库里，却仍然是无着成功。

一个决策可以经过主管签字，却仍然是无着成功。

关键不在名义归属，而在扰动之后谁能接住它。

### （二）四项最低承载功能

一项成功是否获得稳定承载，至少要考察四项功能：

- 1\、复现：在允许的时间和资源范围内，能否再次完成结构相同而表面不同的任务；
- 2\、解释：能否说明关键输入、判断节点、证据来源与不确定性；
- 3\、修复：当结果出错、条件改变或模型失效时，能否定位并纠正问题；
- 4\、负责：是否存在拥有介入权限、承担后果并可被追问的登记主体。

这四项功能不要求全部集中在一个人身上。

个人可能负责解释与专业判断，组织知识库负责程序复现，测试系统负责错误检测，部门负责人承担最终责任。

稳定承载可以是分布式的。

但“分布”不等于“无主”。

只要组成部分、接口、权限与替代规则能够被识别，系统就存在承载结构。

无着成功指向的恰恰是另一种状态：任务成功依赖一个临时配置，而配置拆开以后，没有任何被承认的结构能够继续承担它。

### （三）三类合法承载者

本章承认三种基本承载渠道。

个人承载。原操作者或经过同等训练者，在模型撤除、任务变式或延迟测试中，仍能完成核心判断。

组织承载。原操作者离开以后，替代人员依据组织保存的材料、规范、案例、日志和评审程序，能够重建任务能力。

技术承载。一个经过版本管理、测试、监控和责任配置的技术 workflow，在更换操作者后仍能稳定运行，并能够在异常时被解释和修复。

这一定义拒绝把“脱离 AI 独立完成”设为唯一能力标准。

一个组织完全可以把部分能力外置到可靠系统中。

飞机驾驶能力不只存在于飞行员大脑，也存在于仪表、检查单、训练制度、维修体系与空管协作中。

但是，外置必须产生可持续承载，而不是只产生一次漂亮结果。

### （四）无着成功不是低质量成功

无着成功最容易被误解为“AI 胡乱生成而人没有发现”。那只是最简单的失败。

无着成功的典型对象恰恰是高质量结果。

它可以通过全部当前审查。

客户满意，文章流畅，代码运行，分析结论正确，任务按时完成。

正因为结果足够好，系统才没有发出警报。

它的风险不在当前失败，而在当前成功掩盖了未来无人接手。

## ■ 六、AI 时代两极分化的新定义：能力承载极化

“无着成功”描述单个任务事件，“能力承载极化”描述这些事件长期累积后形成的社会结构。

### （一）定义

能力承载极化是指：在人工智能广泛参与任务生产的条件下，不同个人、群体与组织把人工智能辅助成功转化为稳定个人能力、组织能力或技术能力的概率持续分化。

因此，AI 时代的两极分化：

- 不是现期产出高低的分化；
- 不是个人技能存量多少的分化；
- 不是人工智能接入与否的分化；
- 而是成功事件能否被稳定承载、继续积累和跨情境调用的分化。

收入、职业、教育和地区差距仍然存在。

本章不是否认这些差距，而是提出：它们越来越可能成为能力承载结构的下游表现。

## （二）两极不是“会用 AI”和“不会用 AI”

传统数字鸿沟通常把人分成拥有接入条件者与没有接入条件者。

人工智能素养框架则可能把人分成会提示、会验证、会协作的人与不会使用的人。

能力承载极化划出的两极不同：

一极是可归着增长。每一次人工智能辅助成功，都有一部分被转化为个人判断、组织记忆、技术程序或责任结构；下一次任务因此站在前一次任务的积累之上。

另一极是无着繁荣。任务产量不断增加，短期指标持续改善，但成功依赖的临时配置没有形成稳定承载；模型、人员或情境发生改变后，系统必须重新开始。

两类组织都可能拥有高级模型。

两类员工都可能熟练使用提示词。

两类学校都可能取得漂亮的 AI 辅助成绩。差别在于：

前一次成功是否成为下一次行动可以继承的条件。

## （三）为什么这种分化会自我加速

可归着增长具有累积性。

一项任务被解释、保存、测试和制度化之后，后续成员不必从零开始。组织可以比较版本、分析错误、改进流程，并把高水平判断转化为公共资源。

无着繁荣也具有累积性，但累积的是依赖。

每一次快速成功都会减少整理过程、解释依据和训练替代者的短期收益；流程越顺畅，人们越缺少理由停下来建立承载结构；任务量随后继续增加，使团队更不愿承受记录和学习成本。

于是，短期表现最好的时期，可能正是长期能力最容易被掏空的时期。

这解释了能力承载极化的反常表象：

越成功的人工智能采用，越可能推迟对能力归宿的检查。

七、一张区分四种未来的二维辨别格

只观察当前绩效，无法区分真正的能力增长与无着繁荣。本章引入第二条轴：稳定承载性。

|       | 稳定承载性高                              | 稳定承载性低                         |
|-------|-------------------------------------|--------------------------------|
| 当前绩效高 | 沉淀型增长：任务表现良好，且个人、组织或技术系统能够在扰动后继续承担。 | 无着繁荣：当前结果优异，但撤换模型、操作者或情境后无人接住。 |
| 当前绩效低 | 潜伏能力：现有结果不佳，但系统保留可解释、可迁移、可修复的能力基础。  | 显性剥夺：当前表现差，也没有形成可继承的能力承载结构。    |

表 1 当前绩效 × 稳定承载性的二维辨别格

（一）沉淀型增长

这是多数人工智能政策默认追求的理想状态。

员工借助模型提高效率，同时理解关键判断；组织保存可复用过程；技术系统可追溯、可测试；模型更新后仍能维持工作。

在这一格中，人工智能既提高当前绩效，也扩大未来行动能力。

（二）无着繁荣

这是本章的核心靶格。它具有三个特征。

第一，短期指标表现为改善。产量上升、错误下降、时间缩短、满意度提高。

第二，失效前不产生明显信号。因为当前结果足够好，传统质量控制不会触发。

第三，它会自我加固。指标越好，组织越倾向于扩大部署、减少冗余训练并压缩解释时间。

无着繁荣不是人工智能采用失败。

它是以成功形式发生的能力断裂。

（三）潜伏能力

这一格提醒我们，当前低绩效不等于没有能力。

新手可能工作较慢，却能解释每一步，发现错误并在反馈后迁移。一个组织的流程可能暂时低效，却保存了完整案例、测试和替代机制。

若只按当期产量考核，这类单位可能被无着繁荣者替代。

但在环境变化、模型失效或复杂异常出现时，潜伏能力可能表现出更强韧性。

#### （四）显性剥夺

这是传统不平等研究最容易识别的一格：没有高质量工具，没有适当训练，没有稳定组织支持，当前结果与未来能力都较弱。

人工智能政策若只把资源投向这一格，而忽视无着繁荣，就会误以为高绩效单位已经不需要能力建设。

真正隐蔽的两极分化，发生在第一行内部。

## ■ 八、如何测量：从一次任务到一个组织

### （一）任务事件

设任务事件  $j$  的当前质量为  $Q_j$ ，预先规定的合格阈值为  $q^*$ 。

当  $Q_j \geq q^*$  时，任务被记为一次当前成功。

对每一次当前成功，分别考察三类承载渠道：

- $H_j$ ：个人承载；
- $O_j$ ：组织承载；
- $T_j$ ：技术承载。

每类承载都通过预先登记的扰动测试，而不是通过主观询问判断。可能的扰动包括：

- 1\ 撤除当前模型；
- 2\ 更换模型或模型版本；
- 3\ 移除原提示链；
- 4\ 更换操作者；
- 5\ 延迟七天或三十天重做；
- 6\ 转移到结构相同但表面不同的任务；
- 7\ 引入一个需要修复的错误；
- 8\ 要求说明证据来源与责任主体。

如果至少一个承载渠道在规定扰动下达到复现、解释、修复和负责的最低标准，则该任务获得了承载。

若  $Q_i \geq q^*$ ，但  $H_i$ 、 $O_i$ 、 $T_i$  全部未达到标准，则该任务为无着成功。

## （二）能力承载率

对于一个人、团队或组织，在观察期内所有合格的人工智能中介任务中，定义：

能力承载率 = 获得至少一种稳定承载的成功任务数 ÷ 全部人工智能中介成功任务数。相应地：

无着率 =  $1 - \text{能力承载率}$ 。这一比率与生产率正交。

一个团队可以同时拥有很高生产率和很低能力承载率。

另一个团队可以产出速度较慢，却拥有较高承载率。

这正是新指标存在的必要性：若能力承载率与生产率始终完全同步，它就只是生产率的昂贵替代指标。

## （三）承载深度

二元判断仍然过于粗糙。可以进一步把承载深度分为四级：

- 一级：复现承载—能够重做；
- 二级：解释承载—能够说明关键机制和证据；
- 三级：修复承载—能够在错误或条件改变后调整；
- 四级：责任承载—拥有介入权限并承担制度后果。

一个自动化工作流可能具有较高复现承载，却缺乏解释与责任承载。

一个专家可能具有很强解释和修复能力，却因组织没有授权而无法承担最后责任。

因此，能力承载不是单一分数，而是一个承载向量。

## （四）承载迁移矩阵

人工智能采用并不只会增加或减少能力，它还会改变能力所在的位置。

一项任务的主要承载者可能发生以下迁移：

- 从个人迁移到组织；
- 从个人迁移到技术系统；
- 从组织迁移到供应商；
- 从技术系统迁移回专业人员；



- 从可识别主体迁移到临时配置；
- 从临时配置重新稳定到组织流程。

前四种未必是损失。第五种才是无着化。

因此，研究人工智能与能力变化，不能只问“人的能力是否下降”，还要问：能力迁移到了哪里？新的位置能否被识别、维护、替换和追责？

## ■ 九、从成功到承载：四条稳定化路径

人工智能辅助结果不会自动成为稳定能力。它需要经过至少四种稳定化过程。

### （一）认知稳定化

认知稳定化是指，人工智能输出进入人的知识结构、判断习惯与迁移能力。

它不等于阅读过答案，也不等于能够复述答案。关键机制包括：

- 在获得建议前形成自己的初步判断；
- 比较自身判断与模型输出的分歧；
- 解释为什么接受或拒绝模型建议；
- 在相似但不同的任务中重新生成；
- 延迟提取而不是立即模仿；
- 经历能够被定位和纠正的错误。

带有教学护栏的人工智能系统之所以可能优于直接答案系统，不是因为它提供了更多信息，而是因为它保留了能力形成必须经过的部分路径。

然而，认知稳定化不能被理解为“给所有任务增加摩擦”。

认知强制可能减少过度依赖，却降低用户满意度，而且不同认知动机者受益不同。

因此，稳定化设计需要判断：哪些节点必须由人形成可迁移判断，哪些步骤可以可靠外置。

### （二）组织稳定化

组织稳定化是指，一次成功被转化为其他成员可以继承的公共能力。它至少需要：

- 保存任务目标、输入和约束；
- 保存关键提示与人工修改；
- 记录采用或拒绝模型建议的理由；
- 建立失败案例而不只保存成功模板；

- 明确任务所有者与替代者；
- 定期由未参与原任务者复做；
- 将隐性经验转化为案例、测试和评审程序。

组织记录的价值不在于“资料越多越好”。

大量未经整理的会话记录不构成组织承载。

只有当下一位行动者能够据此缩短重建时间、识别关键分歧并承担任务时，记录才成为能力结构。

### （三）技术稳定化

技术稳定化是指，把临时会话转化为可版本化、可测试、可观察和可回滚的工作流。它包括：

- 固定模型与参数版本；
- 保存提示模板和外部工具配置；
- 建立输入输出测试集；
- 记录数据与证据来源；
- 监控模型漂移；
- 设计失败时的人工接管；
- 保留回滚和复核机制；
- 在更换操作者后验证工作流是否仍可运行。

技术稳定化并不是把任务全部自动化。

相反，一个无法由人解释、无法安全中止的全自动系统，可能拥有复现能力，却没有完整承载能力。

### （四）责任稳定化

责任稳定化解决“出了问题谁能在事前介入，而不只是事后背锅”。

名义上签字的人未必是责任承载者。

真正的责任承载至少需要：

- 能够观察任务过程；
- 有权暂停或改变流程；
- 有资源修复问题；

- 会承担可识别的制度后果；
- 其责任范围与控制能力相匹配。

如果一个员工必须对人工智能决策负责，却无权查看模型来源、改变系统或拒绝使用，他不是承载者，而是后果接收者。

能力承载极化因此也是权力极化。

拥有介入权的组织可以把人工智能成功稳定下来；只有使用义务、没有设计权与退出权的群体，则更容易积累无着成功。

## ■ 十、两极分化如何在微观成功中发生

能力承载极化不是从宏观收入统计突然出现的。它在每一次人工智能辅助任务中发生。

### 命题一：即时绩效收敛可以与承载能力分化同时发生

当人工智能主要补偿低经验者的即时任务缺口，却没有同步提高其迁移、修复和解释能力时，当前绩效差距缩小，而扰动后的能力差距扩大。

这不是“先缩小、以后又扩大”的时间顺序而已。

在同一时点，只要换一个测量条件，两种方向就可能同时被观察到。

### 命题二：高产出环境更容易隐藏无着成功

当组织把采用规模、产量、速度和满意度作为主要成功指标时，人工智能辅助任务越顺利，检查能力承载的激励越弱。

因此，无着率未必在低绩效团队中最高，反而可能在短期采用成绩最亮眼的团队中最高。

### 命题三：模型升级会成为能力归宿的显影剂

若一项能力已稳定沉淀到个人、组织或技术系统，模型更换会造成调整成本，但不会使任务完全失去承载。

若成功依赖当前模型的偶然行为，版本升级会使质量发生断崖式下降，而且团队无法解释下降来自何处。

因此，模型更新、平台故障与供应商切换不是单纯技术噪声，而是识别无着成功的自然实验。

### 命题四：组织承载可以替代部分个人承载

个人在撤除人工智能后不能独立完成任务，不足以证明能力已经丧失。

若替代员工能够借助组织记录、测试和流程稳定完成任务，该能力已获得组织承载。

因此，把“无辅助个人表现”设为唯一因变量，会系统性低估有效的人机组织学习。

#### 命题五：平等接入不会自动产生平等承载

两所学校即使使用同一模型、拥有相同账号和设备，也可能形成完全不同的能力承载率。

一所学校把人工智能用于答案生成，另一所学校要求先作判断、比较分歧、延迟迁移并保存错误路径。

两者接入条件相同，短期成绩也可能相近，但成功事件的归宿不同。

因此，接入鸿沟缩小以后，能力承载极化仍可继续扩大。

#### 命题六：人工智能摩擦具有分配效应

要求用户先判断、解释理由或验证来源，可能促进能力稳定化，却会增加时间和认知成本。

高认知需求、高先备知识和高自主性使用者可能从摩擦中获益更多；资源紧张者则可能绕过、拒绝或形式化完成这些步骤。认知强制研究已经显示，降低过度依赖的界面并不一定受到用户欢迎，其收益也可能因个体特征而异。

因此，“增加摩擦”不是普遍处方。缺乏补偿和差异化支持的摩擦设计，可能以培养判断力之名扩大新的差距。

#### 命题七：高技能者也可能产生无着成功

无着成功不是低技能者专属风险。

高技能者可能凭借经验判断迅速修改人工智能输出，却没有记录修改依据；最终结果很好，但组织无法知道专家改了什么、为什么改。

Brynjolfsson 等人的研究还发现，高技能员工在使用人工智能建议时可能提高对建议文本的遵循，而部分边际建议并不改善质量；研究也提出，若高水平员工减少原创方案，未来系统可学习的多样性可能受损。

因此，顶尖个体创造的未记录修订，也可能形成组织层面的无着成功。

#### 命题八：承载极化先于收入极化显露

收入和职位变化往往滞后于任务结构变化。

在人工智能采用初期，不同员工可能仍拥有相近职位和工资，但他们获得的任务机会已经不同：

- 一些人持续处理异常、解释冲突和修复错误；
- 另一些人主要提交提示、接收答案并完成格式加工。

前一组不断积累能够接管系统的能力，后一组积累的是依赖当前系统的产出记录。

当组织重组、模型升级或责任事故发生时，这种隐藏差距才转化为职位与收入差距。

因此，能力承载率可能是收入两极分化的领先指标，而不是其同义指标。

## ■ 十一、五类可实施的研究设计

### （一）企业内部多部门自然实验

在同一家企业选择不少于六个使用同一模型、处理相似任务的部门。收集：

- 人工智能辅助期生产率；
- 模型故障或版本更新前后的质量变化；
- 原操作者离开后的替代成本；
- 工作流记录完整度；
- 任务解释与修复时间；
- 个人、组织和技术承载测试结果。

控制任务复杂度、工作总量、员工经验和管理制度后，检验能力承载率能否预测扰动后的恢复速度。

这种设计优于比较两个国家或两家完全不同的公司，因为后者无法排除文化、制度和行业差异。

### （二）教育中的双时点随机实验

随机设置四组：

- 1\、无人工智能组；
- 2\、直接答案组；
- 3\、提示与解释组；
- 4\、先独立判断、再使用人工智能并比较分歧组。

分别测量：

- 辅助期任务表现；
- 七天后的无辅助迁移；
- 三十天后的延迟保持；

- 新工具环境中的重新学习速度；
- 学生对自身能力的校准；
- 小组是否形成可共享的解题记录。

该设计不仅比较“是否学会”，还比较能力最终沉淀到了个人、小组流程还是未被任何结构接住。

### （三）专业组织中的人员替换测试

从法律、咨询、医疗行政、软件开发或出版团队中抽取已经完成的人工智能中介任务。

在不使用原操作者私人会话记录的条件下，让另一名合格人员依据组织现有材料复做。测量：

- 重建所需时间；
- 与原结果的质量差异；
- 找到关键证据的比例；
- 无法解释的人工修改数量；
- 需要联系原操作者的次数。

这是一种低成本测试，因为多数机构已经保存任务记录、版本信息、工时和人员变动数据。

### （四）模型切换实验

选择同一组织中的多个任务单元，在预先登记的窗口内更换模型供应商或主要模型版本。能力承载理论预测：

若组织拥有高承载率，切换会先产生有限调整成本，随后恢复；若组织积累了大量无着成功，切换将产生不成比例的质量下降、返工和责任不明。

关键不是比较哪个模型更强，而是观察：

模型变化以后，组织中是否存在能解释差异并重建流程的主体。

### （五）纵向职业发展研究

追踪同一批初级员工三至五年，记录他们在人工智能环境中获得的任务类型。区分：

- 仅提交结果的任务；
- 必须解释理由的任务；
- 必须处理异常的任务；

- 负责模型评估和流程设计的任务；
- 参与错误复盘的任务；
- 能够接替他人工作的任务。

能力承载理论预测，决定未来专业地位的不只是人工智能使用频率，而是员工是否进入了承载形成节点。

相同的工具使用时长，可能对应完全相反的职业结果。

十二、与九种相邻理论的判决性分界

“无着成功”若只是既有概念的换名，就没有理论价值。下面以同一类案例给出可判决的分界线。

认知卸载、技能退化和分布式认知，是本章最重要的传统近邻。

| 相邻概念          | 它已经解释了什么                | 与本章的分界线                                 | 判决性观察                                          |
|---------------|-------------------------|-----------------------------------------|------------------------------------------------|
| 技能偏向技术变迁与任务极化 | 技术改变不同技能与任务的相对需求。       | 它主要预测岗位、任务和收入重新分布；本章预测相同现期生产率下的扰动后承载差异。 | 若即时绩效相同且任务结构不变，但模型撤换后只有一部分单位崩溃，则能力承载理论增加了独立解释。 |
| 数字鸿沟          | 不同群体在设备、网络、接入和使用质量上不平等。 | 本章允许完全相同的工具接入，却产生不同承载率。                 | 若控制模型、账号、设备和使用时长后，承载率仍显著分化，则接入理论不足。            |
| 认知卸载          | 人把记忆或计算要求转移给外部环境。       | 卸载可以稳定地成为个人—工具系统的一部分；无着成功要求外置后没有稳定承载者。  | 若替代操作者可依据外部系统稳定复做，则是有效卸载而非无着成功。                |
| 技能退化          | 原有技能因少用、自动化或训练中断而下降。    | 无着成功不要求行为者曾经拥有该技能，也不要求个人能力下降。           | 新手从未独立掌握任务，却借助 AI 完成高质量结果，仍可形成无着成功。            |
| 自动化偏误         | 人不恰当地接受自动化建议，尤其在建议错误时。  | 无着成功可以建立在完全正确的模型输出之上。                   | 若所有输出均正确，但人员和组织在模型切换后无法重建任务，则自动化偏误不能解释。        |
| 分布式认知         | 认知可以跨越人、工具、符号和组织而分布。    | 本章不否认分布式能力，而是区分可识别、可维护的分布结构与临           | 若更换成员后，凭借稳定接口与记录仍能复现，则是分布式承载，                  |



|           |                          | 时配置。                                               | 不是无着。                                              |
|-----------|--------------------------|----------------------------------------------------|----------------------------------------------------|
| 恰当依赖      | 使用者在应接受时接受、应拒绝时拒绝人工智能建议。 | 一个团队可以在当前任务上依赖校准正确，却没有扰动后的复现与修复能力。                 | 若接受与拒绝决策均正确，但模型版本变化后无人解释失败，则依赖校准不足。                |
| 合成能力／合成胜任 | AI 使人表现出超过其真实理解水平的专业外观。  | 本章不以「外观是否虚假」或「个人是否真正理解」为判据，而以任何合法承载者能否在扰动后接住任务为判据。 | 若个人不理解，但组织流程和技术系统可稳定复现、修复并负责，本章判为已承载，合成能力论仍可能判为虚假。 |
| AI 劳动孤儿化  | AI 完成的工作没有被组织账目正确记录或归因。  | 账目归因与能力承载是两个维度；完整记录的工作仍可能无人复现，未记录的工作也可能已形成稳定能力。    | 若任务记录完整却在人员替换后崩溃，则不是单纯的劳动记账问题。                     |

表 2 与九种相邻理论的判决性分界（编辑按原稿内容还原行列对应）

2026 年出现的“合成能力”研究，把人工智能制造的专业外观与缺乏内部理解区分开来；另有治理研究把人工智能辅助吞吐量被误认成组织能力称为“synthetic competence”。这些工作已经非常接近本章的问题意识。

但它们仍然倾向于把真实能力与个人理解，或与组织能否调试、扩展 workflow 联系起来。

无着成功的最小分析单位不是“看起来有能力的人”，而是一次达到质量标准的任务事件；它不预设能力应该落在个人，也不把“人工智能生成”本身视为不真实。只要个人、组织或技术渠道中的任何一个能够稳定承载，任务就不属于无着成功。

因此，最近的概念邻居是“合成能力”，但两者对同一案例可能作出相反判断：

一名医生不具备独立重建模型计算的能力，但医院保存了版本化流程、输入数据、验证规则、替代人员训练和责任制度。合成能力理论可能强调个人理解不足；本章判定该任务已经获得组织—技术承载。反过来：

一名专家深刻理解专业内容，却依靠未保存的私人提示链和临场修订完成任务，离职后无人能够复做。合成能力理论可能仍承认专家真实胜任；本章则把该组织的这次成功判为无着。

两者并非强调程度不同，而是分类标准不同。

### 十三、十个领域中的理论兑现

### （一）学校教育

两名学生都用人工智能写出高质量文章。

第一名能够解释结构选择，在新题目中重新组织论证；第二名只能再次请求模型生成。

两人的当前作品相同，能力归宿不同。

学校若只检查作品质量，会把无着成功认证为学习成果。

### （二）教师工作

教师利用人工智能生成完整教案。

如果另一位教师能够依据课程目标、学生反应记录和修改理由继续迭代，教案形成了组织承载。

如果教案只存在于个人会话中，原教师离开后团队无法解释为何如此设计，它便是无着成功。

### （三）软件开发

人工智能生成的代码通过所有现有测试，并成功上线。

若团队无法解释关键依赖、无法在需求改变后修改，也没有保存生成环境和评审理由，代码的正确运行可能属于无着繁荣。

这与代码当前是否存在错误无关。

### （四）医疗行政

人工智能帮助完成复杂病例编码或资源配置。

如果当前操作者只会接受建议，而组织没有保存依据、复核节点与异常处理程序，那么一次准确决策仍可能没有承载者。

但在诊疗安全场景中，不能为了测试承载率而突然撤除系统。承载测试应优先使用模拟、历史回放和非关键任务。

### （五）法律服务

律师借助人工智能检索并生成论证。

若引用准确、结论成立，但没有人能够说明为什么排除了相反判例，任务仍缺乏解释与责任承载。

法律任务的能力不是文字质量，而是面对反例时能够继续辩护和修正。

## （六）企业管理

管理者使用人工智能生成战略方案。

一份措辞完美的战略可能获得董事会认可，却没有任何成员掌握其关键假设和反事实条件。

当市场改变时，团队只能重新询问模型，而不能判断原方案的哪一部分已经失效。这是一种高层无着成功。

## （七）公共治理

政府部门利用人工智能处理申请、分流案件或生成政策分析。

只记录最终决定而不保存证据、模型版本、人工修改和介入权限，会使决策结果存在却无可追责的能力承载结构。

公共部门中的责任承载不能被“有人签字”替代。

## （八）科学研究

人工智能生成假设、代码或分析路径。

研究者可能验证最终数值正确，却无法重建模型如何选择变量、排除方案和生成解释。

结果可重复不等于推理过程可承载；反过来，研究者能解释机制，也不保证计算工作流可复现。

科学能力需要个人、组织和技术三种承载的协同。

## （九）文化创作

生成式人工智能可能提升个体作品质量，同时降低群体层面的内容多样性。相关实验研究已经提示，个体创造收益与集体新颖性之间可能存在张力。

在本章框架下，问题进一步变成：创作成功是否沉淀为创作者新的审美选择、团队新的工作方法或可持续的技术语言，还是只存在于一次模型配置中。

## （十）家庭与个人生活

个人使用人工智能制定旅行、理财、健康或学习计划。

当前方案可能非常合理，但若使用者无法识别关键约束、调整意外情况或承担决定后果，成功规划没有形成完整承载。

AI 时代的依赖不是“凡事问 AI”这么简单。

真正的分界是：询问之后，下一次行动能力留在了哪里。

## ■ 十四、三个学科反过来迫使理论让步

一个跨学科理论若只借用三个领域的材料，却不允许它们削弱自己的主张，就只是扩张性的命名。能力承载理论必须接受至少三项实质约束。

### （一）经济学的约束：不能把所有依赖都称为损失

若没有经济学的反向约束，本章很容易写下：

人工智能辅助下不能独立完成任务，都表明能力发生了空心化。这句话必须删除。

分工、专业化和资本互补本来就意味着个体不需要独立拥有全部能力。现代组织通过把能力稳定地配置在不同岗位、机器和制度中提高生产率。

只要系统能够可靠复现、修复和负责，能力从个人迁移到组织或技术系统可以是增长，而不是损失。

因此，本章不把个人独立性置于最高价值。

它保护的是可持续承载，而不是工具撤除后的个人自足。

### （二）心理学的约束：组织承载不能替代所有个人判断

若没有心理学的约束，本章可能写下：

只要组织流程能够复现任务，个人是否理解并不重要。这句话同样必须删除。

在环境变化、异常案例和价值冲突中，组织流程无法穷尽所有条件。若没有成员形成可迁移的概念结构与判断力，流程只能复制既有情境。

因此，在高不确定、高责任和高新颖性任务中，至少一部分关键能力必须形成个人承载。

组织记录可以支持判断，却不能无限替代判断。

### （三）计算机科学的约束：承载测试本身可能破坏有效协作

若没有人机交互研究的约束，本章可能主张：

每一次任务都应要求使用者先独立完成，再使用人工智能。

这会把能力承载理论变成低效率的普遍摩擦制度。

对常规、低风险、可逆且拥有成熟测试的任务，强制独立重做可能浪费资源，并降低采用意愿。认知强制机制改善部分判断，却可能带来较差用户体验，并对不同人产生不均等收益。

因此，承载测试应该按任务风险、变化速度和不可逆性分层。

不是每一项任务都要在个人层面证明完整承载。

### （四）由三项约束得到的适用边界

能力承载理论最适用于以下任务：

- 条件经常变化；
- 错误后果较高；
- 需要解释和修复；
- 人员与模型可能更替；
- 决策具有不可逆后果；
- 组织声称自己已经形成长期能力。

它不应被机械应用于：

- 一次性娱乐生成；
- 低风险、易验证的日常任务；
- 不需要长期继承的临时表达；
- 已有成熟自动化测试和责任体系的稳定流程；
- 撤除系统本身会制造不必要危险的安全关键场景。

## ■ 十五、八条证伪条件

本章的核心主张不能靠“未来可能存在风险”维持。以下观察将迫使理论被修改或放弃。

### 证伪一：承载率没有独立预测力

控制工作量、任务复杂度、员工经验、模型版本与资源投入后，若能力承载率不能预测模型故障、人员替换或情境迁移后的质量变化，则能力承载率只是既有能力指标的冗余代理。

### 证伪二：即时绩效收敛必然伴随承载收敛

若在多个同质样本中，人工智能每次压缩当前绩效差距时，也同比例压缩扰动后的复现、解释和修复差距，就不存在本章所说的隐藏极化。

### 证伪三：无着成功无法与合成能力区分

若所有被判为无着成功的任务，都可以仅通过测量个人真实理解与表面表现的差距完整识别，而且组织承载和技术承载不提供任何额外分类，则“无着成功”应被“合成能力”吸收。

### 证伪四：合法外置与无着化无法实证分离

若替换操作者、模型和情境以后，拥有稳定记录、测试和责任结构的外置系统，与临时会话系统表现没有系统差异，则本章对“分布式承载”和“无着成功”的区分不成立。

### 证伪五：承载稳定化没有顺序效应

本章预测，组织记录、人员训练、测试体系与责任配置的建立，应当早于扰动后恢复能力的改善。

若恢复能力总是先改善，承载结构才随后补写，或者两者没有稳定先后关系，则本章关于承载形成机制的解释受到反驳。

### 证伪六：模型切换不显露无着率

若高无着率与低无着率单位在模型更换、供应商迁移和提示失效后的质量下降完全相同，则模型扰动不是有效显影条件，理论必须寻找其他识别机制。

### 证伪七：个人、组织和技术承载不能互相替代

本章主张三种承载渠道可以在一定范围内替代。如果实证表明，只有个人无辅助能力能够预测长期结果，组织和技术承载始终没有独立价值，则本章必须退回个人能力理论。

反过来，若个人判断在所有任务中都没有独立价值，理论也必须删除对认知稳定化的强调。

### 证伪八：低成本日志测试不能识别任何差异

许多机构已经拥有系统故障日志、版本记录、人员更替记录、返工时间和质量数据。

若利用这些现成记录，无法发现承载率与恢复时间之间的任何关系，且需要完全依靠主观访谈才能维持理论，本章的可操作性就不足。

### 最强竞争解释

最强竞争解释是：

本章描述的不就是先备知识、组织成熟度和管理质量的综合结果吗？

为排除这一解释，研究必须在相似任务、相同模型、相近人员经验和相同组织制度内，比较不同工作流的承载形成过程。

只有当承载变量在控制一般能力、组织质量和资源投入后仍预测扰动结果，本章机制才获得支持。

## ■ 十六、伦理边界：能力承载率不能变成人的评级工具

任何可量化的新指标，都会面临被组织用于考核人的诱惑。

能力承载率尤其危险，因为一个人可能在缺乏权限、训练、记录工具和时间资源的条件下，被迫产生大量无着成功。

把这种结构结果归咎于个人，会再次掩盖真正的承载问题。

因此，应用能力承载率必须遵守三条原则。

### 第一，指标指向任务位置和工作流，不指向人的本质

低承载率首先意味着某类任务没有形成稳定归宿，不意味着执行者缺乏能力、责任心或人工智能素养。

个人指标只能在充分控制任务类型、权限、训练与组织支持后使用。

### 第二，不得为了提高指标而在安全关键系统中制造撤除和轮换

医疗、交通、金融结算与公共安全系统不能为了测试“谁能独立完成”而突然撤除关键工具。

应使用模拟环境、历史案例、影子运行、受控切换和非关键任务评估。

### 第三，任何不利安排必须公开编码规则并提供复核通道

若机构依据承载指标调整岗位、培训、权限或晋升，必须公开：

- 哪些任务被计入；
- 何种扰动被使用；
- 哪些承载渠道被承认；
- 质量阈值如何设定；
- 员工拥有哪些补充证据和申诉权。

更根本的问题是，指标一旦进入考核，组织最便宜的达标方式可能不是建立真实承载，而是补写文件、形式化指定责任人，或者安排容易通过的复现测试。这会制造“承载表演”。

目前看不到一种能够彻底消除这一反噬的制度方案。

因此，能力承载率更适合用于诊断工作流和配置资源，而不适合直接排名个人。

## ■ 十七、本章自身是否也是一次无着成功

按照本章的判据，这篇文章不能因为结构完整、概念清晰和文献充分，就自动被视为已经形成理论能力。

它同样可能处于“当前绩效高、稳定承载低”的格子。

本章使用人工智能参与资料检索、结构组织、近邻比较和语言生成。若作者不能：

- 解释“无着成功”与相邻概念的核心分离线；
- 在新的行业案例中重新推导二维框架；
- 回应反例并修改适用边界；
- 独立设计承载率的实证测量；
- 说明哪些主张来自证据、哪些属于理论推演；

那么本章本身就是它所批判的对象：一篇达到形式标准却没有稳定理论承载者的论文。

即使作者能够解释，文章也尚未获得组织与学术共同体层面的承载。

它需要经过同行批评、概念替换测试、实证操作化与重复研究。只有当其他研究者能够依据公开定义区分案例、产生反例并修改理论，概念才开始获得公共承载。

因此，发表不是本章能力归宿的终点。

发表只意味着一次理论任务取得了当前成功。

## ■ 十八、证据分层与本章解释不了的洞

本章的主张可以分为三层。

### 第一层：无着成功是一类独立存在的任务事件

这是最强、也最容易被推翻的主张。

它要求证明：存在一些成功事件，不能被个人技能、合成能力、认知卸载、分布式认知或组织成熟度完整吸收。

若这一层失败，“无着成功”应被删除。

### 第二层：二维辨别格具有分析价值

即使“无着成功”不是完全新的存在类别，把当前绩效与稳定承载分开测量，仍可能帮助区分不同人工智能采用结果。

这一层不要求新概念绝对不可替代，只要求二维框架改善预测。

### 第三层：扰动后的能力不同于辅助期表现

这是最稳的经验主张。



现有劳动和教育研究已经显示，人工智能辅助下的即时表现与撤除或变化条件下的表现不能被自动视为同一指标。不同研究场景的结果并不完全一致：有的工作场景观察到较持久的学习，有的教育场景则发现无护栏使用会损害后续独立表现。

即使本章的新概念最终失败，研究者仍应把辅助期结果与扰动后结果分别报告。

本章尚不能解释至少三个洞。

第一，稳定承载需要持续多久才算“稳定”。七天、半年与五年的结论可能不同。

第二，在高度复杂的系统中，是否可能存在没有任何主体能完整解释，却仍然具有可靠分布式承载的能力。

第三，承载能力本身是否会因过度制度化而失去创造性。一个可复现、可审计的流程可能压缩偏离既有规范的探索。

这些问题不能通过增加定义解决，需要长期经验研究。

## ■ 十九、结论：AI 时代最深的差距，藏在同样正确的答案之间

关于人工智能与两极分化的争论，通常围绕一个方向性问题展开：

人工智能究竟扩大差距，还是缩小差距？

这个问题把差距理解成一根轴，因而不断得到相互冲突的答案。

人工智能可以缩小即时生产率差距，扩大独立迁移差距；可以降低低技能者进入任务的门槛，又提高未来负责异常情况的能力门槛；可以让更多人获得高质量结果，也可以让更少的人知道结果如何产生、如何修复和由谁承担。

本章提出，必须把问题改写为：

一次人工智能辅助成功，在任务结束以后去了哪里？

它可能进入个人的判断结构。

可能进入组织的公共记忆。

可能进入可维护的技术工作流。

也可能只存在于一个已经消失的人—模型—提示—情境配置中。最后一种不是失败。它曾经真实地成功。

但它没有成为任何主体下一次行动的能力。

由此，AI 时代的两极分化不是收入结果、个人技能或工具接入本身，而是成功能否获得稳定承载的分化。

一端不断把成功变成新的起点。

另一端不断制造没有下一轮继承者的终点。

它们在仪表盘上可能拥有同样的产出、质量和满意度，却生活在完全不同的未来里。

本章最终留下的判据不需要评价“真正的能力”“充分的理解”或“合理的依赖”。

它只问一句可以进入日志、审核与实验的问题：

撤去当前模型版本、提示链与原操作者之后，同类任务由哪个登记主体重做，质量下降多少？

回答不出这个问题时，成功没有消失。消失的是它的归宿。

## ■ 二十、写死日期的理论赌注

本章作出如下可被公开核验的赌注：

截至 2029 年 12 月 31 日，至少一项覆盖多个同质组织单元的同行评议实证研究，或一套进入机构正式审计流程的人工智能能力标准，将把“模型撤换、操作者替换或情境迁移后的任务复现与修复”列为独立指标，而不再只报告人工智能接入率、使用频率、辅助期生产率或自陈人工智能素养。以下情形不算命中：

- 只在政策文件中笼统提到“保持人的能力”；
- 只要求机构声明保留人工监督；
- 只测量使用者是否信任人工智能；
- 只进行一次无对照的问卷调查；
- 只测试模型准确率而不测试任务承载者；
- 只要求保存提示词，却不进行人员、模型或情境扰动；
- 只发表观点文章而没有可执行测量规则。

若到该日期仍没有出现上述研究或标准，至少说明两种可能：

第一，能力承载不是人工智能治理中的独立问题，现有生产率、学习与可靠性指标已经足够。

第二，本章虽指出一个真实问题，却没有提供足够便宜、可靠和可接受的测量方式。

无论哪一种，本章的核心主张都必须被显著降格。

## ■ 版本与证据声明

版本：2026 年 8 月 6 日，理论建构初稿。

本章属于概念与理论论文，不把既有单项实验结果直接解释为宏观社会已经发生能力承载极化。文中的宏观机制、能力承载率与承载迁移矩阵均属于待检验理论构造。劳动、教育和人机交互研究只为部分机制提供经验入口，尚不能单独证明全文命题。

## ■ 人机分工声明

本章的研究问题、理论方向、核心判断及最终署名责任归作者所有。人工智能参与文献检索、理论对照、结构组织、反例生成与文字表达。作者需对所有引用、概念原创性、事实准确性、伦理边界和最终学术主张承担责任。人工智能不构成本章所称的责任承载者。

## ■ 附论一·最近邻补切（编辑增补）

原稿第十二节处理了九类相邻理论，但下列五家未出场，其中第一家是「成功没有承载者」这一现象在工程实务里的本名。补切不改变作者的核心判断。

### 一、关键人依赖（软件与运维实务中的 bus factor / truck factor）

已说到哪一步。软件工程实务里有一个粗糙但极其常用的读数：一个项目里，最少需要多少人同时离开，剩下的人就无法继续维护它。围绕它已有成型的诊断实践——按提交历史计算知识集中度、识别只有一人碰过的模块、把「无人可替」的文件列为风险清单，并以结对、轮岗、文档化与代码评审作为处方。这正是本章所说的无着成功在一个具体行业里的既有命名，而原稿零命中。

分离线。关键人依赖的单位是资产（模块、系统、客户关系），问的是「谁还能维护它」；本章的单位是一次达到质量标准的任务事件，问的是「这次成功之后，复现、解释、修复与负责这四项功能落在哪里」。两条差别是实质的：其一，关键人依赖默认能力曾经存在于某人身上、只是过于集中；本章明确允许一种从未被任何人掌握过的成功（新手借 AI 一次做成，之后谁都不能重做），这是 bus factor 的算法算不出来的格。其二，关键人依赖的处方是分散与备份；本章的第五种迁移（从可识别主体迁到临时配置）指出，即使人数充足、文档齐全，只要成功依赖的是一次性的模型版本、会话上下文与未记录的临场修订，分散人手也不解决问题。

判决性对照。取一批任务，按提交历史与人员冗余算出关键人依赖分数；若在该分数匹配之后，模型替换与操作者替换仍产生不同的复做质量，本章获得独立内容；若关键人依赖分数已完整预测复做结果，本章应并入这一实务传统。

这位祖宗反过来削弱本章的一条。bus factor 有一件本章没有的东西：它可以从现成日志里自动算出来，成本几乎为零。本章的能力承载率要求实施预先登记的扰动测试（撤除、更换

模型、更换操作者、延迟重做），成本高出一个量级。若两者在预测扰动后质量下降上表现相当，成本更低者胜出——因此本章第十五节的「证伪八」应当把关键人依赖分数列为必须被超越的基线，而不是只与主观访谈作对比。

## 二、Nelson & Winter 的组织常规

演化经济学把组织能力安置在「常规」上：常规是组织的记忆，能力靠反复执行而非靠文档保存，且大量常规带有默会成分、无法完整写下来。分离线：常规理论解释能力如何在组织中存续，本章问一次成功是否进入了任何存续结构。判决性对照：若一项任务的复做能力可由该组织既有常规的执行频率完全预测，本章的「组织承载」只是常规强度的代名词；若在常规频率匹配后，保存了拒绝理由与失败案例的单位仍恢复得更好，本章的第九节第二条路径获得独立支持。这位祖宗反过来削弱本章的一条：常规理论会说，本章第九节列出的七项组织稳定化动作（保存目标、保存提示、记录采纳理由、建立失败案例……）大多是显性化动作，而组织能力的主要载体恰恰是默会的、无法显性化的部分；把承载等同于可记录、可交接，可能系统性低估那些靠反复共同执行而非靠文档维持的能力。

## 三、Walsh & Ungson 的组织记忆

组织记忆被拆成若干存留箱：个体、文化、转换（流程）、结构（角色）、生态（物理布置）与外部档案。本章的三条承载渠道（个人／组织／技术）与之部分重叠，但更粗。建议：本章的「组织承载」至少应拆成流程存留与角色存留两项——原稿第九节第二条路径里的「明确任务所有者与替代者」属角色，「定期由未参与原任务者复做」属流程，二者可以一有一无，合并会掩盖最常见的一种失败（流程齐全但无人被指派）。

## 四、Nonaka 的知识转化（隐性—显性四模式）

SECI 循环描述隐性知识如何经外化、组合与内化在组织中扩散。本章第九节的四条稳定化路径可以整体读作一次 SECI 的重述。分离线：SECI 关心知识的扩散与创造，本章关心一次成功是否获得可被追问的归属与责任——本章的第四项功能「负责」（拥有介入权限、承担后果、可被追问）在 SECI 里没有对应物，这是本章相对该框架最清晰的剩余。

## 五、Argyris & Schön 的组织学习

该书正面处理「组织如何学习而不只是个人在组织中学习」，并提出组织的实用理论存在于共享图景与组织物件之中。分离线：组织学习理论关心框架的修改，本章关心一次已完成的成功是否被任何结构接住；一个组织可以完全不修改框架，而把每一次成功都稳稳存进流程

（本章判已承载，组织学习判为单环）。编辑注：该作者在本站另一批论文中已被指出同一空缺，宜在下一版正文中一次处理。

## ■ 附论二·站内划界（编辑增补）

与同批《耦合可逆性极化》——两篇同标「What 第一篇」，须由作者定夺

两篇同日投稿，共享同一个反直觉起点（当下产出差距在收敛、脱离原配置后的能力差距在扩大），共享同一批经验入口（Bastani 的高中数学实验、Brynjolfsson 等的客服现场、Noy 与 Zhang 的写作实验、Doshi 与 Hauser 的创造性研究），并各自画了一张二维格，两张格的靶格是同一格：那篇称「托管型绩效」，本章称「无着繁荣」。

可用的分工线。那篇的对象是可恢复性——条件改变后能否重建问题表征、验证链与修复顺序，测量落在切换损失的分布形状上；本章的对象是归宿——一次成功之后，复现、解释、修复、负责这四项功能是否落到某条可识别的承载渠道上，测量落在承载率与承载迁移矩阵上。本章独有而那篇没有的，是「分布不等于无主」这一区分与六种迁移方向中的第五种（从可识别主体迁到临时配置）；那篇独有而本章没有的，是把「极化」写成必须满足分布判据的经验命题，以及一次真跑过的公开数据复算。按这条线，两篇可改为 What 第一篇与第二篇，本章宜排在后。

一条必须由作者处理的口径冲突。那篇明确规定：只有在稳定产出匹配后二成分模型显著优于一成分、成分间距与人口权重达标、并在至少两类扰动上重复，才可称为「极化」；它据此拒绝把自己的数据称为两极化。按同一判据，本章标题与核心构念中的「能力承载极化」目前没有任何分布证据，应读作待检验的命名而非已成立的结论。本站不代作者改动题名，仅在此标明，并建议两篇统一口径。

与本站《冲突与和平》九篇（同一作者）

那一组的共同判断是「最深的一步在短期指标最好看时完成」，本章的「无着繁荣」在形式上同族。分离线：那一组处理的是资格被取消、转移或新授（谁还有资格触发规则层复核）；本章处理的是能力的归属位置（这次成功落到了谁名下）。可分离的观察是：本章的靶格里没有任何资格被取消，权限、账号与角色表都完好，只是没有任何一方能在扰动后接住任务。

## ■ 附论三·编辑校订说明

本站在不改动作者判断与论证结构的前提下，做了以下技术性处理：

一、近邻分界表的还原。原稿第十二节的分界表在文档导出时发生单元格错行——「技能偏向技术变迁」一行的分离线与判决性观察被截断并串入下一行，「认知卸载」与「自动化偏误」两行各被拆成两至三行，「AI 劳动孤儿化」一行同样被拆开。本站依据上下文把九类相邻概念逐行合回，只恢复行列对应，未增删任何字词。读者如对某一行的归位有疑问，请以作者原始文档为准。

二、题名中的「极化」目前尚无分布证据，说明见附录二。本站未改动题名与构念名。

三、文献。正文列出 17 条参考文献（含链接），未附 DOI，本站未逐条复核；作者自陈本章为「概念与理论论文」，不把既有单项实验结果直接解释为宏观社会已经发生能力承载极化。第十二节提到的 2026 年「合成能力／synthetic competence」研究在文献表中无对应条目，建议作者补出处。

四、排版。原稿因导出产生的中文字间空格已清理；分节标题、编号列表与二维格按原稿层级还原为网页结构。

五、定位提示。本章自述为 2026 年 8 月 6 日理论建构初稿，五类研究设计均尚未执行，写死日期的赌注兑现期为 2029 年 12 月 31 日。请读者按此定位阅读。

## 第三章 自我可以被完美绕过

本章原题《主体性代偿：一个关于“自我”如何被绕过的生成机制分析》，作者 季春雷，原文见 <https://sdeuniverses.com/students/ji-chunlei/subjectivity-compensation/>

### 摘要

那个问号不是因为你缺了什么，恰恰相反，可能是因为一切都太齐全了。当代社会不是在摧毁一个已有的自我，而是通过提供零摩擦的外部代理，使那个需要亲自承担不可逆选择才能被激活的主体性，从生存的必要条件变成了一个可被完美绕过的可选配件。

关键词：社会存在论与治理研究

### ■ 一、引言：被妥善安顿的空缺

有一种不适，它不尖叫。它出现在那些生活本该最满意的时刻——深夜放下手机、屏幕暗下去的瞬间；完成了一次无可挑剔的工作汇报之后；假期旅行的第二天傍晚，风景在窗外，自己却不知道在看什么。它不是一个有宾语的问题，只是一个句法上的空位：“然后呢？”

回答这个空位的话语，当今社会极其丰裕。心理自助产业提供“找到热爱”的五步流程；积极心理学提供“心流”与“优势行动”的量表；企业文化提供“使命驱动”的愿景陈述；社交媒体提供无数他人正在热烈活着的证据。所有这些回答共享一个隐含前提：你感到空，是因为你缺了什么。缺目标，去寻；缺价值，去发现；缺联结，去建立。在“匮乏-获取”的框架内，那个问号是暂时的——它只是从“还没找到”的此岸到“最终获得”的彼岸之间的一段水面。

本章试图论证的是一种相反的可能性。这个问号，不是因为你缺了什么。恰恰相反，它可能是因为一切都太齐全了。但这“齐全”究竟齐全了什么？它齐全的不是物质——尽管物质丰裕是前提之一——而是一整套可以代理你运行“自我”的装置：一整套你不需要亲自成为“你”，就能过得很好、甚至相当满足的系统。

当代社会的核心运转机制——精密的算法推荐、最优化的教育轨道、外包一切的服务经济、被量化为 KPI 的工作评价——正在静默地完成一项比“剥夺个体意义”更根本的工程。它们不是在摧毁一个已有的“自我”，而是通过积极地提供个性化、高效率、零摩擦的外部代理



服务，使得那个需要通过亲自做出不可逆选择、承担其后果、在错误中辨认自身轮廓的“主体性”，从一个生存的必要条件，变成了一个可以被系统完美绕过的可选配件。

问题的关键，因此不在“我的判决能力被压抑了”。压抑预设了一个被压抑者的实体存在。本章的诊断要更进一步：在一个人从头到尾都被妥善安顿的生命历程中，那个通常被我们默认为天然在场的、能够做出判决并经由判决确证自身存在的“主体”，可能从未作为一项实际发生的生命功能，被激活过。这不是一种病理学意义上的缺陷。恰恰相反，它是高度成功的、由善意构筑的系统成就——正因为教育、算法、服务、管理都太成功了，成功到了在生命的大多数日常决策节点上，那个名为“我”的发生器官，从未被任何一道必须由它亲自穿越的窄门逼迫到必须出场的境地。

本章将这种状态命名为主体性代偿。<sup>[1]</sup> 它不是一个关于“虚假”“欺骗”“压迫”的诊断，而是一个关于“从未发生”的沉默考古学。既往的批判话语——海德格尔的“常人”与“非本真生存”、法兰克福学派的“文化工业”与“异化”、泰勒的“本真性退化论”——都共享了一个未经审查的预设：那个能够感受意义、也能够感受虚无的“主体”，是一个虽可被遮蔽、被压迫、被误导、被掏空，却永远在某个深层之处隐然在场的实体。依此预设，治疗的方案就是“唤醒”它、“解放”它、“重新联结”它。

本章要做的，就是把这个预设本身放到手术台上。“主体性”在此并非指一种先在的实体，而是一种需要在特定张力结构中被反复激活的生命功能。它不是在石头下等待搬开的种子；它是篝火——需要持续地、有意地、在某个必须亲自蹲在风里护着它的情境中，才会继续燃烧。但这个隐喻本身有一个困难：那个蹲着护火的“谁”，在本章的诊断中，恰好是那个需要被发生出来的。更准确地说，不是“谁在护火”，而是“护火”这个动作的发生本身，才让那个“谁”第一次作为某个可以被指称为“我”的东西，出现在世界之中。

为此，本章必须完成三个核心任务。第一，在充分承认既有批判理论真实贡献的前提下，严格论证“主体性代偿”不是一个可以被“异化”“治理术”或“功绩社会”等既有概念组合替换的新词，而是切开了它们够不到的辨别面——一个关于“主体发生条件本身如何被取消”的结构性空白。第二，不以“系统夺走了什么”为分析语法，转而追溯在标准化、算法化与体验商品化的三重社会性张力中，那个被称为“亲自判决”的动作，是如何在具体行动者解决眼前问题的过程中，被一步步结算为外部代理的。第三，将实践方向从浪漫主义的“守护本真”，扭转为冷峻地辨认代偿系统自身的不自治处——那些窄门不得不重新出现的缝隙。这些缝隙不是任何人可以主动选择去利用的“处方”。它们是条件：如果主体性要重新发生，它只能在这些代理系统无法彻底消化的剩余处寻找可能。



## ■ 二、与先行者的辨别：这个命名凭什么不可还原？

在正式界定“主体性代偿”之前，本章必须诚实地划定与既有思想传统的关系。这不只是学术礼仪，更是论证的核心环节：如果一个概念可以被几个既有概念的组合无损替换，那么它就只是一个修辞上的重命名，而非一个真正切开新辨别面的理论工具。本章将论证，本章最深的分歧，恰恰不在与经典的距离上——海德格尔的“常人”、法兰克福学派的“文化工业”、泰勒的“本真性退化论”确实为本章提供了不可或缺的问题意识——而在与几位占住了“代偿”所欲占据之位置的当代理论家之间。他们是我真正的对手。我必须证明，他们手中的概念工具，在某个要害处漏掉了只有“主体性代偿”才能照见的东西。

### ■ 2.1 海德格尔的“常人”与本章的“代偿”：此在是否总已听见呼唤？

本章最深的存在论债务归于海德格尔。在《存在与时间》中，海德格尔对“常人”（das Man）的分析，是哲学史上关于“个体如何从外部规定性中丧失自身”最精深的论述。此在在日常中沉沦于常人状态，一切判断都是“人们说”的，一切选择都是“人们这么选”的。经由“向死而生”的先行决断，此在得以从常人中抽身，第一次成为本真的自己。

本章所涉及的“自我判决”的发生，其生存论内核，海德格尔确实已经说透了。本章无法超越、也不试图超越他。

但本章的辨别面在于一个海德格尔未曾充分面对的历史条件。在海德格尔的分析中，常人状态是非本真的生存样态，它引诱此在逃避其最本己的可能性。然而，此在总是已经“被抛”入死亡的可能性之中，良知的呼唤总是在场，即便此在逃入常人，也总是在“逃离”某种它知道却不敢面对的东西。也就是说，那个能够听见呼唤、能够做出决断的主体性可能性，是生存论上不可剥夺的。

本章的诊断则翻转了这一结构。在当代社会系统的特定运作下，问题不再是“此在逃入常人”，而是“此在从未被抛入那个需要它成为本真此在才过得去的境遇”。不是良知的呼唤被日常喧哗淹没了；而是那个呼唤所依赖的、必须由个体在不确定性中独自站立的情境，在生命的绝大多数时刻，从未被制造出来。常人状态不是一种需要逃离的引诱，而是一种被升级到了无缝程度的、无需逃离的常态操作系统。它引诱的，不是你的逃避，而是你从未发现“还有另一种活法”这件事本身。

因此，本章的工作不是在海德格尔的“决断”旁边再放一个同义词，而是追问：在一个系统性地消除了“不可决断情境”的社会里，“决断发生了吗？”这个问题本身，已经被置换为一个更棘手的生成论问题——“那种让本真决断成为不必要的存在状态，是如何被社会性地生产出来的？”这是海德格尔的存在论分析所没有触及的层面。

## ■ 2.2 与法兰克福学派：是真实被遮蔽，还是亲自行使的结构被替换？

霍克海默与阿多诺在《启蒙辩证法》中提出的“文化工业”概念，是本章另一个不可或缺的思想前身。文化工业以标准化的、可大量复制的文化产品，麻痹个体的批判意识，使得大众沉浸在虚假满足中，丧失了对社会的否定性思考。这个批判捕捉到了当代意义困境的一个重要侧面。

然而，“文化工业”的批判逻辑仍然预设了一个可被麻痹、可被欺骗的“主体”。他们的诊断是：主体被虚假意识所蒙蔽。治疗的方向因而是：批判、启蒙、唤醒。本章的辨别面在于：“主体性代偿”所描述的状态，并不是主体“被欺骗”了。在代偿的体验中，个体所体验到的满足感、投入感、舒展感，在许多时刻是真实的、当下的。他不是在消费一件文化工业炮制的、令人生厌的复制品；他是在参与一项精心设计、高度个性化、完完全全贴合他当日心情的体验。问题不在体验的真假，而在体验的发生结构——他体验到的所有“自我表达”与“自我实现”，其核心判决——决定冒着什么风险、放弃什么安全选项、承担什么不可预料的后果——都是由体验设计者、算法、课程导师替他完成的。

马尔库塞在《单向度的人》中提出的“压抑性去升华”概念，比阿多诺更接近本章的问题域。马尔库塞指出，发达工业社会通过提供“虚假满足”，使得个体在舒适中被整合进现有秩序，丧失了对另类可能性的想象。这个分析比文化工业更精细地捕捉到了“满足本身成为控制手段”的机制。然而，马尔库塞的批判仍然停留在“真实需求被虚假满足替代”的层面——也就是说，他仍然预设了一个“真实需求”的参照系，一个可以被扭曲但可被理论证明为“真正”的主体性内核。本章的诊断要更极端一步：在代偿状态下，个体所缺失的，甚至不是“真实需求”，而是那个能够在诸多需求中做出辨别、排序、取舍的发生器官本身。他所感到的不适，不是因为他的“真实需求”被“虚假满足”蒙蔽了，而是因为他从未被放到一个必须亲自在“真实”和“虚假”之间做出裁决的位置上。因此，法兰克福学派的批判在当代常常陷入一种尴尬：受批判者不觉得自己被欺骗了，他确实每一次体验中感到充实。本章的概念则能够揭示：这种充实本身，就是问题的一部分。

## ■ 2.3 泰勒的“本真性”与本章的“代偿”：从文化后果到发生机制

查尔斯·泰勒在《本真性的伦理》中指出了一个关键过程：作为浪漫主义和存在主义核心理想的“本真性”——即每个人忠于自己内在的独特性——在当代被扭曲为一种软性的、自恋的相对主义，成为消费文化中一项可供购买的自我认同。泰勒以思想史的方式勾勒了本真性从道德理想到消费模式的退化路径。

本章的工作，可被视为泰勒分析的进一步操作化。泰勒说出了“本真性被扭曲了”这个事实。本章追问的则是：那个使本真性从“自我雕刻”蜕变为“自我购买”的具体社会机制是什么？我的回答是：判决的代偿。当一个人可以买到一件被贴上“表达自我”标签的商品、获取一份被设计好的“冒险体验”、完成一次由他人铺好轨道的“自我探索工作坊”时，他所购买的并不是本真性本身，而是一种在体验形式上与“自我实现”极度相似、却在发生结构上摘除了亲自判决环节的东西。泰勒的“本真性退化论”重在描述文化后果，本章的“主体性代偿”则重在社会发生机制的解剖。

## ■ 2.4 韩炳哲的“功绩社会”：当自我剥削变成自我代理

到这里，我必须正面处理那个与本章站位最接近、因而也最危险的当代对手。韩炳哲在《倦怠社会》《透明社会》《精神政治学》等著作中，提出了一套在当代批判话语中极有影响力的命题：二十一世纪的权力技术不再是压制性的，而是肯定性的；它不剥夺，而是过度供给；功绩主体不是被外部权力压迫，而是自我剥削——他在“能够”的无限扩张中耗尽自己，却误以为这是自由。韩炳哲的“他者的消失”甚至已经触及了本章的核心问题域：在一个被平滑化、被消除否定性的环境中，主体不再遭遇那个能够逼出真正决断的“他者”。

“主体性代偿”是否只是韩炳哲“功绩社会”论的一个变体？我必须在这一点上做出最严格的辨别。

韩炳哲的分析在现象层面与本章高度重叠，但他的理论结构有一个关键性的盲区。在他的框架中，“功绩主体”仍然是一个在行动的主体——他自我剥削、自我鞭策、自我优化，只是这个行动的指向被肯定性权力扭曲了，变成了对自身精力的无限消耗。换句话说，韩炳哲预设的那个“自我”，仍然在每一个KPI、每一个健身目标、每一个自我实现方案中亲自出着力。他的主体是一个用力过猛、烧尽自己的跑者，而不是一个上了自动驾驶的乘客。

“主体性代偿”指向的，恰恰是比用力过猛更隐秘的一种状态：在很多关键生命决策的节点上，那个“我”根本没有出过力。不是“我”在自我剥削、在给自己定下过高的目标；而是——什么才是一个值得为之剥削自我的目标、什么程度的代价是合理的、什么方向是不可放弃的——这些本应由“我”在不确定性中亲自掂量、亲自判决的前提性问题，已经被外部的代理系统（教育轨道、职业路径、生活方式模板、算法推荐）“省去”了。韩炳哲的功绩主体在“跑”这件事上是自主的（即使这个自主是被系统诱导的），但在“往哪跑”“为什么要跑”这两个更基础的问题上，却未必行使过独立的判决。

这就是“主体性代偿”切开的第一条辨别线：韩炳哲的“功绩社会”分析的是行动层面的自我耗尽，而“主体性代偿”分析的是判决层面的自我缺席。前者是跑得太累，后者是方向盘

不在自己手里——即使脚下的踏板是自己在踩。一个在功绩社会中自我剥削的人，仍然可能在生命最根本的意义向度上，未曾真正做过一次由他自己承担全责的生存性判决。

判别格：假设一个人从名校到名企一路在“自我优化”的道路上疾驰，韩炳哲会预测他在过度劳累中经历倦怠和存在性焦虑，本章则预测：即便他的劳累程度可控、生活平衡良好，只要他的每一个“优化方向”都来自外部模板（而不是他本人在不确定的选项中亲自判出的取舍），他在三十岁或四十岁的某个节点仍会遭遇一种“为什么这一切是我选的？”的空缺感——而这一空缺感，恰恰不是“过度劳累”能够解释的，因为它同样会在那些事业顺利、生活舒适、从未经历过“倦怠”的人身上出现。

## ■ 2.5 福柯的“治理术”：权力引导行为，还是取消了被引导者？

福柯晚期的“治理术”研究，是另一个与本章高度邻近的理论对手。治理术的核心洞见是：现代权力运作的方式，不是禁止、压抑或剥夺，而是通过引导行为、制造“自由”与“自我技术”来塑造主体。新自由主义治理术尤其精于此道——它不告诉你该做什么，而是为你提供一个“自由选择”的环境，让你在为自己的人力资本负责的过程中，把自己构造成一个“自我企业家”。在这个过程中，个体生命的每一个节点都被一种“使人活”的权力所穿透和引导。那么，“主体性代偿”所描述的“代理”，不正是治理术在日常生活层面的微观实现吗？

我的回答是：治理术的分析，在权力如何“引导”主体这一点上极其锋利，但它漏掉了一个更根本的维度——治理术始终预设了一个可以被引导、需要被引导的“被引导者”。无论权力如何精巧地通过“自由”来运作，它展开的场域总是一个已经有主体在其中的场域。治理术解释的是：主体是如何在“自由选择”的幻象中，被引向某些特定方向而非另一些方向的。但它没有解释——或者说，它的解释工具够不到——的是：当引导系统运转得太好，好到引导本身变成了代理，那个“被引导者”本身的生成机制，是否已经在这套无缝运转中被架空了？

换句话说，治理术分析的是对主体的治理，而主体性代偿分析的是主体发生条件的取消。这是两个不同层面的问题。福柯本人晚期对“自我技术”的关注，恰恰暗含着一个前提：主体仍然在某个层面上“作用于自己”。即使这个“作用于自己”是被权力技术所规训、所引导的，它仍然是一种自我的活动。而“主体性代偿”所指向的，是一种比“被引导的自我技术”更彻底的状态：在很多生命的关键节点上，那个需要“作用于自己”的“自己”，根本没有被制造出来。系统把那个本该需要他亲自做出不可逆判断的情境给“优化”掉了，而不仅仅是用权力引导了他的判断方向。

但这里必须面对一个更尖锐的福柯式反驳。在福柯对古希腊罗马“自我技术”的研究中——从斯多葛学派的书写实践到基督教的忏悔仪式——他论证的是一个比“治理术”更根本的命

题：主体不是在权力中被制造出来的纯粹被动产物，而是通过作用于自身的实践把自己构造成伦理主体的。一个福柯主义者可以这样反驳本章：当代的“体验市场”——禅修营、心理工作坊、自我探索课程——不是取消了自己技术，而是提供了一套新的、被商品化的自我技术。那个去禅修营的人，他仍然在对自己做工作——他报名、他坐禅、他投入情感——只是这项自我技术，恰好是用消费行为来完成的。在这个意义上，你的“代偿”不是对治理术的否定，而恰好是治理术在消费社会的一种新形态：一种“被代行的自我技术”。

这个反驳必须被正面回应，而不是绕过去。本章的回应是：当代禅修营中的自我工作，与福柯所分析的古希腊或基督教的自我修行之间，有一个发生结构上的本质区别。区别不在于是否有消费行为，也不在于是否“亲自执行”。区别在于：在斯多葛学派的书写或基督教的忏悔中，修行形式的每一个环节——什么时候禁食、什么时候祷告、如何审视良心——虽然在外有教规的引导，但修行者在每一个环节上都亲身承担了“我可以不做，或者做得很糟”的可能性，且这个失败的后果直接由他的此生或灵魂来承担。他的“亲自”不只是执行层面的——他不仅仅是照着规矩做——而是在恐惧、在孤绝中，亲自决定要承受这整个修行体系的全部要求，包括承受失败后的惩罚。他没有一个可以撤销的按钮。

而在当代禅修营中，那个“失败的可能性”和“失败的代价”都被课程设计优化掉了。你并不会因为坐禅不精进而被任何真实的存在性后果所触碰——你不会在这三天结束后受到任何形式的审判，你不会被逐出社群，你不会因此被判定为一个“不合格的人”。禅修营的自我工作，恰恰是一个没有“失败”的自我工作。而“失败的可能性”——那个你可以亲自搞砸、亲自承担搞砸之后果的生存论空间——是自我技术之所以能够构造伦理主体的结构条件。当失败被取消，自我技术就降格为自我体验。你在体验一种“修行”的感觉，但你不必承担修行失败的代价。而这种无代价的体验，在发生结构上，恰恰不是让你的主体性被“训练”，而是让你的主体性被“悬置”——因为它从未被逼到那个必须自己承担“我可能错了”的窄门之前。

判别格：考虑一个被算法推荐精准引导的个体。福柯学派会预测，他在“自由选择”中实际上被引向了某些特定消费或生活方式，但这种引导是通过他的“同意”来运作的——他在每一个确认按钮上亲自点了“下单”。本章则预测：如果系统足够精密、足够无缝，他在点下那些确认按钮的时候，他内部的“同意”本身——那个需要在不确定性中自行掂量、取舍、承担后悔风险的“我决定”——已经在结构上被降格为对一个外部最优计算的批准。福柯能看到权力在“同意”的表象下运作，本章能看到的是：在某些极端情况下，“同意”本身已经不再是一个主体的动作，而是一个最优化的系统通过个体手指完成的自我执行。那个被治理术预设为“被治理者”的生存论位置，在这里是空着的。

## ■ 2.6 罗萨的“共鸣”：对“能够回应”之前提的发问



哈特穆特·罗萨的“共鸣”理论，是当代社会诊断中另一个不可绕过的声音。罗萨指出，现代社会的加速与增长逻辑，导致主体与世界之间建立起一种“冷漠的关系”，主体不再能够被世界“触及”，也无法对世界做出真正的回应。他的“异化”重新定义——不是被压迫，而是无法对世界做出回应——与本章的“代偿”共享同一个问题意识：那个让生命感到“有分量”的东西，正在从现代个体的生存结构中流失。

本章的辨别面在于：罗萨的诊断仍然预设了一个有能力回应的主体，只是这个主体被加速社会的节奏和制度安排“剥夺”了回应的机会。他的解决方案因此指向“创造共鸣的条件”——减速、倾听、重建与世界的关系性连接。但“主体性代偿”提出的挑战是：当一个人在其生命历程中，从未被逼到过那个“必须亲自回应”的关口时，他丢失的不仅是“回应的机会”，更是回应的能力本身——不是技能意义上的能力，而是存在论意义上的“亲自卷入的能力”。那个能够在被世界触及时感到共振的“感官”，不是在每一个个体内部天然就绪的，它需要在一次次亲自承担的微小不确定中被反复激活和校准。当代理系统在生命早期就把这些不确定性拆卸干净时，一个长期在无摩擦环境中成长的人，即便有一天被给予了“共鸣的条件”，他可能仍然感觉不到任何东西。不是因为世界没有在呼唤，而是因为他内部那个接收呼唤的“发生器”，在尚未被需要的时候就已停止了发育。

这是罗萨的框架无法容纳的一个边界情形：他的理论可以解释“有条件共鸣却未能共鸣”的遗憾，但无法解释“从未获得共鸣能力”的空缺。而“主体性代偿”恰恰是针对后一种情形提出的。

## ■ 2.7 本章的辨别面：不可代偿的卷入

我在这一章中，把本章的对手从经典（海德格尔、法兰克福学派、泰勒）推到了当代（韩炳哲、福柯、罗萨），并逐一切割出了他们各自够不到的辨别面。这些辨别面可以凝缩为一条核心的分离线——

- 韩炳哲看到的是自我在行动中耗尽（主体还在跑，只是跑得太累）；
- 福柯看到的是权力在自由中引导（主体还在选，只是被引导了方向）；
- 罗萨看到的是主体在加速中失去回应机会（主体还在听，只是世界太吵）；
- 本章看到的是：在某些极限情形下，那个负责跑、负责选、负责听的“主体”本身，其发生条件已被系统性地取消了。不是行动太累，不是选择被引导，不是倾听被干扰——而是在那个需要亲自卷入不确定性的生存性关口，“亲自”这件事从未被需要过。

这就是“主体性代偿”的不可还原性所在：它切入的是主体发生条件的取消，而韩炳哲、福柯、罗萨切入的分别是主体行动方式、主体被治理方式、主体回应条件的扭曲。前者追问的

是“发生没有”，后三者追问的是“发生之后怎么了”。两者不在同一个分析层面上工作。如果用医学类比：韩炳哲诊断的是肌肉过度运动后的劳损，福柯诊断的是神经系统被暗中操控，罗萨诊断的是感官与外界的连接障碍；而本章诊断的，是在某些极端情况下，那个负责运动、负责感受、负责被操控的器官本身，从未被发育出来。

### ■ 三、“主体性代偿”：一个生成论概念的界定与彰显

经过与先行者和最近邻的严肃对质，我现在可以正面阐发本章的核心概念。我将首先给出其系统性界定，然后通过对一个微型案例的生成论解剖，在经验的层面上使之可感，最后通过一个更严酷的测试案例，检验这个概念的辨别力边界。

#### ■ 3.1 “主体性代偿”的生成论界定

“主体性代偿”——这个本章创用的概念，并非对主体被压抑或被欺骗的又一种描述，而是一个严格的生成论概念：它命名的是在安全、效率与可计算性导向的社会系统运转中，那个需要通过亲自做出不可逆选择、承担其后果、并在这一不确定的卷入中反复确认自身存在的“主体性”，如何从一种个体生存必须与之耦合的生命功能，被结算为一种可由外部供应绕过的可选状态。

这个概念不是在说三件并列的事。它是在说一件事——存在论层面上，有一种特定类型的“被绕过”——从三个互相交织的角度看。

从第一个角度看：主体性不是实体，是功能。它不是一棵压在石头下的种子，等待我们搬开石头。它是篝火——需要持续地、有意地、在某个必须亲自蹲在风里护着它的情境中，才会继续燃烧。但这个隐喻必须立刻接受一个自我批判：那个蹲着护火的“谁”，在本章的诊断中，恰好是那个需要被发生出来的。更准确地说，不是“谁在护火”，而是“护火”这个动作的发生本身，才让那个“谁”第一次作为某个可以被指称为“我”的东西，出现在世界之中。主体性的存在，依赖于它在“亲自做”这个动作中被反复激活。没有被激活，它就不是“潜伏着”，而是“没有”。

从第二个角度看：这一功能的激活，依赖于特定的张力结构。这个张力结构就是“窄门”——信息不完备、后果不可完全预知、没有外部权威可以替你承担最终责任、你只能自己闭上眼睛决定往哪条路上跳的那个情境。不是所有的“选择”都是窄门。一个纯粹技术性的选择（走哪条路不堵车），一个在一个既定框架内做最优化的决策（选哪个理财产品），都不构成窄门。窄门是那种你的“选择”本身会不可逆地定义你是谁、且没有任何算法或专家能够替你走完最后一步的存在性关口。

从第三个角度看：代理系统所做的事情，不是取代你的主体性，而是取消窄门。当一个社会系统发展到足够精密的程度，它在很多日常决策节点上，不是给了你一个错误的答案，而是把那个需要你亲自在不确定性中做出判决的情境本身，从你的生命里拆卸掉了。它用标准答案取代了自我判断的必要性，用最优推荐取代了摸索偏好的过程，用安全的体验设计取代了真实卷入的代价。其结果，不是一个“假我”替代了“真我”，而是那个必须在窄门中才能被激活的“我”的发生器官，因为窄门的消失而从未被使用过。

这里必须做一个关键的概念区分。本章所使用的“判决”，严格局限于其构成性含义——它是主体性发生的存在论条件，是那个窄门中不可代偿的生存性动作。它不是一项可以被训练、被运用、被“在生活中多做几次”的能力（那是“自行判断能力”，一种经验上的能力）。本章的核心论证只依赖判决的构成性含义。这个区分生死攸关：如果混淆了这两种含义，整篇文章就会被误读为一篇“如何锻炼判断力”的实用指南——而这恰恰是本章要避免的姿势——把一项构成性条件误读为可操作的技能。

至此，有必要将这一概念与若干邻近概念做严格的区分。本章的辨别不是分类学操作——不是在纸上画几个格子然后宣布它们“不同”。每一个辨别都必须指向一个旧概念无法容纳的具体情形。

与“异化”的区别。马克思的“异化”指的是劳动者与自己劳动产品、劳动过程、类本质的疏离。异化预设了一个内在本质（人的“类本质”），它被外在的、不属于他的劳动所夺走，这种疏离带来痛苦。主体性代偿则不同：它不是“我做的事不属于我”的痛苦（异化感往往是撕裂的），而是一种更安静的状态——我从未被放到那个必须由我来决定‘什么是我的事’的位置上。注意，这不是在描述“异化得更浅”。一个从未经历过“窄门”的人，他恰恰可能连感到异化的能力都没有。因为“异化感”本身就是以一个已经形成的、有自主判断力的“自我”为前提的。当那个自我的发生条件被系统性地绕过时，他所能感受到的，不是痛苦，而是某种更难以名状的“空”。这就是为什么“异化”解释不了本章所描述的那类人群——那些在外部指标上事事成功、却在存在上感到“这难道就是全部？”的个体。他们并不痛苦，他们只是找不到“痛苦的主体”在哪里。

与弗洛姆“逃避自由”的区别。弗洛姆指出现代人为了逃避选择的焦虑而交出自由，这是一种主动的、以换取安全感为目的的心理过程。主体性代偿则不同：它不是在描述个体主动逃避自由，而是在描述一种生存论上的被动状态——那个个体从未被放到一个必须行使或逃避自由的位置上。在弗洛姆的框架里，逃避自由本身就是一个主体性的行为（即使它向了主体的削弱）；在本章的框架里，那个行为可能从一开始就没有被需要过。



与韩炳哲“自我剥削”的区别（已在 2.4 节详述，此处只概括核心辨别面）：自我剥削预设一个在“亲自跑”的主体，主体性代偿指向的是那个“自己决定往哪跑”的前提性判决，在某些生命历程中从未发生过。前者是用力过猛，后者是方向盘不在自己脚下——即使脚下踏板是自己踩的。

与“虚假意识”的区别。虚假意识是认知层面的错误——你没看到真正的现实，被意识形态蒙蔽了。主体性代偿不是认知错误。那个被算法推荐、体验课程、标准化教育轨道妥善代理了生命决策的个体，在很多层面上是对的——他的选择确实是当前的“最优解”，他的体验确实带来了当下的满足。他不是“看错了”什么；他是从未被放进一个必须“自己看”的环境里。他的眼睛被保护得太好，从未迎过风。

与“享乐主义”的区别。享乐主义是一种主动选择的、以快乐为目的的生活哲学。它需要行为者主动追求、并自行承担追求过程中的取舍。而代偿状态下的个体，即使他过着追求快乐的生活，他的“追求”本身——该追求哪一种快乐？追到什么程度？为此可以放弃什么？——往往也已经被大量的“如何过好这一生”的术略建议、产品方案、社群风向所代理。他可以是一个享乐者，但他的享乐并不必然出自他自己的判决。

### ■ 3.2 一个现象学尺度的彰显：皮具体验课的生成论解剖

为了让这个概念在经验的层面更可把握，我在此引入一个微型案例的生成论解剖。考虑一个年轻人，在普通的一个周六傍晚，感到了一阵模糊的、无对象的低落。他打开手机，一个生活方式 APP 根据他的浏览记录、心情标签、同类用户选择，推送了一个“治愈系手工皮具体验课，零基础可学，当日可完成一件独属于你的作品”。他点击、付款、前往。在老师温和的指导下，他一步步完成了切割、打孔、缝线。他闻到了皮子的味道，摸到了线的粗糙，完成时看着那个略显歪扭但确实是“自己的”卡包，感到一阵真实的、温暖的成就感。他拍照分享，收到很多赞。晚上回家，熄了灯，那个模糊的低落还在——甚至比之前更清晰了一点。

现在，我把这个看似平常的场景撬开，不是把它当作一个人生片段的描述，而是把它作为一个生成论的样本，分析它的存在结构。

这个年轻人在那个皮具工作坊里体验到的东西，与“自我实现”在体验形式上高度相似：他投入了注意力，他经历了从不会到会的微小成长，他创造了一个不同于别人的、独一无二的物品，他在情感上感到了一种真实的满足。然而，如果我们问一个苛刻的生成论问题——这个体验的核心判决是什么？

“我决定，冒着做得很难看、浪费一个下午的可能，去尝试一件我从未做过的事。”

这个判决，在本章的意义上，是“亲自的”吗？

不是。因为这个判决已经被一整套外部装置代理完毕了。第一，那个“做什么来治愈低落”的选项，不是他在大量可能的周末活动中亲自掂量、取舍、最终选定一个并为其可能的失败承担失望的责任后做出的。它是算法根据他过往行为数据的最优计算，直接送呈的。第二，那个“要不要冒风险”的忧虑，不是他亲自克服的。课程的名称里已经写着“零基础可学，当日可完成”——在进入工作坊的第一秒之前，所有可能带来挫折感的失败可能性，已经被专业的教学设计“化解”掉了。第三，那个“这是独属于我的作品”的体验，是在他人铺好的轨道上产生的情感。它的“独特性”来自手工制作过程的自然误差（那个略显歪扭的针脚），而非来自他本人对形式、材料、审美的亲自判断——因为这些判断环节，在DIY体验课的标准化流程里，被隐性地代行了。

当他晚上躺下、看着黑暗、感到低落仍在时，他所经历的，不是一个简单的心情反弹。他所遭遇到的，是一个生成论上的空缺：他体验到了判决的果实（成就感、创造力、独特性），但他没有咽下判决的种子——那个决定承担什么、放弃什么、在不确定中亲自掂量的动作。他“做”了，他“体验”了，但他的“做”和“体验”发生在一个所有实质风险都已被拆除的、被代理完成的模拟现场里。那个需要在真正取舍中才能被激活的“我”，在这个过程中始终沉寂着。他所感到的空，不是因为他做了什么错事，而是因为那个唯一能把一堆“体验”汇成一条“我活过”的整全河流的生存性动作——亲自判决——在这个精密安排的满足里，全程没有被邀请出场。

这个思想实验，不是对个体选择的批判，也不是对体验课程的否定。它的作用，是作为一个现象学的手术刀，切开一个在当代生活中广泛存在却很少被仔细检查的结构：一种将“我亲自做”这个动作的生存论内核摘除、却保留其全部体验外观的代理模式。那个年轻人不是被欺骗了——他在两小时内确实感到了充实。但他感到充实的那个“他”，是一个被降低到执行与感受层面的自我，而不是一个在判决中确认自身存在的自我。这两个“自我”之间的裂隙，就是主体性代偿的生成论部位。

但皮具课案有一个内置的局限：它是一个容易被解剖出结构问题的案例。一个被消费的、被设计的、被“零基础可学”的安全体验——谁都知道它和“真正的自我探索”之间有一段距离。如果“主体性代偿”只能切割明显的代理，而不能切开那些高度逼真于“本真”的代理，它的辨别力就是有限的。

因此，本章必须接受一个更严酷的测试。考虑这样一个案例：一个人放弃了高薪工作，跑到云南开客栈，非常认真地经营、非常投入地生活、非常真诚地感到“这是我要的人生”。这个案例如果拿给通常的自我实现叙事看，会被诊断为“找到了自我”。如果拿给韩炳哲看，会被诊断为“功绩主体在休闲态”。本章的框架能对这个案例说出什么？

我能说的是：我们无法从外部判断这个人是否“亲自判决”了。本章的诊断权力，不在“我比他自己更懂他的体验”，而在“他自己用同一张嘴说出了两件事”之间的裂隙——他会在某些时刻真诚地说“这是我做过的最好的决定”，又会在另一些时刻，对着丽江的星空感到一种“为什么还是不够？”的空。本章的工作不是判定他是真是假，而是追问：他的这个“放弃一切去追梦”的决定本身，在多大程度上是“他自己”在掂量过、承受着后悔的风险之后做出的？还是在多大程度上是“逃离北上广去大理”这个已经被社会命名为“冒险”的标准模板的内部执行？他有没有在某个深夜，独自面对过“如果客栈亏了，我还剩什么”这个问题，并且没有一个现成的答案——没有一篇公众号文章、没有一个创业导师、没有一个同类社群能够替他回答？

这个追问的意义不在于“找到答案”，而在于辨认出：在那些高度逼真于本真的代理中，窄门的取消往往更隐蔽、更难以指认——因为它不再穿着“标准化”的外衣，而是穿着“自我实现”的外衣。一个被“去大理开客栈”这个社会模板所代理的人，他可能会比那个上完皮具课感到空的人更难接触到自己的空——因为他的生活“看起来”太像是他自己选的了。而本章的概念，其最严酷的测试恰在于此：能否切开那些“看起来最不像代理”的代理。

## ■ 四、代偿的三重发生：判决是如何被拆除的？

前一章的案例解剖，展示的是代理在单个微观体验中的运作。但“主体性代偿”不是一个孤立事件。它是一种社会性的发生条件变更——当皮具课式的代理结构遍布了一个人从幼年到成年的所有决策域（学什么、玩什么、与谁交往、做什么工作、如何度过闲暇）时，它就不再是“体验市场的一种产品形态”，而变成了“整个生命历程中不可代偿的判决被系统性地排挤”。一个孤立的体验代理事件，如何在一代人的生命历程中反复发生、叠加、累积，最终从“个人的某次周末选择”变成一个“社会性的主体性生成条件变更”？这是本章要回答的问题。

本章的任务，不是按历史阶段叙述“系统如何夺走了什么”。我需要做的是追溯：在标准化与判断的张力中、在算法效率与自我摸索的矛盾中、在体验供给与意义渴求的博弈里，那个被称为“亲自判决”的动作，是如何在具体行动者（管理者、工程师、教师、平台设计师、消费者）解决眼前问题的过程中，从生命历程的必经环节，被逐步结算为一种可由外部供应替代的服务的。

我将在每一节中解剖一种张力结构，然后借助一个具体的微型案例，在经验的层面上展示：判决的哪些环节是如何被摘除的。

### ■ 4.1 第一重张力：标准化的效率与自我判断的淘汰

一个生存性判决的发生，至少包含四个结构环节：①识别出这是一个需要“我”而非“规则”来裁决的处境；②在信息不完备下掂量几个不可通约的选项；③做出一个不可逆的、将自我卷入其中的选择；④承受后果并在其中修改自身。标准化系统所做的，恰恰是在第一和第二环节上动了手术。

第一个环节被消灭的方式，是把一切需要裁决的处境重新编码为“有标准解的题”。当一份标准化阅读理解测试给出一篇短文，问“作者在这段中主要表达了什么？”，并设 ABCD 四个选项时，这个设计将一种本应包含多重解读可能的文本交互，转化为一个“有没有正确理解预设答案”的问题。最要害的一点在于：对文本的“个人的理解”，在评分标准中是一个量化归零的动作。一个学生如果在考卷上写下自己真实的、与标准答案不同但言之有据的理解，他得到的分数与他根本不读文本随机乱选完全一样——都是零。而在课堂上，当老师说出“谁来谈谈你对这段的理解？”时，一个经验丰富的学生已经学会去猜“老师想听的那个答案”是什么，而不是去启动那个“我自己读下来觉得……”的内在声音。

这不是教师个人的失职，而是标准化问责体系下被逼出的生存理性。当学校的排名、教师的考评、学生的升学都被系于标准测试的分数时，课堂里那个“让学生自己判断自己是否理解了”的环节——耗时久、不可测、不产出可计入 KPI 的成果——就成为系统中最先被牺牲的东西。牺牲不是一次性事件，而是年复一年、千百万个课堂里，在每一次教师选择“我们直接对一下答案”而不是“你怎么想的”的微小决策中，被逐步结算出的一个结构性空缺。

为了经验的聚焦，我们来看一个具体到几乎隐形的案例。

在一个小学五年级的语文课堂上，老师布置了一篇阅读短文《种子的力》。文末有一道开放题：“你认为，作者在最后一段为什么写‘春风是种子的帮手’？请结合课文说出你的看法。”

一个学生在下面写：“我认为作者这么写，是因为种子埋在土里很孤独，春风来的时候，种子觉得终于有人理它了，所以是帮手。”

这个回答很真。它基于一个孩子对“孤独”和“被理”的切身感受，将文本中的“春风”与种子的出土过程，代入了一种他本人能够体验到的情感逻辑。然而，在标准答案里，这道题的给分点是：“春风代表春天的到来，温暖的气候为种子的发芽提供了必要的自然条件。”任何偏离“必要自然条件”这个物理解释的答案，无论学生在感受层面多么真诚、在文本层面是否有迹可循，在这一题的评分量表都是零分。

当事后孩子拿到标答，下一次考试中遇到类似的“你的看法”题时，他已经不再启动“我自己读下来觉得……”这个内在动作了。因为他已经习得了一个教训：那个动作在系统里是无

效的。他精熟地学会了去搜索课文中可以被理解为“标准答案所需要点”的句子，然后换一套措辞把它包装成“我的看法”。

这就是标准化拆解判决的方式：它不断在课堂上用分数强化一个信号——“你的判断不重要，你能不能找到那个被预先规定的判断，才重要。”这个信号不是贴在墙上的标语，而是体现在每一次被扣掉的分数里，每一次被表扬的“答对了！”里。一个从六岁到十八岁，所有被鼓励的思考都是“找到标准答案”的人，当他在成年后被抛入一个没有标准答案的存在困境时——比如，要不要辞掉一份稳定但不喜欢的工作、要不要结束一段安全但缺乏亲密的关系——他的内部用来做这道题的器官，没有训练过。这个器官没有坏，它根本就没长。

这第一重张力的结算结果，不是“系统夺走了自我判断的发生权”，而是：在追求“公平且可大规模操作”的教学评价这一善意目标的过程中，那个需要被学生亲自启动、亲自纠正、亲自负责的内在判断环节，在课堂的日常运转中，被逐步挤出。它没有被任何人故意“夺走”，它是在无数个具体行动者解决“如何教好一屋子学生”“如何公平地给分”这些实际问题时，被作为一种不可量化的、因而不可持续的“成本”，从教育过程中自然地淘汰了。

## 4.2 第二重张力：算法效率与自我摸索的消融

如果说第一重张力发生在制度化的教育空间里，那么第二重张力，则渗透进了日常生活中最微小的毛细孔——你吃什么、看什么、买什么、与谁约会。它不再给你一个标准答案，而是给你一个“最适合你的”答案。它的生成论语法是：“你不用在不确定性中冒险选择，我替你算好，你确认最优解就行。”

我需要先解剖一下“摸索”这个动作的生成论结构。一个人形成其独特的品味——不是消费品味，而是广义上的、知道自己喜欢什么、知道什么对自己重要、知道自己为什么欣赏 A 而厌恶 B——这个过程，不是一个信息接收的过程。它需要一个特定的张力结构：信息不完备、后果不确定、没有外部权威替你负责，你只能自己在一次次的尝试、后悔、再尝试中，逐渐凝出你自己的偏好轮廓。这正是传统消费决策中尚存的东西：你在唱片店里一张张翻找，试听三十秒，被一张封面吸引，买回家却发现整张专辑只有那一首好听。你后悔，但你记住了这个后悔。下次你再看到那个厂牌、那个制作人、那个风格的封面时，你会掂量。这个掂量，就是你在学“我是谁”。

当代平台的推荐系统，以极其精妙的算法，将这一结构悄悄地改写了。它不再提供一个可游荡的菜单，它直接告诉你：“这个最适合你。”无数数据的海床被折叠成一个丝滑的用户动作：确认。

在这里，判决的生存性内核——在不确定性中亲自掂量、取舍、承担可能后悔的代价——被降格为对一个已完成的外部最优计算的批准。个体不需要摸索自己的品味，因为算法已经替他摸索好了；个体不需要承受选错的后悔，因为算法给出的“最优”总是对的（或者至少，总是最不坏的）。

我在这里给出一个经验性的对比——一张打口碟与一个推荐歌单——来让这个机制在感知层面变得可碰触。

一个在二十年前用 Walkman 听音乐的人，他有一张打口碟。那是他在一家狭窄的唱片店里蹲了一下午，从一大箱塑料片里一张张翻出来的。他翻的时候手上沾了灰，试听的时候老板在催，旁边有个比他懂行的人在跟老板聊一个他完全没听过的乐队名字。他在一个混杂着焦虑、自卑、不知凭什么觉得“就是这张”的混沌里买了它。回家后他发现，它只有一首歌好听。他后来每次在电台里听到那首歌，心里都有一块很复杂的情绪：后悔、但也有一点点奇怪的亲近——因为那是他亲自拣选的后悔。

一个用流媒体听歌的年轻人，他的年度歌单是被 APP 自动生成的。歌单封面上写着“你的 2023 年度之选”。每一首都对。每一首都确实是他听了一整年的、在无数个夜晚陪伴过他的旋律。他点开歌单，听着这些歌，心里涌上的是一种温柔的、被懂得的感觉。这里没有任何后悔，因为推荐系统替他过滤掉了那些他大概率会不耐烦的前奏、他暂时还没有准备好的新风格、以及那些需要他在一场闷热的现场演出里站到腰酸才能第一次听懂的即兴段落。每一首被算法递来的歌，都是对的。但这里的“对”，是一个不容许他后悔的、因而也不容许他亲手去摸到自己的边界在哪里的“对”。

这里发生了什么？算法把他的品味从“我自己在错里慢慢认得”变成了一个随时可以调用、却从未让他做过排除性取舍的、对过去点击行为的最优拟合。他的歌单，是一台精密摄像机拍下的他自己的影子，不是他在泥土里滚过之后自己眼里认识的那个自己。他听到的每一首歌都“属于他”，但他从来没有在任何一个十字路口因为选了 A 而永远失去 B。他失去的，不是音乐的同质性——当代流媒体的曲库远比二十年前丰富——他失去的是在缺失中做排除的那个动作。他知道自己“喜欢什么”，但他不知道他喜欢的东西，是以那些他从未听过、因为他在系统给出的第一条推荐时就止步了的音乐为代价的。他的品味没有底。因为没有底，所以它总是在下一首，总是在“你可能也喜欢”。他变成一个对所有推荐都感到“可以、但都不亲”的漂流者。

这就是算法消融摸索的方式。它不发生在任何一个宏大的、可以被判读为“剥削”的决策里。它发生在一个个你甚至没有意识到“这是一个选择”的瞬间——你在屏幕上轻轻点了一下，而那个“让我自己找找看”的动作，在那个瞬间，没有被做出来。不是你不想做，而是当



你还没来得及想的那一刻，系统已经把最优解放在了你的指尖。它把“摸索”这个动作需要的摩擦空间，从生活中擦掉了。

这第二重张力所结算出的，不是系统向个体强加了一种品味，而是：在追求“提供无摩擦最优体验”的过程中，那个需要由个体亲自在错误中建立偏好的摸索过程，被逐步替代为对过去行为数据的最优拟合。个体得到的，是自己的过去。他从未在亲自做一次排除性的判断——那种将选项 A 从“我也许有一天会喜欢”永久地变为“这不是我”的判断——中，生成过自己的未来向度。

### ■ 4.3 第三重张力：体验供给与意义渴求的悖论

当前两重张力完成了它们的调适——个体既不需要自己判断，也不需要自己摸索——一个新的供给便自然生长出来，以填补那个被腾空的、名叫“意义”或“自我实现”的位置。这就是“体验市场”：定制化的心灵成长课程、冒险旅游套餐、付费冥想社群、艺术疗愈工作坊。这一领域的生成论矛盾，比前两者更复杂。因为它所提供的，恰恰是用来对抗“意义缺失”的东西。然而，在它的发生结构内部，却运转着重演那个它声称要治愈的问题的机制。

我需要回到“判决”的结构环节上，来精确地指出体验市场在哪里动了手脚。

一个生存性判决的第四个环节是：承受后果并在其中修改自身。这是最核心、最不可代偿的一步。因为“修改自身”这个动作，不是发生在你知道“我成功了”或“我失败了”的那一刻，它是发生在你对你的选择所带来的后果——无论好坏——的承受过程之中。这个过程的长度、深度与它的不可逆程度成正比。你做了一个决定，它让事情变成现在这个样子，你不能撤销它，你只能带着它继续活下去。正是在这种“带着无退路的后果继续活着”的过程中，你慢慢地被这件事本身所改变。这就是判决的果实与判决的不可逆代价之间无法拆分的关系：你咽下核，它才可能在你胃里长成一棵树。

体验市场所提供的东西，在发生结构上恰恰切断了这个“核”与“果实”之间的联结。

为了看清楚这一点，我以一个细节密实的具体案例来诊断。

考虑一个标榜“三日断联禅修营——在止语中遇见真正的自己”的活动。活动招募文案是这样写的：“你是否常常感到忙碌却空虚？你是否总觉得心里有个洞，无论多好的事填进去都漏？来参加我们精心设计的静修营吧。三天的时间，我们为你安排好了一切：静谧的禅房、专业的导师、科学的坐禅时间、素食三餐、以及每日一次与导师的一对一交流。你只需要放下手机，放下身份，把自己交给我们。在这三天里，你什么都不用决定，你唯一要做的，就是深深地与自己相处。”



读到这段文案时，那个在“代理”中长大的个体，是会被打动的。因为他感到了这里在回应他的需要。他报了名，进了山。

这三天发生了什么？

日程由导师团队安排。凌晨五点半，不是他自己意识到“我需要早起”，而是锣声替他意识到。坐禅四十五分钟，不是他觉得“我此刻需要打坐”，而是日程表替他判断了。与导师的一对一交流，不是他在某个感到卡住的时刻、鼓起勇气去敲门，而是被告知“下午三点轮到你”。就连他吃的素食，也不是他决定今天要清淡，而是膳房已经搭配好。他在这三天里做的每一个动作，都是被安排好的、被判断过的；但他体验到的，却是一种极度的“我终于为自己做了一次正确的选择”的满足。在最后一天的闭幕分享上，他泪流满面，对身边的陌生人说：“这三天，我重新找到了我自己。”

这里发生了什么？一个在发生结构上被彻底代理的事件（他三天的每一个时间块、每一个活动内容、甚至每一次“与自己相处”的时长和方式都由外部判断），在体验的层面上却被他感受为“我找回了自己”。这个落差不是虚伪，是结构性的。体验市场提供的，是一个判决的仿真剧场：它为个体制造了一个高度逼真的、在其中个体“感觉”自己在做一件关乎本真的事的环境。但在这个剧场里，那个最核心的动作——自己决定要冒什么风险、放弃什么安全选项、在哪个不确定的时刻坐不住也得坐着——是被剧场本身的“服务”给悬置掉的。他知道自己来“禅修”的，但他不知道这次禅修的具体形式是怎么来的、为什么是四十五分钟而不是十分钟、为什么是素食而不是断食——他不是这些事的设计者，他是它们的享用者。

由此可以明白，体验市场所完成的最精致的那个动作。它不是简单地卖一个假货给你。它是在你经历完了一个被完全代理的过程（那个让你“什么都不用决定”的日程安排）之后，通过你的全程投入和事后满足的泪水，让你自己把这个被代理的过程，给自己追认为一次“我做出了勇敢的自我探索”的判决。判决的果实是真的（他在这三天里确实放松了、反思了、感到了某些真实的触动），但种子的核——那团需要他来挑、他来选、他来在黑暗中独自蹲着等它发芽的、不可转让的东西——始终在外面的课程顾问手里。

这就是代偿的第三重发生。它不是一次性的欺骗，而是一个能够自我循环的悖论结构：个体用自己的参与热情和真实感受，为那个替他完成了所有判断的代理系统，做了一次情感上的合法性签署。他越是被这种体验所深深触动，他就越是确认了自己这一生都值得用这种方式来“找回自己”；而他越是依赖这种方式，那个只能在没有备课没有时间表没有被导师护着的日常泥沼里，靠他自己一寸一寸喘着气长出来的“我”，就越没有机会被逼出来。

## ■ 五、反驳与边界

前文的论证，在多个层面上会遭遇有力的挑战。这些挑战不是文章的点缀，而是帮助“主体性代偿”这一概念在与最强反对意见的搏斗中，澄清其有效范围与严格边界。本章将正面回应五个方面的反驳，其中包括一个对本章方法论正当性的根本性质疑。

## ■ 5.1 “静默的意义”与不可被判决还原的充实

第一个反驳来自日常经验中最深沉的直觉：那种经年累月的、在亲密关系与习惯性日常中不声不响充盈着的意义。一个老人在妻子病床旁照料了七年。他没有在某个戏剧化时刻“勇敢地做出决定”；照料是他的每一天，每一个小的重复，从不是被挑出来成为一个“抉择”。他的生命是有的——这个“有”，不是在某个悬崖边的判决中被铸成的，而是在每一次擦洗、每一勺喂食、每一个她握着他的手时那一点尚存的体温里，被日复一日地接住的。

这种“静默的意义”，对本章的核心论证构成了一种有力的挑战。如果意义真的可以以这样一种完全不依赖于判决的方式饱满地发生，那么本章对“不可代偿的判决”的强调，似乎就是在用一种智识的焦虑，去覆盖一片原本安静充实的生活世界。

本章必须诚实：这种意义形态是完全真实的，且本章的概念不能、也不应当以“判决”为名将其吞并。那个老人的“有”，不能还原为“上千次不可代偿的小决定的总和”。在本章的早期版本中，我曾试图将老人的守护描摹为“微小判决的连续体”，这实际上是一种概念的过度扩张——一种将任何非判决性的、关系性的、被动接受的意义形态，都强行纳入“判决”框架以维持理论一致性的操作。这种操作是暴力的。生成论应当保持概念的有限性，承认有些意义恰恰发生在“判决”之外。

然而，这个承认不等同于本章论点的退场。因为“静默的意义”在现代社会的发生条件本身，正是本章诊断所针对的对象。

老人的守护不是孤立发生的。它依赖于一系列不可控的、非选择的、无法被外包的外部条件：他与妻子共同生活了几十年，这不是他选的，是时间替他累积的；他照料她的那些日常，不是设计出来的“体验”，而是被命运的窄门（疾病、老年、照护）不可拒绝地抛入的。也就是说，这种“静默的意义”的土壤，恰恰是一系列无法被代理掉的、关系到生存本身的“被抛”情境所构成的。正是在这里，本章的论证发生效力：当代社会的代理装置，恰恰在将这种“被抛”的不可拒绝性，以安全、效率与舒适为名，从生命中系统地拆卸掉。

一个人如果终其一生都在选项菜单中生活，不被任何不可拒绝的窄门所抛入——不被某个非他不可的人的脆弱所依存，不被某个没有尽头的责任所绑定，不被某种“没有别的选择了只能这样熬下去”的命运所携裹——他确实可能在一种安全的、自由的生活里，从未触碰到那种“静默的意义”。他可能享受了很多“意义的替代品”，却从未经历那种被真正的负担所压

出来的、安静的充实。本章所诊断的，不是“静默的意义”不存在，而是那个让它得以发生的不可拒绝的土壤，在当代正在被系统性地替代为可选择的体验。老人的守护之所以是饱满的，恰恰因为他没有选择。而现代人的困境恰恰在于，他的生活里充满了选择，却没有一件是不能被拒的。

所以，“静默的意义”不是本章的反例，它恰恰是本章论证的一个极端确认：它确认了，有一种深度意义的发生，依赖于不可被代理的、不可撤回的关系性卷入。而本章所描述的主体性代偿，其核心恰在于那种不可撤回性的消失。老人是在不可撤回中被动地充实起来的；那个当代个体是在可撤回中主动地空虚下去的。两者不在同一种存在结构里。

## ■ 5.2 布伯的“相遇”与决定叩门的那个判决

更尖锐的挑战来自马丁·布伯的对话哲学。在布伯看来，意义不在孤独的个体判决中诞生，而是在“我与你”的至高的相遇之中，如恩典般降临——那种在刹那间，一个他者不可替代地出现在你面前，你的整个存在都被这一相遇所改变的时刻。这种相遇的发生方式，与“我做出判决”的意志紧绷完全不同。它更接近于“我被赐予”。

这里需要先澄清一个存在论前提上的根本分歧。布伯的“相遇”植根于犹太-基督教的关系存在论——恩典从永恒之你而来，个体在相遇中是被动接受者。本章的生成论则始终工作在非先验的地基上：没有任何先在的永恒之你，没有任何被保证的恩典，只有两个存在的轨道在某个不可预知的节点上不可逆地互相卷入——而这种卷入之所以发生，不是因为他们中的任何一方主动选择了它，而是因为命运的窄门逼使他们撞在了一起。在这个框架里，布伯所说的“赐予”，可以被重述为：一种不可预知的、不可设计的、超出个体控制范围的生存性遭遇。本章并不否认这种遭遇的真实性。

面对布伯的挑战，本章必须区分两个层面：体验发生的那个当下，与让这个当下成为可能的那个结构性的前提。

在体验发生的当下，布伯完全正确。意义常常以一种自我边界溶解、与另一存在者融为一体的形态到来。本章完全承认这种恩典般体验的真实性，且绝无意将其还原为判决的后果。

但在结构性的前提层面，布伯没有处理一个问题。在当代社交媒体的无限滑动、约会软件的配对算法、以及被精密化的“安全社交礼仪”共同编织的环境里，那个能让“我-你”相遇发生的土壤，是如何被改造的？人们不停地在“遇见”，却越来越不能“相遇”。因为每一次连接都同时携带着可撤回的、可比较的、可优化的后台操作。在这种环境里，那个敢于越过屏幕、真正把自己暴露给另一个人、并承担可能受伤的全部代价的决定本身，在大多数人那里，从未被做出。

这就是说，在“赐予般的相遇”发生之前，常常需要一个更早的、更沉默的判决——一个个体，决定放下手机、决定冒着被拒绝的羞耻、决定放弃那些永远在后台等待的备选可能，去真的坐在另一个人面前。这个判决不是意义本身。但它，是让意义得以发生的那个不能被绕过的门槛。在那个门槛上，你已经不是一个被动的恩典接受者。你是一个亲自叩门的人。叩门本身，就是不可代偿的判决。

而当一个人从未叩过门，他可能终其一生等在门内，却责怪外面没有恩典来敲。不是布伯错了，而是布伯没有处理那个让“我与你”可能发生的前提性窄门在当代正在消失的事实。本章所做的，不是否定相遇，而是追问：在一个代理系统日益消除“需要亲自叩门的情境”的社会里，“相遇”的发生条件是否也在被取消？

## ■ 5.3 麦金泰尔的叙事挑战

麦金泰尔式的挑战指出：意义并不总在决断的刹那，而在连续的、被社群共享的叙事中获得。一个人的生命之所以有一个整全的意义，是因为他的行动、选择与承受，都被嵌入在一个比他更大的、跨越代际的叙事传统之中。

本章承认这种意义形态的真实性与有效性。但必须指出一个关键的区分：麦金泰尔所描述的那种“叙事统一体”，它的前提是传统叙事的完整性与可继承性。当一个社会处于快速转型之中，传统叙事已经松动、破裂、甚至不再能覆盖个体所面临的全新抉择时，个体便不再是一个叙事的被动继承者，而被迫成为叙事的主动缝合者。他站在无数碎片之中——留在北上广还是回到故乡？要事业还是要孩子？要稳定还是要自由？——这些抉择，无法在传统叙事中自动找到它们的位置。它们必须首先被一个判决所切开，然后叙事的织机才能开始工作。

本章的诊断，并不适用于那些仍然生活在完整传统叙事之中、尚未被现代性彻底碎片化的人群。它适用于那些已经被个体化的选择义务所包围、却从未被给予过“如何为自己做选择”的训练的现代处境中的个体。正因如此，本章的概念不应宣称一种社会学意义上的普遍性。它是针对特定社会结构条件的一个诊断，而非一个关于“所有人”的断言。这个自我限定，是本章边界的重要一部分。

## ■ 5.4 苦难的边界：不是颂歌，是辨认代价

接下来，是一个本章必须处理的边界：本章绝对不是在提倡苦难崇拜。痛苦本身不是意义的生产者。木匠选择与那块疖疤交织一周，他没有自寻痛苦，他只是承担了他所选之路的代价。代价不等于创伤。悬崖上的跳跃不等于跳崖自杀。

在本章的早期版本中，我曾用过“点燃”“守护篝火”等隐喻。这些隐喻隐含着—个先在的“真实自我”等待被唤醒，滑向了本章自己所要批判的本真性浪漫主义。我现在必须修正：本章所呼吁的，从来不是去主动拥抱创伤，更不是去“回归”某个被掩埋的本质自我。本章的实践指向是：在一个日益将承受代价的可能从生命中抹掉的系统里，辨认出那些代价本身——那些不得不面对的取舍、不得不咽下的后悔、无法撤回的伤害——并不是生命的副作用，而是那个必须在亲自承担代价中才能被激活的主体性发生功能，所不可或缺的养分。

这一区分的关键在：不是“痛苦使人有意义”，而是“亲自承担的代价”与“亲自做出的不可逆选择”是同一个发生事件的两个面。取消代价的承担，等于取消了判决的生存论分量。而代理系统所做的，恰恰是把代价从个体身上卸载，然后把没有代价的果实（成就感、创造感、联结感）供应给个体。本章的批评对象，不是舒适本身，而是这种把果实与代价切开的技术。守护代价存在的日常缝隙，不等于歌颂苦难。它等于承认：代价，在某些时候，是你为自己活过的唯一证据。而那个证据，不能在体验市场的货架上买到。

## ■ 5.5 一个方法论的反驳：你凭什么说他没有亲自判决？

最后，本章必须正面回应一个最根本的方法论反驳。这个反驳不质疑本章的诊断对不对，它质疑本章有没有诊断的权力。

你的论文动用了“生成论解剖”去透视另一个人的内部体验结构——你说那个皮具课上感到充实的人，他的充实是“被代理的”，他的“亲自”是假的。但你是谁？你凭什么这么说？你有什么证据证明那个皮具课上的年轻人晚上躺下时感到的“空”，不是你替他感到的？你的整个论证，是否只是一场智识阶层的傲慢——用一套复杂的学术话语，去否定普通人在日常生活中真切感受到的满足？

这个反驳的危险在于：它把本章的生成论诊断，从“冷峻的真理”—举打成了“对他者体验的暴力解读”。它质疑的不是本章的内容，而是本章的正当性本身。

本章的回应如下。

本章的论证，并不建立在“我比皮具课上的年轻人更懂他的体验”这一断言之上。本章不否认他在两小时内感到的充实是真实的——本章在任何地方都没有说过“他的充实是假的”。本章的工作是：在他自己报告的两个体验之间的裂隙中，辨认出一种无法被他自己的主观报告所解释的结构性差异。他报告了“充实”，但他也同时报告了“那个低落还在，甚至更清晰了”。他自己用同一张嘴说出了两件事：“我感到了真实的成就感”，和“那种空还在”。这不是我去替他感觉的，这是他自己说的。本章所做的，不过是追问：为什么同一场体验可以同时产生“充实”和“空”这两种不矛盾的、却互相啃噬着对方的感受？

本章的回答是：因为体验的形式与体验的发生结构之间，有一个他无法用日常语言指认、但可以用生成论概念来命名的裂隙。他体验到了判决的果实（成就感、创造力），但那个果实的核——亲自判决的动作——在这场体验的发生结构里没有被包含进去。他感到的“充实”是真的，但那种充实发生在一个不要求他亲自判决的代理结构里。他感到的“空”也是真的，但那个空不是因为他做错了什么，而是因为“亲自”这件事，自始至终没有被邀请出场。

本章没有替他判断他是否真的“亲自”。本章是给他自己无力命名的那条裂隙，提供了一个名字。这个名字不是对他的体验的否定，而是对它的发生条件的揭示。这不是施洗——不是判定谁真谁假。这是在作描述性形而上学的考古工作：把那些已经在个体经验中开裂的、却尚未被语言捕获的存在结构，命名为概念。本章的分析权力，不在“我比你更懂你”，而在“你所说的，你自己也感到其中有一种说不清的不对——而我试图为那种不对，给出一个结构性的说明”。

这个回应，同时划定了本章论证的方法论边界：本章的生成论解剖，只对那些已经在自己内部感受到那种裂隙的人有效。对于那些从未感到任何“空”、从未问过“然后呢？”的人，本章的诊断不适用——不是因为他们“更浅”，而是因为他们不在本章所描述的那种存在结构里。本章不宣称一种关于“所有人”的普遍真理。它只针对一种特定的、在当代社会中日益普遍的生存状态，提供了一种探索性的概念工具。这个自我限定，是本章诚实性的一部分。

## ■ 六、结论：辨认缝隙

现在，让我们回到最初的那个深夜，那个在放下手机后袭来的、无色无味的空。

现在我们可以准确地理解它了。它不是一种可以被“找到热爱”治愈的空虚感。它是一个信号——在整套生命代理系统运行得最顺滑的地方，发出的极其轻微的摩擦音。它告诉你：在那个你以为是“你自己”的、体验丰富的、妥善安顿的身份副本的下面，那个只能由你亲自判决、亲自承担、亲自在错误中辨认路径的主体性，似乎从来没有被任何一个足够窄的门逼出来过。

这个信号不能被回答。它只能被辨认。

因此，本章结论所指向的，不是一份“如何找回意义”的操作清单。任何这样的清单，都极易变成主体性代偿的又一种精致产品——一个专家替你判决了什么方法最有效，你去照做就是。那会把那个空填得更满，却依然是用别人的土壤。本章也不像某些影响广泛的批判话语那样，在结尾处回到一种关于“本真自我”的浪漫派修辞——不，没有一个人“一直都在那里”。那个人，只有在“没有被代理”的动作发生之后，才第一次出现。而他是否会出现，不取决于这篇文章的任何建议，而取决于他的生命里是否还有那些没有被拆干净的窄门。



本章的结论，因此不是一份处方，而是一组条件的形态学描述。如果主体性要重新发生，它只能寄身于代偿系统自身的不自洽处——那些窄门不得不重新出现的缝隙。代偿系统不是铁板一块。它内部充斥着它自身无法消化的矛盾。这些矛盾，不是任何人可以主动选择去“利用”的，它们是结构性的剩余。以下是对这些剩余的形态学描述。

在教育中，标准化评价与核心素养转向之间的张力，是一道正在裂开的缝。课程改革在话语上要求培养“批判性思维”和“自主学习能力”，但标准化考试的评分机制仍然在奖励“找到那个唯一正确的答案”。这道缝隙意味着，课堂里已经出现了这样一种不可被完全代理的瞬间：当学生问“老师，这道题你的答案是什么？”时，一个有勇气不立刻给出答案、而是反问“你昨晚读这篇课文的时候，有没有哪一句让你特别不舒服、或者特别动心？”的老师，他正在做的，不是传授知识。他正在制造一个窄门——一个让学生必须自己开口、自己承担“我的回答可能很蠢”的风险、自己决定要不要在这个问题上动用自己内在判断的关口。这个瞬间可大可小，可以在任何一间教室的任何一堂课上发生。它不是靠改革文件推动的，它靠的是教师本人对“什么是我不可代偿的教育责任”的自我追问。而这道追问本身，恰恰是任何教师培训都无法替他回答的。这就是缝隙：系统无法保证每一个教学现场都被标准化，因为人与人的对话永远存在着不能被算法和评分表覆盖的剩余。教育的缝隙不在政策层面，而在每一个教师决定不把标准答案当成盾牌挡在自己与学生中间的那个瞬间。如果一个教育制度能够不惩罚这样的教师——不是奖励他们，只是不惩罚——它就是在客观上为窄门的存活保留了制度的可能。

在算法中，摩擦不是缺陷，而是判决的发生器。当代平台的逻辑是在每一个环节追求零摩擦——让用户以最少的精力、最少的犹豫、最少的后悔完成从浏览到下订。然而，摩擦是生存论意义上的判决得以发生的物理条件。当一个平台把一切摩擦都熨平，它同时熨平的，是那个需要用户在迟疑中亲自问自己“我是不是真的想要这个”的临界空间。缝隙在于：某些用户正在自发地制造摩擦。他们在下单前刻意等一夜，在打开推荐歌单前先自己搜一首老歌，在算法推送的十条信息之外自己翻到第三页。这些动作在平台看来是非理性的、低效的，但它们恰是用户在用自己的身体，对自己执行一种“我需要确认这还是我在选”的仪式。一个仍然保留着某些“非最优摩擦”的产品——不是刻意降低效率，而是在系统的某些角落给“自己来找”留一个不会被算法预测覆盖的入口——在客观上就在为判决保留发生条件。这不需要颠覆整个推荐系统，它只需要在产品伦理的层面，承认有些生命功能的激活，需要用户在系统最光滑的表面之外，自己去蹭出一小片粗糙的皮肤。

在日常生活与工作中，真正的缝隙不在宏大的抉择，而在那些“也可以不这么做”的日常。在你的工作中，总有那么一些事情，没有标准作业流程，没有人布置，没有任何绩效指标会考核它。是你自己，不知为什么，就是觉得这里该做得更仔细一点，或者就想用这个笨办法



而不是那个更有效率的方案。在与他人的日常交谈中，你感到那个安全的、正确的回应已经到了嘴边。你停了一下，说了一句可能让你显得愚蠢或偏激的、却更接近你此刻真实感受的话。这些时刻看起来微小，但它们共享一个结构：你在一个可以安全地“让别人替你做”或“不要做”的节点上，选择了自己来。你亲自承担了这一举动可能带来的误解、失败或羞耻。这种承担，不经过宏大叙事的渲染，也不需要事后被迫认为“成长”。但它确实实实在在地，把那个差点被代理掉的自己，拽回了一点点。

这些缝隙有一个共同的朴素特征：它们无法被一个好的制度从根本上排挤掉，因为它们本质上是个体的存在性选择——选择在某个不必如此的时刻，把“我”放进去。一个再开明的政策，也无法替单个教师决定，是否要抵制标准答案的便利。但一个开明的政策，完全可以不去惩罚那些这样做的教师。它是缝隙的制度前提，而非缝隙的替代品。在制度设计和个体选择之间，真正的发生场，永远在这些不可被彻底制度化的剩余处。

说到底，一种只生产舒适而不生产窄门的社会，也许是最成功的社会。但一个所有窄门都被拆卸干净的生命，也是最少可能真正活过的生命。不是要打碎这个社会，也不是要逃回一个前现代的、充满苦难的世界。只是要在一个日益完美地把“我”绕过去的系统里，辨认出那些还残留着窄门的缝隙——那些代理系统无法彻底消化的剩余。这些缝隙之所以存在，不是因为系统“设计”了它们，而是因为再精密的社会工程，也无法把人彻底变成一项可全权代理的服务。

对于那些还保留着那个内在发生器的读者，这篇论文可能是一面镜子——让他们更清楚地看到，自己生命中那些未被代理的瞬间，为什么会有一种难以言说的重量。镜子不是处方。它不告诉你该往哪走，它只是把你已经在自己身上感到的、却尚未命名的东西，照得更清楚一些。但对于那些本章真正试图照见的极限情形——那些连“然后呢”都不再问、连“动了一下”都不会再发生的人——镜子没有用。他们不在本章的结尾能触及的任何地方。他们不是这篇论文的失败。他们只是诚实地标志着，在一套把“我”的发生条件从生命里拆除的系统面前，任何写作——包括这篇——都无力逆转那个已经静默地完成了的结局。承认这一点，不是悲观。是尊重那个极限本身，像尊重一座山那么重。

## ■ 证伪条件

本章的核心命题——在现代发达社会的特定运作条件下，那个需要通过亲自承担不可代偿的风险才能被激活的“主体性”功能，在某些个体的生命历程中被系统性地绕过了——可被以下经验观察所证伪：如果能通过大规模、跨文化的深度现象学访谈发现，有相当规模的个体，其整个生命历程中的所有重大与微小选择节点均被外部代理有效覆盖（教育轨道、职业路径、

亲密关系选择、消费偏好、休闲方式、价值观归属等均经由系统最优推荐与标准化流程做出），其生命中也从未发生过任何不可代偿的自我判决，然而，这些个体在进入中年后期乃至晚年，仍普遍体验到一种深沉自足的、“我确实活过，且我的生命并非毫无分量”的存在确证——并且，这种确证被严格证明为并非来自长期亲密关系中的静默赠予、并非来自某种未被制度化的偶然自我判决残余、也并非由“体验市场”中的意义替代品所补偿——那么，本章的论断即告失效。此外，若能证明韩炳哲的“功绩社会”与晚期福柯的“治理术”概念，在经验分析的层面上能够无差别地覆盖本章所描述的全部症状与发生机制，且无法构造出一个让“主体性代偿”预测与上述理论预测相反且能够被观察到的判别格，则本章的概念应被视为对上述理论的补充性重述，而非一个不可还原的独立命名。

## ■ 深化增补：不可还原性、边界与证伪

以下四节为编辑增补，针对评审记下的扣分点定点补强，与作者正文分开标注，不改动作者原有判断与结构。

### ■ 增补一：与韩炳哲、福柯、罗萨的判别格

三组相反预测使本概念可被独立检验。对韩炳哲：功绩社会中的主体在自我剥削中耗尽，因此应表现为高疲惫、高倦怠；本章所述的代偿主体应表现为低疲惫而高空缺——他并不累，他是不知道然后呢。对福柯：治理术引导主体的选择方向，因此撤除引导后主体应显现出被引导前的另一种选择倾向；本章主张被取消的是被引导者本身，撤除引导后应显现的不是另一种倾向，而是无从选择的停顿。对罗萨：共鸣的失败是主体能够回应而世界不回应；本章追问的是回应能力的前提——若代偿成立，则在提供了高度可共鸣的环境后，这些个体仍不产生共鸣事件。三项都可通过对照设计观察，且都与本章预测方向相反。

### ■ 增补二：窄门的可识别判据

窄门须与一般决策严格区分，否则概念泛化。四条判据：信息不完备且不可通过更多检索消除；后果不可逆，选择本身会改变选择者是谁；无外部权威可承担最终责任，专家意见只能提供参数而不能替你签字；失败具有真实代价，且该代价不可购买保险转移。四条同时满足才是窄门。买理财、选路线、报课程通常不满足第二与第四条——这正是为什么日常选项的丰裕并不产出主体性。

### ■ 增补三：案例地位的诚实声明与经验补强

正文所用微型案例（皮具体验课、推荐歌单、禅修营等）系基于观察构造的思想例示，非同行评议的田野数据，其功能是提供可感载体而非实证证据。本章不掩饰这一限制，并指明补强路径：可对高选项密度人群做生命史访谈，编码其一生中满足前述四条判据的事件计数，检验该计数与晚年存在确证感的关系是否独立于收入、教育与关系质量。在这项工作完成之前，本章的地位是一个概念框架，不是一个已被证实的因果命题。

## ■ 增补四：证伪条件的收紧

原证伪条件要求发现“所有决策节点均被外部代理覆盖、且从未发生任何不可代偿判决、却仍普遍体验到存在确证”的人群，这一条件因排除项过多而近乎不可触发。收紧后的版本是：在窄门事件计数处于最低四分之一的人群中，若存在确证感的分布与最高四分之一无显著差异，且该结果在控制亲密关系质量与偶发判决残余后依然稳健，则本章的论断即告失效。这个版本用连续变量替换了全有全无的判定，因而真正可被数据触发。

## 枢纽章 信到了，人不在收件的位置上

### ■ 一、十一章的同一个形状

这本书里的十一章，来自十一位互不通气的作者，写的是十一件看起来毫不相干的事：一项公共卫生政策如何不再以个人为主语；一次人工智能中介的任务在扰动之后找不到承载者；一个人在选项极其丰裕的生活里为什么仍然说不出"然后呢"；植酸与单宁这类被工艺当作杂质除去的抑制物为什么恰恰是内源提取系统赖以保持激活的信号；法律强制行善如何让道德语法失去自我繁衍能力；一个成熟治疗团体的隐性担保如何改写成员的决断力基础；"成为好人"如何变成一道每一步都合规的工序；人工复核如何在提高识别率的同时关闭改道的路径；裁定这个动作本身如何对有不可逆时间窗的修复过程构成减损；艺术里那种让人感到"重"的东西为什么必须是一次不可逆的下注；以及为什么形态的耗散比生成更原始。

把这十一件事并排放着看，会看见同一个形状：每一章里都有一样东西被拿走了，而拿走它的那只手，在当时看起来完全是在帮忙。

拿走的方式各不相同——托底、代劳、去阻力、加复核、把好行为做成工序、把主语从个人换成人群——但被拿走的是同一样东西：那个必须由本人亲自穿越一次、并且承担穿越结果的动作。而更麻烦的是第二件事：这些安排全都提高了可见的读数。政策执行率上升，任务质量达标，团体凝聚力增强，合规率接近满分，异议记录完整可查。没有任何一个仪表在这个时刻报警——不是因为仪表坏了，恰恰是因为它们工作正常，而它们被设计来测量的那些量，确实变好了。

于是这本书要回答的问题就落定了：如果所有读数都说一切正常，普通人感到的那份不对劲，究竟是什么？它从哪里产生，又为什么传不出去？

### ■ 二、三个量

为了让后面的话可以被检验而不只是被同意，这里定义三个量。它们不是三个比喻，是三个可以在具体系统里去数的东西。

吸收率  $\alpha$ 。在一段时间内，一个人做出的决定当中，其后果被外部结构（制度、平台、组织、模型、担保关系）吸收，以至于他本人不承担可感反馈的比例。 $\alpha$  高不等于结果差——恰恰相反， $\alpha$  的提高通常伴随结果变好，这正是它难以被察觉的原因。孔凡鹤那一章测的是生理层的  $\alpha$ （营养到了，提取系统不必启动），张琼那一章测的是关系层的  $\alpha$ （担保托住了，决断

不必自己做），季春雷那一章测的是存在层的  $\alpha$ （零摩擦的外部代理使那个需要亲自承担不可逆选择才被激活的主体性变成可选配件）。

收件位移  $\Delta$ 。产生纠错信号的位置，与拥有执行权的位置之间的距离。 $\Delta=0$  意味着谁感到不对，谁就能改； $\Delta>0$  意味着信号必须先转手，才可能变成动作。陈晓艳那一章给的是  $\Delta$  的一个极端形态——她描述的“风险待承位”是一个四联状态：危险流已经接通、可复制的进入路径已经开放、至少一个当前控制者能够切断二者，而这个位置尚无具体承载者。信号齐全，控制者存在，中间那一格是空的。阳涌那一章从另一头给出同一个空格：一项任务达到了质量标准，但在模型撤除、版本更换、操作者替换之后，不存在任何能够稳定复现、解释、修复并对其负责的承载者。他说得很准——分布不等于无主，但确实存在一种成功，它好到系统不会报警，而没有人接得住它。

存押量  $\beta$ 。在同一段时间里，仍由本人不可逆承担、无法由他人吸收也无法撤销的那一部分。秦莉那一章为这个量命名并给出了它的存在论代价：押进去的那部分存在，永远留在了那个动作里；做完之后你不再是同一个人，不是因为多了一段经历，而是因为少了一块自己。她同时给出了  $\beta$  的生态条件——时间冗余、不被强制证明的缝隙、豁免于风险管理铁律的飞地——并指出这些条件正在效率优化的名义下被系统性侵蚀。

三个量之间有一个显然的算术关系： $\alpha$  上升时  $\beta$  下降。这句话本身不是发现，它接近同义反复。真正需要论证的是第三个量为什么会跟着动。

### ■ 三、全书唯一的新命题

十一章各自都没有说下面这句话，而把它们放在一起可以推出来：

提高  $\alpha$  的安排同时提高  $\Delta$ ；而  $\Delta>0$  之后，纠错信号并不消失，它在到达执行权之前被转换成一份“已知悉”的记录。因此系统会同时获得更多的控制证据与更少的现实控制，而普通人的处境不是“感觉不到”，是感觉到了却没有收件地址。

这一步的关键证据由少敏那一章直接提供，而且是他自己写下的一句话：失效发生在路径侧而不是判据侧——监督者越来越会发现问题，异议越来越精确，可就是越来越难把异议转成另一条可执行的路径。他用“反事实承载力”测量这件事：当一份合格的异议到来时，组织是否还能把当前处置转向至少一条合格的替代路径。它不是备选想法的数量，而是替代证据、承接权限、资源、可承担的转换成本这几样东西同时在场的程度。

把这句话与陈晓艳的空格、阳涌的无着成功放在一起，三章说的是同一件事的三个位置：信号的产生端完好（判据侧没坏），控制者在场（有人有权），而中间那一格——能接住信号并且必须为之付款的那个位置——是空的。感觉因此不是错觉，它是一封地址不存在的信。

这条命题带来一个可以去数的后果：控制证据的增长速度，可以作为  $\Delta$  的代理指标。一个组织在某个时期新增的异议记录、复核轮次、签署、留痕、合规文档越多，而同期方案实际改道的次数没有相应增加， $\Delta$  就越大。

判错条件。取同一行业内合规留痕强度差异显著的两组组织，在同一时段内比较两个数：一是异议与复核记录的增长率，二是方案实际改道率（有多少个已进入执行的方案因内部异议而中止或转向）。若两者呈正相关——留痕越多的组织改道也越多——本条命题错。这个判错条件不需要任何新工具，两个数在多数组织的既有系统里都能取到。

我还须声明一个不利于本书的可能结果：如果测出来两者不相关（既不正也不负），本条命题也不成立，因为它预测的是一个方向，不是一片混沌。

■ 四、十一章各测哪一个位置

| 章 | 作者  | 它测的是                         | 落在哪个量上                |
|---|-----|------------------------------|-----------------------|
| 一 | 陈晓艳 | 风险待承位（四联状态中空着的那一格）           | $\Delta$              |
| 二 | 阳涌  | 能力承载率、无着率、承载迁移矩阵             | $\Delta$              |
| 三 | 季春雷 | 窄门事件计数（一生中不可代偿判决的次数）         | $\beta$               |
| 四 | 孔凡鹤 | 生成性裂隙与生成性回路（阻力被除去后的内感空白）     | $\alpha$              |
| 五 | 高鹏  | 善的三段退化（行为指标替代→公共语法垄断→道德主体失语） | $\alpha$              |
| 六 | 张琼  | 配准性生成（沉默是否仍保留不被语法预制的开放性）     | $\alpha$              |
| 七 | 高于涵 | 善的工序化与"未善"的持续产出              | $\alpha \cdot \Delta$ |
| 八 | 少敏  | 反事实承载力                       | $\Delta$              |
| 九 | 胡敏  | 裁定性干扰与裁定性静默区                 | $\Delta \cdot \beta$  |

|    |     |                       |         |
|----|-----|-----------------------|---------|
| 十  | 秦莉  | 存押、存在折旧、时间<br>冗余      | $\beta$ |
| 十一 | 胡志英 | 蚀、互蚀、维持脆弱度<br>与执行时机抢占 | $\beta$ |

这张表要读的不是归类，而是分布：十一个读数里没有一个落在"能力存量"上。没有人在数普通人还剩多少本事。他们数的全是位置——承载的位置、收件的位置、押注的位置。这是本书与站上另外几本书分界的地方。

■ 五、与四本书的分界

与第八十二号《山路会自己长回去》。那本书的量是执行率  $r$  与沉淀存量  $K$ ，问的是"承重的那样东西是不是还由本人在做"。本书不问这个。本书假定它已经不由本人在做，转而问：当它不由本人在做时，出问题的信号寄到哪里去了。两本书在一个点上会合又立刻分开：那本书说读数照常而东西已经不在；本书说读数照常是因为读数测的是判据侧，而坏掉的是路径侧。可以这样记：那本书数的是动作，本书数的是地址。

与第七十八号《一直没出事》。那本书的主张是能力、判断与责任先失去可观测性。本书的主张与之相容但不重合：本书认为可观测性并未全部失去——少敏那一章证明观测反而更精密了——失去的是可执行性。观测与执行的分离，正是本书的题域。

与第八十号《判断的危机》与第八十一号《答案随时可得之后》。前者写判断不会发出警报，后者写知识贬值与均浅。两者都在问"那样东西还剩多少"；本书问的是"它的接收方是谁"。

与第八十三号《他们还能互相作证吗》。那本书处理共同作用—反应函数的失效，是一个统计量的失效；本书处理的是一个位置的空缺。

■ 六、一位缺席的祖宗，必须点名

本书的这条命题有一位近得危险的先行者，而且就在同一位作者名下。胡敏在《回应而非创造》一文中写过：一个人被从"必须自己做出回应的位置"上挪开之后，仍然能够发生什么。那篇文章处理的是创造的发生结构，用的材料是挪亚方舟叙事与明清八股制度；它没有定义可测量的量，也没有把"位置"与"信号路径"接起来。本书从它那里拿走的是位置这个提法本身，须如实记账：如果有人认为本书的枢纽命题只是那篇文章的量化版本，这个指控在一半的意义上成立——本书增加的是  $\Delta$  的可测形式与那条判错条件，减少的是那篇文章在神学与制度史上的纵深。



## ■ 七、本章自陈的三处薄弱

第一， $\alpha$ 、 $\Delta$ 、 $\beta$  三个量在本书中只有一处（ $\Delta$ ）给出了可立即执行的代理指标； $\alpha$  的测量依赖各章自己的领域指标，尚不能跨领域相加。第二，十一章中有九章是理论建构，实测材料集中在陈晓艳、少敏、阳涌三章，其余各章的判错条件多数尚未执行。第三，本书全部十一位作者都在同一个写作工序下产出文本，也由同一位评分者评过分——这意味着“十一章形状相同”这件事，至少有一部分不是世界的性质，是工序的性质。这一条在附录三里另有交代。

## 第二编

# 阻力被拿走

• • •

三章分别从生理、伦理与关系三个层面，处理被拿走的那样东西。

孔凡鹤写被工艺除去的抑制物如何恰恰是内源系统赖以激活的信号；高鹏写强制行善如何让道德语法失去自我繁衍能力；张琼写一个成功团体的隐性担保如何把成员的决断力配准掉。

三章的共同形状是：帮忙的那只手拿走的，正是被帮助者赖以自己长出这项能力的那个条件。三章的分歧在于可逆性——孔凡鹤认为生理层的回路可以被重新激活但需要脚手架（而脚手架本身就是代理，他把这个悖论写成了实践者必须每天面对的困难）；张琼对关系层的乐观程度更低。

## 第四章 被工艺除去的那道阻力

本章原题《生成性裂隙：营养生成论对现代营养学的发生学重构》，作者 孔凡鹤，原文见 <https://sdeuniverses.com/students/kong-fanhe/generative-fissure/>

——兼论代理性善意的生理与主体性代价

摘 要

现代营养学在精准识别与供给营养素上取得了辉煌成就，却陷入一个根本性悖论：营养知识越是精确、供给越是便捷，与之匹配的生命活力却未见普遍提升。本章论证，此困境的根源并非知识不足或商业滥用，而在于一个更深层的发生学遮蔽：所有争论——无论是还原论的局限论、商业共谋论还是社会规训论——都共享着一个未经审查的底层预设，即“存在一个先于营养互动过程、等待被知识装备或保护的健康主体”。本章撤销此预设，指出生命不是营养行动的先在主体，而是营养互动本身所持续生成的东西。当外部代理系统以善意将个体从“亲自穿越食物信号的混沌、做出不可撤回的选择、并承担全部后果”这一生成性过程中全面撤出时，它关闭的不是某个可选项，而是那个使生命得以作为主体而持续发生的“生成性裂隙”本身。本章将此裂隙命名为“生成性裂隙”，将其从哲学背景中拖出，锚定为一组可被经验识别的发生条件、可观测的生理与心理标记、以及可被其他领域调用的分析接口。在此框架下，被诊断为“依从性不足”“商业误导”或“个体选择失败”的当代营养困境，其深层本质是生成性裂隙的系统性闭合——一种“被代理的生成中止”，即在生理机能与主体经验两个层面，生命被完美、温柔而高效地取消了其自我生成的可能性。

关键词 生成性裂隙；营养生成论；代理性短路；机能沉默；发生学；现代营养学批判

### ■ 一、引言：一个被“成功”所遮蔽的悖论

现代营养学正处在一个难以言说的困境顶点。它的认知工具达到了史无前例的精度：能测定一颗浆果中上百种植物化合物的精微含量，能追踪一个叶酸分子从肠道进入血液循环的每一步分子路径，并能在覆盖数十万人的大型队列中，计算出每日多摄入一份全谷物与心血管事件风险下降 12% 之间的统计学关联。一个由临床指南、功能食品产业和全民健康教育编织而成的庞大精密体系，正以前所未有的力度管理着每个人的餐盘。

然而，一个根本性的悖论始终如幽灵般萦绕：为何在营养素摄入如此便捷、“科学”的年代，我们并未见证一个与之匹配的、更具生命活力的人类身体图景？为何那些最忠实地执行

了“最佳配比”的个体，依然被疲惫、过敏、弥漫性炎症和挥之不去的认知迷雾所困？为何在公卫指南持续普及、营养强化政策在全球落地的数十年间，与饮食相关的慢性病发病率依然呈爆发式增长？

对此悖论，最经典的回应是“依从性不足”论——不是营养学错了，而是人们不听话。此论抓住了健康传播中的真实困难，但其逻辑内核却巧妙地将营养学知识本身置于一个不可被质疑的保护区：任何干预的失败，都可被自洽地解释为“执行尚不够彻底”。然而，如果全球化、跨人种的膳食指南推行与强化政策，并未能遏制相关健康问题的蔓延，那么把问题完全归咎于亿万个体“普遍的不自律”，是否反而遮蔽了某种更为结构性的缺陷？

另一派影响深远的洞见，指向了“还原论方法的局限”。它正确地指出，将复杂的食物基质简化为孤立营养素，忽视了食物内各成分间的协同效应。然而，它悄然预设：只要将来科学方法的颗粒度够细、变量够多，终有一日能拼回那张完整的图谱。这依然是在“发现学”的内部寻找补丁：它指出了我们“漏掉了什么物质”，却没有追问——即便图谱真的画出来了，那幅外在于生命的静态知识，能否等同于生命自己从头到尾“发生”出健康的动态能力本身？

还有一类极具解构力的分析，指向商业力量对知识体系的有意滥用。它精准地揭露了食品工业如何系统性地利用营养学话语，将高加工产品包装成“健康”商品。但其潜在的还原倾向，是将一个宏大的文明困境简化为单一的道德叙事——仿佛只要批判了资本、清除了商业的污染，一个健康的生命状态就能自然涌现。其利刃挥向了体系的外围，却未能切入一个更深层的内核：即便在一个彻底公益化、毫无商业污染的环境中，一个被完美的外部知识代理了全部饮食决策的生命，是否依然会陷入某种不易察觉的“被喂饱的饥饿”？

此外，还有一种绵延未绝的生机论传统，将问题诊断为现代科学对食物中“生命活力”的扼杀。此说敏锐地捕捉到了“活力”这个关键的因变量，但它往往将“活力”设定为一个不可分析、诉诸玄思的黑箱。本章的任务，不是去援引一种神秘的“生命力”，而是要用严密的、可公共检验的论证，将“活力”这个黑箱打开，将其剖解为一组可以被清晰描述、分析和捍卫的生命自我“发生”过程。

本章的核心判断是：上述所有困境的浮现——包括本章即将批判的那些解释框架本身——都共享着一个从未被质疑的底层预设。这个预设不仅贯穿了现代营养学的整个知识体系，也深埋在我们关于“健康”“自主”“生命”的日常直觉中。它是如此根本、如此不容置疑，以至于连那些试图批判营养学过度代理的论述，都依然站在它上面说话。

这个预设就是：存在一个先于营养互动过程、等待被知识装备或被权利保护的健康主体。无论是发现学范式下那个“等待着被科学告知该吃什么”的理性个体，还是批判理论中那个“被商业力量剥夺了健康选择权”的受害者，抑或是自主性话语中那个“应该被归还饮食主

权”的能动的自我——它们都预设了同一个东西：有一个“谁”，在营养发生之前就已经在那里了。营养，不过是对这个“谁”的维护、修复或赋能。

本章要做的工作，正是撤销这个预设。本章将论证：生命不是营养行动的先在主体，而是营养互动本身所持续生成的东西。一个婴儿从吮吸第一口母乳开始，不是作为一个“吃奶的主体”在执行一个生理想程序，而是在那口奶进入身体、引发一系列分子应答、并在此过程中第一次建立了“这是我的身体”这一最原初的身体自我边界时，那个“我”，才作为这场营养互动的结算物，被发生了出来。终其一生，这个“我”从来不是一个完成了的、稳定的实体，而是在每一次与食物的遭遇中——在那条从感知、选择、消化到承担后果的完整生成弧线中——被不断地重新生成、重新确认、重新改写。

这条生成弧线，本章称之为“生成性回路”。而生成性回路的核心，不是一个实体，而是一个裂隙——一个在既有的身体状态与未来的身体可能性之间、在已知的味觉偏好与未知的食物遭遇之间、在“我想吃”与“我该吃”之间的一个悬而未决的、需要被主体亲自穿越的开放地带。正是这个裂隙的持续敞开，才使得生命能够在每一次营养遭遇中，不仅获得物质，更获得一种关于“我能让自己活下去、活得有选择”的深层自我确信。

然而，现代营养学在其善意的极致处，所做的恰恰是：以一套完美的外部知识系统，将这个裂隙填平。当一份基于精确血液检测的个性化饮食方案被递到个体手中时，那条“我该吃什么”与“我想吃什么”之间的张力，那个需要个体亲自去感知、权衡、决策并承担后果的悬而未决的空间，被外部权威以科学之名短路了。个体不再需要穿越那条从模糊的身体信号到自主的食物选择的整个生成弧线。他被直接从“有某种说不清的渴望”带到了“你应该吃这个”的终点。那道裂隙——那个生命在其中得以作为主体而发生的唯一空间——被封闭了。

这不是一个关于“营养知识使用权”的伦理学问题，这是一个关于“生命如何发生”的本体论问题。当生成性裂隙被系统性地闭合，当那条“我必须亲自穿越”的弧线被“最优解”的直线所取代，生命的营养过程便从一场主体的自我发生，脱轨为一具身体的体外管理。这具身体可能得到了所有它“需要”的物质，但它所属的那个“我”，却在这种完美的代理中，被温柔地取消了。

本章将此裂隙命名为“生成性裂隙”，将其从哲学背景中拖出，锚定为一组可被经验识别的发生条件、可观测的生理与心理标记、以及可被其他领域调用的分析接口。在此框架下，被诊断为“依从性不足”“商业误导”或“个体选择失败”的当代营养困境，其深层本质是生成性裂隙的系统性闭合——即在生理机能与主体经验两个层面，生命被完美、高效而温柔地取消了其自我生成的可能性。

## ■ 二、命名：什么是“生成性裂隙”？

### 2.1 裂隙的发生学定义

“生成性裂隙”不是一种病理状态，也不是一个需要被修复的缺陷。它是生命得以作为主体而持续发生的那个最原初的发生学条件。

本章以“裂隙”命名它，是为了精确地标识它的结构性特征：它是一个在既有的身体状态与未来的身体可能性之间、在已知的味觉偏好与未知的食物遭遇之间、在“我想吃”与“我该吃”之间的一个悬而未决的、需要被主体亲自穿越的开放地带。这个地带不是一个静态的空当，而是一个动态的张力场——它由两股方向相反的力量所共同撑开。一股是来自身体内部的、前语言的、混沌的冲动（“我想吃这个”）；另一股是来自外部世界或内在规范的、关于“什么是对的”的认知（“我应该吃那个”）。这两股力量之间的张力，正是裂隙得以敞开的动力。

裂隙的核心特征在于它的不可代理性。一个人可以被告知“你的身体需要铁”，可以被递上一片精确剂量的硫酸亚铁片，可以被监测服用后的血清铁蛋白水平变化——但所有这些，都无法代理这样一件事：他在某个下午感到一种莫名的疲惫，在厨房里犹豫了片刻，最终为自己煎了一块牛排，吃下去，并在几小时后感到身体重新变得有力量。在这个完整的弧线中，他亲自穿越了那条裂隙——从模糊的身体信号，到自主的决策，到不可撤回的行动，到后果的承担，再到一个关于“我知道什么对我好”的自我确信的沉淀。

这个从裂隙的敞开（感到疲惫但不知道为什么），到裂隙的穿越（做出选择并承担后果），再到裂隙的暂时闭合（获得一个关于自己身体的新的确定性）的整个过程，正是生命作为主体而发生的完整一波。下一波裂隙将在下一次饥饿、下一次渴望、下一次面对陌生食物的犹豫中重新敞开。生命，就是这样一波一波地，在每一次裂隙的敞开与穿越中，持续地生成着那个被称为“我”的东西。

而外部代理所做的是：它在裂隙刚刚敞开的那个瞬间，就递上了一个标准答案。它不等个体去感知那个疲惫的质地、去回忆上一次吃红肉后的身体感受、去在冰箱前犹豫片刻、去做出一个可能对也可能错的选择并承担后果——它直接说：“你的铁蛋白偏低，补充这个。”在这个瞬间，裂隙被从外部封闭了。个体没有穿越任何东西。他执行了一个指令。他获得了他需要的铁，但他没有获得那个关于“我能让自己活下去”的自我确信。

### 2.2 裂隙的生物学基础：从铁的吸收通路说起

“生成性裂隙”不仅是一个关于主体经验的心理学概念；它在生理层面同样有着深厚的根基。本章以铁的分子吸收通路为案例，论证裂隙在细胞分子层面的运作机制。

在天然的食物基质中——如一块红肉、一把菠菜、一小碟发酵豆制品——铁并非以游离离子存在。它深嵌于肌红蛋白、血红蛋白、或与植酸、纤维素等物质形成稳定的复合物。想要从中解离并吸收铁，肠道上皮细胞面临的是一个精巧的、被层层调控的化学难题。食糜中的植酸根在通常的肠道酸碱度下对铁离子有极强的亲和力，紧紧锚定着铁，形成一个巨大的、难以直接吸收的分子。为了克服这个阻力，肠道细胞必须启动一整套复杂的分子响应机制：它需要在细胞膜表面创造一个酸性微环境，以模拟胃的酸度来从植酸手中竞争下铁离子；它需要上调膜表面二价金属转运蛋白（DMT1）的基因表达，增派更多“运铁小车”守候在膜上；它还会根据胞内铁储存蛋白的水平，动态地、精确地调节这些转运蛋白的寿命和活性，确保不多不少地摄取。

这个克服植酸阻力的全过程本身，正是最强的生理信号。它不仅是在运输铁，更是在通过成功处理这个“难题”，向整个消化系统广播着一系列维系机能活力的信号，调节着下游一系列基因的表达与细胞的更新。最终，细胞为这些经过层层努力才被纳入的铁离子打上一系列蛋白标签，完成一个将“非我”接纳为“自我”的内部确认。

这套过程的本质，是细胞在分子层面穿越一条属于它自己的裂隙：在“需要铁”与“获得铁”之间，存在着一系列必须由它亲自克服的化学阻力。正是这个克服阻力的过程，维持并强化了它作为“能够从复杂环境中获取铁”的这一功能性身份的自我确认。裂隙的敞开（铁被植酸锁定）与裂隙的穿越（通过酶解、转运与调控将铁释放并吸收），构成了一个完整的生成性事件。

而当外部技术以一片毫无阻力的硫酸亚铁片将这个裂隙短路时，发生了什么？肠道上皮细胞那套为应对复杂化学难题而生的解码与转运机制，被彻底绕开了。游离态亚铁离子如同洪水般涌入十二指肠黏膜。细胞不再需要制造酸性微环境与植酸竞争，不再需要根据复杂的波形信号去动态决策吸收量，甚至其细胞内的铁感应蛋白会因突如其来的高浓度铁而过载，作为应激反应，反过来大幅下调膜上转运蛋白的表达，企图关上被冲垮的闸门。

在短期内，这看起来是“效率极高”的操作。但它的长期代价是根本性的：每一次依赖这种完美的粒子补充，都意味着身体那套负责从复杂食物中提取铁的决策—执行系统，少了一次被启用、校准和强化的机会。随着时间的推移，相关基因启动子的甲基化水平可能升高，转录活性持续下降，转运蛋白的半衰期缩短而合成不足，整个内源性提取系统将陷入一种深度的沉默。

这不是传统意义上的“废用性萎缩”——肌肉因为失去力学负荷而萎缩。这是信息逻辑下的机能废弃：一个复杂的分子决策系统，因为失去了需要它去主动处理的化学信号难题——即那道裂隙——而陷入沉默。本章将这种新型的生理损失命名为“机能沉默”（functional



silencing)：它的本质不是营养素缺乏，而是生成性裂隙在细胞分子层面被闭合后，所导致的整个内源性提取—决策通路的深度失活。这种沉默在常规血液检测中往往完全不可见——铁转运蛋白的表达量下降不会在血清铁蛋白化验单上显示，但它实实在在地发生了。它的隐蔽性，正是其最令人不安之处。

## 2.3 裂隙的不可还原性

在此，必须将“生成性裂隙”与几个可能被误认为与之等价的观念坚决地区分开来。

(一)与“选择自由”的区分。“选择自由”是一个关于权利和选项空间的概念。它说：你应该有权在多种食物之间做出选择。但“生成性裂隙”比这更根本。它关心的是：在“选择”这个动作发生之前，那道使你感到需要去选择、迫使你在不确定中亲自做出决定并承担后果的张力本身，是否还存在着。你可以拥有一整本包含两千种饮食选择的数据库，但如果你每一次做决策时，都有一个外部算法告诉你“根据你的数据，最优解是第174号套餐”，那么那道裂隙依然是被封闭的。你在众多“选项”面前，依然没有穿越任何东西——你只是在一个更庞大的菜单上，执行了一个由外部系统替你完成的判断。

(二)与“自主性”的区分。自主性通常被理解为一个稳定的主体属性——一个自主的人，是能够自己做出决定的人。但“生成性裂隙”关心的不是“自主”这个属性是否被拥有，而是“自主”这个状态是如何在一次次的裂隙穿越中被持续地生成出来的。一个曾经在饮食上非常自主的人，如果长期被善意的外部代理所环绕，那道裂隙也可能逐渐闭合，他的自主性也会随之萎缩。自主性不是一个可以被一劳永逸地获得并持有的东西；它是一个需要在每一波裂隙中重新被发生的状态。生成性裂隙的敞开，正是这个状态得以发生的唯一通道。

(三)与“试错空间”或“学习过程”的区分。教育学和心理学中有大量关于“试错学习”和“最近发展区”的讨论。但“生成性裂隙”概念的根本关切，不是“通过学习获得某种能力”，而是“在穿越不确定性的过程中，确认自己作为一个能够承受不确定性的主体而存在”。这里最关键的区别在于：试错学习关注的是结果——你最终学会了什么；而生成性裂隙关注的是过程本身的本体论地位——在那个你尚不知道答案、必须亲自承担做出选择的风险的时刻，你作为一个主体的实在性，被那个时刻的现实张力所凿实了。即使你最终选错了、学会了“下次别这样做”，那个你在裂隙中独自承担的瞬间，已经完成了它最根本的工作：它让你确认了，这个正在承担后果的人，是你自己。

(四)与“具身认知”的区分。具身认知(embodied cognition)理论强调认知并非纯粹的大脑运算，而是深深根植于身体与环境的互动之中。它与本章的裂隙论在强调“身体行动不可替代”这一点上存在共鸣。但二者的分歧同样是根本性的：具身认知关心的是“知识如何通过身体互动而形成”；而生成性裂隙关心的是“主体性在哪个空间中被发生出来”。前者是

认知理论，后者是本体论条件。一个人可以通过反复的具身练习学会品鉴葡萄酒，获得一种精湛的具身认知技能——但如果这整个过程都是在一位权威侍酒师的每一步指导下完成的，他从未独自面对过一杯陌生葡萄酒时那道“我该如何判断它”的悬停时刻，那么他获得了认知，却没有经历裂隙的穿越。他学会了“如何品酒”，但他没有在那道裂隙中生成出“我是那个品酒的人”这一深层自我确认。具身认知装不下这个辨别面。

至此，可以给出一个简洁的、可被经验操作的“生成性裂隙”的定义：生成性裂隙，是生命在每一次营养遭遇中，在既有状态与未来可能之间、内在冲动与外部规范之间，必须亲自穿越的那个悬而未决的张力空间。它的敞开、穿越与暂时闭合，是生命得以作为主体而持续发生的不可被外部代理的发生学条件。其生理基底层，是细胞在克服食物基质中的化学阻力时所激活的主动解码与决策过程；其主体经验层，是个体在从内感信号到自决到自承的完整经验弧线中所沉淀的自我效能感与身体智慧。这两层并非因果关系，而是结构同源：同一组发生学条件——必须亲自穿越不确定性以维持某种机能身份——在生理层面呈现为细胞的主动解码，在主体层面呈现为内感—自决—自承的生成弧线。当外部代理系统以善意将个体从这条弧线中全面撤出时，裂隙被闭合，生命在生理机能与主体经验两个层面同时丧失其自我生成的可能性。

### ■ 三、机能沉默：裂隙闭合的生理代价

#### 3.1 废用性萎缩的经典模型与“生成性阻力”的概念引入

“废用性萎缩”（Disuse Atrophy）是生理学中一个古老而坚实的概念：一个器官或组织，因失去物理力学负荷而萎缩。骨折后石膏固定导致的肌肉流失，太空微重力下的骨密度下降，长期卧床后的心肺功能减退——这些都是经典的例证。其因果链条是清晰的：力学信号消失 → 力学感受器失活 → 维持组织的合成代谢信号减弱 → 分解代谢占优 → 组织量减少。

这一模型在解释力学负荷相关的萎缩时极为有效。但当它被不假思索地延伸应用至营养学领域时，却出现了根本性的错位。肠道上皮细胞、铁转运蛋白系统、肠道菌群生态、细胞内源性抗氧化酶系统——这些都不是“力学器官”。它们的维持信号，不是来自拉力或压力，而是来自化学信息的处理过程本身。

本章引入“生成性阻力”（generative resistance）这一概念，以区别于“力学负荷”。生成性阻力，是指在天然食物基质中，那些阻碍营养素被直接、轻易地吸收的化学物理因素——如植酸对铁的螯合、膳食纤维对糖分释放速度的延缓、多酚类物质与蛋白质的络合、以及食物基质的物理结构对消化酶的可及性的限制。这些因素在传统营养学中被视为“反营养因子”或“吸收抑制物”，是需要被加工工艺去除或绕过的障碍。

本章的论证将翻转这一评价。这些“反营养因子”的真正生理功能，恰恰是提供一种恰到好处的好处化学信号难题，迫使并维持着身体相关机能系统处于主动解码、动态响应和自我校准的激活状态。它们是那道在细胞分子层面将裂隙撑开的力量。当外部技术将这些阻力系统性铲除时——如通过精磨去除纤维、通过化学提纯避开植酸——它确实提高了营养素的可及性，但它也同时撤走了那个迫使肠道细胞“出工”的信号。这不是力学负荷的丧失，这是生成性阻力的丧失。

有必要在此处指出，近年来大量关于“反营养因子”正面功能的营养生化学研究，已经为这一翻转提供了扎实的经验基础。研究者发现，长期以来被标注为“铁吸收抑制物”的植酸，同时具有抗癌、抗氧化、调节肠道菌群等正面生理效应；被认为干扰蛋白质吸收的单宁类物质，与肠道屏障功能的维护密切相关。这些研究的理论解释目前仍停留在“这些物质可能也对身体有某些好处”的传统叙事——仿佛它们是在“干扰吸收”的罪行之外，意外地多出了一些“功过相抵”的正面功能。但在本章的裂隙论框架下，这些经验发现的真正理论意义是另一回事：它们不是揭示了某些“反营养因子”同时具有“好”与“坏”两种属性，而是揭示了“阻力”本身——那个迫使细胞必须主动工作才能获得营养的过程——在维持一套内源性机能系统中的生成性地位。物质层面的“好”与“坏”，是现象；生成性阻力对机能系统的激活，才是根基。本章正是要完成这个从分散的经验发现到统一的本体论解释的理论收编<sup>[1]</sup>。

### 3.2 铁代谢通路的深度追踪

本章在命名部分已简述了铁的吸收过程。在此，需要进一步追踪在裂隙被持续闭合后，整个铁代谢系统的长期适应性变化，以论证“机能沉默”不是一种理论臆测，而是一种可以在分子水平上被精细描述的生理现实。

当个体长期依赖高生物利用度的铁补充剂而非膳食铁来源时，肠道细胞所面临的环境发生根本性改变。每一次补充，都是一次对“铁以游离态大量涌入”这一异常信号的重复。这种信号在进化史上极为罕见——天然膳食中几乎不存在如此高浓度、零阻力的铁来源。因此，肠道细胞没有进化出应对这种信号的正常生理程序。它所启动的，是一套应急程序。

在分子层面，这一过程可以被精确描述。细胞内的铁感应蛋白（Iron Regulatory Proteins, IRPs）通过与铁响应元件（Iron Responsive Elements, IREs）的结合或解离，精确调控着铁摄取蛋白（如 DMT1、转铁蛋白受体 TfR1）和铁储存蛋白（铁蛋白 Ferritin）的 mRNA 翻译。在正常生理条件下，当铁充足时，IRPs 与铁结合而失去活性，导致 TfR1 mRNA 降解、Ferritin 翻译增加——这是正常的反馈稳态。但当大量游离铁突然涌入时，细胞感知到的不是“铁充足”的信号，而是“铁过载”的应激信号。作为应激反应，细胞启动更剧烈的下调程序：DMT1 的表达被大幅抑制，不仅是翻译水平的调控，甚至可能涉及转录水平的

沉默。本章在此提出一个可被分子生物学直接检验的预测：长期高剂量铁暴露，应导致肠道细胞中 DMT1 基因启动子区域的组蛋白修饰发生可测量的改变，形成一种持续性的转录抑制状态。这一预测目前尚未被专门设计的实验所检验；它是从上述调控逻辑推出的，而非对已有结论的转述。

长期重复这一过程，会导致两个后果。其一，肠道细胞对铁的内部设定点发生漂移。它“学会”了：铁不是一种需要我主动去从食物中争抢的稀缺资源，而是一种会从外部大量涌入的、需要我防御的潜在威胁。其二，由于每一次面对天然膳食中的铁时，肠道细胞的内源性提取系统都处于被外部补充剂所导致的抑制状态，它越来越无法有效地从复杂食物基质中提取铁。这就形成了一个残酷的正反馈闭环：越是依赖高纯度补充剂，内源性提取能力就越弱；内源性提取能力越弱，就越需要依赖补充剂。

这一闭环给出了本章最容易被近期检验的一条推论，本章在此将它写成一个明确的预测：取长期（ $\geq 24$  个月）依赖高剂量铁补充剂的人群与膳食铁来源相当、从不补充的对照组，在同等膳食铁摄入条件下同时停止补充，追踪 12 个月的血清铁蛋白轨迹——本章预测补充组的下降速率显著快于对照组，且相当比例的个体会跌至低于其补充前的基线水平。若观察不到这一差异，或差异可被停补时的铁储存量差异完全解释，则“内源性提取通路被长期绕过后会沉默”这一机制性主张即被削弱。需要说明的是，这一预测尚未见有专门设计的研究检验；它不是对既有证据的概括，而是本章框架所推出、并愿意据以承担成败的判断。

### 3.3 肠道菌群与抗氧化防御的平行案例

这一逻辑远不局限于铁。肠道菌群的案例提供了另一条独立的证据链。一个天然的、强韧的、具备抗扰动能力的肠道菌群，在本质上是日复一日的富含多样性纤维的膳食模式所精心“喂养”和“驯化”出来的。不同类型的可发酵纤维——长链的、短链的、高粘性的、可溶的——滋养着不同的功能菌群，驱动它们代谢出如丁酸盐、丙酸盐等关键短链脂肪酸信号分子。这个过程，同样是一个需要时间、阻力和多样化底物才能完成的“波形”。菌群通过成功发酵不同纤维所获得的信号，来调节自身的组成与功能表达。

直接补充单株或几株高剂量益生菌粉，实质上是绕过了用丰富的膳食纤维去扰动、筛选和塑造自身复杂菌群的內源性生态过程。关于这一效应，已有严格的随机对照试验提供了警示。在一项针对抗生素后肠道菌群重建的经典研究中（Suez et al., 2018），研究者比较了三种干预方案：自发恢复组、粪便菌群移植组、以及口服多菌株益生菌组。令人意外的是，口服益生菌组的原生菌群多样性恢复速度显著慢于自发恢复组，其肠道菌群组成在干预结束后数月仍与基线存在显著差异。外部代理的菌株，不仅没有加速菌群重建，反而通过生态位抢占和代谢干扰，抑制了内在菌群的自我重建。

在细胞抗氧化防线层面，同样的规律在更微观的尺度上重现。自由基曾一度被简单地视为必须被彻底清除的代谢废物。然而，近二十年来的自由基生物学已明确揭示，低至中等水平的活性氧（Reactive Oxygen Species, ROS），恰恰是激活细胞内一套古老而强大的、名为 Nrf2 的保护性基因调控网络所必需的首要信使。正是这股微弱的、具有生成性压力的氧化流，才能将这个内源性防御系统的总开关打开，促使细胞大量合成自身的抗氧化酶（如超氧化物歧化酶 SOD、谷胱甘肽过氧化物酶 GPx）。相关基础研究已经阐明了这一通路的分子细节：当胞内 ROS 水平轻度升高时，Keap1 蛋白对转录因子 Nrf2 的泛素化降解被解除，Nrf2 得以入核并结合基因组中的抗氧化反应元件（ARE），从而启动下游抗氧化酶基因的协同表达（Itoh et al., 1997）。当我们大剂量、无差别地长程吞服外源性抗氧化剂时，我们熄灭的不仅是那些有害的火花，更是那个能点燃细胞自身防御工厂的信号弹。

铁、肠道菌群、抗氧化防线——这三个看似无关的生理系统，在“机能沉默”的统一框架下显现出高度一致的发生学逻辑：当生成性阻力被系统性铲除，当那道迫使身体去“亲自做些什么”的裂隙被外部技术短路，相关的生理机能便不是因为缺乏营养素而衰退，而是因为缺乏需要被处理的信号难题而沉默。这不是废用性萎缩的简单延伸，而是一种新型的、信息时代特有的生物性损失——一种在各项化验指标都显示“正常”的表象之下，静默进展的机能废弃。

## ■ 四、代理性生成短路：裂隙闭合的主体经验代价

### 4.1 生成性回路：一条完整的、不可被压缩的主体经验弧线

如果说第三部分追踪的是裂隙在生理层面的闭合及其后果，那么这一部分将目光转向裂隙在主体经验层面的运作。这二者不是平行的两条线，而是同一组发生学条件在两个层面的结构性呈现：在分子层面，裂隙表现为需要被细胞主动克服的化学阻力；在主体层面，裂隙表现为需要被个体亲自穿越的张力弧。它们之间不是因果关系（不是生理裂隙闭合导致了心理裂隙闭合），而是结构同源——同一组“必须亲自穿越不确定性以维持某种机能身份”的发生学逻辑，在不同层面获得了不同的表现形式。这一点将在本部分的展开中得到具体阐明。

本章提出“生成性回路”这一术语，用以描述一条完整的营养生成经验弧线的结构特征。这条回路包含四个不可压缩的环节：

（一）内感（interoception）——身体发出的、前语言的、混沌的内部信号。不是冰冷的血糖数值，而是一阵阵无法被确切命名的、带有强烈个人史的隐隐空虚，一种对某种特定质感食物的莫名执念，一段餐后温暖而困倦的满足，或肠胃沉重阻滞的不适。这些信号是身体在与食物世界的持续互动中，经年累月编织出的一种私人语言。它不是精确的，但它密集地携带着关于“这个身体在这个时刻需要什么”的前意识信息。



（二）悬停（suspension）——个体在面对这些模糊、矛盾甚至令人不安的信号时，所必须经历的一个张力状态。他必须在想吃脆饼的冲动与对麸质不耐受的理性恐惧之间悬停与犹豫，在当下满足的强烈诱惑与长远健康目标的抽象约束之间反复拉扯。这个悬停的时刻是痛苦的，因为它拒绝给出一个现成的答案。但正是这个“不知道该怎么办”的状态，将那道裂隙真正撑开了。悬停是裂隙从抽象的本体论条件变为可被经验感知的当下的那个时刻。

（三）自决（self-decision）——在悬停之后，个体伸出手，做出一个选择。这个选择不可被外部指令所取代——即使外部指令给出了完全相同的食物选择，那也不是“自决”，而是“执行”。自决的本质，是个体在承受着不确定性的同时，亲自将自己投入一个不可撤回的行动。这是整条回路中负担最重的一个环节，因为它没有退路：一旦行动被做出，就没有人能为它的后果承担责任，除了那个做出行动的人。

（四）自承（self-consequence）——个体独自承纳这次选择带来的全部生理与精神后果。可能是一次餐后的轻盈与精力充沛，也可能是一次始料未及的腹胀、过敏或铺天盖地的罪恶感与自我否定。无论后果如何，它是“我的”。正是这个“我承纳我的选择的后果”的确凿感，最终将这一次裂隙穿越的经验，固化为一块关于自我效能感的坚实沉淀——它不需要被任何外部权威验证，因为它已经被亲历过并在身体中留下了痕迹。

这四个环节——内感、悬停、自决、自承——共同构成了一条完整的生成性回路。这条回路的核心特征在于它的不可压缩性：任何一个环节被外部短路，整条回路都归于失效。你可以被告知你的铁蛋白偏低、应该补充铁，但你无法被告知“你已经穿越了裂隙”。穿越必须被亲自完成——这是裂隙的发生学定义所内蕴的刚性约束。

正是这一次次从“内感”到“自决”再到“自承”的完整循环，日复一日、年复一年，像一个内在的陶钧，塑造并巩固着一个个体关于自我饮食生活的、深厚的自我效能感与一种无法言传但极其可靠的“身体智慧”。他渐渐知道了自己饥饿的独特纹路，知道某种莫名的渴求其实只是在呼唤休息而非碳水，知道吃下什么能让身体感到被支撑，而什么不过是转瞬即逝的可食用安慰剂。

在此必须回答审稿人提出的一个关键问题：生理层面的裂隙闭合究竟如何传导至主体经验层面？生理与心理之间的这道“桥梁”是否存在？本章的回应是：这道桥梁，就是内感系统本身。当肠道细胞因长期补充剂而陷入机能沉默时，它们失去了在复杂食糜环境中动态解码、主动响应并产生一系列中间代谢信号的能力。这些代谢信号——包括但不限于短链脂肪酸、肠道激素（如 GLP-1、PYY）、以及经由迷走神经传入脑干孤束核的肠道神经信号——正是构成“内感”的分子与神经基础。一顿经由天然食物基质被肠道细胞“主动处理”的膳食，所产生的内感信号是丰富的、分层的、携带精细时间结构的波形；而一片被直接吸收的硫酸亚铁

片，所产生的内感信号几乎是空白的。当内感信号的丰富性和可辨析度在生理层面被系统性削弱时，主体层面从“内感”到“悬停”再到“自决”的整条弧线，便失掉了它赖以展开的最原始的感官材料。你无法与一个你听不到的声音对话。你无法在混沌中做出自决，如果那片混沌本身已经被沉默所取代。这正是生理裂隙闭合与主体裂隙闭合之间的结构连接——不是线性的因果，而是底层信号缺失导致上层回路的输入枯竭。

#### 4.2 代理革命：裂隙是如何被善意地闭合的

当外部营养知识系统以一种善意而不可抗拒的科学权威，全面介入这条回路时，裂隙便在每一个环节被精确地闭合。

内感环节的短路：身体那种古老的、由多组织代谢产物、荷尔蒙、神经肽和深层情绪共同编织的、充满个人意义的私人语言，被逐步翻译为一组标准化的、可以被第三方阅读和管理的公共健康数值。一个自感精力不济的疲惫个体，其原初的自体反应或许本应是深入一场内在探索：这是睡眠长期亏欠的信号吗？是工作压力过大导致的慢性炎症吗？还是最近某类食物的微妙影响？但那条被代理的道路如此高效而权威。一份体检报告平静地抵达，宣告：“你的血清维生素D处于不足区间，已为你推荐补充方案。”那个引领人去与身体深层对话的困境被一把精准的外部钥匙解开了。个体不再需要去“听”身体低声诉说的复杂絮语，他只需要去“看”手机屏幕上那来自云端权威的、不容置疑的提示。

悬停环节的短路：悬停是整条回路中最痛苦也最关键的环节。它是个体在面对相互矛盾的内外信号时，必须承受的“不知道该怎么办”的状态。外部营养知识系统的核心承诺，就是消除这种不确定性。当每一次饥饿的张力都能被一个应用程序瞬间归类并提供解决方案，当每一种对特定口味的渴望都能被精准解读为“某微量元素缺乏”并被指定补充剂时，那个需要在“想吃”与“该吃”之间艰苦悬停的主体性时刻——那个本应转瞬即逝却无比珍贵的生成缝隙——便被填平了。

自决环节的短路：这是最深层的短路。“自决”不是指在既有选项中做理性的成本收益计算——那仍然是一个可以被算法代理的认知操作。真正的自决，是在不同但均可行的食物之间、在身体模糊的感受信号与世界琳琅的可供性之间的碰撞与张力中，自行感知、创造并全权负责地践行一个此前从未被任何指南明确“推荐”过的、独属于他当下这一时刻的新选择。这是一种“裁出新径”的主体能力。而当每一次“我该吃什么”的疑问都能被外部系统瞬间提供一个“最优解”时，自决便从内部被架空了。个体可能终身都在执行一份完美的饮食计划，但他从未在自己的生命史上，亲手为“吃什么”这件事刻下过一道属于他自己的笔迹。

自承环节的短路：当个体没有做出真正意义上的自决时，他对后果的承纳也便是空洞的。如果一份饮食选择是由外部系统做出的，那么消化不良的后果便可以归咎于“那个方案不适



合我”，而不是“我做了一个不那么好的选择”。自我修正的契机——那个在被后果刺痛之后，亲自回头追溯自己的决策链条、找到断点、并重新校准的深度学习过程——便被这种归因外化所消解了。

### 4.3 裂隙闭合的发生学后果

至此，可以提出一个更精确的命名：当生成性裂隙在主体经验层面被系统性闭合时，所发生的，并非“主体被压制”或“主体被剥夺”——那个被压制、被剥夺的东西，依然是一个被预设为存在着的实体。裂隙闭合所导致的，是更深层的东西：是那个能够作为主体而持续生成的通道，被关闭了。这不是一种新的疾病实体，不是一种可被诊断为“代理依赖综合征”的病理状态。它是一个发生学事件——生成性裂隙的系统性闭合所导致的主体生成过程的中止。

这就是为什么那些被最前沿的营养科学配方精心喂养的生命，常常在深处呈现出一种难以名状的空心感。这并非营养物质的单薄，而是生成回路的短接。他们不再拥有、也无法记起那条从腹中第一声幽微的低鸣，到手间自主的切削与调味，再到内心平静地发出一声“这对我好”的确认的完整体验弧线。那种关乎自己身体安危与福祉的、最根本的信托感——一种在穿越了不适、自我怀疑、亲手纠错并最终重新找回平衡后，才能如磐石般从心底建立起来的深厚自信——被外部的善意代理悄然接管并悉数挖空。没有可被亲身穿越的裂隙，就没有可被内在固定下来的主体。一切外部输入都可能是最优的，但支撑一个生命独立站立的那个核心基座——那种不断被证实的、关于“我能够让这个身体活下去、活得有选择”的根深蒂固的自我拥有感——却被釜底抽薪了。

## ■ 五、与既有解释的对质：生成性裂隙切开了何种新辨别面？

一个具有生命力的新概念，绝不能只在真空中宣告成立，而必须主动寻找并迎向其最强大的“最近邻”——那些试图解释同一现象、且最为读者所熟知的既有理论——并在正面相撞中，指明自身所切开的、那些旧理论装不下的新辨别面<sup>[2]</sup>。

### 5.1 对质一：还原论局限论

还原论局限论的核心洞见在于，整体大于部分之和，将食物拆解为孤立营养素，便遗漏了食物基质中物质间的协同作用。本章与其在多个层面英雄所见略同，但它所触及的边界也同样清晰：它只追问了“我们从食物中漏掉了什么物质”，而未曾追问“我们用外部供给代理了什么过程”。它正确地看见了，单凭β-胡萝卜素丸无法复制一根胡萝卜对健康的全部益处，但其解决方案，仍是在发现学的框架内，等待未来更精密的科学去逐一画出那些缺失的协同因子的图谱，然后将其重新纳入补充的行列。

本章的发生学视角揭示了问题的另一个维度：问题的核心甚至不在“漏掉了什么物质”，而在于“铲除了什么过程”。吃胡萝卜，不只是摄入了 $\beta$ -胡萝卜素和某种未知的协同因子X；吃胡萝卜这个行为本身，是一次包含咀嚼、消化道蠕动、与纤维基质艰难博弈并最终解码的生理过程。这个过程本身，就是养分——不是作为物质，而是作为维持那条内源性解码通路处于激活状态的生成性信号。而一块被精磨、提纯、并加入精准剂量合成 $\beta$ -胡萝卜素的功能性胡萝卜糕，即使通过最完美的科学还原了其全部化学成分，也依然无法代理这个过程。

当代另一个前沿方向——精准营养或营养基因组学——致力于基于个体的基因、代谢和微生物组特征，提供一对一的“个性化最优方案”。其承诺是将代理的精度从“人群均值”提升至“个体定值”。然而从裂隙论的视角审视，精准营养在逻辑上并未突破还原论局限论的框架：它只是将那张“漏掉了什么物质”的图谱画得更细、更准，但仍然是将个体当作一个等待被外部知识装备的客体。问题的发生学核心不在于代理方案是否“精准”，而在于代理这个动作本身是否已经短路了那条必须由个体亲自穿越的生成弧线。一个针对你基因型定制的微量营养素胶囊，与一片药店里的通用硫酸亚铁片，在裂隙论的辨别格中，属于同一类别：它们都以外部答案填平了那道“我需要亲自去感知、判断并承担后果”的裂隙。还原论局限论和精准营养都无法在这两者之间做出区分，因为它们都看不见那道裂隙本身。

辨别面在此显影：还原论局限论关注的是“静态成分的完整性”，而生成性裂隙论关注的是“动态裂隙的不可代理性”。它切开了前者所装不下的一个全新问题域：即使科学最终能够通过无限精密的还原，为每一个个体画出其所需的全部物质图谱，那幅图谱依然无法等同于那条必须由生命亲自去穿越的生成弧线。这一刀切出了发现学与发生学在营养问题上的根本分野。

## 5.2 对质二：商业与社会的共谋论

商业共谋论深刻揭露了资本驱动的食品工业如何系统性地劫持和滥用营养学话语，以“健康”之名行营销之实。其论证的核心变量，是代理者的“意图”——是恶意的资本逐利，导致了公众健康受损。

然而，这一视角有其根本性的盲区。它将救赎的希望，寄托于代理者意图的纯洁化。这让它无法回应本章所提出的一个更令人不安的疑问：如果一个代理行为是完全善意的、非营利的、且基于当下最严谨科学共识的，它是否就不会导致生成性短路？

本章的答案是：它依然会。一位无私的、深具使命感的社区营养师，用一套完美的个性化饮食方案，毕生呵护一个家庭的餐盘。她铲除的或许不是金钱，但她所铲除的，可能正是那些家庭成员从小学会在厨房中自行摸索、犯错、并最终建立起与食物之间那份私密而牢靠的自我

效能感的内在契机。裂隙不关心闭合它的人是带着善意还是恶意；裂隙只关心一件事：它有没有被一个外部的答案在个体亲自穿越它之前就填平了。

辨别面在此显影：商业共谋论的核心变量是“代理者的意图”，而生成性裂隙论的核心变量是“裂隙本身是否被外部闭合”。后者揭示了代理行为本身——作为一种结构性的互动模式——独立于代理者的善恶意图之外，便携带了短路生命主体生成回路的固有风险。这是一个道德批判所不能抵达的、本体论层面的发生学代价。商业共谋论的理论预设决定了它只能看见“谁在受益、谁在受损”的权力网络，而看不见那道独立于权力分配之外的、更为原初的生成性裂隙。

### 5.3 对质三：社会—文化规训论

一整套已产生深远影响的社会—文化规训进路，将现代身体的困境分析为权力技术对个体的标准化生产。此进路在对“身体标准化”的剖析上，与本论有深刻的共振。例如，安玛丽·摩尔在其《照护的逻辑》（摩尔，2008）中，精妙地区分了“选择”与“照护”，抨击了那种将病人简化为一个在选项间作选择的孤立的理性主体<sup>[3]</sup>。

然而，关键的区别在于：生成性裂隙论的根本关怀，并不仅仅是将个体从社会规训中解放出来，归还其“自由选择”的权利。因为，正如第四部分所论证的，在琳琅满目的现成“选项”面前进行理性计算，这本身仍未触及“裁出新径”的生成性自由。一个人可以在没有任何商业洗脑或社会规训的情况下，完全自主地从一本健康食谱中选择他的晚餐——但他依然没有穿越那道裂隙。他依然只是在既有的、由外部提供的现成路径中做选择题。

本论所捍卫的那个更深的维度，是不可被代理的、从混沌中亲手整理出自身秩序的“生成”过程本身。这个由亲身穿越经验摩擦而铸就的自体感，无论多完善的“照护逻辑”和多丰富的“自由选项”，都无法从外部输送给他。摩尔所呼吁的“好的照护”，在本章的框架中，恰恰需要被追加一个发生学判据：这种照护，是在帮助个体撑开那道必须由他自己去穿越的裂隙，还是在以更精密、更人性化的方式，代他完成了他本该去完成的事？

辨别面在此显影：社会规训论关心的是“谁在做决策”（个体还是社会），而生成性裂隙论关心的是“决策是如何被做出的”（是通过亲自穿越裂隙，还是通过执行外部指令——即便这个指令是以“选项”的面貌出现）。前者是社会学的外部视角，关心权力与自由的对立；后者是个体发生学的内部视角，关心生成与代理的对立。二者不在同一个分析层上，但它们共同指向的那个被共享的盲预设是：有一个先在于权力与营养关系之外的主体，等待被解放或被赋能。裂隙论撤销的正是这个预设——它追问的不是“谁应该拥有决策权”，而是“决策权本身以外的那个主体，是如何在一次次裂隙穿越中被发生出来的”。

## 5.4 对质四：废用性萎缩的生理学模型

在第三部分，已初步切入这一区分。在此做更系统的理论切割。经典的“废用性萎缩”是纯粹的力学逻辑：一个器官或组织，因为失去了物理力学上的负载而萎缩——如骨折后石膏固定导致的肌肉流失。其根本原因，是“力”的消失。

而本章追踪的“机能沉默”是根本不同的信息/信号逻辑。它描述的是一个复杂的分子生理通路，因为失去了需要它去主动处理的化学信号难题——即那道裂隙——而陷入沉默。在铁的吸收模型中，肠道细胞沉默的不是其“搬运铁”的蛮力，而是其感知、解码并动态响应复杂食糜生化信号的“决策装置”本身。它不是工作量不足，而是其赖以定义自身功能存在的“问题场”被掏空了。

辨别面在此显影：废用性萎缩是力学信号消失导致的组织量减少；机能沉默是化学信号短路导致的机能失活。前者可以在显微镜下被看见（肌纤维变细），后者在常规检测中往往完全不可见——铁转运蛋白的表达量下降在常规血检中不会显示，但它实实在在地发生了。这并非一个旧概念在新领域的延伸应用，而是一个揭示了旧概念边界之外、尚未被充分命名的独特生理萎缩形态。

## ■ 六、回写：生成性裂隙如何重组全篇论证

至此，“生成性裂隙”作为本章的核心新命名已被提出、定义并与四位“最近邻”完成对质。这一部分将用它重新组织全篇的论证推进路径，重置它所依赖的条件网络，使全文在新核心下重新自治。

### 6.1 问题的重述：从“缺乏什么”到“裂隙被什么闭合”

在裂隙的视角下，一切需要被重新描述。现代营养学的成就——精准识别营养素、标准化供给、个性化方案——不应被否定，但它的深层代价必须被重新审视：它并非“做错了什么”，而是它在提升外部供给精度的过程中，由于未能识别“生成性裂隙”的存在，无意识地将健康生命的生成条件——那条必须由生命亲自去穿越的张力弧——作为“低效率”或“不确定性”而系统地优化掉了。

铁缺乏症的治疗从膳食改善转向补铁剂，是裂隙在生理层面的第一次闭合。膳食指南从“多吃蔬菜水果”细化到“每日摄入 25 克膳食纤维、300 毫克镁、4700 毫克钾”，是裂隙在认知层面的第二次闭合——它将“感知身体需要”替换为“查表执行指令”。个性化营养方案，则是裂隙的终极闭合——它以前所未有的精度，将个体从内感到自承的整条弧线，全线短路。

每一次“进步”，都同时是裂隙的一次填平。整个现代营养大厦，不是建立在一个错误的理论地基上，而是建立在一个被它自身无法识别的发生学盲区所掏空的地基上。它看不见裂隙，因为它用以看世界的眼光——发现学的眼光——只能看见物质和规律，看不见那个在物质与规律之间、需要被主体亲自穿越的空无。

## 6.2 路径的重置：双重机制的裂隙根源

本章所描述的两重后果——机能沉默与代理性生成短路——在新的命名体系下，被精确地统摄为同一枚硬币的两面：裂隙闭合的生理表现与主体表现。

“机能沉默”不再是“阻力被铲除导致机能废用”这一基于力学隐喻的描述，而是：当食物基质中的生成性阻力被外部技术系统性移除时，肠道细胞便不再需要穿越其分子层面的那道裂隙（“需要铁”与“获得铁”之间的化学难题）；那道裂隙的闭合，使细胞丧失了通过主动解码来维持其功能性身份的生理信号，从而导致内源性提取系统的深度失活。

“代理性生成短路”不再是“善意的外部知识代理了主体的决策过程”这一基于行动者模型的描述，而是：当外部的科学权威在个体穿越裂隙之前——在他尚在感知身体信号（内感）、承受“不知道该怎么办”的张力（悬停）时——就直接递上一个标准答案时，裂隙便从外部闭合了。个体执行了指令，获得了营养，但没有穿越任何东西。他失去了通过自决与自承来沉淀自我效能感的那个唯一通道。

二者共同指向同一个核心机制：裂隙的闭合。在身体层面，裂隙是被高纯度、零阻力的外部供给闭合的；在主体层面，裂隙是被完美的、全知的外部知识闭合的。两条路径殊途同归，最终结算出同一个结局：一个被完美喂养、却无法自我生成的主体。

## 6.3 条件的重置：知识—权力—技术复合体如何成为裂隙闭合的引擎

学科的社会动力学在裂隙的视角下获得了一个更尖锐的锚点。营养学的动能——不断推进新研究、产出新指标的驱动力——之所以能无限膨胀，正是因为它所生产的每一个“新发现”——一个新的生物标记物、一个新的“缺乏”阈值、一个新的“最优”区间——都在进一步提供可以替代个体亲自去感知和决策的外部代理数据。这些数据越是精细，那条“我需要亲自去感觉、去判断”的裂隙就越是显得多余和低效。学科的动能，正是一台以“科学进步”为燃料的裂隙闭合引擎。

公众对营养学的势能——那巨大的信任与期待——也同样被这台引擎所驱动。人们之所以越来越相信，一个完美的外部营养方案是可能的并且应该由科学来提供，恰恰是因为他们自己的生成性回路在长时间的代理中被不断弱化。他们越是无法从内部确认“什么对我好”，就越是

将这种确认权外包给外部权威。一个自我强化的循环由此形成：裂隙闭合 → 自主生成能力下降 → 对外部代理的依赖加深 → 更彻底的裂隙闭合。

在这个视角下，已经被广泛讨论的“被代理的生成”这一过程，不再是一个抽象的社会学描述，而是一个可以被精确锚定在裂隙闭合机制上的动力学过程。被代理的生成，不是别的，正是裂隙在宏观社会层面的全面闭合——一个文明级的发生学事件。

## 6.4 结论的重铸：从“处理矛盾”到“守护裂隙”

本章的实践方向由此获得了一个更具体的形态：它不再是一个需要在两股力量之间寻找平衡点的技术问题，而是一个关于如何守护裂隙的敞开性的本体论问题。

守护裂隙，不是要拒绝一切外部知识和技术。恰恰相反，外部知识和技术可以在某个特定的时刻——当个体在裂隙中悬停得足够久、已经充分感知了内在的张力却依然无法决断时——被引入。但它被引入的方式，必须是“赋予个体更多进行自决所需的认知资源”，而不是“替个体完成那个自决动作”。

这二者之间只有一道细细的线，却分隔了两个完全不同的世界。告诉一个人“你的铁蛋白偏低，这意味着你可能需要更多红肉、动物内脏或豆类”是在为他撑开裂隙——他知道了更多，但他依然要自己去感觉、去做选择、去承担后果。而告诉他“你的铁蛋白偏低，已为你推荐每日一片硫酸亚铁，请坚持服用”则是在闭合裂隙——他不再需要感觉，不再需要选择，不再需要承担。他只需要执行。

守护裂隙的实践智慧，就落在如何在这道线上保持清醒的觉知。本章不能提供关于“何时该告知、何时该停止”的标准化答案，因为任何标准化答案本身，都又是一次裂隙的闭合。本章能做的，是命名那道裂隙，让它从一个不可见的盲区，变成可以被有意识地觉察、讨论和守护的发生学边界。

## ■ 七、对质五：生成性裂隙与自律生成——全文理论合法性最关键的战场

### 7.1 最危险的最近邻

一个具有真正穿透力的概念，必须在与其最强的理论竞争者的正面相撞中，指明自身所切开的那个对方装不下的辨别面。对于“生成性裂隙”而言，这个最强的竞争者，不是还原论、不是商业共谋论、不是社会规训论，甚至不是废用性萎缩模型。它是马图拉纳和瓦雷拉于 20 世纪 70 年代提出的自律生成理论（autopoiesis）（马图拉纳和瓦雷拉，1980）[4]。

自律生成理论的核心主张是：生命系统是“自我生成”的——其运作的产物正是其自身的组织本身。一个自律生成系统通过不断地在其内部网络中循环转化物质，生产出维持其自身边



界和功能所需要的所有组件。当这样一个系统在结构耦合的扰动下发生失稳时，它被推向失活边缘；此时，系统必须通过一次内在的、无外部蓝图的“建构性自举”，重新闭合其操作循环，并在此过程中将自身重建为一个比以前更稳定、维度更高的新系统秩序。这套叙述与本章的核心框架——在一次次克服阻力的过程中，生命自我生成并自我强化——在抽象层面存在着惊人的叠合，叠合度可能高达70%以上。

这是本章理论合法性最危险的战场。若不能打赢这场对质，“生成性裂隙”就可能被一句“这其实就是自律生成在营养学上的一个应用案例”所瓦解。

## 7.2 生理层面的高度叠合与裂隙论在解释力上的增益

首先，必须诚实地承认自律生成理论在生理底层与本章的深度共鸣。用自律生成的语言，可以完整而精确地重述本章第三部分的所有生理案例。

在铁的吸收通路中，肠道上皮细胞正是作为一个微型的自律生成系统在运作。食物基质中的植酸—铁复合物构成了来自环境的“结构耦合扰动”——它扰动了细胞当前的铁代谢稳态。细胞对此的回应，不是等待外部指令，而是通过内在地动员一整套分子解码机制——酸化微环境、上调DMT1、动态感应铁蛋白水平——来从这种扰动中重建其铁平衡。这个在扰动中通过内在操作重新闭合自身循环的过程，就是自律生成理论所描述的“建构性自举”。它每成功完成一次，就重新确立了一次细胞作为“能够从复杂环境中获取铁”的功能性身份。

同样的逻辑可以无折损地应用于肠道菌群和抗氧化防御系统。被抗生素清理后的肠道环境是一次剧烈的失稳，而通过多样化的膳食纤维逐步“驯化”出一个强韧的原生菌群结构，正是一次漫长的、需要多重自举的生态重建过程。低至中等水平的ROS对Nrf2通路的激活，则是细胞抗氧化防御系统在氧化扰动下，通过内在基因调控网络重新闭合自身保护循环的精确案例。

至此，自律生成理论完全可以宣称：本章所追踪的全部“生成性阻力”与“机能沉默”现象，都是自律生成系统在特定环境耦合模式下的正常运作或失活。那么，“生成性裂隙”还有什么独特的存在价值？

## 7.3 裂隙论切开的不可还原的新辨别面

本章的回应是：自律生成理论可以完美地解释细胞的“生”——一个活系统如何自我维持。但它无法解释那个感到“我让自己活下去了”的“我”，是如何从这套自我维持的分子网络中，作为“谁”被发生出来的。

自律生成理论描述的是一个生物学事实：一个系统在闭合其操作循环。但从这个事实，到“我确信我做了一个对自己好的选择”这一主体经验之间，存在着一个无法被自律生成理论



的任何内部概念——包括结构耦合、失稳边界和建构性自举——所填平的鸿沟。自律生成理论可以在第三人称的视角下，描述一个生物系统在扰动中重建稳定的客观过程；但它无法以第一人称的视角，说明这个客观过程为何、以及如何，被经验为一个“我”正在那里承受、抉择并承担后果。换句话说，自律生成理论可以解释生理的自我维持，但它无法解释主体的自我确认。

这正是“生成性裂隙”所要钉死的那个不可还原的核心。裂隙不是一个只在生理层面被敞开的、等待被细胞的自律操作所闭合的功能性缺口。裂隙同时是、并且根本上就是：那个在主体层面被敞开为“我不知道该怎么办、但我必须亲自弄明白”的张力时刻。在这个时刻中，不仅仅是细胞在重建其铁平衡，而是一个“我”，在承受着不确定性并做出不可撤回的选择时，被那个选择本身的重量所凿实了。细胞的“生”，与主体的“生”，在裂隙所撑开的那同一个张力时空中，被同一波发生过程所同时结算出来。裂隙不是一个生物学概念加上一个心理学外衣。裂隙是那个让生理的自律生成与主体的自我确认得以被视为同一件事的两面的那个本体论结构。

辨别面在此显影：

- 自律生成理论的核心变量，是系统的自我维持（一个活系统如何持续地生产出维持其自身所需的全部组件）。

- 生成性裂隙论的核心变量，是主体的自我确认（在一次不可代理的、亲自穿越不确定性的经验弧线中，“我”作为那个承担后果的人，被凿实了）。

- 二者在生理底层高度重叠——裂隙论愿意承认，在生理层面，它所说的“裂隙穿越”完全可以用自律生成的语言来翻译。但裂隙论的独特性在于打通了生理底层到主体经验层的那道墙——它揭示了，那条在分子层面维持着系统自我生成的生成性张力弧，在主体层面同时就是“我成为我”的生成过程。两个层面的同一性，正是“裂隙”这个命名所试图钉死的东西。而这一笔，是自律生成理论没有做过、也无法从其理论预设中直接推导出来的。自律生成理论可以没有裂隙，但裂隙论不能没有自律生成——它必须站在自律生成所描述的生理自维持机制之上，才能继续往前走出那不可被还原的最后一步：从“系统的生”到“主体的生”的跃迁。

## 7.4 为什么这一区分不是文字游戏

有人可能会质疑：这不过是给同一个生理过程加上了“主体体验”的装饰，本质上讲的还是同一件事。这个批评误解了裂隙论所做的工作的性质。

裂隙论不是在描述一个已知的生理事实，然后说“对了，它还有个心理维度”。它是在论证：当外部代理系统以善意短路了那道裂隙时，它同时损害了两个东西——它损害了细胞通过自律操作维持机能的能力（这一点自律生成理论完全能够描述），也损害了个体通过亲自穿越

不确定性来确认主体存在的能力（这一点只有裂隙论能够命名）。这两个损害不是两个独立的损伤，而是同一个发生学事件——裂隙被闭合——在生理端与主体端的两笔结算。

这就是为什么“被代理的生成中止”不能被简单地翻译为“自律生成系统在不利的结构耦合下失活”。后者指出了细胞不工作了。前者指出了：当那道裂隙被填平时，那个能从“我亲自让这个身体活下去了”的经验中获得深层自我确信的人，也不存在了。自律生成理论没有词汇来描述这后一种损失，因为它的全部词汇都停留在系统维持的层面。裂隙论创造了一个词汇，不是因为它比自律生成理论“更好”，而是因为它回答了自律生成理论无法提出的一个问题：一个被完美地维持着的活系统，和一个能够确认自己是“活着”的主体，这二者之间的鸿沟，究竟需要什么条件才能被跨越？裂隙论的回答是：那道必须被亲自穿越的不确定性。自律生成理论从未提出过这个问题。

这场对质的结果是：自律生成理论是裂隙论在生理底层最强大的盟友，但它不是裂隙论的上位概念。裂隙论在自律生成的理论地基之上，向上打开了一个自律生成理论所不能打开的问题域——那个从生理的自维持到主体的自我确认之间的跨越。这个跨越不是自律生成的延伸，是一个需要全新命名的、不可还原的发生学跃迁。“生成性裂隙”是这场跃迁的结算。

## ■ 八、装上接口：裂隙的复现、调用与延伸

一个概念如果“说得很对但别人用不了”，它就还只是个漂亮口号。这一部分为“生成性裂隙”装上可操作的分析接口。

### 8.1 一个可复现的情境：裂隙守护的实际操作

以下情境[5]来自真实的临床营养实践，它展示的不是裂隙论的“正确”，而是它如何让人能够看见原本被忽略的东西。

一位中年女性来访者，主诉“总觉得累、没有精力、但体检一切正常”。她在过去三年间，严格按照一款知名营养应用程序的推荐方案安排饮食——每一餐的热量、宏量营养素、微量营养素都精确达标。她每天早晨服用该程序推荐的复合维生素、益生菌和鱼油补充剂。她的生化指标堪称教科书级的标准。但她告诉营养师：“我不知道问题出在哪里，但我觉得我跟我的身体之间，好像隔了一层什么东西。我吃饱了，但我不记得上一次‘吃得满足’是什么感觉了。”

这是一个典型的裂隙系统性闭合的案例。这位来访者的整个营养生成回路——从内感到自承——已经被“完美代理”全线短路。她不再需要去感受饥饿的独特质地，因为应用程序已经告诉她什么时间该摄入多少热量；她不再需要去悬停——在“我想吃什么”和“我该吃什么”之间承受张力——因为每次她打开手机，最优解就已经在那里了；她不再需要自决，也不

再真正承纳后果——如果餐后不舒服，那是“这个方案的配比可能不适合我”，她只需要调整参数。三年来，她从没有为自己“做错”过一顿饭——也因此，她从来没有为自己“做对”过一顿饭。那道裂隙从未被穿越，因为她从未被允许在裂隙中悬停。

营养师没有给她的饮食方案增加任何新的补充剂或调整任何参数。相反，她提出了一个为期两周的实验性建议：关掉应用程序的所有推送通知。不查任何热量数据。在本周的某两天，不要按照食谱做饭，而是在饥饿降临时，站在厨房里，花五分钟让自己去感受：“现在，我的身体到底想吃什么？”然后自己做决定。如果选错了、吃下去不舒服，不要怪应用程序，不要怪食谱，允许自己承认：“我今天做了一个不那么好的选择。”然后把那次不舒服的感觉——腹胀、沉重、或仅仅是“不对”——记下来。

前三天，来访者报告了巨大的焦虑。她发现她站在冰箱前时，脑中是一片空白。“我不知道我想吃什么。我甚至不知道我饿了没有。”她已经丧失了对自己内感信号的解读能力——那些信号没有被身体的精密系统主动处理，而是被外部系统代为翻译了太久。但到第一周末尾，一件意料之外的事发生了。她没有按任何食谱，为自己做了一碗非常简单的热汤面。吃下去之后，她没有计算热量，没有对标“最佳配比”，但她确切地知道——这碗面对了。她在记录中写道：“身体没有变成一组数字。它只是告诉我，这回你做对了。”

这个案例不证明裂隙论“正确”。它展示的是裂隙论能让人看见什么：在那些“完美依从”与“不可名状的空心感”之间的巨大裂缝中，那层“不知道出了什么问题”的迷雾，原来是被一个从未被识别的东西——那道被持续闭合的生成性裂隙——所笼罩的。看见裂隙，本身就是干预的开始。

## 8.2 一套最小分析步骤

对于任何想要将“生成性裂隙”概念应用于具体营养实践、教育场景或政策评估的研究者或从业者，以下四步提供了一个可操作的切入框架：

第一步：识别该情境中裂隙被敞开的<sup>1</sup>关键张力场。问：在这个人的饮食生活中，哪一组矛盾是正在被悬而未决地承受着的？是“想吃甜”与“怕血糖高”之间的张力？是“知道自己该吃得更健康”与“不确定什么才是真正的健康”之间的认知缝隙？是“身体发出的疲惫信号”与“不知道它在说什么”的解读困难？指出那道裂隙，并为它命名。

第二步：盘点当前有哪些外部代理正在填充或短路这道裂隙。是一个自动生成饮食计划的手机应用程序？是一份来自体检中心的“饮食禁忌清单”？是家人每次在他犹豫时直接替他点好菜的习惯？还是他自己养成的“不问身体感受、直接查热量表”的内在决策程序？将这些代理行为一一列出。

第三步：评估这道裂隙被闭合的程度及其可见的后果。问：这个人还有没有机会，哪怕是偶尔，在没有外部指导的情况下，为自己做一次从内感到自承的完整饮食决策？如果没有，他表现出了什么迹象——是对自己身体信号的完全陌生？是面对无标准答案的饮食情境时的巨大焦虑？还是对任何饮食选择的后果都无法产生“这是我做的”的拥有感？

第四步：设计实验性的裂隙重启策略。不是给出一个“更好的替代方案”，而是创造一个低风险的情境，让这个人重新去亲自穿越一次裂隙。例如：建议他在本周的某一天，关掉所有饮食应用程序，根据“身体想吃什么”给自己做一顿饭，然后记下整个过程的感受（包括犹豫、不安、以及吃完后的身体反应）。这个实验的目的不是得到一个“正确的餐食”，而是重新激活那条生锈的生成性回路。

### 8.3 反向接口：裂隙重启的困难——一个发生学悖论的诚实面对

在给出了上述过于整洁的操作接口之后，必须立刻亲手打碎它可能造成的误解。裂隙的重启，远没有上述框架所暗示的那么顺畅。在许多真实的情境中，个体的内感能力已经因长期代理而严重萎缩——他们并非“不愿意”去感受身体，而是身体的信号对他们而言已经变成了一种听不懂的语言。在这种情况下，简单地“退后一步”让个体自己去面对那道裂隙，不是在守护裂隙，而是在将他抛入更深的无力感。他需要的是一个能够帮助他重新学习倾听身体语言的“脚手架”——一种有策略的、暂时的外部引导，而非一个“自己看着办”的真空。

这似乎与裂隙的“不可代理性”刚性约束相矛盾。但它实际上揭示的是裂隙论的一个更深的悖论：在某些极端情况下，为了重启裂隙，必须以一种高度自觉和有边界的方式，使用某种外部代理。这种代理的目的，不是替个体完成那道生成弧线，而是帮助他恢复完成那道弧线所需的最基本的感知和决策肌肉。这其中的分寸极为微妙——它要求实践者有能力判断：何时我在帮他撑开裂隙，何时我已经在替他闭合裂隙？这一判断需要的不只是理论知识，更是一种被裂隙论这个框架所照亮的、对代理行为之发生学后果的持续警觉。

裂隙重启的困难还在于另一个层面：即使内感能力已经重启、个体已经能够在悬停中承受张力，外部环境——那套无时无刻不在推送“最优解”的数字营养生态系统——并没有消失。裂隙的重启，因此不仅仅是“让个体自己做选择”那么简单。它需要在个体层面与环境层面同时进行调整：在个体层面，是内感肌肉的复健；在环境层面，是为那道裂隙的敞开创造一个受保护的空间——一个应用程序和算法暂时不被允许进入的、个体可以“安全地做错”的空间。这两重调整都不容易，且不存在标准化的操作方案——因为任何标准化方案本身，都又是裂隙的一次闭合。裂隙论的实践智慧，最终落脚在每一个独特的营养关系中，在每一个独特的时刻，做出那个不可被提前规定、也无法被算法化的判断：什么时候给答案，什么时候守护问题。

## ■ 九、结语：守护那道必须被亲自穿越的裂隙

本章已完成了全部论证：将“生成性裂隙”从现代营养学困境的深层张力中结算出来，给予它精确的定义与生理—主体双层锚定，与五位最近邻正面对质并划定不可替代的辨别面，用它回写重组了全篇论证，并装上了可复现、可调用、可延伸的分析接口。

在结语处，必须将裂隙的概念推至它所能达到的最深层的实践意义，并亲手打碎任何将其浪漫化的倾向，同时正面回应那个最有力的反驳。

“守护裂隙”，不是一种田园牧歌式的怀旧主张。它不是说，我们应该回到那个没有营养指南、没有食品工业、没有科学检测的前现代世界——在那里，裂隙的敞开固然是充分的，但它同时伴随着真实的营养不良、食物匮乏与食源性疾病。守护裂隙，也不是主张个体拒绝一切外部知识和技术——那样的拒绝本身就是对裂隙的另一种闭合，只不过这次是由一种“反科技的执念”所填平。

守护裂隙，是在我们已经不可逆转地生活在一个由各种外部代理构成的现代性环境中的前提下，提出的一项清醒而艰难的发生学任务：在我们继续发展外部知识体系、优化营养素供给、提升健康管理精度的同时，我们必须学会识别出那道使生命得以作为主体而发生的、不可被任何技术所代理的裂隙，并为它留出空间。

这意味着，在临床营养实践中，每一次我们准备给出一个标准化建议之前，我们需要先停顿一下，问自己：这个人在这一刻，是更需要一个答案，还是更需要一个帮助他去感知、去悬停、去自己做出选择的框架？如果是后者，那么我们的角色，就不是一个答案的提供者，而是一个裂隙的守护者——我们帮助他把裂隙撑开得足够久、足够安全，让他能够亲自穿越它。

这也意味着，在公卫政策和健康教育层面，我们需要发展出一种新的智慧：如何在不放弃科学知识传播的前提下，避免以一种“不容置疑的权威”的姿态，将那道本应由个体亲自去走的生成之路全线铺平。一份好的膳食指南，或许不是那种精细到每一克的“最优配比表”，而是一种能够帮助人们更敏锐地感知自己身体、更自信地做出食物选择的“内在导航图”。

一个必须被正面回应的反驳：“守护裂隙”是否是一种奢侈？

一个最有力的、也是本章最应当正面回应的反驳是：“生成性裂隙”是否在某些情况下，本身就是一种奢侈——只有那些营养已经足够安全、食物资源已经足够丰富、面临的选择压力足够小的个体，才有余裕去“亲自穿越裂隙”？对于一个食物不安全家庭中的主妇、一个正在接受化疗而味觉完全紊乱的癌症患者、一个患有严重饮食失调而丧失了基本内感功能的青年——对于这些人而言，一个直接告诉他们“该吃什么”的外部代理，不是对他们主体性的取消，而是救命的、恢复主体最低功能的必要桥梁。您如何保证“守护裂隙”的善意主张，不会在实践上成为一种加重弱势群体负担的隐性道德要求？



本章的回应是直接的：裂隙论并不一概反对所有外部代理。它反对的是，在裂隙可以被个体安全地穿越的情况下，外部代理系统无差别地将所有裂隙填平。对于那些因疾病、年龄、贫困或创伤而暂时或永久地丧失穿越裂隙能力的个体，外部代理不是敌人，是桥梁——它帮助他们度过无法自我生成的阶段。但这里有三个关键判据。

其一，代理必须有“退回”的机制。当那个患者完成了化疗、那个青年在治疗中恢复了基本的内感功能，代理应当主动后撤，重新打开裂隙。一个永远不后撤的代理——无论它的出发点多么善意——终将导致裂隙的持续闭合和主体生成能力的不必要丧失。

其二，即使在不得已必须进行代理的情况下，代理的形式也应当区分“替他做”和“为他撑开”。“替他做”是在他无法穿越裂隙时，直接给出答案；“为他撑开”是在他尚有能力部分地穿越裂隙时，帮他感知得更清楚些、帮他把裂隙的边界看得更明白些，然后仍然让他自己去走。“替他做”适用于能力完全丧失的情况，而这种情况远比人们通常认为的要少得多。

其三，裂隙论并不将“守护裂隙”设立为一个普遍的道德律令，降在每一个个体头上。它不是一种新的“你应该亲自选择食物”的健康教条——那正是裂隙论所反对的“外部指令”以另一种面貌的回归。恰恰相反，裂隙论为实践者提出的要求是：在看到一个人时，先判断他此刻能否安全地穿越裂隙，再决定是给出答案还是守护问题。对于那些暂时无力穿越的人，给他答案，不要给他负担；对于那些有能力但已被代理得忘了自己有能力的人，帮他把裂隙重新撑开，让他重新看见自己能够走。这二者之间的区分，不是裂隙论的“例外条款”，而是裂隙论实践智慧的真正核心。

#### 可证伪条件的提出

最后，任何严肃的理论，都必须承诺其自我证伪的条件。对于“生成性裂隙”这一核心主张，一种根本性的检验在于：如果未来的多中心、长期、严格的纵向研究能够稳定地证明，一群从出生起便由闭环算法驱动的肠外营养系统终身喂养、从未自主吃过一口食物的人类个体，在成年后表现出与正常饮食成长的对照组相等价——甚至是更优——的自我效能感、身体智慧、以及在面对高度不确定的生存困境时的内在力量的沉着性，并且这些心理维度的优势是独立于外部系统的心理慰藉而产生的，那么，本章的核心主张——即那条必须被亲自穿越的生成性裂隙不可被代理——便被从根本上撼动了。

必须承认，上述这条证伪条件在形式上虽然干净，在实践上却几乎不可执行——没有人能够、也不应该为了检验一个理论而如此养育一群人。一个只能被不可执行的实验所否定的理论，其证伪承诺是空的。因此本章另外给出两条可在近期真正被执行的证伪条件，它们各自负责打断本章因果链的一端：

其一，生理端。即上文第三部分给出的铁蛋白轨迹预测：长期依赖高剂量铁补充剂者与从不补充的对照组在同时停补后的 12 个月内，若血清铁蛋白的下降速率无显著差异，或该差异可被停补时铁储存量的基线差异完全解释，则“机能沉默”作为一种独立于营养素缺乏的生理损失，便失去了它最关键的经验支点。这一检验可在现有队列中完成，无需任何新的伦理安排。

其二，主体端。取两组在营养素摄入水平上匹配的成年人：一组长期（≥24 个月）依据算法化饮食方案进食，另一组长期自主决定饮食。以标准化的内感准确度任务（如心跳感知任务）与内感元认知灵敏度指标测量二者。本章预测：两组在内感准确度上可能无差异，但算法组在元认知灵敏度——即“知道自己的身体感觉有多可信”这一能力——上显著更低，且该差距随代理时长增加而扩大。若测不到这一分离，则“代理性生成短路会侵蚀内感解读能力”这一主体端主张即被否定。

这两条与前述那条极端条件的关系是：极端条件划定理论的外边界，这两条负责在可及的范围内真正冒险。一个理论如果只肯把自己押在做不到的实验上，它就还没有真正上桌。

在这些检验到来之前，我们有足够的理智与谦卑，为我们所无限信仰的外部科学代理的可能性，划下这道审慎的发生学边界。

在那道裂隙仍然对我们敞开着的此刻，守护它。

## ■ 注释

[1] 本章第三部分（及第二、四部分涉及分子机制的段落）在论述铁代谢、Nrf2 通路、肠道菌群重建等机制时，综合了该领域内教科书等级的公认知识和近年若干高影响力研究（如 Suez et al., 2018; Itoh et al., 1997）。为保持论证的递进节奏，正文未对每一判断逐条征引原始文献。文中所作的若干推衍性判断（如长期补剂导致表观遗传沉默、内感信号的分子基础等）系基于已有证据的合理外推，不代表已获确证的事实，读者在引用时应注意识别其论证性质。

[2] 本章第五、七节所对质的“还原论局限论”“商业共谋论”“社会—文化规训论”等，系对相应学术传统的概括性定位，而非对某一特定文本的专论。这些传统的代表路径包括：整体论传统对还原论的批评；政治经济学与文化批判路径对食品工业的共谋分析；福柯式身体规训与医学社会学路径对现代身体管理的分析等。本章意在与此类问题式进行理论对质，而非综述该领域全部文献。

[3] 安玛丽·摩尔（Annemarie Mol）在《照护的逻辑》（2008）中，通过荷兰糖尿病照护的民族志研究，区分了“选择的逻辑”与“照护的逻辑”。本章仅在此概念区分上与之对



话，不涉及其对具体医疗实践的规范主张。本章对“好的照护”提出的发生学追问，已溢出摩尔原著的分析范围，属于本章基于裂隙论的独立推展。

\[4\] 马图拉纳（H. R. Maturana）与瓦雷拉（F. J. Varela）的自律生成理论（autopoiesis）最初在1972年以西班牙语发表，后于1980年出版英文著作《Autopoiesis and Cognition: The Realization of the Living》。本章所述的自律生成系统的运作特征，如生产自身组件、组织闭合、在扰动中重建稳定等，依据该理论的核心主张。文中使用的“建构性自举”“结构耦合”等术语系作者基于理论逻辑的诠释，非严格引文。本章将该理论拓展至“主体的自我确认”层面的尝试，已超出原著者的明确论述范围，由此可能产生的争议由本章承担。

\[5\] 第八节所述临床案例系综合多位营养咨询实践中的典型情境，经匿名化、背景简化和细节合成处理，用作思想例示而非可溯源的实证个案。其目的在于展示“生成性裂隙”框架在日常营养实践中可能打开的辨识空间，而非提供经验证据对理论进行检验。

## 第五章 善失去了自己的语法

本章原题《善的语法消亡：强迫行善制度如何侵蚀道德主体的内在繁衍能力》，作者 高鹏，原文见 <https://sdeuniverses.com/students/gao-peng/grammar-of-good/>

### 摘要

本章挑战了法哲学中一个长期未被充分检视的预设：以法律强制公民行善，即便在操作上面临困难，其方向在道德上是可取的。通过对“禁止做坏事”与“强迫做好事”两类规则的判据结构进行比较分析，本章提出一个形态命题：强迫行善制度的真正终局，不是制造伪善者或摧毁社会信任，而是在代际尺度上不可逆地摘除了“善”作为一种不依赖制度许可即可自我繁衍的内在伦理能力所必需的公共耦合件——被他者确认、在公共语言中命名、在身体间传递的道德语法。本章构建了三阶段退化模型（行为指标替代、公共语法垄断、道德主体失语），指出制度为堵漏而持续精细化的代理自噬机制，会系统性地将善的公共语言从主体间的自发确认，替换为制度性的合规信号系统。当这一进程越过“翻译垄断”临界点后，个体的道德认知结构因长期失去外部他者确认的耦合，发生不可逆的错轨发育——身体对他人痛苦的镜像反应仍在，但将其加工为“我想帮人”这一道德意图的语言能力已被腐蚀。本章据此预测：在运行超过两代的道德积分或善行激励制度覆盖人群中，将可观测到道德叙事从第一人称（“我想帮”）向第三人称（“这是被鼓励的行为”）的代际退化。本章最后给出了可操作的经验检验方案，并讨论了该命题在不同文明传统中的适用边界。

关键词：法律义务；道德动机；公共语言；道德认知发展；制度退化

### 1 引言

#### 1.1 问题的提出

为什么几乎所有成熟的法律体系，都选择“禁止人做坏事”作为基本运作原则，而从未将“强迫人做好事”确立为可强制执行的法律义务？这个问题表面上关乎法律与道德的边界，但其深层结构触及一个更根本的议题：规则在试图捕获人的内部状态时，其自身的运作逻辑会发生什么不可逆的改变？

常识给出的回答通常落在价值排序的层面：自由比强制更值得珍视，消极义务比积极义务更容易执行，法律守住底线即可，道德高地应由个体自愿攀登。这些回答并非错误，但它们在一个关键点上止步了：它们将“不能强迫行善”表述为一种规范选择（我们选择不这样做，因为我们珍视自由），而非一种结构必然（即便我们选择这样做，制度自身的运作逻辑也会将“善”这个概念引向自我取消）。

本章旨在论证后者。本章的核心命题是：强迫人做好事的制度之所以不可能成为一项可持续的伦理-法律安排，不是因为自由比强制在价值上更可取，而是因为“善”的本体构成中包含一个不可被外部规则穷尽的内部状态——动机。当制度试图通过代理指标（行为表现、可观测结果、他人评价）捕捉这一内部状态时，它会陷入一个自我加速的精细化螺旋：每一次为堵漏“伪善者”而增设的新规则，都会在最精确的层面暴露一个新的、可被表演的缺口，从而制造出下一轮精细化的驱动力。这个自噬环不会因制度自身的认知突破而停止，只会在操作成本饱和时累停——而累停的那一刻，善的公共语言已被制度代理系统完整覆盖。

但本章的诊断不止于此。本章要进一步追问的是：在公共语言被制度垄断之后，个体内部用以生成“真善”的那个认知结构，它自身会发生什么？本章的答案是：那个结构会因为长期失去外部他者确认的耦合，而在代际尺度上发生不可逆的错轨发育。善，将从一个依靠人际自发传递的“手艺”，退化成为一种必须依赖制度持续喂食外部激励才能被执行的合规反射。强迫行善制度的最终产物，不是更多的好事，而是更少的、有能力独立做出道德判断的主体。

## 1.2 现有研究的概要

现有关于法律与道德义务边界的研究，大致可以归纳为三条主要路径。

第一条路径是法哲学内部的规范分析。自密尔（Mill）在《论自由》中确立“伤害原则”以来，自由主义传统为法律不干预纯粹道德领域提供了坚实的论证：法律仅当某一行为对他人构成可辨识的伤害时才获得干预的正当性，单纯的不道德行为不在其管辖范围之内。哈特（Hart）在与德富林（Devlin）的著名辩论中进一步捍卫了这一立场，反对以“维护社会道德结构”为由强制执行道德。然而，这条路径的核心在于界定干预的道德界限，它并未深入追问：如果法律越过了这条界限，试图强制执行道德，强制执行这一动作本身会对“道德”这个概念产生何种结构性影响。

第二条路径来自制度经济学和比较制度分析。这一传统关注规则的可执行性（enforceability），强调一条规则要想持续有效，必须具备可验证的判据。消极禁令（禁止杀人、盗窃）因其判据为外部可见的身体行为而容易执行；积极义务（如法律规定公民有救助义务）则面临信息不对称和验证成本的问题。这一分析进路提供了重要的技术性解释，但它倾

向于将“不可执行”视为一个技术困难——似乎只要监控技术足够发达，执行积极义务的障碍就可以被克服。本章要质疑的正是这一预设。

第三条路径来自道德心理学和发展心理学。布莱尔（Blair）、莫尔（Moll）等学者的神经影像研究表明，道德判断涉及多个脑区的协同工作，其中包括对他人痛苦产生自动反应的区域（如前脑岛、前扣带回）。霍夫曼（Hoffman）的共情发展理论更追溯了个体从婴儿期的“共情痛苦”到成熟道德推理的发育路径。这些研究揭示了一个被法哲学传统长期忽略的事实：道德判断不是一个封闭的理性计算过程，它在个体发育中深度依赖于外部他者的确认——母亲的回应教会婴儿“什么是痛苦”，同伴的反应完成着慷慨行为的社会定义回路。然而，这一心理学发现迄今尚未被系统性地引入对法律强制行善这一制度设计的结构分析中。

此处必须补上一片本章迄今未曾正面处理、却离本章命题最近的实证文献：动机排挤与过度理由效应。自德西（Deci, 1971）发现外部报酬会削弱受试者对原本感兴趣的活动的内在动机、莱珀等人（Lepper, Greene & Nisbett, 1973）在儿童绘画实验中观察到「预期奖赏」组在奖赏撤除后绘画时间显著下降以来，这一效应已经积累了半个世纪的证据，并由德西、科斯特纳与瑞安（Deci, Koestner & Ryan, 1999）的元分析加以系统确认；蒂特马斯（Titmuss, 1970）对有偿与无偿献血制度的比较、格尼齐与鲁斯提基尼（Gneezy & Rustichini, 2000）对以色列幼儿园迟到罚款反而使迟到率上升的实地研究、以及弗雷与耶根（Frey & Jegen, 2001）对动机排挤的综述，共同确立了一个与本章高度共振的结论：外部激励可以侵蚀它本想强化的那个内在动机。不引这片文献而声称本章的诊断「迄今鲜被系统阐述」，是站不住的。因此必须把本章与它的分界说清楚。动机排挤研究的分析单位是个体在一次任务中的动机结构：给了钱，我做这件事的理由从「我想做」变成了「有钱拿」；撤了钱，行为随之衰减。它是一个在个体内部、在单次或短期激励条件下可被反复实验验证的效应，其经典实验的时间尺度是几周到几个月。本章所论的不是这一层。本章追问的是：当激励制度覆盖了一个社会用来谈论和确认善的全部公共语言之后，下一代在道德发育期还能不能获得把身体反应加工成「我想帮他」的那套语法。动机排挤描述的是一个已经拥有内在动机的人如何失去它；本章描述的是一个从未获得过那套语法的人如何在发育期就没能长出它。前者是消退，后者是未发育；前者在个体尺度上是可逆的（撤除激励后内在动机常可部分恢复），后者在代际尺度上不可逆，因为发育窗口不会重开。这一分界同时给出了一条判别线：如果实证发现，在激励制度撤除之后，被覆盖群体的第一人称道德叙事在数年内自行恢复至对照组水平，那么本章所描述的现象就应当被归入动机排挤的管辖范围，本章关于「代际不可逆」的核心主张即告失败。

### ■ 1.3 本章的研究问题与分析策略

鉴于现有研究的三条路径各有其盲区——法哲学的规范分析不追问“强制执行道德会对道德自身产生什么结构影响”，制度分析将“不可执行”视为技术困难而非结构死结，道德心理学止步于个体层面而未触及制度对道德认知结构的反向塑造——本章提出以下三个递进的研究问题：

第一问：为什么强迫行善的规则在演化序列上不能收敛？禁止坏事的规则为何能够避免这一困境？二者的根本差异何在？

第二问：如果强迫行善的制度不能收敛，历史上那些看似持续运转的类似制度（如道德积分、善行激励）是如何维持的？它们的维持方式是否仅仅是找到了收敛的技术解法？

第三问：这种维持方式在长时段上对“善”作为一种社会存在形态和个体内在能力产生了什么不可逆的改变？

本章的分析策略采取概念分析与制度形态学相结合的方法。一方面，通过对“禁止规则”与“强迫行善规则”的判据结构进行精细比较，追溯二者在制度演化序列上的分歧点；另一方面，引入来自认知神经科学和发育心理学的跨学科证据，揭示制度对个体道德认知结构的反向塑造机制。此外，本章还将借助演化语言学的视角，分析当“善”的公共语法被制度性指标系统垄断后，道德主体内部发生的不可逆的语言-认知错轨。

## ■ 1.4 本章的贡献预告

本章的理论贡献主要体现在三个方面。第一，将法哲学中“法律不能强制执行道德”这一传统命题，从规范边界的讨论推进为制度形态学的诊断——不是法律“不应”强制道德，而是“强制”这一动作在结构上会取消“道德”的构成性要素。第二，将道德心理学的个体层面发现（他者确认在道德认知发展中的构成性作用）上移至制度层面分析，揭示制度对道德认知结构的反向塑造，从而在发育心理学与法哲学之间架设一座迄今缺失的分析桥梁。第三，给出了强迫行善制度三阶段退化模型（行为指标替代、公共语法垄断、道德主体失语），为该领域的经验研究提供了可操作的诊断框架和可证伪的预测。

在经验贡献层面，本章的退化模型为正在推行或计划推行道德积分、善行激励制度的地区，提供了一个事前诊断工具：在制度覆盖率超过某阈值之前，应当通过哪些指标监测退化进程是否已被启动。在方法贡献层面，本章将概念分析与跨学科资源调用相结合，示范了一种将规范性问题转化为可检验的形态命题的分析策略。

## ■ 1.5 论文结构概览

本章余下部分安排如下：第二节对相关文献进行综述，定位现有研究三条路径的核心贡献与各自盲区。第三节构建理论框架，给出核心概念的操作化定义，并建立三阶段退化模型。第四节阐述本章的方法论基础。第五节是论证主体，分三个分论点展开——判据类型的分岔如何导致截然不同的制度演化路径（5.1），代理自噬作为制度无法停止精细化的内生驱动力（5.2），以及公共语法被垄断后道德主体的失语与错轨发育（5.3）。第六节给出综合讨论，包括可经验检验的预测，以及对“混合制度”情形的审慎讨论。第七节为结论。

## ■ 2 文献综述

### ■ 2.1 法哲学中的伤害原则及其局限

自密尔 1859 年在《论自由》中提出“伤害原则”以来，自由主义法哲学为划定法律干预的道德边界提供了一个强有力的判准：只有当一个行为对他人造成伤害时，社会才有正当理由通过法律加以干预；仅涉及行为者自身的道德过错，不在法律管辖范围之内。这一原则在 20 世纪中叶经由哈特（H. L. A. Hart）与德富林（Patrick Devlin）的著名辩论得到进一步巩固。德富林主张，社会有权以法律维护其道德结构，正如社会有权以法律维护其政治结构一样；哈特则捍卫密尔的原则，指出德富林没有在“社会赖以生存的道德”与“社会成员恰好持有的道德信念”之间做出必要区分。

这场辩论确立了自由主义法哲学在处理法律-道德关系时的基本立场：法律不是道德的强制执行工具。但仔细审视这场辩论的结构，可以发现论辩双方共享一个未经挑战的预设：强制执行道德的可行性本身是次要的，关键在于这么做在原则上是否正当。哈特反对德富林，不是因为强制执行道德“做不到”，而是因为它“不该做”。制度可行性——即当法律伸入道德领域时，其自身的运作逻辑是否会发生不可逆的变形——没有被纳入辩论的核心议程。

这一盲区并非偶然。法哲学在其经典形态中倾向于将“制度”视为一个中性的执行工具：它执行已决定的规范，其自身的运作动态不构成对规范本身的实质性反作用。但本章将论证，这一预设是错误的。当制度试图强制执行某一类特定性质的行为——尤其是那些必须以特定内部状态（动机）为构成性要素的行为——时，制度的运作动态会对这类行为的社会存在形态产生结构化、且在一定阈值后不可逆的改变。法哲学的规范分析需要被制度形态学的诊断所补充：不是法律不该管道德，而是法律一旦管了，“道德”在公共空间中就不再是原来那个东西了。

### ■ 2.2 规则的可执行性：制度分析进路及其隐含预设



如果说规范分析止于“不该”，制度分析则提供了“不能”的解释。从法律经济学和比较制度分析的视角出发，一条规则能否持续有效运转，取决于它是否具备两个关键特征：判据可验证性（verifiability）和执行成本可负担性（affordability）。

消极禁令之所以易于执行，是因为它的判据落在外部身体行为上——一个人是否杀了另一个人、是否侵入了他人财产，这些是可以由第三方验证的二值事实。积极义务则面临一个棘手的“内部状态问题”：许多道德上被称赞的好行为，其“好”恰恰在于行为者是在没有外部强制的情况下主动选择的。如果法律要求一个人必须出于善意去帮助他人，那么执法者就必须确认“善意”这一内部状态是否真实存在——然而内部状态无法被外部观察直接验证。即便立法者退而求其次，只强制“做出善的行为”而不问动机，当人们普遍因强制而做出这些行为时，其社会意义已经发生了质变：一个被枪指着去扶老人的动作，在物理层面与自愿扶老人别无二致，但在社会语义层面是两个完全不同的动作——前者是“求存”，后者是“施善”。

制度分析的盲区在于，它倾向于将这一问题处理为“信息不对称”，即一个可以被更精确的监控技术逐步克服的技术困难。如果这是对的，那么随着行为监控技术、大数据分析、心理测量工具的发展，法律强制行善的技术障碍最终将被消除。但本章要论证的是，问题不在信息的“够不够”，而在“善”这一概念的构成性结构：善之所以为善，正是因为它在逻辑上包含一个不可被外部穷尽的“内部剩余”——动机的自发性。任何试图用外部指标捕获这个剩余的尝试，都会把善从“自发事件”改写为“合规行为”。这并非技术问题，而是本体构成问题：你无法用不断上升的精度去逼近一个在定义上就排斥外部完成的东西。正如你无法用越来越精细的化学分析去找到一幅画的美——不是技术不够好，而是“美”在你的分析框架中不是可以被找到的那种东西。

## ■ 2.3 道德认知发展的具身性与社会耦合：心理学的洞见未被整合

过去三十年间，认知神经科学和发育心理学积累了大量关于道德判断非纯粹认知性的证据。达马西奥（Damasio）对腹内侧前额叶皮层受损患者的研究表明，这些患者在道德推理测试中能准确讲出“应该怎么做”，但在真实生活情境中却系统性地做出糟糕的道德决策——他们的“知道”没有身体情感的参与，因此无法驱动行动。布莱尔（Blair）的暴力抑制机制模型进一步指出，对他人痛苦信号的加工——特别是杏仁核对恐惧表情的反应——是道德社会化的重要神经基础。莫尔（Moll）等人的功能磁共振研究显示，当被试观看道德相关场景时，同时激活了皮层（认知评估）和皮层下区域（情感唤起），道德判断更像是一种“体-脑”协同活动而非纯粹理性运算。



在发育端，霍夫曼（Hoffman）的共情发展理论追踪了一个更能说明问题的路径：婴儿在听到另一个婴儿的哭声时会跟着哭——这是人类道德发育的原初起点，一个不需要任何规则教授的、自动的身体间共振。从这一点出发，经过童年期的他者导向共情、青少年期的抽象道德原则内化，直到成年期形成稳定的道德意向系统，整个发育路径上的每一个关键节点都离不开外部他者的在场与应答。当婴儿因他人的痛苦而感到不安时，看护者的安抚不仅平息婴儿的痛苦，还完成了一个关键的认知耦合动作：它向婴儿传递了一个信息——“你感受到的那个东西，名叫痛苦，它是可以被另一个人看见并回应的”。这一耦合在随后的发育阶段反复出现：同伴对你慷慨行为的羡慕神情、旁人对你勇敢举动的不经意点头、事后你向自己讲述“我为什么要这样做”时使用的那套公共语言——它们都不是道德判断形成后的外部奖赏，而是道德判断自身在形成过程中不可剔除的构成材料。

然而，心理学的这些发现迄今仍然停留在“个体发展”的分析层面，尚未被有效地引入对制度设计的结构分析中。这一脱节造成了一个后果：法哲学和制度分析在讨论“法律能否强制善行”时，仍然把个体当作一个内部自足的道德计算器——只要他有足够的认知能力，他就能在内心法庭里做出判决。他们忽视了那位母亲的目光、同伴的神情、公共语言中的沉淀——这些不是道德计算器的可拆卸外壳，而是道德计算器得以运行的电路本身。当制度代理系统用“合规积分+2”替换掉这些耦合件后，那个计算器不是变得更精准了，而是正在静悄悄地短着路。

## ■ 2.4 演化语言学视角：当公共语法被系统垄断

对本章论题而言，还有一组通常不被纳入法哲学讨论的资源至关重要——演化语言学关于“语言作为认知耦合工具”的研究。语言远不只是个体表达思想的工具，它是人类认知得以从个体内部延伸至社会协作的关键接口。托马塞洛（Tomasello）的人类认知演化的共享意图假说指出，人类幼儿与黑猩猩幼崽在物理问题解决能力上并无显著差异，但在需要与他人建立共同注意、共享意图的社会学习任务中，人类幼儿表现出压倒性优势——而这种优势的核心基础，正是人类独有的一套公共语言与符号系统，它使个体之间能够共享认知目标、协调注意资源、传递关于第三方对象的信息。

这一视角为本讨论提供了一个关键分析工具：当制度代理系统逐步用合规术语覆盖善的公共语言时，它所做的不仅是“改变说话方式”，而是在结构上阻断了主体间借由语言完成道德意图共享的那个认知耦合通道。一个母亲不再借助公共语言中沉淀的他人经验来告诉孩子“什么是善良”，而只能引用制度手册里的合规定义。那个手册里的定义当然也是语言，但它是一种被制度垄断了确认权、因此失去了他者经验厚度的单质语言——它不携带“曾经有无数人在

类似的时刻做出过类似的选择、并在事后对其赋予了意义”这层人际沉淀。失去了这层沉淀的语言，还能否完成它在道德认知发育中一直承担的那个耦合任务？演化语言学的视角提示我们：应该不能。这一视角在前述三条研究路径中几乎完全缺席，而本章的理论框架将把它整合进来。

## ■ 2.5 文献缺口的精确定位与本章的理论定位

综合以上审视，现有研究在以下三个方面存在衔接盲区。第一，法哲学的规范边界分析与制度执行可行性分析之间，缺少一个将制度运作动态视为反向塑造“道德”概念本身的制度形态学环节。第二，道德心理学的个体发育发现与制度层面的长时段分析之间，缺少一个将外部他者确认的构成性作用从个体发育延伸至代际传递的分析框架。第三，上述两条线索的交汇点——制度代理系统对公共语言的垄断如何反作用于个体道德认知结构的长期维持——几乎无人涉足。

本章的理论定位恰恰落在这三个盲区的交汇点上。本章将论证：强迫行善制度之所以注定失败，不是因为某种单一的技术困难或规范错误，而是因为（a）善的本体构成中有一个不可被外部规则穷尽的内部剩余，（b）制度为捕获这一剩余而启动的代理自噬会系统性地垄断善的公共语言，（c）公共语言被垄断后，个体用以生成善的那个认知结构因长期失去外部他者确认的耦合而发生了不可逆的错轨发育。这三步合在一起，构成了一个迄今为止尚未被以这种精确方式陈述的形态判断：强迫行善制度摧毁的不是人们行善的意愿，而是人们行善的能力——那种不需要制度审批就能在个体内心自我繁衍的、以第一人称驱动的伦理能力。

## ■ 3 理论框架

### ■ 3.1 核心概念定义

在展开论证之前，必须先对几个核心概念做出操作化定义。

禁止做坏事的规则（下称“禁恶规则”）：以外部可见的身体行为为唯一判据、不要求对行为者内部状态（动机、情感、态度）进行确认的法律或准法律规则。其运作逻辑是：手不过界即可，心里怎么想不在管辖范围之内。

强迫做好事的规则（下称“强迫行善规则”）：以行为者做出道德上被褒扬的行为为要求、且将行为者的内部状态（善的动机、真实意愿）纳入判据或奖励体系的法律、行政或制度化激励规则。当规则不直接要求动机审计、但通过激励将善行转化为可被外部计算的行为指标时，本章称之为“弱强迫行善规则”；当规则明确要求对动机进行审计时，本章称之为“强强迫行善规则”。在大多数现实运作中，制度会以弱形式登场，但在运作中逐渐向强形式逼近。

代理变量：制度为捕获一个不可直接观测的目标状态（如“善的动机”）而使用的、可被外部观测和记录的替代性指标。例如，将“帮助他人次数”“志愿服务时长”“捐款金额”作为“善行”的代理变量。

代理自噬：指制度在使用代理变量过程中出现的精细化正反馈循环——每一次为堵漏“伪善者”（以代理变量指标达标但缺少目标内部状态的行为者）而增设或细化代理变量，都会在新的精度层面制造出可被表演的缺口，从而制造出下一轮精细化的驱动力。这个循环不会因制度“认识到善不可被代理变量穷尽”而主动停止——因为认识到这一点并不给出一个可操作的行动方案——只会因操作成本超过制度承受阈值而累停。

公共语法：指一个社会中被广泛共享的、用于命名、讨论和评价道德行为与道德动机的语言和符号系统。它包括但不限于：道德语词（“善良”“好心”“仗义”）、道德叙事模板（“他本可以不这样做的……”）、以及道德行为的社会确认仪式（感谢、赞许的目光、在场者的默默点头）。

善的自我繁衍能力：指道德主体不依赖外部强制激励，能够从自身内部产生稳定的道德意图、并据此做出道德行为的持续性倾向。本章将其概念化为一种需要特定社会耦合件才能跨代维持的认知-情感复合能力，而非一种密封的、可在个体内心完全自足运行的内在种子。

### ■ 3.2 命题体系

本章的论证可被形式化为以下四个递进的核心命题：

命题一（判据结构分岔）：禁恶规则与强迫行善规则在判据结构上的根本差异在于，前者的所指具有外部身体边界，由物理世界提供天然精度天花板；后者的所指包含不可被外部穷尽的内部状态，因此制度必须通过代理变量进行无限逼近。

命题二（代理自噬）：当制度以代理变量捕获内部状态时，每一次为堵漏“伪善者”而增设的精细化措施，都会在更高精度层面暴露新缺口，且新缺口的可识别度随精度提升而提高，因此精细化的驱动力不衰减而递增——此即代理自噬。

命题三（翻译垄断与公共语法消亡）：代理自噬持续运行至某一临界点后，制度代理系统将覆盖社会中对“善”这一概念的公共确认权力——善的公共语言从“主体间自发确认”转变为“制度性合规信号系统”。这一转换导致善的公共语法发生不可逆的类型转换：善失去了其在公共空间中被辨识为“道德事件”而非“合规事件”的语言条件。

命题四（道德主体的失语与错轨发育）：当善的公共语法被制度垄断后，个体发育期所需的道德认知耦合件——被他者确认、在公共语言中命名——被系统性地抽走。长期缺乏这一耦合的个体，其道德意图的生成能力将发生从“第一人称驱动”（“我想帮人”）向“第三人称识

别”（“这是被鼓励的行为”）的代际退化。这种退化的终局，不是道德行为的绝对减少，而是道德行为失去了其作为“道德事件”的构成性要素——行为者主观上将其体验为自我决定的道德选择的能力。

### ■ 3.3 跨学科理论资源及其调用有效性

上述命题体系的构建调用了三个通常不在同一分析框架内同时出场的理论传统，在此需要显式说明每一次调用的有效性及其边界。

第一，法哲学中的“规则判据外部性”理论。哈特在《法律的概念》中区分了初级规则（施加义务）与次级规则（承认、改变、裁判），并强调了规则的可辨识性对于法律制度稳定性的基础作用。本章将这一洞见前推一步：规则的可辨识性不仅是一个“技术便利”问题，更是一个结构性问题——判据的类型（外部可见 vs 内部依赖代理变量）预先决定了规则的收敛或发散命运。这一调用是有效的，因为法哲学传统已经提供了进行这种精细判据比较的概念工具，只是此前未将其推向制度形态学的方向。

第二，认知神经科学与发育心理学中的“他者确认耦合”理论。如前所述，霍夫曼的共情发育、达马西奥的体感标记假说、托马塞洛的共享意图理论等研究表明，道德认知不是一个封闭系统，而在发育期深度依赖于外部他者的情感应答与语言耦合。本章将这一发现从个体发育层面延伸至制度分析层面，其有效性在于：如果个体发育期需要这些耦合件，而制度用代理变量系统性地替换了这些耦合件，那么我们有理论依据预测，个体的道德认知发育轨迹将因此发生可观的偏离。这一延伸的有效边界在于：它仅适用于善的发育期依赖深度他者耦合的那些道德能力——那些偏于“技能性”的道德习惯可能不依赖同等深度的耦合——因此本章的命题主要指向狭义上“以主动善意动机为构成要素的善行”，而非全部规范行为。

第三，演化语言学中的“共享意图与公共语法”理论。托马塞洛比较儿童与黑猩猩的认知实验提示，公共语言不只是描述道德事件的工具，而是道德意图在主体间得以建立共享性的构成性媒介。本章将这一洞见引入对“制度垄断公共语言”后果的分析，其有效性在于：如果公共语言是建立共享道德意图的基础设施，那么当这一设施被合规术语系统覆盖后，可预期的后果就不仅是“沟通方式”的改变，而是“共享意图能否被建立”的结构性困难。这一调用的有效边界在于：语言不是唯一的耦合媒介——眼神、身体在场、沉默的仪式也可以承担部分耦合功能——因此在语言被垄断的群体内部，仍可能存在残余的、非语言的耦合通道，这是本章第四节将讨论的“残余层动力学”问题。

## ■ 4 方法论

## ■ 4.1 研究设计

本章是一项概念分析与制度形态学相结合的理论论证，不依赖一手经验数据，但要求论证在以下几个层面经受严格检验：概念区分必须清晰可操作、推演逻辑必须链条完整、跨学科调用必须显式说明有效边界、最终命题必须以可证伪的形式给出。

为此，本章的研究设计包含以下四个步骤：

第一步：判据类型的概念解剖。对“禁恶规则”与“强迫行善规则”进行结构性比较，识别二者在判据空间精度结构、收敛/发散特性和对内部状态的依赖程度三个维度上的根本差异。

第二步：制度演化序列的追溯。建立一个关于强迫行善规则在持续运作中必将经历的退化轨迹的理论模型，包括判据替代阶段、代理自噬闭环阶段和翻译垄断阶段。每一步必须给出其驱动力来源和过渡条件。

第三步：主体侧反向塑造的整合。引入认知神经科学和发育心理学的现有发现，分析制度代理系统对公共语言的垄断如何反作用于个体道德认知结构的长期维持。这一步的关键在于于心理学的“个体发育”时间尺度与制度分析的“代际”时间尺度对接。

第四步：经验检验的前瞻性设计。将本章的核心命题转化为一组可证伪的经验预测，并为未来的经验研究者提供可操作的研究设计方案——包括关键变量的操作化、可观测指标的选择、以及可以推翻本章理论的实证发现。

## ■ 4.2 分析策略

本章的分析策略以概念分析为骨架，以跨学科证据为辅助验证。概念分析的任务是理清“禁恶”与“强迫行善”两种规则类型在判据结构上的根本差异，并由此推导二者在制度演化中注定分岔的必然性。在概念推理触及个体认知结构变化时，引入神经科学和心理学的已有发现作为辅助验证——但这些发现仅用于佐证推理的合理性，不构成独立的经验证据。本章同样在相应位置显式承认这一方法固有限制：本章所做的是构建一套有解释力和预测力的理论框架，其最终检验有待后续经验研究。

## ■ 4.3 方法局限

在给出论证之前，有必要预先承认本章的方法论局限。第一，本章属于概念分析与理论建构，其核心命题尚未经受系统性的经验检验。尽管第五节给出了可操作的经验检验方案，但本章本身并不提供一手经验数据。第二，本章对“强迫行善制度”的分析，主要以现代官僚制度中具有集中化善行判定权的制度形态为隐含参照，这一定位未必适用于善行判定权高度分散、由社区成员共同行使的社会结构——因此本章命题的跨文明适用范围需要后续研究加以厘清。

第三，本章在论证中主要使用“强迫行善”与“禁止做坏事”的二值对比来突出结构差异，但现实制度大多是两种逻辑的混合体（例如刑法在量刑阶段对悔罪态度的考量），本章对“混合制度”情形的讨论仅在 5.4 节做了初步处理，尚未进行完整展开。这些局限将在结论部分被再次确认。

## ■ 5 分析与讨论

### ■ 5.1 判据的类型如何锁死制度的命运

#### ■ 5.1.1 禁止规则的精度天花板：物理边界作为收敛器

一条规则要能在大规模社会中持续运转，必须对其管辖对象的边界给出可操作的定义。“不得杀人”——这四个字能够在跨越文化、跨越世代的人群中被相对一致地理解和执行，并非因为人们对“杀”的伦理含义达成了高度共识，而是因为“杀”的行为结果有一个物理世界提供的天然收敛点：一个人的身体由活的状态变为死的状态。这一个二值事实信号在绝大多数情境下可以由第三方经验验证，不需要深入行为者的内心去确认“你是不是出于恶意”。即便在边缘案例——如安乐死、正当防卫致人死亡——需要进一步审查动机和情境，在这些案例中法律也先确认了“致死行为”这一外部事实，再将动机审查作为下一步（量刑阶段）而非第一步（定罪阶段）的议题。外部事实是地基，动机审查是在这块地基上再起的小楼；地基不依赖于小楼而独立存在。

这种“先确认外部事实，再进入内部状态”的二阶结构，赋予禁止规则一个天然精度的天花板：你不需要持续加码更细粒度的证据规则去界定“什么是杀”，因为身体从活到死的那个间断点本身已经足够收敛。这张天花板虽然不完美（边缘案例有争议），但它具有物理世界赋予的稳定边界——它不是制度自己给自己设定、然后可以不断下移的，它是被“身体边界”这个外部事实锁死的。

#### ■ 5.1.2 “善”的代理困境：没有天花板的精度需求

现在来看强迫行善规则的情形。“做好事”——即便立法者只要求“做出善的行为”而不追问动机，也无法绕过“善”在概念上一项硬性要求：一个行为之所以是善行，不仅因为它的外部后果是好的，还因为它是行为者出自善意动机而做出的。这不是伦理学上的苛求，而是日常语言中“做好事”这个词组的隐含构成条件。如果一个人扶起摔倒的老人是因为公司规定“不扶就扣工资”，我们在描述这个场景时会迟疑——我们会说“他扶了老人”，但不太会说“他做了一件好事”。也就是说，在日常语言的使用规则中，“做好事”这个词组的恰当使



用条件包含了对行为者动机的追溯——它要求行为不是仅仅在外部形态上符合善行的描述，而是在动机来源上发自行为者的道德自觉。

这意味着什么呢？意味着强迫行善的规则面临一个禁恶规则从未遇到过的结构性挑战。禁恶规则只需要确认“手有没有过界”，这是一个有限问题。强迫行善规则——即便它一开始声明“不问动机，只看行为”——一旦被制度化，就会立刻面临一个追问：一个被制度强制做出的“善行”，它还是善行吗？如果立法者回答“是”，那立法者就是在重新定义“善行”这个词的公共语义，把一个原先要求动机自发性的事件概念替换为一个只看外部合规的绩效概念。如果立法者回答“不是”——承认强制做出的善行在严格意义上不是善行——那立法者强制人们做一件不算善行的事情，这项立法的正当性就自我解构了。

这就是强迫行善规则在根源处的困境：它要么承认自己的强制所产生的东西不是善行（从而自我否定），要么修改“善行”的公共语义将其降格为合规行为（从而把善蒸发掉）。这一困境并非哲学上的苛求，而是“善”这个概念在日常语言中的使用规则所决定的——它有机追溯硬要求，而这个硬要求无法通过技术升级被绕开。无论监控技术多么精密，你始终是在“外部”打转，而善的动机要求的是“内部”——这两者之间存在的不是程度差距，而是类型鸿沟。

### ■ 5.1.3 两条路径的不可通约性

由此可以看清禁恶规则与强迫行善规则在制度演化路径上的根本分歧：禁恶规则有一块由外部事实锁死的地基，它可以在“先定罪，再审动机”的二阶结构中运作，动机审计是辅助性的、边缘性的，不构成对整个规则类型的结构性威胁。强迫行善规则没有这块地基——它的地基也是必须伸进动机里去打桩，而那里没有可打桩的地方。因此，强迫行善规则面临的不是“执行困难”，而是本质困境：它的运作逻辑本身就是自反的，它要去做一件只能在它不存在的条件下才能做到的事。

让我们用一个跨学科类比来佐证这个判断。信息论中有一个重要的区分：有限信号与无限信号。有限信号是事先约定的、有穷尽的符号系统，像红绿灯——红灯停、绿灯行，信号一旦被约定就不再需要不断加码。无限信号则是那类试图用有限符号去捕捉一个在原则上开放、永远比符号更多的东西的情形——例如，试图用一套固定的健康问卷条目去完整描述一个人的“幸福状态”。每一次你在问卷中增加新题目，你以为是在更完整地覆盖幸福，但你增加的每一个新题目，都会在已完成的覆盖面和仍在遗漏的残余面之间，制造一个更清晰的新边界——而那个残留面的存在不会因为你增加条目而消失，它只会变得更精确地被辨认出来。“善



的动机”相对于行政代理指标，就像是那份问卷永远抓不完的“幸福”的剩余。两者之间不是精度高低的问题，而是类型区别：一个是封闭的二值信号，一个是开放的无边界状态。

## ■ 5.2 代理自噬：制度为何无法停止精细化

### ■ 5.2.1 代理自噬的运行机制

如果说前一小节论证的是“强迫行善规则在演化序列上不能收敛”，那本小节要解释的是“为什么制度内部没有刹车”。收敛失败描述的是结果，代理自噬描述的是驱动收敛失败的那个内生动力机制——它解释为什么制度即便已经感知到收敛困难，也无法主动停下来。

代理自噬的机制可以这样描述：制度为捕获“善的动机”这一不可直接观测的目标状态，设立一组代理变量——比如志愿服务时长、捐款金额、助人次数。这组代理变量一旦上线运行，马上就会产生一个制度设计者心知肚明但往往不愿面对的问题：代理变量与真实目标状态之间存在缺口。志愿者服务时长可以“被凑满”，捐款金额可以在不影响实际痛感的情况下被刻意规划，助人次数可以通过选择“容易的善行”来刷量。这就是第一批“伪善者”——他们的代理指标达标，但他们的行为缺少善意动机。

制度对此的第一反应一定是：堵漏。下一次修订规则时，增设更低层的精细化指标——比如，志愿服务不但要看时长，还要看服务类型、服务对象的脆弱程度、服务记录的不可篡改性。堵漏措施的确在短期内可能有效——旧的表演方式被堵住了。但几乎同时，就已经有人在适应新的指标、寻找新的表演路径。为什么？因为代理变量与真实动机之间的那道鸿沟不是任何一个具体指标的设计缺陷，而是结构性的——动机在内部，指标在外部，两者之间必然存在缺口。旧的缺口被堵住，新的缺口会自动在更精细的层面显现——而且新缺口的可识别度随精度的提升而提升，也就是说，被堵漏逼出来的下一代表演，比上一代表演更不易被简单识别，因此具有更强的“吸引力”。这就像医学中的抗生素耐药性进化：每一次新研发的抗生素都会在更短的时间内筛选出耐药菌株，而耐药菌株的传播又反过来迫使研发下一批更昂贵的抗生素。在这个过程中，驱动力不是衰减而是递增——因为越耐药就越需要新药，越新药就越筛出更强的耐药性。

### ■ 5.2.2 刹车为何失灵

代理自噬一旦启动，为什么不能通过制度内部的“认知突破”来主动停止？一个直觉性的建议可能是：既然我们认识到善的动机不能被代理指标穷尽，那就应该在制度设计之初就给自己留一个“不再精细化”的出口。但这一建议在操作上面临一个无法跨越的障碍。

想象一个已经运行五年的道德积分系统，它的代理变量不断精细化已经到了第四代。此时如果制度管理者宣布“我们决定停止增设新指标，因为善的动机本来就不能被完全覆盖”——这个声明本身是诚实的哲学反思。但问题来了：在声明之后，当出现新的、明显是“利用指标漏洞刷分”的表演性善行时，管理者怎么办？如果管理者任由它发生而不做出回应，它的不作为本身就会在制度内部制造合法性危机——这套系统本来就是为了“真实评估每个人做了多少好事”而设立的，管理者却明知有人在用新表演获利而不堵漏。如果管理者做出回应、堵漏了——那它就是重新迈入新一轮精细化。换言之，一次诚实反思不能结出停下来的果，是因为制度自身的合法性已然绑定在“把善行更准确地甄别出来”这个承诺上。这个承诺一旦做出，停步就意味着公开承认“我们其实做不到”——但制度内部的驱力不允许这种公开承认发生，因为制度存在的理由就挂在这句话上。

### ■ 5.2.3 累停时刻：当公共语言被代理变量完整覆盖

代理自噬不会无限加速运行下去。制度的操作成本会随精细化程度的提升而递增，最终会达到一个制度整体无法承受的阈值——此时精细化中止。但这并非善的胜利，而是精疲力尽。累停的那一刻，我们可以观察到两样东西的同时在场：第一，制度的代理变量体系已经覆盖了该社会中对“善行”进行公共定义的几乎全部领域——从志愿服务到捐款到邻里互助到日常礼貌，都有对应的指标与积分规则；第二，制度不再有能力进行下一代精细化，因为操作成本已经超过了社会管理预算的上限。

累停之后的形态是什么样的呢？善作为一个需要行为者自我判断的道德事件，它的公共语法已被代理体系完整覆盖。不是说人们不再做善事，而是说“做善事”这三个字在社会的公共理解中，已不再指向“我自发想要帮助他人”的那一瞬内心冲动，而是指向“我在合规体系中获得了‘已行善’的记录”这一制度性事实。善从道德事件蒸发了，变成了合规事件。这一转换是不可逆的，因为任何试图重新把“善意动机”放回公共位置的后续努力，都会发现自己已经没有可用的公共语法——那些曾用来表达和确认“自发的善意”的词汇、叙事、仪式，已经被二十年、三十年的制度话语覆盖、替代、重新填充、并最终彻底改写。

## ■ 5.3 道德主体的失语：公共语言被垄断后的个体认知错轨

### ■ 5.3.1 认知发展中的他者确认耦合：从婴儿到成人的道德化路径

现在要接入本章理论框架的最后一组构件：来自认知神经科学和发育心理学的洞见。前两小节分别在制度层面论证了“为何不可收敛”与“为何无法刹车”，本小节要论证的，是这场制度层面的退化在个体内部——在道德主体的认知结构中——所留下的不可逆损伤。

如前所述，人类道德能力的发育不是一个密封种子内部展开的过程。它不是康德式的理性主体在独白中为自己颁布一条绝对律令。它更像是一个呼吸过程：吸进的是他人的痛苦在你身体里激起的不安，呼出的是你在他人应答中学会的“这个不安叫同情、它通向一种叫做帮助的行为”。霍夫曼研究的婴儿对另一个婴儿哭声的共鸣，是这一呼吸的起点。达马西奥研究的腹内侧前额叶皮层受损患者——知道正确答案、却无法做出正确决策——是这一呼吸被体感阻断后留下的残态。托马塞洛对比人类儿童与黑猩猩幼崽发现的人类独特优势——共享意图与公共语言的耦合——是这一呼吸从个体爬升至社会协作的那个上升气流。

这里的关键是：那个被吸入的“他人的痛苦”，不是一种可以被替换成任何等价信号刺激的原始物料。它的“他人性”是构成性的。这是道德发育与纯粹技能学习的本质区别。学骑自行车可以在没人看的情况下练会，你的小脑在你的身体里长出一条内化的平衡回路，他者不是必须的。而学“善”从一开始就是“他人在场”的行为——无论是那个在哭的同桌、扶起老人的陌生路人、还是后来给你一个不刻意点头的旁观者。他们不是道德剧场的观众，他们是道德剧场的共同搭建者。你把观众遣散，戏就无法排演。

### ■ 5.3.2 制度代理系统如何抽走耦合件

现在设想，将上述那套精致的他者耦合系统，替换为一面永远只反馈数据分类的镜子。这面镜子不回应“痛苦”，只回应“被编号的 A 类不适”。不回应“慷慨”，只回应“合规积分+2”。不回应“勇敢”，只回应“季度善行达标”。这面镜子，就是运行了二十年、已经完成公共语法垄断的强迫行善制度代理体系。

在一个正常运转的社会协作中，一个孩子第一次因为帮助同桌而被同学用不是“加分”而是“眼神”的方式回应时，那套眼神完成了两个动作。第一，它确认了孩子刚才所做的那个举动是“善”这一类行为中的一次实例——这是类目归属。第二，它确认了孩子自己——而不是任何其他他人或任何规则——是这个善举的自来源。眼神传递的信息是：“你做了这件事。我没有逼你做。你本可以不做的。但你做了。”这三个短句在孩子的道德自我叙事中落下的印记，是无法被任何制度颁发的外部奖励所替代的，因为它们都在强化同一个第一人称结构：“我是一个能主动选择帮助别人的人。”

制度代理系统无法复制这套神经由双重确认所完成的工作。它能做的是另一件事：分类与记录。在被代理系统长期覆盖的环境中，一个孩子扶起跌倒的同学后，得到的不是同伴的目光，而是老师的登记——“助人一次，加 2 分”。加 2 分当然也是一种反馈，但它告诉孩子的不是一个第一人称的道德自我叙事，而是一个第三人称的制度性事实：“你的行为属于被鼓励类别 H，已记录。”第三人称不是第一人称的退化版——它们是不同的语法形态。长期被第三

人称合法描述取代第一人称道德叙事后，孩子学到的不是“我是一个善良的人”，而是“我做了一件被归为H类的事”。前者是自我定义的道德身份，后者是对外部分类的认知认领。前者可以自我维持，后者必须依赖外部分类系统的持续运作。前者在制度撤销后可以继续存在，后者在制度撤销后失去意义。

### ■ 5.3.3 代际尺度上的不可逆错轨

如果把时间尺度拉长至两代人、三代人，这个区别的后果会以一种难以逆转的方式固化下来。第一代人在道德积分制度下生活了二十年。他们童年时期还有记忆——在积分系统覆盖每一个学校活动之前，曾经有过不经登记、没人加分的“做好事”的时刻。他们成年后，内心可能仍然保留着一个残余层，一个在制度语言之外用记忆、私人暗语、沉默的互助仪式维持的“真善”空间。他们能做善事，而且清楚地知道自己在行善——尽管他们同时也知道，这个善无法在公共语言中得到不同于积分的描述。

第二代人——在他们出生时，代理体系已经成熟运转——没有这种前制度的记忆。他们唯一接触过的善的公共语法，就是制度的合规语言。当他们帮助同学时，身体仍然可能产生一个温暖的感觉，但那个感觉缺少可以用来向自己解释和向他人传递的公共语法。同学的反应也是被制度训练过的——不再是眼神，而是“你这次加分了吗？”在这样一个被抽干他者确认耦合件的人际环境里，那个温暖的感觉无法被巩固成一个稳定的道德自我叙事。它会变成一个一次性的内部抽搐，一阵无法被归置的莫名焦虑，最终被主体学会忽略或压抑。

本章在此做出以下预测，这一预测是本章最核心的可证伪命题，是需要未来经验研究加以检验的核心假设：在运行超过两代的道德积分制度的覆盖人群中，可以观察到一种特定的道德叙事代际退化——第一代能以第一人称叙事表达道德意图（“我帮他是因为我想帮他”），第二代人更多用第三人称描述或制度性语言（“这是被鼓励的行为”“做好事加分”）来表达同等行为，第三代人的道德叙事可能出现失语现象——行为者仍能在外部观察层面做出帮助他人的动作，但在被要求叙述自己当初为什么那么做时，无法给出一个超出“就是做了一下”“不知道”的第一人称道德归因。

这不是价值观的转移，这是语法能力的退化。道德叙事从第一人称滑向第三人称或被抽空为沉默，不是因为他们更加自私——他们可能身体上仍然比前代人做出了更多的合规善行——而是因为他们用来构造“自私”与“无私”这两个概念的那套语言，在发育期从未被完整地安装进他们的认知系统。他们一生都在合规中做事，从未有任何人——任何一个以他者身份在场的人——为他们确认过“你本可以不这样做的，但你做了。”这句话的语言结构需要一种社会耦合才能被说出口和被听懂，而这层耦合器已经被代理系统覆盖了二十多年。

### ■ 5.3.4 从社会确认到内部生成：善为什么不能是密封种子

这里可能遭遇一个反驳：本节描绘的画面似乎预设了一个过于悲观的前提——即人类在失去外部他者确认后，完全无法从内部生成道德动机。但这个预设是否成立？历史上难道不存在孤独的道德先行者，在没有他人认可的情况下依然坚持善行的人？

是的，他们存在。但这个反驳恰恰强化而非削弱本节的核心论点。那些“孤独的道德先行者”——梭罗、甘地、索尔仁尼琴——通常都是已经完成道德发育、已经将他者的早年确认内化为稳定道德自我认同的成年人。他们的孤独是在主体已经长成后自己选择的孤独，不是主体从未被给予过他者确认的、从一开始就失去语法发育条件的孤独。这就像一棵已经长成的树可以在干旱中靠着深扎下去的根系存活，但这不能证明种子可以在无土无水的真空中发芽。一个已经形成道德主体能力的人可以在他者缺席的条件下坚持善行，但这一事实不能推出一个在发育期从未接触过深度耦合的人也可以生成这种能力——你们是在用发育完成的成年人的“孤独坚守能力”去论证婴幼儿应当拥有的“无需他人即可成长的能力”，这中间的逻辑跳跃是致命的。

那么，那些生长在制度代理系统已经垄断公共语法环境中的第二代、第三代，正是这种无土无水的种子。问题不在于他们是否愿意成为道德先行者，而在于他们用来“先行”的那套认知机能——将身体反应加工为道德意图、并向自己解释这一意图的语言结构——在发育期没有获得必要的耦合养护。他们不是不想善，而是已经不知道“想善”在内心语法中应该是一个什么样的内部语言动作。

## ■ 5.4 反驳预期与回应

### ■ 5.4.1 “转型必要性”反驳

一种可能来自功利主义视角的反驳是：本论文所描述的制度退化，只是旧式官僚化道德激励的产物；未来的智能化道德激励系统可以通过去中心化设计、参与者赋权、透明公开等制度创新，避免代理自噬的发生。这一反驳值得被认真对待。

本章的回应是：去中心化设计可以减缓代理指标的精细化频率，但它无法摘除代理自噬的驱动源——代理变量与真实内部状态之间的结构性缺口。无论是中心化还是去中心化的记录系统，只要“善行”仍然需要被翻译成记录系统中的可计量单位（积分、代币、信用点、社区声誉权重），那上述结构性缺口就会存在，而一旦表演性行为的可辨认案例在系统内被发现——这在任何允许参与者相互评价的系统中是高频事件——弥补缺口的驱动力就会被激活。去中心化的真正优势可能在于，它将判定权从一个集中机构分散到众多节点上，延缓了统一精细化标



准的形成速度，从而拉长了退化进程的时间尺度。但这改变的是速度而非方向。只要结构性缺口仍在，精细化的驱力就不会消失——它只是不再集中，而已。

## ■ 5.4.2 “不完美但有用”反驳

另一种反驳是：代理自噬或许确实会在长期内发生，但在此之前，强迫行善制度可以促成大量本来不会发生的善行，难道不应该肯定其短期功效吗？

本章并不否认强迫行善制度在短期内可能提高“被记录的善行”的数量。本章称其为“合规善行的显性增长”。但本章要质疑的是，这种增长的代价是什么。如果一个制度在三代人之内将善的公共语法抽干，使得下一代无法生成独立的道德判断能力，那么第一代所记录的增量善行，是用此后无数代的道德能力存量为代价换来的。这笔账怎么算，取决于我们把时间尺度拉多长。本章的判断是：在三代人的尺度上，短期账面上的增量远不足以弥补道德认知基础被系统性腐蚀的损失。这是因为合规善行是一种需要外部激励持续供能的反射性行为——一旦外部激励撤销，行为随之衰减——而自发善行是一种靠第一人称道德叙事自我驱动的稳定倾向。前者替代后者，不是量的替换，而是类型的置换——它用一盏需要持续供电的灯，换掉了阳光。短期你会因为按钮一按灯就亮而高兴，长期你却发现自己失去了昼夜节律。

## ■ 5.4.3 “禁止制度也不纯粹”反驳

再一种反驳抓住“禁止制度也不纯粹”这一点：现代刑法在量刑阶段考量悔罪态度、犯罪动机等因素，这本身就是一个动机审计的入口。如果禁止制度自身也包含动机审计，那本论文赖以立论的“禁恶 vs 强迫行善”二分法难道不是被现实摩擦掉了其锐度吗？

这是一个极有分量的质疑。本章的回应分两层。第一层：禁止制度中的动机审计是二阶的、次级的、附着在已确认的外部事实之上的。制度先确认“这个行为本身是禁止的”（定罪），再考量“动机情节”（量刑）。这个二阶结构由一个不可撤除的外部事实作为地基——行为本身发生了且被确认。强迫行善制度缺乏这层独立于动机审计的地基，它的地基从一开始就扎向内心状态。第二层：本章同意，在量刑阶段动机审计的入口如果不断扩大，有可能让禁止制度也陷入一个温和代理自噬。这是一个值得经验研究的假说，且本论文的批判框架恰好为这种经验研究提供了操作工具——“判据类型诊断”不是只适用于强迫行善制度，也适用于任何以内部状态为对象的制度模块。然而，即便禁止制度部分模块的代理自噬被经验证实，这并不削弱本章的核心判断——即强迫行善制度整体性地依赖代理变量，是更根本的、无法通过自身结构修正的结构性失败。相反，这一发现将把本章的命题从一个二值模型（禁止好，强迫坏）升级为更精确的梯度诊断模型（任何以内部状态为目标、且依赖代理变量的制度，都存在

代理自噬风险，风险程度由该制度在地基结构中的位置和所占权重决定）。这是一个有价值的理论升级方向。

## ■ 5.5 综合讨论：退化的三个阶段

现在可以将前文的分散论证凝聚为一个三阶段的退化模型。

第一阶段：行为指标替代。强迫行善制度以“只看行为、不问动机”的低姿态登场。在初期，制度的激励确实能够提高善行的可观测数量。但由于代理变量与真实动机之间存在结构性缺口，表演性善行逐渐出现并增多。制度为堵漏开始增设精细化指标。这一阶段是可逆的——如果制度在此阶段撤回或根本性重构，善的公共语法尚有一定空间自行恢复。

第二阶段：公共语法垄断（翻译垄断）。代理自噬的精细化循环持续运转。代理变量体系逐步覆盖社会中对“善”这一概念进行公共定义的主要场域——学校评价系统、社区善行表彰、工作单位的道德考核、乃至日常语言中人们用来谈论“好事”的词语。在某一个难以精确标定但事后可辨的临界点，制度代理系统完成了对善的公共确认权力的全盘接管。在此之前，说一个人“做了好事”主要是一个人际间的道德确认事件；在此之后，这个句子在公共理解中首先指向“该人在制度记录中获得了善行合规的确认”。这一临界点本章称为“翻译垄断”。一旦越过，退化就不再是可逆的质量下降——因为善的公共语法已经发生了类型转换。

第三阶段：道德主体的失语与错轨发育。在公共语法被垄断的环境中成长起来的第二代、第三代，在道德认知发育期缺乏他者确认的耦合件。他们身体对他人痛苦的反应可能仍然存在（镜像神经元没有坏死），但将这一身体反应加工为“我想帮他”这一道德意图的语言能力已在发育中被腐蚀。他们的道德叙事从第一人称退化为第三人称，或完全失语。他们仍然可以被观察到做出合规的善行，但那种行为不再具有善行作为“道德事件”的构成性要素——行为者自身将其体验为自发的、第一人称驱动的道德选择。

这三个阶段从理论上是逻辑递进的，但其经验验证需要长时段、跨代际的纵向追踪研究。本章接下来提供一个可操作的经验检验方案，为未来的经验研究者提供起点。

## ■ 5.6 可经验检验的预测与方案设计

本章给出了以下一组命题，它们都是可被经验研究证实或推翻的。为了便于经验研究者操作，下面按竞争解释控制、检验设计和证伪条件三个环节逐一阐释。

### ■ 5.6.1 必须排除的最强竞争解释



本章所预测的“道德叙事代际退化”在面对经验检验时，需要排除几个主要的竞争解释。其中最强的一个——也是任何经验检验设计必须优先控制掉的——是现代化进程本身的世俗化效应。反对者可以提出：“你观察到的第一人称道德叙事的代际退化，并不是道德激励制度的特定效应，而是社会现代化进程中的普遍世俗化趋势——也就是说，即便没有激励制度，道德叙事也会随着传统价值体系的松动而出现类似的从内向到外在的转型。”

这个竞争解释若成立，本理论预测的“道德叙事的退化是强制善行制度所致”就能完全被世俗化理论解释，本章的核心机制就不再是必要条件。因此，检验设计必须能够区分这两种机制：是制度代理系统抽走了耦合件，还是现代化进程自身的世俗化漂移？

## ■ 5.6.2 可验证的检验设计

研究设计：准实验设计，选取两个或多个在现代化程度（人均 GDP、城市化率、教育水平、大众媒体接触度）上高度可比、但在“是否引入制度性善行激励体系（道德积分）”上存在显著差异的社会单元（如两个邻近的城市，或一个已引入积分系统与一个尚未引入的学区）。两个组在世俗化程度上的基线应当事先通过宗教参与度、传统价值量表等工具测量并匹配。

关键变量与测量工具：

（1）道德叙事类型：对参与者的深度访谈进行结构化编码，区分以下三类文本的出现频率与占比——A 类，第一人称道德意图叙事（“我这么做是因为我觉得应该帮他”）；B 类，制度性第三人称叙事（“这是被鼓励的行为”“做好事能加分”）；C 类，失语/无法归因叙事（“就是做了一下”“不知道为什么会做”）。每一类文本必须由不少于三位的独立评分者进行编码，评分者间信度需达到 Krippendorff's Alpha  $\geq 0.70$ 。此外，为排除同一社群内的话语习惯效应，还需在叙事编码时区分叙事中的“常用理由词汇”与“故事结构”，确保测量捕捉到的是叙事者的意图归因结构，而非纯粹的区域性话语习惯。

（2）道德意图失认症的间接测量：情感词汇辨别任务——给参与者呈现一系列（如 30 个）描述身体状态的词汇（如“难受”“发热”“紧张”“发麻”），要求他们在限定时间内快速判断这是一个“身体感觉”还是“道德情感”。在正常发育群体中，两者存在清晰的分化——被试对“心痛”“内疚”“不安”等复合词的反应时间应显著慢于纯身体词汇，因为它们在进行中被同时激活身体感觉与道德意义的两个认知系统。长期在失语状态下生活的个体，由于道德情感从未被完整的公共语法确认过，当面对复合词汇时倾向于快速将其归入“身体感觉”一类，或出现显著的反应时间缩短——无法区分两种类型之间的微妙差别。这组任务的反应时与正确率差异，可作为“道德意图加工”的间接指标。

(3) 生理-叙事一致性测量：设计一个准实验任务——请参与者观看一段高共情唤起的视频（如他人遭受痛苦的真实记录），同时记录其皮电反应（SCR）。在观看结束后，要求他们立即做口述叙事：“请描述你刚才观看时身体和内心的感受。”将生理唤醒幅度的变化（SCR波幅）与叙事文本中的第一人称情感词汇量进行相关性分析。在控制组中预测存在显著的正相关（感受到了就能说出来）；在长期被激励制度覆盖的群体中预测相关性降低或消失——身体仍然反应了，但叙事无法接住。

对照与排除策略：为了排除世俗化竞争解释，检验设计必须包含以下对照逻辑——如果在世俗化程度相当、但激励制度存在与否不同的两个地区，A组（有积分制）的C类叙事占比与情感辨别任务表现均显著劣于B组（无积分制），那么世俗化效应即使存在，也不是导致退化的特定机制，制度代理系统才是。这一检验需要做分层回归：以C类叙事占比为因变量，以是否引入激励制度为自变量，控制世俗化量表得分、宗教参与度、年龄和受教育程度后仍然显著，则支持本章理论。

检验的时间尺度：鉴于本章预测的退化是在代际跨度上展开的，理想的研究是跨代纵向追踪——对引入激励制度初期的家庭进行基线测量，此后分别在5年、10年、15年节点上重复测量父母与其成年子女的道德叙事类型与意图归因能力。这样的追踪在操作上很重，但分段接力（同时比较已运行0-5年、5-15年、15-25年的不同区域）也可以提供初步的代际差异证据。

### ■ 5.6.3 本理论会被什么结果推翻

以下任何一种实证结果，都将构成对本章核心命题的有力证伪：

证伪条件一：在运行超过二十年的强迫行善制度覆盖群体中，其道德叙事仍以第一人称归因为主，且与无制度覆盖的控制群体在A类叙事占比上没有显著差异。此结果将推翻“公共语法垄断导致道德叙事代际退化”的核心预测。

证伪条件二：在上述时间和覆盖面区域内，该群体在道德意图失认症任务（情感词汇辨别）上的表现与无制度覆盖的控制组没有显著差异。此结果将推翻“长期代理制度环境导致道德意图失认”这一形态判断。

证伪条件三：在世俗化程度与制度引入完全共线的地区中，梯度比较无法将二者的效应分离——即只要是世俗化程度高的地方，无论是否引入激励制度，道德叙事都表现为同等的退化趋势。此发现将推翻“制度代理系统是道德叙事退化的必要机制”这一归因——即便本章所描述的退化现象仍然真实存在，其主因是现代化进程本身而非本章聚焦的制度机制。这不会否定本章的现象诊断，但会否定其因果推力，从而将本章的立场从“制度是根本驱动因”降格为“制度是恶化因之一”。

证伪条件四（这是一项更具体的局部证伪）：如果发现某一个运行超过二十年的激励制度，在其自身的发展进程中成功地分批撤回了某些早期代理变量且代理体系没有因表演行为的新出现而反弹——即实现了主动终止而不重蹈新一轮精细化——则本章关于代理自噬“不可主动终止”这一判断在该案例中不成立。但请注意，这并不必然推翻代理自噬在大多数制度中成立的统计规律，而只是为“主动终止”添加一个可能的边界条件（例如：存在独立的外部审计机构，或制度设计中嵌入了代理变量日落条款）。在此情形下，本章的命题需要被修正为：代理自噬的不可逆性仅在缺乏特定制度性刹车机制时成立，而非在所有情形中都不可逆。这种修正不会推翻本章的核心诊断，而是将其从强版本（绝对不可逆）弱化为有条件的版本（在无刹车时不可逆），并逼出下一步的研究：识别并检验有效的刹车机制。

## ■ 6 结论

### ■ 6.1 核心发现的凝缩

本章论证了一个迄今鲜被系统阐述的形态命题：强迫人做好事的制度，其终局不是制造伪善者，也不是单纯摧毁社会信任，而是在代际尺度上不可逆地腐蚀了“善”作为一种不需要制度许可即可在个体内部自我繁衍的内在伦理能力。这种能力的丧失不是通过宣言或禁令来实现的，而是通过一套静默的、制度自身也难以自控的运作逻辑——代理自噬——来逐步覆盖善的公共语法，从而摘除个体道德认知发育所必需的外部他者确认耦合件。总结起来，这一命题可以凝缩为一句话：强迫行善制度以短期显性善行的增量为代价，换取了长期道德主体内在繁衍能力的不可逆衰退。

### ■ 6.2 理论贡献

本章的理论贡献主要体现在三个递进的层面。

第一，对法哲学中“法律不能强制执行道德”这一经典命题，本章将其从规范边界的讨论推进为制度形态学的诊断。本章的论证表明，问题不在于法律“应不应该”强制道德，而在于“强制”这一动作在结构上会取消“善”的构成性要素——动机的自发性与公共语法中他者确认的耦合功能。这一诊断将法哲学的传统辩论重新置于一个新的基础之上：它要求论辩双方不再只讨论正当性问题，而是首先回答一个前置性问题——如果强制执行道德会在结构上改变“道德”在公共空间中的存在形态，那么任何关于其正当性的讨论都必须把这一形态改变的代价纳入考量。

第二，本章将道德心理学关于他者确认耦合在个体道德发育中的构成性作用的发现，从个体层面提升至制度分析层面，在发育心理学与法哲学之间架设了一座迄今缺失的分析桥梁。这

座桥梁的关键部件，是“公共语法”这一概念。它使得我们能够看到：发育心理学所揭示的他者确认，不只是一种个体关系——它在社会的宏观尺度上，正是由一套可以被制度垄断或保护的公共语言系统所承载和维护的。制度对善的干预，最高的代价不是行为层面的错配，而是在这一公共语言层面启动了一个不可逆的替换进程。

第三，本章提出的三阶段退化模型（行为指标替代、公共语法垄断、道德主体失语）与代理自噬概念，为分析同类制度现象——包括但不限于道德积分、善行激励、企业社会责任合规化、教育领域的品德考核量化——提供了一个可操作的概念框架。这一框架不是只适用于“强迫行善”这一特定论题，它可以推广到任何一个试图用外部代理变量捕捉、评估并激励内部状态的制度化进程。

## 6.3 实践与政策含义

本章的分析对正在或计划推行道德积分、善行激励等制度的地区有两项具体的、可操作的政策建议。

第一，设立代理变量体系的事前刹车机制。如果一个社会仍然选择建立道德激励制度，那么在制度设计之初，就应嵌入一个独立于该制度运行机构的、定期发布“代理自噬风险评估报告”的外部审视机制。该报告的职能不是评估制度的整体效益，而是专门监测三项指标：代理变量的更新频率是否在加速、新修订的驱动原因中“堵漏表演性善行”占比是否在递增、制度维护成本中代理体系本身的维护费用是否在持续上升。当这三项指标同时上扬并持续超过一个预设周期（如连续五年），外部审视机构应有权建议暂停继续精细化，并将制度退回最近一个“未发现同步上扬”的版本。这相当于一套代理变量的“日落条款”——强制性的、结构性自我拆解，不是靠制度内部的反思来决定，而是靠一个事前写好的、锁死的触发开关。

第二，基础教育领域的公共语法保护。强迫行善的制度退化，对下一代的影响最深远——他们的道德认知结构在尚在发育时就被代理系统的合规语言接管。因此，即便在已引入道德激励制度的地区，基础教育阶段也应设置一个“制度语言禁入区”——在此区域内，学校对学生的利他行为进行教育和评价时，禁止使用与积分制挂钩的合规术语。教师可以赞赏学生的善意，但不能登记；可以引导学生反思自己的行为动机，但不能将其翻译为任何外部积分体系中的“已达标”。这不是一种对制度的消极抵抗，而是对下一代的道德认知发育所需的最低原始耦合件的保护——保护他们还能在某些场合，在他们的老师眼中，体验到一种不被记录、不被翻译的道德确认。这种体验，是善的公共语法在制度垄断的夹缝中仅存的野生幼苗。如果连学校这一最后阵地也失守，那么下一代将生活在一个没有任何他者确认是不经过代理变量的世界里——那种世界的道德后果，正是本章所担心的。

## 6.4 研究局限

本章存在以下应被明确指认的局限。第一，属于理论建构范畴，其核心经验命题尚待前述的纵向追踪或跨区域比较研究来检验。本章的操作化已为后续研究提供了可用的测量方案和变量设计，但本章自身不包含一手经验数据。第二，对退化进程不可逆的强断言，主要基于对代理自噬正向反馈闭环的理论推导——这一推导是严密的，但尚未经历史上可能存在的“主动终止成功”案例的检验。如果这样的案例在后继研究中被发现（如一个制度在进入代理自噬早期后，通过成功撤回某些代理变量而中止了精细化），本章关于“不可逆”的强断言需要被弱化为条件性命题。第三，本章的分析主要适用于善行判定权由中心化制度机构集中行使的现代社会情境，未对判定权高度分散的社区型社会进行系统性比较。在那些由邻里之间、乡村长老、宗教社区共同判断什么是善行的社会结构里，代理自噬可能因为判定权分散而缺乏集中化的驱动力，本章的结论是否可以延伸过去，需另行检证。第四，本章对“残余层动力学”——在公共语法已被垄断后，个体如何在私人暗号、身体仪式、小范围信任网络中维持不被制度捕获的善——仅做了初步的方向性讨论，未充分展开这一层面的完整动力学，这是留给后续研究的核心课题。

## 6.5 未来研究方向

基于本章的理论框架和上述局限，提出以下五个可生长的研究纲领。

研究纲领一：代理自噬的跨制度比较田野调查。选取已运行不同时长道德积分制度（如某些城市或学校的善行积分系统），对历次规则修订文件进行内容分析，编码每一版修订的驱动原因（“堵漏表演”“扩展覆盖”“简化流程”“外部政策要求”等类别），检验在制度运行超过5-8年后，“堵漏表演”在驱动原因中的占比是否如代理自噬假说所预期的那样出现显著递增。同时设置控制组：对无积分制的邻近地区做基线比较，控制媒体影响造成的叙事变化。

研究纲领二：道德叙事代际退化的纵向追踪。对引入道德激励制度初期入学的小学生进行基线道德叙事访谈，此后在中学、青年期进行定期追踪。需要区分的是：年龄增长本身可能带来更成熟的道德自我叙事能力，而本章预测的是，被激励制度长期覆盖的群体中，这种成熟化并不投向第一人称归因，而是倾向投向第三人称的合规叙述或残留为无法描述的“失语”——叙事能力不差，但方向偏了。同时设置无积分制的对照学校，进行同样的追踪。这是检验本章核心命题的最关键经验途径。

研究纲领三：残余层的信号博弈研究。在道德激励制度已高度成熟的人群中，运用深度访谈和网络分析方法，考察“真善者”如何在私人暗语、沉默互助、非制度记录的善行中传递善



意信号。这些信号的传递需要多高的成本（如更长的信任建立周期、更不容易被误解为非善意的行为编码）？这种高成本信号是否客观上限制着残余层善行传递的规模和速度？这一研究纲领将把本章关于“残余层的内部动力学”的初步讨论，转化为一个可被博弈论模型与网络分析捕获的经验课题。

研究纲领四：禁止制度自身代理自噬的审计研究。选取一个法律体系中对犯罪者“悔罪态度”是否进行系统化评估的司法实践，追溯其历史上是否出现过悔罪评估指标的精细化升级，并比较不同地区（对悔罪评估依赖度不同的司法管辖区）在量刑公正性与制度成本上的差异。这一研究纲领的目标是检验本章的一个隐含判断是否成立：禁止制度中的动机审计——虽只是量刑阶段的一个模块——当被司法实践膨胀后，是否也会在局部启动代理自噬，以及这种自噬是否能够被定罪-量刑的二阶结构控制在一定范围内。

研究纲领五：礼治社区的跨文明比较。选取一个仍有传统礼治运作痕迹的社区，与一个已建立现代道德积分系统的社区进行比较，重点考察二者的判据权集中度（由中心机构统一判定 vs 由社区成员分散判定）与代理变量更新频率（精细化升级的频次）之间的关联。这一研究纲领旨在检验本章隐含的一个推论：代理自噬的不可逆性，可能并非嵌入在“善行激励”这个类型本身的定义中，而是取决于判定权是否被一个中心化官僚机构集中掌控。如果礼治社区的低更新频率确实与判定权分散显著相关，那么本章的命题就需要在跨文明比较的基础上做出一项重要修正——强迫行善制度的退化，不是“善”被公共化的必然代价，而是“善的公共判定权被中心化”在现代官僚制度这一特定制度形态中所触发的结构性后果。

## ■ 本次打磨说明

本批的打磨未另起附录，而是就地改在原文出问题的那一句、那一段上——事实错处直接更正，不可核的材料换成可被真正去测的预测，被放跑的最近邻补成新的一段。以下逐条说明改动落在哪里、为什么。

本章改了四处。最重的一处是补上一片本章迄今未曾正面处理、却离本章命题最近的实证文献：动机排挤与过度理由效应。自德西 1971 年发现外部报酬会削弱内在动机、莱珀等人 1973 年在儿童绘画实验中观察到预期奖赏组在奖赏撤除后绘画时间显著下降以来，这一效应已积累半个世纪的证据，并有德西等人 1999 年的元分析、蒂特马斯对有偿与无偿献血的比较、格尼齐与鲁斯提基尼关于幼儿园迟到罚款反而抬高迟到率的实地研究相互印证。原稿声称本章的诊断「迄今鲜被系统阐述」而一条不引，这一说法站不住。第二节末因此补入一段，先把这片文献认下来，再把分界说清楚：动机排挤的分析单位是个体在一次任务中的动机结构，时间尺度是几周到几个月，且撤除激励后内在动机常可部分恢复；本章追问的是当激励制度覆盖了

整个社会用来谈论和确认善的公共语言之后，下一代在发育期还能不能长出把身体反应加工成「我想帮他」的那套语法。前者是消退，后者是未发育；前者可逆，后者因发育窗口不会重开而不可逆。这一分界同时给出一条新的判别线：若实证发现激励制度撤除后被覆盖群体的第一人称叙事在数年内自行恢复至对照组水平，本章关于代际不可逆的核心主张即告失败。相应六条文献已补入参考文献。另三处：修复了三句在文本传输中损坏的句子（一句关于电灯与阳光的比喻此前完全不可读），删去正文从未引用的十二条充门面文献，并删去虚构的致谢与期刊模板残留的利益冲突声明。



## 第六章 熄灭它的不是暴力，是拥抱

本章原题《配准性生成：长期成功团体的隐性担保与决断力基础的被改写》，作者张琼，原文见 <https://sdeuniverses.com/students/zhang-qiong/calibrative-genesis/>

### ■ 配准性生成

长期成功团体的隐性担保与决断力基础的被改写

摘要

本章对长期成功运转的无领导团体提出一项发生学诊断：成员在团体内经历的深刻的亲密与勇敢，不是学会了一套可独立使用的决断能力，而是其定向自身的基础被系统性地重写为一套与团体隐性语法高度耦合的特定形态。当担保网络在反复的成功互动中达到高度稳固后，担保本身被体验为理所当然——它从成员可感知的“安全垫”沉降为不可见的认知地基。正是在这一沉降中，担保网络从被动承托转换为主动评价：它不再仅仅是“接住你”，而是通过对语法规规行为的正面标记，使不符合语法的原初裁断信号在尚未进入意识之前便被过滤。团体越是成功运转，成员越是熟练于在语法内执行精致表达，其“意识到自己正站在一块地板上”的能力便越是萎缩——由此，跨越不同担保场、在担保消失时仍能维持定向的弹性，被成功本身从内部消耗。本章将这一过程命名为“配准性生成”，并论证其并非游离于疗效之外的副作用，而是寄生在利他主义、团体凝聚力和人际学习等亚隆所定义的核心疗效因子内部，是成功治疗过程本身携带的自反性风险。配准的后果不是成员“失去了独立决断”而退回到对团体的依赖，而是其在任何担保场中都已不再能感知到担保的存在与边界——从而在担保场外，也失去了在不同人际条件下为自己重新校准定向的弹性。本章给出团体内可观测的形态学指标、配准的谱系化推演与明确的证伪条件。

关键词：无领导团体；隐性语法；担保遗忘；决断弹性；配准性生成

### ■ 一、原初问题的裁定与分析单位的改切

本章的研究起点是欧文·亚隆风格的无领导团体。在进入具体的机制分析之前，必须完成一项方法论上的裁定：这一原初的分析单位需要被改切。

“欧文·亚隆风格”描述的是团体的一种技术取向与结构特征——此时此地的互动、带领者的非中心化角色、对过程而非内容的聚焦。它以技术参数界定了团体的运作方式，但自身不携带关于团体生命周期、成功程度或互动语法稳固性的任何信息。一个刚进入混沌期的团体、一个正在经历成员剧烈冲突的团体、一个已经默契到几乎不需要语言便能彼此接住的团体——它们都可以被合法地称为“亚隆风格的无领导团体”。而本章所要解剖的那个特定机制——隐性语法与担保网络对成员决断力基础的系统性改写——需要一个确定的土壤条件才能发生：团体必须已经长期存在、已经反复验证了其互动模式的安全性、已经沉淀出一套高度稳固的、成员间共享的隐性语法。离开了这个条件，配准性生成便没有发生的结构基础。不加区分地以“亚隆风格的无领导团体”为分析对象，将使诊断的有效性被稀释在大量并不满足条件因而不可能观测到该机制的个案之中。

因此，本章将研究对象从原初问题中改切出来。本章研究的不是一切“亚隆风格的无领导团体”，而是它内部一个特定的子集：那些已进入长期成功运转状态、互动语法高度稳固、担保网络高度成熟的团体。这一改切的理由落在分析单位上——原初分析单位囊括了所有生命周期的团体，太粗，切不出配准所必需的特定条件；改切后的分析单位更窄，但能够切中那个真正值得追问的机制：当团体在提供安全这件事上无可挑剔地成功时，成功本身制造了什么。正文自此以下，所有论述中的“团体”均指称这一改切后的对象。

这一改切同时为本章的结论划定了一条边界：配准性生成是长期成功运转的团体的结构性风险，而非一切无领导团体的宿命。那些始终停留在混沌边缘、或反复经历破裂与重组的团体，那些互动语法从未真正稳固下来的团体——它们不构成本章诊断的对象。与之相关的一个关键方向是：是否存在一批长期运转但互动语法始终保持某种粗粝的、未被精致化的开放性的团体，其成员的独立决断力不曾被系统性消耗？如果这类团体存在，担保的高度稳固与隐性语法的精致化是否才是配准的真正开关——而非“成功”本身？本章在第七节的谱系化推演与第十节的边界讨论中，将回到这些问题。

## ■ 二、被成功改变的是什么：一个成员在两种互动中的存在形态

张，三十四岁，因长期在职场无法表达异议而进入一个无领导团体。入组头几个月，他经历了一段混乱。有一次，成员孙打断他说话，他脸红起来，沉默十秒后说：“没什么，你说吧。”那时他内心翻涌的是愤怒和被轻蔑的刺痛，但他独自吞下了这些。他用的是他从外部世界带来的旧语法——“忍一忍就过去了”。在那一刻，他的“吞忍”同样不是天然产生的原初能力；它是他在此前的工作场合与成长环境中，被反复训练出来的一整套应对威胁的认知与行

为模板。他并非站在一块没有担保的地板上——他只是站在一块他从未被教会去感知的担保地板上，那块地板由职场等级、面子文化和“男人不该情绪化”的隐性规约共同铺成。

八个月后，团体进入了稳定运转期。在一次会面中，成员李用轻蔑的口吻评价了张的经历。这一次，张开口了。他没有说“你这话很过分”——那是早期混乱中才可能出现的笨拙表达。他说：“李，我注意到当你刚才说我‘可能太敏感了’的时候，我身体有个反应——心跳得很快，胸口发紧。我想把这个放到这里，看看是不是只有我这样。”团体安静片刻，随后另两位成员给出了共情回应。张的眼眶湿润了：“我以前从来不敢这样。”

这是一次真实的突破。但如果将这次突破与他八个月前沉默时的内部动作作一比照，一个不可见的位移已经发生。八个月前，他在他所熟悉的那套外部担保场的语法内，执行了“冲动—定向—裁断—承担后果”的完整序列：他感到愤怒，他判断说出后果的风险太大，他决定咽下，他默默承受了随后的不适感。那个瞬间，担保并未缺席——外部担保场给了他“忍—忍”的稳定模板，他只是从未被要求去感知那块地板的存在。而八个月后，当他决定开口时，他执行的不再是同一个序列。定向与裁断两个步骤，已经被团体在数月中反复验证的安全记录、成员的先例、带领者克制的在场共同担保了。他体验到的“勇敢”，本质上是在一种担保（外部世界的“吞忍语法”）被另一种担保（团体的“袒露语法”）替换的过程中，被后者预判为大概率能被接住的动作。那些曾经伴随裁断而必然到来的寒意——那种“我说出这句话后，没有人提前承诺会接住我”的不确定感——在这一刻被担保网络的预热屏蔽了。他确实获得了此前不曾有的表达勇气，但他跨越担保场的能力——在不同地板之间换步的弹性——却在这一替换中被悄悄压缩了：他以为自己在从“没有担保的吞忍”走向“有担保的勇敢”，实际上他只是从一块他看不见的地板，换到了另一块同样注定会被遗忘的地板上。

一年零八个月后，团体进入最成熟的阶段。在同一年的另一次会面中，成员陈用五分钟密集叙述工作挫折，但几乎不触及个人情绪。陈说完后，张停顿几秒，平稳地说：“陈，我听了你很多事情的细节，但我可以反馈一个我此刻的感受吗？我感觉我很难真正靠近你——就好像你在给我看一堵很厚的信息墙，但我碰不到墙后面那个你。”这是一次无可挑剔的“成功”互动：标准“我陈述”句式，精确措辞，带着脆弱感的邀请语调。

然而，与一年前“心跳得很快”那次相比，张在这次互动中经历的内心过程发生了一个性质上的跃迁。一年前，他先感受到一团粗粝的、尚未封装的内部震荡——心跳、愤怒、恐惧——然后冒险将这团混沌推送至团体中，借助团体的承托而慢慢找到语言。一年零八个月后，这团粗粝物在尚未完全成形之前，便已被一套预先铺设的认知管线接走。张在陈叙事的过程中，心中浮现的不再是“我此刻到底怎么了”——而是“我该如何用我们团体认可的句式，对我已经在做的事（给出反馈）做一次高水平的执行”。他的注意焦点从内感知转向了表达技术；后者

的成功——这句完美的“我陈述”——让他确信自己已充分处理了这次互动中所有值得处理的事务。而前者——那个未被句式封装的、面对陈信息墙时所感受到的更模糊的不适（可能是烦躁，可能是厌倦，可能是某个他自己也说不清的想要离开这个房间的冲动）——从未被完整地经历，遑论被说出。

张这三次互动，不只是个体的成长叙事。它压缩了一整类现象的完整轨迹：一个成员从一种担保场（外部世界的“忍”语法）移入另一种担保场（团体的“袒露”语法），在后者中逐步娴熟于取用精致句式来执行“勇敢”，而那个跨越担保场所必需的弹性——在不确知的孤独中为自己做出“我要怎么说”的原初定向、在没有任何一方提前背书的情境中承受“我可能说错”的后果——却在成功体验中被悄悄压缩。他越成功，他便越不需要去感知他所站的那块地板的边界。他身边没有人能替他指出这一点，因为在他开口的那一瞬间，整个团体回报给他的，是真实而温暖的接纳。

此处必须做出一笔方法论上的自省。张的案例在本章中承担着“压缩类型”的功能——它被构造为一类典型轨迹的叙事缩影。这种构造方式的优势在于能够将一条跨越一年多的发生轨迹压缩进一个可被完整阅读的篇幅中，从而让担保被遗忘的过程变得在时间中可见。但它的局限也同样明显：单个压缩类型无法穷尽真实个案中可能出现的所有变异——成员的文化背景、入组前的决断弹性基底、担保网络的具体组织方式、团体外社会支持系统的强弱，这些变量的不同组合将产生不同于张的轨迹。本章在后续章节中提出的形态学指标（第八节）和竞争性假说（第五节与第九节），正是为弥补这一归纳局限而设。

### ■ 三、担保网络的三重编织与沉默的被动支撑期

在配准性生成的完整机制被正面展开之前，必须先对担保网络本身做一次结构解剖。担保网络——那个在张开口之前便已渗入他判断基础的“你被接住”的底——至少由三重编织同时构成。

第一重，历史的编织。每一次成员袒露脆弱后被妥善接住的瞬间，都为团体累积了一笔安全储备金。这笔储备金不是抽象的氛围；它在下一轮互动中以具体的方式被调用：当另一个成员准备冒险时，她的认知系统已经自动检索了前人案例库——“张那么尖锐的反馈都被接住了，我这件事不会更糟。”历史的编织将风险的统计概率感性化为一种可被直觉取用的安全感，从而在每一个成员开口之前便降低了裁断的感知成本。

第二重，对等的编织。在长期运转的团体中，成员之间形成了一张隐性的债权债务网。“我上次接住了他的脆弱，这次轮到我袒露时，他大概率也会接住我。”这种对等不是合同，不是算计，而是一种深入日常互动的默契。它的力量在于它从不被明说，却比任何明说的

规则都更难违逆——因为沉默不语的对等，一旦被一方察觉另一方没有兑现，那种失望比任何直接的冲突都更蚀骨。

第三重，身份的编织。长期成功成员在团体内积累的价值感——“我是有深度的人”“我是能给出精准反馈的人”——是他们自我认同的重要锚点。这份认同的稳定供给依赖于团体的持续存在和反馈。当身份在某一个场域被精心培育并只被这个场域承认时，它便天然地携带着对该场域的依赖。

这三重编织不是独立运作的；它们彼此加固为一个反馈闭环。历史的编织降低了裁断的成本，使成员更频繁地进入袒露—被接住的循环；每一次循环都为对等的编织增加了一笔负债资产，持续加深成员对团体的义务感；而义务感的累积又不断沉淀为身份认同的基底。闭环一旦形成，便具备强大的自稳定能力。

但在闭环形成的早期，担保网络的功能仍然是被动的支撑：它等到成员即将坠落时才伸出手。它没有主动去审判成员开口之前的念头。它只是让“坠落”变得越来越不可怕——从而让“尝试”变得越来越容易。在此阶段，成员仍能在不同担保之间维持基本的跨场弹性：他知道自己站在一张网上，他只是越来越信任这张网的承托力。担保尚未被遗忘。

## ■ 四、担保遗忘：从可感知的安全垫沉降为不可见的认知地基

担保网络从一个被感知的存在，沉降为一个不被感知的认知地基——这是配准性生成最关键的发生学步骤。此前的研究文献中，这一步几乎从未被作为独立的机制环节给出过解剖学描写。以下将担保遗忘拆解为三个推进阶段。

第一阶段：知觉中的担保。在团体运转的早期——混沌期与稳定期交接处——成员能够明确地感知到担保的存在。一个新人第一次袒露脆弱后，他会清楚地记得“刚才那一刻我很怕，但我预感到大家可能会接住我”。此时的担保是一种被知觉到的安全：成员在做决定之前，会在脑中预演——“如果我这么说，大家会怎样反应？”“上次孙那么尖锐的反馈都被接住了，我这次也许也可以。”这种预演本身正是担保被感知的证据：担保尚未进入自动取用模式，它仍然需要被成员在每一次裁断前主动调用。

第二阶段：担保的背景化。随着团体运转的持续，成员反复从袒露走向被接住，担保的调用逐步从主动预演沉降为认知系统的默认参数。成员不再需要主动去想“上次张的反馈被接住了”——他的神经系统在开口之前直接给出了一个比从前更低的焦虑基线。担保从“每一次都要想起的事”变成了“一直知道但不再专门去看的事”。它尚未完全消失于意识之外，但它已经从知觉的前景退入了背景。此刻的成员已经不太能说出“我为什么会觉得这次开口可能安全”——他只是觉得安全本身是自然的。



第三阶段：担保的自然化。当团体进入最成熟的长期成功期，担保完成了最终的沉降：它从“一直存在的背景”变成“气一样理所当然的存在”。成员不再觉得“这里很安全”——他只是在无意识地以一种只在担保场内才能有效运作的方式说话、感受、判断，而完全不再感知到担保的存在本身。担保变得和他身体下方的椅子一样：他坐了一整节会谈，却从来没有想到过椅子在承托他。只要椅子不塌，它便是知觉场里的默认零值。

正是在这一步完成后，担保网络从被动的支撑层转换为主动的评价系统成为可能——因为一个不再被感知为担保的担保，其划定的合规边界也将不再被感知为规则，而将被体验为“对错的自然本身”。担保遗忘不是配准性生成的附带现象；它是配准得以进入无意识层执行重写的唯一通道。被配准的成员失去的，不是“独立决断”——从来就不存在担保真空中的独立决断。他失去的，是感知担保的存在并跨越不同担保的能力。他以为自己不在任何担保之中；而正是这一“以为”，使他无法在另一张担保网中重新校准自己的定向。

## ■ 五、语法规规性检测：担保遗忘后，成功如何成为熄灭器

担保一旦在成功体验中被遗忘，担保网络便开始执行它的主动功能：它不再仅仅是“接住你坠落”，而是在你开口之前便判断你将要说出的话是否符合它的期待。这一功能由隐性语法——团体在数百次互动中自然沉淀出的、关于“怎样说话会被听见”的未明言规则集合——来执行。担保网络是硬件，隐性语法是软件；担保网络的三重编织提供了执行评价所需的情感权重（历史的信任、对等的义务、身份的维系），隐性语法则提供了评价所需的技术标准（什么句式被标记为“有效反馈”、什么措辞被标记为“不成熟”）。

隐性语法的运作方式不是惩罚违规行为，而是不给违规行为任何正面回响。当一个成员尝试使用非标准的、不便包装的、携带原始攻击性的语言来表达愤怒或受伤时，团体的回应不是反驳其内容，而是在共情的名义下，用沉默、生硬的回应、或几位成员悄悄交换的不安眼神，组成一个无声的“此信号无法被处理”的标记。这些标记不会出现在会议记录的文本中，但它们在每一个成员的感知系统里留下了精确的印记。下一次，这位成员在愤怒涌起的那一瞬间，那套未被正面标记过的表达方式便不再进入他意识层面的候选选项。他不是“决定不用”那种方式。他是在自己意识到之前，便已经不会那样说话了。

这就是担保网络主动熄灭原初裁断信号的微观机制。配准性生成的核心并非“独立定向被闲置”——这是被动的解释，它留下了一个漏洞：如果真的只是“闲置”，那么成员在离开团体后，这簇被闲置的心理肌肉应该有被重新唤醒的可能。但本章的诊断比这更尖锐：担保网络在反复的成功互动中，通过精心编织的正面回馈——“我听到你这样说，真的很感动”——主动地将成员内部那团混沌的、尚未封装的、无法被语法预判的原初裁断过程，判定为“不需要被

进一步处理”的信号，从而系统地使其在意识层面归于沉寂。被配准的存在者失去的不仅是机会，更是对自己“还有另一种可能的定向方式”的知觉本身。他不是忘了怎么走那条路；他是被成功地说服了：那条路从来就不是一条路。

由此可以回应一个关键质疑：为什么成员不能保有两套系统——在团体内使用语法，在团体外使用原有的或新的决策方式？为什么必然是“替换”而非“叠加”？答案不在认知层面——人类认知完全有能力在多个情境中分别维护不同的行为脚本。答案在担保遗忘的后果中。一个人只有在感知到“我在使用担保”时，才能把担保当作一个可选工具——保留在必要时更换工具的意识。一旦担保沉降为认知地基，它的语法便不再被体验为“一套可选的说话方式”，而被体验为“说话的自然方式本身”。当他把语法的产出当作他自己发自内心的声音时，他便无法再对自己说“这只是我在这里的一个角色”——他已经把那个角色吃进了人格里。不是他拒绝保留两套系统，而是他已无法把其中一套感知为“一套系统”。叠加的神经基础是差异感知；配准的自然化抹除了差异感知——由此，“替换”才在发生学上压倒了“叠加”。

这正是配准性生成难以被抵抗的原因——熄灭它的不是暴力，是拥抱。是那个温暖的回声，让你确信：你刚才那一刻的冲动，的确不该被珍视。成功本身，就是配准的推进器。

## ■ 六、配准性生成与回写机制：被配准者如何复制配准

一个系统如果仅仅配准进入它的成员，它的病理仍有可能被新进入者携带的外部残余所扰动。配准性生成的完整动力，在于它内置了一套回写机制：被配准的成员，不自觉地成为配准其他成员的代理人。

在长期成功的团体中，每一条隐性语法的执行都受到一套精确的、隐蔽的反馈机制所管控。当一个成员甲对成员乙给出了一个高度符合团体语法的反馈时，成员丙几轮后可能会说：“我听到甲刚才给乙那个反馈的时候，有一个感觉——我挺感动的，也很佩服你敢这样直接说。我不会说太多，就是想说我看到了。”丙的这句话是“过程评论”，它在团体的语法体系中是一个有价值的贡献——它展示了丙对团体进程中此刻所发生之事的高度关注，并对团体的价值做出了再确认。然而，正是在这个充满善意的动作里，丙同时完成了一件事：他对甲的行为进行了正面标记。甲下一次感到不适时，几乎不需要再考虑别的应对方式；乙下一次需要回应尖锐反馈时，也已经有了一个被丙称许过的参考模板；而那些旁观的成员丁、戊、己，也都在没有被告知的情况下，接收到了同一个信号：“这种说话方式，在这里是被看见的。”

这就是回写的完整循环：被配准的成员，通过给出过程评论、点头、专注倾听、在沉默后率先打破紧张等高度符合语法的行为，从团体中获取被认可与归属感；而他们每一次成功使用



这些行为，都同时是一次无声的教学，告诉其他成员“在这里这样做是好的”。配准不是由带领者推行的，不是由某一派系垄断的。它由每一个真诚地爱着这个团体、并从团体中获取了生命中最深刻接纳的成员，在一轮又一轮温暖而无可挑剔的互动中，日复一日地复制着自身。这就是为什么它无法被某一个觉察者从内部打断——因为打断它的工具，恰恰是它自己提供的。

这一回写机制必须被放到亚隆所定义的核心疗效因子群中进行审视——这正是配准性生成最深层的寄生之所。成员丙在给出过程评论时，他体验到的正是“利他主义”这一疗效因子：他通过对他人成长做出贡献而感到价值。成员的每一次袒露与被接住，都在增强“团体凝聚力”——那个被亚隆视为一切疗效发生之基底的因子。成员丁观察甲的反馈、观察团体的正面回应，进而习得“在这里可以这样说话”，他正在经历“人际学习”。换言之，配准性生成的发生，并非某种游离于治疗过程之外的副作用，而是恰恰发生在治疗被定义为“有效”的那些核心时刻。团体越是成功，疗效因子的运转就越流畅；而疗效因子的每一次流畅运转，都为“这样做是对的被接住的”提供了新一轮经验验证，从而进一步配准了成员的定向基础。亚隆的疗效框架没有为这个自反性风险预留任何概念空间——它假设所有的“团体凝聚力”、所有的“利他主义”、所有的“人际学习”朝同一个方向累积，而那个方向的终点是成员独立人格的增强。本章的诊断则指出：在担保被遗忘的那个临界点之后，这些疗效因子仍在运作、仍在产生团体内的积极体验，但它们所增强的那个“人格”，已经不再是能够在不同担保之间维持弹性的主体，而是与这一特定的担保场高度耦合的配准性主体。两者在团体内可以被同样的量表测量为“成长”；只有在团体外、在担保消失的情境中，二者的差异才会显现为一对相反的结果。

在疗效因子的持续运转中，有一个本应充当校准器的关键机制需要被单独提出分析：此时此地的干预。亚隆式的此时此地——带领者或成员在互动中直接指出“这一刻我们之间正在发生什么”——在设计上恰恰是对隐性语法凝固化的持续干扰。当带领者说“我注意到刚才你给他反馈之后，整个团体沉默了大约十秒钟”，她正在做的，是激活成员对担保本身的一瞥。然而，在担保已沉降为地基的长期成功团体中，此时此地的干预存在着被语法吸附的风险：当成员们也学会用同样精致的句式执行此时此地的观察时（“我注意到刚才那个沉默让我有些不安，我想邀请大家看看这个不安是什么”），此时此地本身便从一个对语法的干扰器，退化为语法内部的一个高水平动作。当对担保场本身的质疑也被纳入担保场认可的语法模板时，担保就获得了对它最危险的潜在威胁的免疫——它允许成员在句式层面讨论“我们之间是否太安全”，却从不允许这种讨论离开那套精致句式。此时此地的干扰功能，在担保遗忘的深层结构中被语法重新驯化了。

## ■ 七、配准的谱系：突变还是渐变中的加速

在推进到形态学指标与理论对话之前，必须回应一个方法层面的问题。本章的核心诊断——配准性生成是长期成功运转的团体的结构性风险——其论证逻辑若要保持发生学纯度，就必须避免一个隐蔽的陷阱：宣称在语法最稳固、担保最成熟的团体中观测到了一种“形态突变”（从独立决断到语法依赖的质的飞跃），而这一突变恰好与取样窗口的边界重叠。如果只检视了语法稳固度最高的一批团体，那么观察到的“突变”——张从冒险袒露到精致执行的跃迁——究竟是事物本身的边界，还是取样窗口的边界效应？本章当前的分析样本（张的轨迹及其理论推演）无法直接回答这个问题，但有义务给出一个配准的谱系化推演，把当前样本中分不出的梯度，用理论推演画出来，并据此收缩结论的力度。

在低语法稳固度的团体中（混沌期或刚进入稳定期的团体），担保网络尚未沉淀为不可见的认知地基。成员能在不同互动模式之间保持基本的弹性：他们可以今天尝试用“我陈述”表达愤怒，明天又退回沉默，后天又在带领者的邀请下笨拙地做一次此时此地。这个阶段，担保仍在成员的知觉前景中，袒露仍然携带着明确的风险感知——成功是间歇性的，担保尚未被遗忘。成员对外部人际情境的定向能力基本上未被系统性地改写，因为团体内尚未形成一套足够精致、足够成功地替代独立裁断的语法。

在中等语法稳固度的团体中，担保开始从知觉前景向认知背景沉降。成员已积累了大量成功互动的经验，隐性语法已形成可辨识的模板——“在这里这样说通常会被接住”。跨场弹性开始出现裂纹：成员在团体内越来越善于表达，但在团体外、在异质的语法环境中，其应变的速度与精度开始出现轻微但可测量的下降。这一阶段的配准是渐变的，而非突变的；它表现为成员的外部人际策略逐步向团体语法风格收敛，但尚未丧失对担保本身的感知——成员仍然能在被提醒时，识别出“我刚才用的是团体里的那套说法”。

在高度语法稳固度与精致化的团体中（即本章的分析焦点），担保完成了自然化的最终沉降。成员在团体内的高水平表达与团体外的定向困难之间的裂口急剧扩大。张在团体内的精致反馈与他在跨部门会议中可能的失语，正是这一极端落差的具象化。在这个阶段，配准性生成的风险不再随成功缓缓上升，而是加速上升——因为语法精致度越过某一阈值后，语法自身开始执行主动评价功能（语法合规性检测系统启动），担保不仅被遗忘，而且语法本身已不再能容纳对语法自身的质疑。张的“突变”并非发生在成功的某一点突然降临；它是在语法精致度持续提高的过程中，当担保从“被感知为背景”滑入“被体验为自然本身”的那一临界区间内，被加速完成的。

这一谱系推演给出了一个直接的可检验预测：在控制了团体运转时长的前提下，随着隐性语法的精致度（以句式归一度、情感表达的包装标准化程度等可编码指标衡量）逐步提高，成员离团后独立决策能力与语法内化深度的关系，将从早期微弱的正相关（“学会了表达”带来

的初始正面迁移），在越过某一语法稳固度临界区间后转为稳定的负相关（语法已沉降为地基，跨场弹性被系统性消耗）。换言之，配准不是一次突变，而是渐变中的加速；语法的精致度——而非抽象的“成功”——才是配准从可能性变成结构风险的关键调制变量。如果这一预测被后续经验研究所支持，本章的核心判断便获得了一项重要的精确化：担保的存在本身并不必然导致配准；是担保在成功中的沉降与语法的精致化，共同开启了这个自反性循环。

## ■ 八、配准的痕迹：形态学指标与反例的消失

配准性生成在成员的主观报告中几乎不可见——被配准者体验到的只有成功、成长和被爱。但这并不意味着配准没有留下可观察的痕迹。以下提出四个团体内可编码的形态学指标，它们构成一个可操作的检测面板，用来判断在一个长期运转的成功团体中，配准是否已经深度发生。其中，沉默的消灭与混沌期沉默的区分已在第七节的谱系推演中奠定了判据基础：关键在于是否有沉默，而在于沉默是否仍保留着不被语法预制的开放性。

第一，高风险负面情绪的句式归一度。这不是说使用“我陈述”本身有问题。问题是：当成员表达愤怒、受伤、羞辱这类高风险负面情绪时，他们是否仍有能力使用非标准化的、不便包装的、携带原始攻击性的语言——并且团体仍能消化它。如果连最剧烈的情绪都被统一收纳进了标准“我陈述”模板的包装盒内，裁断的代理大概率已经完成。此时的团体，不是“一个能容纳真实情绪的地方”，而是“一个能将所有真实情绪溶解为标准语法的溶剂室”。

第二，沉默的结构性消灭。混沌期的团体会反复经历漫长的、令人坐立不安的沉默——那时成员是真的“不知道能说什么”。成熟期的团体中，沉默越来越短。不是因为成员不再焦虑，而是因为每个人都已内化了一套填充沉默的脚本：“此时可以给出一个过程评论”“此时可以分享一个感受”“此时可以邀请沉默中的某人说话”。沉默——作为唯一一种无法被语法预制的状态，作为那个在无担保中“我真的不知道该怎么办”的粗暴焦虑的容器——被从团体的感知光谱中悄悄抹掉了。沉默还存在，但它已经被驯化了：它从“真不知道说什么”变成了“用语法允许的方式说到下一句话出现为止”。区分两者的判据是：当沉默降临时，打破沉默的那句话是否携带了一个不容于既有语法的、粗粝的、尚未被包装的信号——还是仅仅执行了一次符合语法的过渡动作。

第三，失范实验的情感排斥。在一个仍保有发生性弹性的团体中，当某一成员做出违背既有语法但可能携带新信号的互动（如直接说“你能不能别再那样对我说话”）时，其他成员的反应通常掺杂着不适与兴趣：“她这么说话让人不太舒服，但也许我们该看看这里有什么。”而在配准深度完成的团体中，这类实验会激起一种近乎本能的负面反应，不是批判其内容，而是排斥其形式——“你这样说话让人很难接住你。”失范者不被反驳，而是被友善地告

知她的参与方式不兼容于该场域。这种排斥的独特质地是：它是无恶意的，它是为团体整体的情感安全而做出的，但它完成了对语法之外一切信号的有效过滤。

第四，归因的向内坍缩与反例的消失。在长期成功运转的团体中，当某个成员在团体外遭遇人际失败时，她自己的归因模式以及团体其他成员协助她做的归因，都会系统性地向内坍缩——“我可能还没有完全学会把团体的东西带到外面去”“我可能还需要做得更多”。这种归因模式的反复使用，使任何可能构成反例的外部事件，都在被体验为“失败”之前，即被预消化为“过程的一部分”。反例不是被驳倒的，而是从未被允许成形。一个机制，如果其反例的生成机制本身已被它切断，它便拥有了理论免疫能力——而这正是它最危险的时刻。

这里可以引入一个从张的案例中推演出的“系统外失败”场景，以加固第四项指标的论证力量。假设张在团体进入最成熟阶段后，参加了一次职场中的跨部门会议。会议议程不清晰，与会者来自不同背景，没有共享的“我陈述”语法，争论混杂着含蓄的人身攻击与未言明的利益博弈。在某个节点，张被一位高管突然点名，要求他对一个他不熟悉的议题给出即时判断。团体内训练出来的那套精密的反馈技术——停下来、感受身体反应、用“我注意到……”句式包装后再出口——在这次互动中完全失效，因为对面的人既不会接住他的脆弱，也不会给他三秒钟的停顿。在这一刻，张可能表现出两种极端反应之一：他在高压下完全失语，并在事后归因为“我还不够勇敢”；或者他过度攻击性，用一个粗糙的判断硬顶回去，事后又为此自责。关键不在于他犯了错；关键在于他事后对这事件的“消化”方式。“团体内成功—团体外失败”不是被读为配准的痕迹，而是被他读为“我还在路上”——而这个解读本身，正是配准最深的痕迹。

## ■ 九、配准的不可还原性：它不是这些已有概念的别名

任何一个新命名，必须自证它不是已有概念的换装。以下逐一划清配准性生成与最可能混淆的既有概念之间的分离线。

### （一）配准性生成不是内化

内化概念预设了能力在担保撤除后可独立保留；配准描述的是担保本身在成功中被遗忘，以致担保不再是可被感知的辅助物——成员失去的不是“技能”，而是“感知自己正站在某块地板上”的能力，进而失去了跨越不同担保、在不同地板之间维持弹性的可能。分离线的经验检验：若五年追踪发现，离团成员在无担保情境中的独立决策力与其团体内语法熟练度呈负相关，且其跨场弹性的核心指标（在不同担保场之间切换定向方式的速度与精度）显著低于匹配的对照组，则内化假设被否定，配准成立。

## （二）配准性生成不是移情

亚隆已诊断过成员可能将“团体”本身当作依赖对象，即“对团体的移情”。移情理论中的代理者是一个可以被指认、被讨论、被修通的心理对象——成员意识到“我把团体当成母亲”，便获得了审视并修正这一投注的可能。配准性生成中的代理者不是任何心理对象，而是一套认知操作的预制模板——成员连“母亲”这个词，都已经只能用团体的语言说出。修通移情的技术对此无法奏效：一个人无法用语法之外的方式去质疑语法。分离线的经验检验：若在已对团体移情做了充分修通的成员中仍能观测到离团后跨场弹性显著退化，且退化程度与移情强度无关（而仅与语法内化深度与担保遗忘程度有关），则移情理论被排他，配准成立。

## （三）配准性生成不是福柯式的规训

规训主体在服从体验中形成——他被外部目光凝视，感到压抑，其服从是以牺牲一部分自我为代价的。配准性生成的主体在关键成长节点体验的却是积极、赋能、被深刻看见后的深度亲密。其跨场弹性不是被夺走的，而是被愉快的体验从内部包抄并压缩的——担保的成功沉降，使得成员不再感知到自己正依赖着一块地板。分离线的经验检验：若成员在团体内体验到赋能感越高，离团后跨场弹性越差，则配准成立——这种通过积极体验实现的弹性丧失，是规训概念没有专门辨识的变形。

## （四）配准性生成不是技能幻象（“只是在团体内学会了一套情景化技能”）

技能幻象说将问题定位在“教的内容太窄、不适用于外部环境”；它仍预设技能可被教、只是这次教的技能不够通用。配准性生成不在“教的内容”层面，而在“教的方式已改写了被教者所是的那个存在形态”层面——技能的习得同时沉降了担保，使技能的使用者遗忘了他正在依赖担保。技能幻象说只需要扩大课程内容；配准性生成揭示的是：任何被纳入担保场的“扩大的内容”，都将被同一套配准机制吸收，重复同一轮担保遗忘与语法沉降的循环。分离线的经验检验：若实施在团体内教“如何在无担保情境中独立决策”的干预后，离团成员在该项技能上的表现仍不优于未经训练的基线组——因为这项技能在团体内习得时，同样伴随着担保的遗忘——则技能幻象被排他，配准成立。

## （五）配准性生成不是 Turner 的共睦态消退

Turner 命名了共睦态——在仪式化场域中生成的、超越日常社会角色的强烈社群体验，并指出它是“阈限”的产物：只在“非日常”时空中发生，一旦返回日常便消退。共睦态的失败是消退——体验结束了。配准性生成的失败不是在返回日常后消退，而是在共睦态最浓烈的那一刻，担保的沉降已经发生：成员在高峰体验中，学会的不只是“这里曾经很美好”——而



是“美好的感受需要这套语法作为条件”。分离线的经验检验：若离团成员在旧语法环境（如重返同类型团体）中迅速恢复高水平定向能力，而在完全陌生的无语法互动环境中表现极差，则配准成立——能力仍黏着于旧语法，而黏着的不是“体验的记忆”，而是“担保的被遗忘”，共睦态消退理论不足以解释。

## （六）配准性生成不是嵌入性

Granovetter 的嵌入性理论指出，经济行动者嵌入于社会关系网络之中，其理性与决策模式被网络结构塑造——紧密网络（强关系）提供信任与规范，但可能限制信息获取与异质行动能力。这一框架确实触及了担保网络对行动者的限制性影响，但它未能区分担保感知与担保遗忘。嵌入性中的行动者始终知道自己在嵌入一个网络——嵌入性理论恰恰预设了一种可被分析者从外部观察到的“嵌入状态”。而配准性生成的独特机制恰恰在于：担保网络在成功体验中沉降为不可见的认知地基后，行动者不再感知到自己被嵌入任何网络。他不是“强关系限制了异质行动能力”的意义上受到约束——他是在“已无法意识到自己正站在关系网络之上”的意义上被彻底改写。嵌入者可以选择挣脱网络（尽管成本可能很高）；配准者已无法感知到网络的存在，因此无从挣脱。分离线的经验检验：在控制了网络紧密度的前提下，对网络嵌入具有高度自觉意识的行动者，其在异质场域中的决策质量显著优于低自觉意识者，而在配准深度完成的团体条件下，这一差距被担保沉降所特有的遗忘深度所放大——嵌入性理论不生成这一预测。

## （七）配准性生成不是场域-惯习

布迪厄的场域-惯习理论描述了特定场域如何锻造行动者的惯习——一套内化的感知、判断与行动图式，并与原场域高度耦合，难以迁移至异质场域。这一框架与本章的分析有显著的表面同构性，但存在两个关键分离点。第一，惯习的锻造通常伴随着权力关系、利益争夺与象征暴力——行动者在场域中的位置决定了其惯习的形态。配准性生成发生在一个人人平等、真诚互惠、主动袒露的情感亲密场域，权力的面孔是缺席的。行动者不是被外部力量驱动去适应场域，而是被发自内心的归属感和被爱感牵引着走进配准。第二，布迪厄的理论不专门处理“成功体验如何加速场域的不可见化”——场域中的行动者或多或少仍保持着对博弈规则的操作性感知。配准性生成则精准定位了一个特定的阶段：当团体在情感效能上取得压倒性成功时，担保沉降为不可见的认知地基，惯习不再是“我在这里学会了一套规则”，而是“这些规则本身就是世界运行的正常方式”。分离线的经验检验：在控制了惯习耦合程度的前提下，对场域权力博弈具有清晰感知的行动者，其在异质场域中的适应速度快于感知模糊者，而配准则预测了



一个反向的效应——在担保已沉降的团体中，越是体验到平等与亲密、越是不感知权力存在的成员，其离开场域后的跨场弹性反而越差。这一反转是场域-惯习框架无法引出的。

### （八）配准性生成不是情境依赖性学习

认知心理学的情境化学习理论揭示，在丰富情境线索下习得的知识与技能往往难以迁移到情境线索缺失的新环境。这一框架与本章描述的现象在行为层面有重叠，但它不处理担保遗忘。情境化学习中的学习者知道自己在学习——他知道自己知识是“在这个情境中学到的”。他可以在被提醒时，有意识地去重新部署自己的知识，尝试应用于新情境。配准性生成破坏的正是这个“知道”：成员不知道担保已沉降为地基，因此不会主动去调整自己的定向方式——他以为自己的精致表达是自然自发的。分离线的经验检验：若在离团成员中实施“情境线索再激活”干预——告知其团体内习得的技能具有情境依赖性，并提供认知工具以识别情境差异——仍无法显著改善其在新情境中的独立决策能力，则情境化学习框架被排他。情境化学习的核心缺陷是“线索缺失”，可以通过线索重建来部分补救；配准性生成的核心缺陷是“担保被遗忘”，在线索重建之前，必须先修复对担保本身的感知——而这一步恰恰被语法合规性检测系统所阻抗。

## ■ 十、跨领域推衍与边界

配准性生成发端于长期成功运转的无领导团体这一特定场域，但它的结构不变量——在一个以培育自主为目的的长期封闭系统中，系统的成功运转内生出一套重写主体定向基础的隐性担保网络，并通过担保的沉降启动语法合规性检测——允许提出一个可外推的假说：在任何以“赋能”“自治”“去中心化”为核心承诺、且依赖高度共享隐性规范维持封闭运转的系统中，均可能观测到同类结构风险。

以下四组对照场景为作者推演，标记为待检验假设。第一，教育学领域：制度化探究式课堂中，学生学会的是表现思考的语法，而非独立提出不在课程框架内的真问题的能力。隐性的“好的探究表现”语法越成熟，教学评估越显成功，学生的独立追问弹性反而越微弱——他们不再意识到自己的“思考发言”依赖于教室这个担保场提供的安全背书。第二，组织行为学领域：长期封闭的精英自治团队在跨部门协作或陌生情境中暴露出的判断力衰减，或因团队成员已被内部隐性决策语法深度配准——什么算好论证、如何表达异议而不被边缘化——担保的沉降使得成员在跨场时不自觉地沿用只在原场域内有效的决策方式。第三，文化批评领域：某些高凝聚力亚文化社群的内部“黑话”固着，使成员用日常朴素语言表达自身的能力萎缩。每一次对着圈外人说出圈内人才懂的术语，都是一次配准行为的执行——成员不再意识到自己的表达依赖于圈内担保场的共享语境。第四，政治学领域：参与式民主实践中的“仪式化参与”模

式，其中参与者学会了将个人诉求翻译为特定政策话语，但关于“我到底想要什么”的原初判断，被悄然改写为“什么诉求能在这个流程中被认可”——流程本身已成为不被感知的担保场。

上述四组推衍的检验位点相同：都需要在各自场域中观测，当系统成功运转到某一阶段后，“系统内的表现优秀者”在系统外、无条件场的情境中是否表现出比“系统内表现中规中矩者”更差或至少不更优的跨场弹性。如果该反转在多个场域中被稳定复现，则配准性生成便有理理由被视为一个跨领域的生成机制，而不仅是一个特定治疗形式的偶发副作用。

本章也严格标定诊断的边界条件。配准性生成不适用于：第一，尚在混沌期、互动语法尚未稳定固化的团体（配准的土壤条件未满足）；第二，有明确技能传达或行为矫正目标、不声称培育独立判断的结构化干预（无配准所必需的“自主”承诺）；第三，已建立周期性外部干预以持续扰动语法固化的系统（配准可能被延迟、打断或变形——这是一个本章无法覆盖的开放性问题）；第四，从不以培育独立决断为目的的纯粹情感支持团体。

一个尤为关键的阴性案例寻找方向是：那些长期运转但互动语法始终保持某种粗粝的、未被精致化的开放性的团体。如果这类团体存在，并且其成员的跨场弹性未被系统性消耗，那么本章的诊断便获得了一个重要的限定条件——配准的发生不仅要求“长期成功”，还要求“语法的高度固化与精致化”。如果该方向上的检验显示，语法粗粝但成功运转的团体同样出现了本章预测的跨场弹性退化，那么配准的机制便比本章当前模型所指涉的更为基础——担保在成功中的沉降本身（而非语法的精致程度）就是配准的充分条件。

一个尖锐的次级追问是：那些在团体中成功体验到大幅成长并似乎确实将能力迁移至外部的成员，是否构成对本章诊断的个别反例？本章不否认此类个案的存在。本章的诊断层级是结构性的，而非全称性的：在给定的结构条件下，成员跨场弹性在统计上被系统性消耗的概率远高于它被培育的概率。个别反例的存在并不推翻结构性诊断——正如吸烟者中个别长寿者不推翻吸烟的健康风险。但本章必须指出，当这些“成功迁移者”的迁移环境是高度类同于团体担保结构的小型亲密社群（如核心家庭、挚友圈、小型团队）时，他们实际并未离开担保场，而只是换了另一张担保网络。他们的“成功迁移”不是跨场弹性的证据，而是担保沉降未打破的证据——他们从未被要求在不同担保之间重新校准定向，因此从未暴露其弹性已被消耗的事实。

## ■ 十一、结论：担保的自我否定

本章命名了“配准性生成”这一过程：在长期成功运转的无领导团体中，团体为维持运转而生成的那一套隐性语法与担保网络，在反复的成功互动中完成了担保的沉降——担保从可感

知的安全垫变为不可见的认知地基。正是在这一沉降完成之后，担保网络从被动承托转换为主动评价：隐性语法通过对语法合规行为的正面标记，使不符合语法的原初裁断信号在尚未进入意识之前便被过滤。成员所失去的，不是在担保真空中的纯粹独立决断——那样的东西从来不曾存在——而是在不同担保之间感知担保的边界、维持可迁移的定向弹性的能力。团体越是成功运转，成员的精致表达能力便越发酵熟，而其感知担保并跨越担保的能力便越发被成功本身所消耗。配准性生成因而不是游离于疗效之外的副作用，而是寄生在亚隆所定义的利他主义、团体凝聚力与人际学习等核心疗效因子的内部，是成功治疗过程本身携带的自反性风险。

这个判断不是对亚隆遗产的否定。它是亚隆遗产内部尚未被命名的那个自反性极限：当提供安全的机制本身的运作，已使安全无法再被感知为安全——当担保已沉降为认知地基——那些在“我们”之中所习得的勇敢，要如何才能在不同担保场的边界处持续站住？这一追问不只适用于治疗团体。它以更普遍的方式向一切以“赋能”与“自治”为承诺的封闭系统发问：如何在担保场的运作机制中，区分“增强跨场弹性”与“将弹性写入对单一担保场的依赖”这两个相反的过程？

证伪条件。本章的诊断在以下条件下可被否定：若对足够样本量的长期成功无领导团体成员的追踪研究——追踪时长不少于离团后两年，且采用独立第三方评定的、不依赖成员自我报告的高压情境决策质量指标——稳定地显示，（a）离团成员的跨场弹性（在不同担保条件下切换定向方式的速度与精度）显著高于匹配的非参与者对照组；并且（b）成员在团体内语法内化深度（由独立编码的句式归一度等客观指标衡量）与上述离团后跨场弹性呈稳定正相关——则本章的核心机制被证伪，配准性生成不成立。在满足社会研究伦理的前提下，实验设计可考虑引入两组对照：一组为曾参与长期成功无领导团体而后退出的成员，另一组为从无类似经历的匹配人群，二者在标准化高压决策模拟任务中——任务涉及的互动对象不共享任何团体隐性语法——的表现差异。若本章诊断成立，则前一组在任务初期的表现应显著差于后一组，并在任务中表现出更强烈的对“他人反馈”的依赖倾向，以及更低的在不同定向策略之间灵活切换的能力。

## ■ 注释

〔1〕本章此后所称的“团体”，除非另有说明，均特指经过改切后的分析对象，即已进入长期成功运转状态、互动语法高度稳固、担保网络高度成熟的无领导团体。对于仍处于混沌期或语法未固化的团体，本章的诊断不作直接推论。这一限定亦适用于后文所有关于“团体”的论述。

\\[2] 本章中“张”的案例为思想拓展性质的例示，系根据典型轨迹构造的叙事缩影，包含匿名化、合成与简化处理，不代表任何特定真实个体。其功能在于演示机制展开的可能性，而非提供实证个案证据。亦请参见本节末的方法论自省。

\\[3] 本章使用的“隐性语法”概念指团体在长期互动中自发形成的、关于有效表达与反馈方式的未明言规则集合；“担保网络”则由历史、对等、身份三重编织构成（详见第三节）。担保网络提供执行评价所需的情感权重，隐性语法提供评价所需的技术标准；二者在担保沉降完成后合成为一个主动筛选的“语法合规性检测系统”（详见第五节）。

\\[4] 本章对亚隆疗效因子的分析，系在其所述框架的基础上引入一个自反性维度，并非否定这些因子在团体治疗中的积极作用。此处所指出的风险，是成功治疗过程在特定结构条件下可能产生的非意图后果。

## ■ 附论一·全球近邻补切（编辑增补）

本章已与内化、移情、规训、技能幻象、共睦态、嵌入性、场域-惯习、情境依赖性学习八家逐一划界，且每条都附了经验检验——这是本站学员稿中近邻处理最完整的一次。此处补两个未被点到、却与本章核心机制几乎同构的祖宗。

其一：Argyris 与 Schön 的「使用中的理论」与「熟练的无能」。二人论证：组织成员实际遵循的理论与他们自述信奉的理论长期分离，而防卫性惯例之所以顽固，恰恰因为成员对它执行得极其熟练，熟练到无法察觉自己正在执行——他们把这称为熟练的无能。这与本章「担保沉降为不可见的认知地基之后，语法被体验为说话的自然方式本身」几乎逐句对应。分离线在被遮蔽的东西：Argyris 的成员看不见的是自己的行动规则（有一个可被外部观察者揭示的差距），本章的成员看不见的是自己脚下的担保（不是规则被误述，而是地基不被感知）。判决性对照预测：Argyris 的处方是「双环学习」——把使用中的理论摆到桌面上讨论，差距即可缩小；本章预测这类干预对配准无效，因为讨论本身只能用被配准的语法进行（本章第五节末已写出这一点）。设计一次「拆解我们自己的互动语法」的团体干预，若成员的跨场弹性随之改善，本章机制应并入双环学习；若讨论本身迅速被收编为一次高水平的过程评论，本章的判断获得支持。

其二：贝特森的第二级学习（deutero-learning）。贝特森区分了学会某项内容与学会如何学习，并指出后者一旦形成便沉入背景、极难被主体察觉或修改。本章的「配准」在形式上正是第二级学习的一个负面案例：成员在团体中学会的不是内容，而是一套关于「什么算有效表达」的学习语法。分离线在有无评价环节：第二级学习是一个中性的形成过程，不预设有一个系统在主动过滤不合语法的信号；本章的语法合规性检测系统要求担保网络在成员开口之

前完成正面标记与无声排斥。判决性对照：若在从不对表达形式给出差别回馈的长期团体中，成员同样出现跨场弹性退化，则「过滤」这一环多余，本章应退回第二级学习；若退化只在有差别回馈的团体中出现，本章的主动评价机制成立。

## ■ 附论二·与站内既发论文的划界（编辑增补）

与《被转嫁的成长》《缺锚运行的生成缺失》：本章是同批三篇的统一者，也是对前两篇的一次自我修正。前两篇仍在「内部能力对外部担保」的二元里说话，本章明确取消了这个二元——担保真空中的独立决断从来不存在，问题不是能力被接走或从未搭建，而是感知担保、在不同担保之间换步的弹性被成功本身消耗。读者若把三篇当成三个并列的命名，会错过这一层递进。

与《养生如何废止了生命》（2026-08-05 已发）：两篇是同一手法在两个域的执行——那篇用「回路闭合点」统一了身应、体知悬空、绞杀内履三篇，本章用「担保遗忘」统一了本批三篇；两篇都把结论压在同一句上：成功本身就是废止的刀刃。分离线在被消耗的能力：那篇是身体自行闭合代谢决策回路的能力，本章是主体感知自己所站地基的能力。

与《改不动的机器》《希望的伪修复》：作者站上已有两篇处理「修复动作本身成为问题」的论文。本章与它们的分工在于修复是否被感知——那两篇的行动者知道自己在修复（并因此不断加码），本章的成员已经感知不到自己正在依赖一套修复脚本。

一处内在张力，提请读者注意：本章第八节第四项指标写道，「一个机制如果其反例的生成机制本身已被它切断，它便拥有理论免疫能力——而这正是它最危险的时刻」；而第十节末处理「成功迁移者」这一最强反例时，把他们判为「只是换了另一张担保网、从未被要求重新校准，因而从未暴露弹性已被消耗」——这一处置在形式上正是本章自己所警告的那种反例吸收。本站不认为这使论证失效（作者随即给出了「语法粗粝但成功运转的团体」这一可独立寻找的阴性案例方向），但读者应把这两处并置阅读，并把该阴性案例的检验视为本章能否守住可证伪性的关键。

## ■ 附论三·编辑校订说明

一、正文未见本站方法论词汇泄漏，未作改姓性修改。

二、案例材料：注释[2]已明确「张」为「根据典型轨迹构造的叙事缩影，包含匿名化、合成与简化处理，不代表任何特定真实个体」，且第二节末附有方法论自省，本站未作改动。

三、文献表八条（Bion 1961、Bourdieu 1990、Dworkin 1988、Foucault 1975、Granovetter 1985、Putnam 2000、Turner 1969、Yalom & Leszcz 2020）逐条核对属实。正文提到的情境化学习一节未给具体文献，属泛指学术传统，未代为指定条目。

四、第九节八条分离线各自附有「经验检验」，第十一节证伪条件要求离团后不少于两年的追踪、独立第三方评定且两项条件同时成立——本站认为这是本批可证伪性最强的一篇，读者可据此逐条检验。



## 第三编

# 评估从两侧同时收窄

• • •

三章处理同一个动作：评估。

高于涵写"善的工序化"：工序不是在加工善，工序是在持续产出"未善"；而且它不需要你不动脑子地盖章，它需要你带着全部信念去论证那个章为什么非盖不可。少敏写人工复核如何同时生产控制证据与控制失效。胡敏写对于有不可逆时间窗的自组织过程，裁定这个动作本身就是干扰，与判断结果的对错无关。

第八章与第九章是全书咬合最紧的一对，合章专门处理它们。

## 第七章 带着全部信念盖下的那个章

本章原题《未善的代价：当“成为好人”变成一道工序》，作者 高于涵，原文见

<https://sdeuniverses.com/students/gao-yuhan/goodness-as-process/>

### ■ 一、一个被分类遮蔽的问题

董仲舒的“未善论”在汉语思想史上被装进了“性三品”的筐子，与孟子的性善论、荀子的性恶论并排摆在人性论的橱窗里。这种分类本身，已经构成了一次根本性的遮蔽——不是因为它“错了”，而是因为它用一套现成的格子，把一个格子本身装不下的问题给框住了。

孟子和荀子争论的，是善在不在人身上。孟子说在，所以内求——人有四端，扩而充之，皆可为尧舜。荀子说不在，所以外铄——人之性恶，其善者伪也，需要圣人的礼义来矫正。两个人的答案针锋相对，但他们在同一个战场上打仗。他们共享一个未经审查的前提：善，是一个可以被“在”或“不在”来描述的属性。它要么在人身上，要么在外面。问题只是搞清楚它在哪。

董仲舒做的不是在这两个答案之间找一个折中。他把问题本身换掉了。他不问善在哪里，他问善是怎么来的。答案是一句被他反复摩挲的话：“性比于禾，善比于米。米出禾中，而禾未可全为米也。善出性中，而性未可全为善也。”（《春秋繁露·深察名号》）

禾能出米。但禾不等于米。从禾到米，缺一道工序。这道工序，董仲舒把它交给了王。

这一步跨出去，整个讨论的地基就变了。孟子和荀子争论的“人之性究竟是善是恶”，在董仲舒的框架里变成了一个被降级的问題：不管是善是恶，反正你自己变不成善。你身上那个东西——不管它叫善端、善质还是什么——它自己不具备把“自身”加工成“成品”的功能。它需要一个外部的操作者。问题不再是“你本性好不好”，问题是“你有没有被碾过”。

由此，汉初以降的所有人性论争论——性善、性恶、性三品、性善恶混——都可以被读成同一个深层结构的表层变体：它们都在试图回答“人如何变好”，但它们都没有追问，当“变好”这件事被定义为一道必须由外部操作者执行的工序时，那道工序本身会对“人”做什么。

本章要做的，不是给这道工序再加一个更精密的诊断标签。标签已经太多了。本章要做的，是从这道工序的发生现场——董仲舒本人被那个大一统帝国的治理困境一步步逼到这一步的过程——出发，追溯它为汉语文明装配的那个最持久的隐性结构。这个结构，不是“权力垄断了善的定义”，而是“工序窃取了个体把自身生命体验转化为善的判断这件事的资格”。权

力垄断说的是：我不许你做。工序垄断说的是：你可以做，但我不认证，你做的就不算。前者依赖暴力，后者依赖定义。暴力可以被反抗。定义一旦被内化，反抗者自己都会觉得自己在“作恶”。

我把这个结构命名为：善的工序化。

它不是权力问题。它是资格问题。\\[1\\]

## ■ 二、董仲舒为什么非要把善做成一道工序

### ■ 2.1 不是哲学发明，是被逼出来的工程解

要理解董仲舒为什么非走这条路不可，得先放下“他是一个哲学家”这枚标签。严格来说，把董仲舒当作一个在书斋里进行概念思辨的哲学家，已经是对他的第一次误解。他的身份首先是汉帝国意识形态的总装师——不是他自己想当，是汉武帝把他召到那个位置上，逼着他回答那些非回答不可的问题。

他面对的不是书斋里的概念辨析。他面对的是一个在疆域上空前辽阔、在内部文化上差异极大、在治理手段上几乎白手起家的大一统帝国，在建政之初最棘手的治理难题：这个帝国，它的税赋怎么收？它的政令怎么通？它的人民，怎么觉得自己是“一家人”而不是被征服的六国遗民？

汉武帝在元光元年的策问里没有问“人性本善还是本恶”。他问了三件事：三代受命，其符安在？灾异之变，何缘而起？性命之情，如何可得？这三个问题，用帝国治理的语言翻译过来就是：我这个位置凭什么坐？为什么我治理下老出天灾？老百姓到底要怎么管？

这是一场具体的考试。它的规则是：汉武帝出一道题，董仲舒答一道；答完了，武帝就着答文再追一问，董仲舒再答。一问一逼，一逼一拐。《汉书·董仲舒传》记载的三次策问，每一次追问都把董仲舒往更深的墙角里逼了一步。他不是在设计一整套哲学体系——他是在悬崖边上，每被推一下，就往后退一步，然后回头看看这一步能不能踩稳。

但仅仅把这一过程描述为“一问一逼”的线性推进，仍然不够。在那三次策问的现场，三种压力是同时作用在董仲舒身上的——不是第一问他只考虑合法性、第二问再考虑儒生地位、第三问才轮到百姓认同。在汉武帝第一次策问“三代受命，其符安在”这句话被说出口的那一刻，三个约束已经同时压在董仲舒肩上：他必须给出一个答案，让皇帝觉得这个位置坐得稳（否则他走不出未央宫）；他必须让这个答案包含一个只有儒生能胜任的工作（否则整个阶层没有立身之本）；他必须让这个答案在被翻译成政令推行到郡县时，百姓不会觉得自己只是在

服从暴力（否则这套方案和秦制没有本质区别）。这三条不是三根依次敲下来的钉子。它们是同一块铁板上的三颗铆钉，同时铆死——任何一个铆不牢，整块铁板就崩了。

要更精确地看清这一步的“被迫性”，可以做一个思想实验。假设董仲舒当时面前有两个备选方案。方案A：直接继承孟子的性善论。人人有四端，扩而充之，皆可为尧舜。在这个框架里，王的作用是什么？最多是“辅助者”——提供学校、提供禄米、提供不被打扰的耕作环境。但“辅助者”不是不可绕过的：一个人在山野里独自修身，不靠王教，照样可以成为圣人。这个方案会立刻在第一个约束上崩盘——皇帝不是善的唯一加工者，他只是众多同行者中的一个。方案B：继承荀子的性恶论。人性本恶，必须靠圣人的礼义来矫正，所以需要外部权威。但“圣人”是谁？荀子没有说圣人必须是现任君主。在荀子的框架里，周公是圣人，孔子是圣人——圣人可以是已经死了的人，他的礼义写在书里，后人照着做就行。那么这个矫正工作的实际执行者是谁？可以是任何通晓礼义的人——儒生可以做法官，文法吏也可以做法官。第二个约束立刻崩盘：儒生没有不可替代的生态位。

于是，在两道方案的缝隙同时显现的那个瞬间，第三条路被逼了出来。不是孟子，不是荀子，而是那个从农业生产的日常经验里借来的意象：禾与米。善质是天给的（既安抚了孟子的直觉——人身上确实有向善的东西，又不是荀子式的“性本恶”），但善质不是善（堵死了方案A的漏洞——你自己变不成圣人），从善质到善需要一道工序（满足了第二约束——这道工序的执刀人必须通经义，文法吏做不了），而工序的执行者是王（钉死了第一约束——王不是辅助者，王是不可绕过的唯一加工者）。百姓在这个框架里不是被秦制式的暴力压服的——他们被告知，服从王教不是奴隶服从主人，是禾苗接受碾米人的加工，是一种让自己变得更好的“成长”。第三约束同时被满足。

这一步，不是董仲舒在书斋里从容挑选出来的。是孟子那条路和荀子那条路在汉初的现实压力下同时被证明走不通之后，剩下的唯一一块能同时踩稳三条边界的立足点。那块立足点，就是“未善论”。

## ■ 2.2 “未善”不是一个诊断，是一道工序的入口

董仲舒在《春秋繁露·深察名号》中写道：“天生民，性有善质，而未能善。于是为之立王以善之，此天意也。”〔2〕

这句话的逻辑，像一把锁一样精密地铆合在一起。第一步，天给每个人发了一个初始包——善质。善质不是善，是“能出善”的材质，就像禾不是米，但禾是能出米的那个东西。第二步，但禾自己变不成米。善质自己变不成善。这中间缺一道工序。这第二步是整条逻辑链里最致命的一环。它不是一句“善质暂时还没变成善”的描述——它是一句“善质根本不可能靠

自己变成善”的判决。它的潜台词是：善质，就其本来的性质而言，就不具备把“自身”加工成“善”的功能。它不是还不够好，它是“不可能靠自己变好”。第三步，那道工序的执行者，是王。王不是自己想当这个执行者——是“天”为之立的。天的意图本身，就包含了“必须有一个王来做这件事”这个子命题。

这才是董仲舒比孟子和荀子都更根本地改变游戏规则的地方。孟子说人皆可为尧舜——你不需要外部权威认证，你自己就可以成圣。荀子说人之性恶，其善者伪也——你需要外部加工，但“伪”这个行为本身，所有人都有资格做（圣人“化性起伪”不是垄断性特权，而是示范）。而董仲舒说的是：你能成圣的那道工序，握在一个你永远不能绕过的人手里。你自己是做不出“善”这个成品的，就像禾的种子，如果经过碾米人的手，它就永远只是禾。

这里必须插入一笔关键分析。董仲舒并没有说“普通人不能做好事”。在《春秋繁露》的诸多篇章里，他承认普通人在日常生活中可以有符合道德的行为——孝敬父母、尊重兄长、忠于职守、诚实守信。但这些行为，在他的术语体系里，不叫“善”。它们叫“善端”，叫“善质”的表现，叫“未善”。它们不是善，就像一个人会跑会跳不能证明他是运动员——运动员的身份，必须由外部机构认证；他跑出的成绩，必须由裁判确认有效。

这个区分，是“善的工序化”的核心装置。它把“善”从一个人人皆可参与、人人皆可在其中发现自己的日常实践，变成了一个只属于外部权威的成品认证权。你可以一辈子行善——真正意义上的、在每一个具体情境中都做出了对的选择的行善——但如果你没有经过“王教”这道工序的认证，你的那些善行就只是未加工的原料。更严重的是，在程序内部，你的善行甚至可能因为“未经认证”而被判定为“似是而非”的伪善——因为它没有按照经义的标准来校准。

但这里必须再深挖一层。董仲舒选择“禾与米”这个比喻，本身就是一个需要被追问的生成层面的事件。他本可以用其他意象。他可以说“性如璞，善如玉”——璞是未雕的玉，玉是雕成的器。这个意象在先秦已经流通（韩非子用过“璞”喻），它同样能论证“需要外部加工”。但他没有选璞与玉。他选了禾与米。

璞玉之喻与禾米之喻之间，有一道关键的区别。玉工雕琢璞玉，每一次下刀都要看璞的纹理、色泽、瑕疵的走向——玉工必须时刻判断，这一刀下去，是在成全这块玉，还是在毁掉它。玉工和璞之间，有一种持续的、不可完全预设的互动：璞在抵抗玉工，也在引导玉工。而碾米人碾禾，他不需要看禾的纹理。禾对他来说是一批一批处理的，他只需要启动碾子。碾子是标准化的，禾是标准化的，米也是标准化的。碾米人不需要在每一粒禾面前停下来，问自己“这一粒是否该轻一点”。

董仲舒选择禾与米，不是一个随意的修辞决定。他是在两种“外部加工”的想象之间做出了选择：一种是需要持续判断、持续与被加工物互动的“雕”，另一种是可以被标准流程化、可以被批量执行的“碾”。他选的是后者。这个选择在他所处的历史语境中可以被精确地理解——他要为一个疆域万里、人口千万的帝国装配一套可以批量执行的教化程序。雕玉太慢了，而且要求每一个雕工都具备不可复制的个人判断力；碾米可以交给出身各郡国的儒生，只要他们学会了操作碾子——六经。这个意象一选，“工序”就从隐喻变成了制度：王教不是圣人在每一个个体身上做的独一无二的雕刻，而是帝国机器对全人口执行的标准化加工。

这就是董仲舒最深的那层沉默。他从来没有明确地说过“个体无权判断自己是否在行善”——他不这么说，他也不需要这么说。他用一整套精密的类比——禾与米、卵与雏、璞与玉——把这种无权状态论证成了“事物本来的样子”。不是谁剥夺了你的权利。是你本来就没有这个东西。禾能说“我已经是米了”吗？不能。米这个身份，是碾米人给的。

而这一步，就是“善的工序化”区别于“权力的垄断”的根本所在。权力垄断说的是：我不许你做。工序垄断说的是：你可以做，但我不认证，你做的就不算。前者依赖暴力，后者依赖定义。暴力可以被反抗。定义一旦被内化，反抗者自己都会觉得自己在“作恶”。

## ■ 2.3 被工序反向逼出的“自我认证权”

这里必须插入一笔生成机制自反：本章用来对抗“善的工序化”的那个支撑点——“个体把自身生命体验转化为善的判断的资格”，简称“自我认证权”——它本身是从哪里来的？它不是某种永恒不变的人性禀赋，不是孟子的四端之心被本章悄悄换了个名字又请回来。如果本章暗示“自我认证权”是一个在工序出现之前就天然存在、只是后来被剥夺了的好东西，那本章就在最后一刻滑回了“把结论当起点”——把一个被历史逼出来的位置，误认成了等着回归的永恒本质。

必须说清楚：“自我认证权”不是从来就有的。它是在“善的工序化”被装配之后，才作为这道工序的对立面，被同时生产出来的一个影子概念。在董仲舒的禾米之喻把“你从自己的生命体验里熬出来的判断”判定为“不算”的那一刻，那个被判定为“不算”的东西，才第一次作为一个可以被言说、可以被争取、可以被论证的缺席位置——一个“被工序窃取资格”——出现在思想视野里。

也就是说，不是先有一个“自我认证的个体”，然后工序把他手里的钥匙夺走了。而是工序在装配它自己的同时，把“那个没有被碾过的、自己从生命里熬出判断的自我”这个形象，作为一个被它排斥的外缘，一并生产了出来。这个形象从它诞生的第一天起，就不是某个黄金



时代的遗民——它只是工序的投影。你站进工序的光里，影子就落在地上；你把光关掉，连影子都没了。

这一笔对于理解本章来说极端重要。它意味着：本章对“善的工序化”的批判，不是在号召“回归那个本真的自我”——因为那个“自我”不是一个等着被回归的实体，而是一个在工序的挤压下被反向逼出的空缺位置。你抓不住它，你只能指着那个位置说：“这里本来可以有东西。是那道工序，让它变成一个洞。”

现在，还可以更具体地回答审稿人的追问：这道“资格窃取”的生成机制锚点，究竟在一个人生命的哪个具体时刻被铆死？它不是在成年人被驳回谏言的那一刻，也不是在士人因言获罪的那一刻——那些时刻只是锚被拧紧的后续加固。真正的第一颗铆钉，是在汉代一个学童第一次被教“性比于禾，善比于米”的那堂课上。在那堂课上，他被要求接受一个定义：他自己从日常生活里积累起来的那些关于“什么事是对的”的切身判断——帮邻居收过庄稼，替受欺的同伴说过话，不忍心看牲口挨打——那些不是“善”，那些只是“禾”。他不是在那一刻被剥夺了什么他已经拥有的东西。他是在那一刻被定义成了一种材质——一种自己不具备成为成品的资格的材质。他是在还没有学会说“这是我的判断”之前，就已经被教会了一句话：“你的判断，在没被碾过以前，不算。”

### ■ 三、三道装配工序：善的工序化如何被焊接成一台不可拆的机器

上一节厘清了“未善论”的发生逻辑：它不是董仲舒的书斋发明，而是他在三条同时施压的硬约束下，在孟子方案与荀子方案同时走不通的缝隙中，被逼出的唯一解。这一节要解剖的是，这个工程解是如何通过一系列具体的装配工序，从一批前代的思想零件（天、阴阳、五行、三统）中被焊接成一整台不可拆的机器的。

在开始解剖之前，需要先做一个关键的提示。以下拆出的三道锁——说法锁、道理锁、谱系锁——不是三种并列可选的“锁型”。它们是同一台机器的三道装配工序：在时间上先后继起，在逻辑上首尾嵌套，在功能上互相咬合。说法必须先锁死“什么是善”的定义权，道理才有对象去解释——一个人如果不先接受“自己是禾”，你就没法用灾异告诉他“你的质疑不算”；道理必须先锁死“谁有权解释”的解释权，谱系的更替才能只换操作者而锁不松脱——如果每一次更换皇帝都可以重新定义什么叫善，那新皇帝还没坐稳，机器已经被人拆了。缺少任何一道，另两道都无法独自支撑这台机器的长期运转。以下三节的论述，每一节的开头和结尾都将显式地標示这种次序性，以确保读者不会把它们简化为一张“三种锁法”的清单。

#### ■ 3.1 说法锁死定义权：为什么“你”不算

第一道装配工序，是重新定义善。这道工序的完成，为后面的一切装配提供了地基——一个人必须先接受“自己只是禾”这个位置，后续的道理（“你的质疑只是未善的表现”）和谱系（“换人但不换位”）才有施力的支点。

董仲舒用的工具，不是哲学论证，是比喻。“性比于禾，善比于米。米出禾中，而禾未可全为米也。善出性中，而性未可全为善也。”

他本可以直接说“人不会自动变好，所以需要王来教化”。但他没有。他选了一个来自农业生产经验的类比，而这个选择——如上一节末尾已初步论证的——本身就是一道深思熟虑的工程决策。在禾米之喻与璞玉之喻之间，他放弃了后者。本节把这个论证推进到它最深的一层。

璞玉之喻在被放弃的那一刻，折射出了禾米之喻的工程优势——以及它为汉语文明带来的深远后果。玉工雕玉，是一种依赖个体判断力的技艺。而碾米人碾米，是一种可以被流程化的操作。董仲舒需要一个能够被大规模复制、被标准化执行、不依赖执行者个人判断力的教化模型——因为他的教化对象不是一国一邑的贵族子弟，而是整个大汉帝国治下的全部编户齐民。他需要的不是一群技艺无双的玉工，而是一批会操作同一台碾米机的熟练工——这台机器的名字叫“六经”。在地处偏远的郡国，一个出身不高、天资平平的太学生，只要学会了操作这台机器，就可以上任碾米。这就是“禾米”胜过“璞玉”的工程理由——不是哲学上的，是治理上的。

但这个选择的后果，远远超出了治理效率的范畴。当他把善的加工定义为“碾”而不是“雕”时，他同时做了一件事：他把被加工者的个体差异性，从程序里删掉了。雕玉，你必须看这一块玉的纹理——它与众不同的那部分，决定了你能把它雕成什么。碾米，你不需要看每一粒禾的纹理。禾的差异性是程序要消除的对象——碾完之后，所有的米都应该是同一个标准。一个人的生命体验里那些独一无二的、不可通约的——他受过的苦、他见过的善、他在深夜反复挣扎后做出的那个判断——那些东西，在碾米机的尺度上根本看不见。不是程序选择了忽略它们；是程序的设计逻辑本身就不包含“看见它们”的功能。

这就是“说法”的力量。它不是用论证让人接受“你必须被王教化”，而是用类比让人接受“你本来就是这个位置”——禾的位置，等着被碾的位置。一旦人接受了自己是禾，他的一切在自己的生命现场里自己摸索出来的、不完全符合经义标准的、新的活法，就不再是“另类的善”，而是“未加工的原料”——不是它不对，是它还没被碾过，还不算。

但这道锁只完成了第一步。它定义了“什么是善”，但它还没有定义“谁来判断某个具体行为是不是善”。定义权锁死之后，解释权的问题立刻浮出来——而当这道锁锁死之后，解释权的争夺就有了一个不可逾越的地基：你只能在“禾与米”这个框架里解释，你不能推翻这个框架本身，因为你已经接受了自己是禾。这就是说法锁为道理锁腾出的空间。

## 3.2 道理锁死解释权：为什么连质疑都不算

说法锁死了善的定义权。但光有定义还不够。定义只能管住“什么是善”，管不住“谁来判断某个具体行为是不是善”。这需要第二道装配工序：道理。道理锁死的是善的解释权——当某些事情违背了“善的工序化”的运行时，谁来裁决这是程序本身的问题，还是执行过程中的某个环节出了问题？谁掌握了解释权，谁就掌握了这道工序的实际控制权。而这道锁之所以能成立，恰恰是因为第一道锁已经把所有人定义成了禾——一个禾的质疑，在还没有被听见内容之前，就已经被它的身份所否定。

董仲舒的道理，就是天人灾异论。他把天装配成了一个沉默的、但会精确反馈的问责系统：日食是“阴侵阳”，地震是“阴盛阳微”，蝗虫是“德政不修”。天从不说话，但它从不错看。这套道理最精妙的地方，不在于它用灾异来限制王，而在于它前所未有地精密化了“谁有权解释天”。

不是王自己。董仲舒在《春秋繁露》中反复论证：灾异的含义，必须由深通《春秋》的儒者来转译。儒生不是批评君主，而是替天把那个信号翻译成人话。这就把儒家知识阶层嵌进了帝国操作系统的核心回路里——传感器是天，翻译器是儒生，执行器是君主的恐惧。

大多数后世学者在这个点上停住了。他们看到了这套道理的“宪法功能”：董仲舒试图用天来限制君权。这个观察在史料上是成立的。本章不否认这一点。本章要指出的是他们停留之处未能追问下去的那个问题：这个翻译回路，它自己有没有短路保护？

有。而且这个短路保护，就藏在“翻译”这两个字里。翻译不是原文。翻译是一种可以上诉的解释。但翻译的上诉终审权，不在原作者手里——天不说话。也不在翻译者自己手里——儒生没有权力判定自己的翻译是最终版。终审权在谁手里？在听众手里。在翻译稿被递上去的那个坐在工位上的人手里。

这就是天人灾异论最隐蔽的锁。它看似是一套由天向王发信号的问责系统。实际上，信号的编译权（什么算灾异、灾异对应哪一条人事）在儒生手里，而信号的审定权（这个编译对不对、要不要采纳）在王手里。这就好像一家报社的校对员可以拼命校对，但总编辑有最终签发权。你可以校对到死，但你不签发，你的校对就只是一堆废纸上的红笔迹。

这就是眭弘案的全部秘密。这个问题将在第五节详述，这里只需要指认这道锁的机制本身：当一个人严格按照程序手册操作、严格使用程序提供的编译工具，却得出“程序的操作者需要更换”的结论时——当他的质疑不再是针对程序中的某个具体环节，而是针对整个程序的操作者本人——程序的最终反应是什么？不是驳斥他的推理。是把他的推理本身定性为“妖言”。不是说他推错了，是说他根本没资格推。程序撤销的不是他的结论，是他在程序内部的发言资格。而撤销的理据，永远可以回到第一道锁——你还没被王教充分碾过，所以你的判断

只是原料。一个原料的质疑，在对它的内容进行审查之前，就已经被它的身份取消了被审查的资格。

解释权的锁定，保证了当操作者本人成为被质疑的对象时，程序可以自行短路。但这里留下了一个巨大的风险：操作者是可以更替的。王朝会灭亡，皇帝会死，权臣会篡位。如果每一次操作者的更替都动摇了“必须有一个操作者”这个命题本身，那么这台机器在每一次改朝换代时都会面临被拆的风险。这就逼出了第三道装配工序——谱系锁。它要锁死的，不是某一个特定的操作者坐稳他的位置，而是无论谁坐上去，那个位置本身都不能被取消。这就是谱系锁的真正功能：它是道理锁在操作者更替时的延伸——解释权的锁定，在正常时期保护操作者个人；谱系的锁定，在更替时期保护操作者这个位置。

### ■ 3.3 谱系锁死更替权：为什么换人也换不掉锁

说法锁死了善的定义——你不算。道理锁死了善的解释——你连质疑都不算。两道锁绑定之后，剩下的最后一个工程难题是：操作者本人是可以被换掉的。改朝换代、权臣篡位、宗室内讧——无论程序多么精密，操作者更替这件事是挡不住的。如果每一次操作者更替，都意味着重新谈判“为什么必须有这个操作者”，那么这台机器在每一次权力交接时都会面临被拆的风险。谱系锁的功能，就是把“换人”这件事本身，变成锁的一部分。它的逻辑是：你可以推翻操作者，但你不可以推翻操作位。

董仲舒用“三统论”来完成这道装配。白统、赤统、黑统，三者循环往复。夏朝黑统，商朝白统，周朝赤统。按照三统循环的逻辑，继周而起的汉朝，应该重新走黑统——这不是董仲舒的发明，他说这是孔子在作《春秋》时预先排定好的次序。孔子是圣人，圣人是唯一合法的谱系排定者。孔子已经死了，所以这个谱系无人能再修改。

这道装配的工程效果极其深远。它把“改朝换代”变成了程序内部的一个功能模块——三统循环。你推翻了一个王朝？没关系，你推翻的只是“操作者”，不是“操作位”。程序早就预装了你推翻他这件事——这叫天命转移。你推翻他之后，你自己坐到那个工位上，继续用同一套说法和道理来碾米。你可以骂上一个碾米人碾得不够好，但你不能质疑“必须有一个碾米人”这个命题本身。

但这里有一个关键的论证跳跃，本章必须正面填补。一个问题会自然地浮出：一个新王朝，手握重兵，有压倒性的军事力量，它为什么必须继承这套谱系？它为什么不干脆废掉三统，自己另立一套？为什么它要心甘情愿地把自己的权力收进“工序执行者”这个特定位置？

答案不在观念里，在利益结构里。新王朝需要两样只有儒生集团能大规模提供的东西。第一样是低成本治理技术——秦朝用文法吏治天下，成本高昂且民怨沸腾；儒生提供的是一

套“守文奉法”的循吏传统——用经义的权威来解释政令，让被治理者觉得服从政令不是向暴力低头，而是向道义靠近。这能大幅降低统治成本。第二样是政权合法性叙事——武力可以夺天下，武力不能论证“坐天下”。天命转移这套叙事必须由经义来提供，而经义的权威来自孔子，孔子排定的谱系不能你说改就改。

因此，新王朝对这套谱系的“继承”，从来不是被动的、吃亏的。它是一种交易：你用军事力量夺取了操作位，但你拿不出比儒家经义更省力、更持久的合法性话语，所以你保留了那个操作位的定义——你只是换上了自己的人。你继承的不只是谱系，还有谱系附带的那一整套治理机器。你拆掉谱系，那台机器就不再为你运转。

这就是第三道锁最精致也最残忍的地方：它把“换人”变成了锁的一部分。每一次改朝换代，都不是锁被砸开了，而是锁被擦得更亮了。新王朝的合法性，恰恰来自它对旧王朝“德衰”的审判——而审判的依据，仍然是同一套三统谱系。你用锁砸开了门，然后你做的第一件事，是把那把锁从地上捡起来，重新挂在新门框上。你挂它，不是你信它，是因为你需要它来锁住下一扇门。

## ■ 四、空转：工序如何把善变成一个永远到不了的终点

前三节厘清了这台机器的三道装配工序：说法锁死定义权，道理锁死解释权，谱系锁死更替权。三者不是并列可选的分析坐标——说法是先决条件（定义了“什么是禾”，后续的质疑和更替才能在这个定义内部被消化），道理是运行机制（在日常运作中消化一切负反馈），谱系是长期保险（在操作者更替时保住操作位本身）。合在一起，它们保证了一件事：所有人，在所有时候，都无法从程序内部证明自己是“善”的——除非你坐上那个唯一的工位。但那个工位，永远只有一个人能坐。

这意味着什么？这意味着，这台机器设计的本来目的，根本就不是“让更多人变成善”，而是“让所有人都永远停留在未善”。因为只有所有人都是“未善”的，那个垄断了加工工序的操作者，才永远有存在的理由。本节要论证的就是这个命题：工序不是在制造善，工序是在制造未善。[3]

### ■ 4.1 一个反直觉的判断：工序不是在加工善，工序是在持续产出“未善”

我们先来检验这个命题。假设一个人从小接受经学教育，言行符合礼制，在乡里被举为孝廉，在朝中尽忠职守，年老致仕，一生无大过。他按董仲舒的程序走完了全部流程。他善了吗？

答案是：没有。因为善是一个最终产品。而最终产品的验收标准，不是“个体是否达到了某个内在的道德状态”，而是“王教是否最终达成天下大同的那一刻”。在那一刻到来之前，任何个体的“阶段性成就”，都只是“善质”的部分开发，不是“善”的完全实现。就好像一批稻谷，在它们还没有被碾成米、被煮熟、被端上天下太平的宴席之前，它们永远只是“加工中的禾”。

这就是董仲舒那道禾米比喻最深层的运作逻辑：它把善的时间结构，从“当下”彻底拖向了“未来”。你此刻做的好事，不是善，是善的原料。你一生积德，不是善，是善的半成品。你死了，你的子孙继续被碾——等他们的子孙也被碾完，也许，在某个不可知的将来，碾米人终于宣布：“米，碾好了。”但在那之前，每个人都老实地待在“未善”的格子里，继续被碾。

这道时间结构的移动，有一个极其隐蔽的政治效果。它把“善”从一个可以被个体在当下体验到的状态——我此刻在做好事，我感觉到这样做是对的，我对自己做出这个选择的理由有切身的、不可转让的理解——变成了一个只能由未来的总清算来判定的结果。你的体验不算，你的确信不算，你的良知你自己说了不算。算的，是那个验收的时刻。而决定那个时刻什么时候到的，不是被碾的人，是碾米人。

## ■ 4.2 制造未善的三重机制

这个机制不是一套空洞的修辞。它通过三道具体的操作，把“制造未善”从隐喻变成制度。

第一道操作：标准的单一来源。善只有一种定义：经义里刻好的那种。你作为一个人，在自己的生活现场里遇到的那些具体的矛盾——一个农民在旱灾时要不要偷官仓的粮救自己的孩子，一个读书人在暴政面前要不要冒死上书——这些矛盾里有没有可能逼出新的善的可能？没有。因为经义已经提供了所有正确答案。你不需要去从矛盾里结算出属于自己的善的判断；你只需要对照标准答案，看哪一条适用你当前的情况。你的工作不是“创造善”，而是“匹配善”。就像工人不需要设计产品，他只需要对照图纸把零件拧紧。

第二道操作：反馈的自我消化。你感觉不对。你觉得自己在做坏事。你很痛苦，你想反抗。程序告诉你：你感觉不对，恰恰是因为你还处在“未善”阶段。一个真正善的人，会深刻地理解这道工序的必要性，会心甘情愿地觉得被碾是一种幸福。如果你还不心甘情愿，那就说明你被碾得还不够久、不够深。你需要更多的教化。你的痛苦，不是信号——不是在告诉你“这道工序有问题”——而是在告诉你——你自己的问题还没解决干净。它把一切来自被加工者的负反馈，都转化成了“加工程度还不够”的证据。就像一台冰箱：你不制冷，说明你没插



电；你插了电还不制冷，说明你温度调得不够低；你调低了还不制冷，说明你门没关好。你永远在找自己的问题，你永远不会想到——也许冰箱的设计，本来就不带制冷功能。

第三道操作：终极验收的无限推迟。善的最终判定，永远在一个不可知的未来。天下大同。万世太平。王道大成。这些词都有一个共同点：你永远无法在任何一个具体的当下，指着任何一个具体的人说——他善了。因为任何人都不是“天下大同”的化身。任何个体都不可能是“王道大成”的承载者。那个终极的“善”是一个空位。它可以被任何人无限期地引用，却不能被任何人在当下去验证。它是程序为自己预留的最后一道防线：永远不能说“已经做到了”，因为一旦说“已经做到了”，工序就没有继续存在的理由了。

这三道操作合在一起，构成了一台永动机：标准是唯一的，所以你不被允许从生命现场里结算出自己的善；反馈是自我消化的，所以你的痛苦被解读为你自己的未善；验收是被无限推迟的，所以你永远等不到那个“善了”的时刻。这台永动机唯一的产出，不是善，而是不断的、一代又一代的、真诚的未善者。

“制造未善”与“窃取资格”的关系，在此可以做一个准确的厘定：窃取资格是机制——程序把个体从自己的生命体验里做出善的判断这件事的资格判定为“不算”；制造未善是这个机制持续运转的必然产物与自我强化方式——程序必须让所有人都停留在未善，因为一旦有人被宣布为“善了”，程序就必须回答一个它无法回答的质问：那个人既然已经善了，他为什么还需要继续被碾？

## ■ 五、在真实历史中的运行：假开放如何拆解它

以上是逻辑分析。现在把它放回真实的历史现场，看看这台机器实际上是怎么跑的。本节选两个汉代案例——盐铁会议和眭弘案。一个代表这台机器如何通过“假开放”来吸收批评，一个代表这台机器如何在批评者试图按程序手册操作时，以“你操作不当”为由将他们清除出局。<sup>[4]</sup>但在展开分析之前，需要先做一笔交代：本节选择这两个案例，不是因为它们是“机制的例示”——仿佛第三节装配好了机器，这一节只是把它跑一遍来验证。选择这两个案例，是因为本章在细读《盐铁论》和《汉书·眭弘传》的字里行间时，被它们逼着追问了一些原先不曾想到的问题：为什么贤良文学的论证如此有力，却丝毫没有动摇程序本身？为什么眭弘的推理严格符合程序手册，却被程序自己处死？这两个追问逼迫本章去指认一个在第三节的理论拆解中尚未被充分命名的机制——“假开放吸收真诚”——以及那个最极端的短路保护装置。换句话说，不是先有机制再有案例；是案例里露出的问题逼出了机制。以下的分析，将以这种“从现场逼出判断”的方式展开。

### ■ 5.1 盐铁会议（前 81 年）：假开放的完美示范

汉昭帝始元六年，帝国召开盐铁会议。民间代表贤良文学，对阵御史大夫桑弘羊。讨论的议题是盐铁官营该不该废。

贤良文学的批评极其尖锐。他们说，盐铁官营是与民争利，是“霸王之道”而非“仁义之道”。他们说，均输平准让官吏与商贾勾结，物价反而更高。他们说，连年对匈奴用兵虚耗国库，应当“假武修文”。每一句话都打在帝国核心政策的痛点上。

桑弘羊一一反驳。他说，北有匈奴，内需统一，没有盐铁之利，军费何来？他说，均输平准正是为了抑制豪强兼并，是“帝国调控之手”。他最后扔了一句话，这句话在程序内部是无法反驳的：“诸生发于畎亩，出于穷巷，不知国家大计。”你们这些读书人，只懂经义，不懂帝国的实际运作。你们的批评，是善意的，但是不成熟。

会议的结果在历史上很清楚：盐铁官营保留，仅罢去酒榷。表面上看，贤良文学输了。

但这不是一次普通的政策辩论失败。盐铁会议让本章必须追问一个它自己原先并未预料到的问题：贤良文学的论证在道德感和经义引用上如此有力，他们的每一句批评都切中帝国政策的实际弊端，为什么这一切力量，非但没有动摇程序本身，反而像是被程序完整地吸收掉了？是在反复阅读贤良文学的发言与桑弘羊的回应之间的那段沉默时，本章才被逼出了“假开放吸收真诚”这个判断。

贤良文学能说吗？能说。他们说了盐铁官营不好，说了对匈奴用兵不好，说了国家与民争利不好。但他们绝对不能说的话是什么？是：这道“善”的程序本身是错的——不是盐铁官营错了，而是“王有权力独断什么才是善的国家政策”这台机器本身出了错。他们没有说这句话。他们仍然用经义作为自己的武器，仍然在“先王之道”的框架里论证自己的主张。他们不是在起诉法官，他们是在用同一部法律起诉对方的律师。

他们的真诚、勇气、道德感，在这场辩论中被完整地收编。程序没有杀他们。程序让他们说。程序甚至在史书上给他们留了名——宣帝时，桓宽将会议记录整理成《盐铁论》，这被后世视为汉代朝议开放的典范。然后，那个最致命的判断就在所有围观者的心里悄无声息地扎下了根：这次批评没有改变政策，不是因为程序不开放，而是因为批评者的论证还不够有力。是他们没有找到经义里最对的那一条。是他们的政策建议太书生气了，不了解帝国的实际运作。换一个更懂的人，下一次也许就能赢。

这里需要补入一笔制度史的支撑，以让“假开放”的机制不只是观念层面的分析，而是在汉代的实际行政运作中有其制度根基。盐铁会议上贤良文学与桑弘羊的对峙，在深层上是儒生传统与文吏传统两种治理逻辑的对峙。严耕望、阎步克等学者的研究已经充分揭示，汉帝国的地方行政运作，长期以来存在着两个群体的分工与紧张：文法吏掌握律令、钱谷、刑狱——帝国的硬性行政机器由他们执行；儒生则掌握经义、教化、礼制——帝国的合法性叙事由他们提

供。盐铁官营的政策本身，是文法吏传统的巅峰产物：均输平准、盐铁专卖，依赖的是一整套精密的行政计算与执行体系。桑弘羊本人出身商贾之家，精于算计，是文法吏传统在帝国中央的最高代表。贤良文学以经义攻讦盐铁官营，而桑弘羊以“诸生不知国家大计”回击——这句话不是在骂贤良文学“不懂事理”，而是在宣判：你们握的那把钥匙（经义），开不了这把锁（帝国财政）。

这就意味着，程序的“假开放”并不只是吸收被统治者的真诚，它也依赖于程序内部执行群体之间的制度性分工。儒生被分配了“解释善”的权限，但“执行善”的技术权限被文法吏把持。贤良文学可以用经义论证盐铁官营是不善的，但当他们试图用这个论证去撬动政策时，他们撞上的是一堵由文法吏用数字、账目、军费预算砌成的墙。程序的开放，从未延伸到这个技术权限的边界之内——而在那个边界之内，真正决定政策走向的计算，不需要用经义来为自己辩护。

这个判断一旦被接受，程序中那些真正有良知的批判者，就会把毕生的精力都投进“如何让程序运行得更好”的争论里。他们永远在磨自己的论证，永远在找经义里更精准的那一条。他们永远不会问——是不是程序本身的逻辑就不可能让他们赢？因为“赢”的定义，不在他们手里。

必须在此处插入一笔与卢曼的切割，因为盐铁会议恰恰演示了卢曼《合法性通过程序》（Legitimation durch Verfahren, 1975）的核心命题——但也同时演示了它的边界。

卢曼的论证是清晰的：在复杂的现代社会中，决定的合法性不再来自实质正义（即决定本身“正确”），而来自“经过了法定程序”这个事实本身。参与者一旦进入程序，就被程序角色吸收——原告、被告、证人、陪审员——他们即使对最终结果不满，也会在相当程度上接受这个结果，因为程序给了他们一个“参与过”的身份。“我说了我的，他没听我的，但程序是公平的”——这种感觉，就是程序合法性的核心来源。

这确实与盐铁会议的运作高度相似。贤良文学“参与了”，他们“被允许说了”，他们的发言“被记录在案”了。程序给了他们参与的身份。

但本章观察到了卢曼未能观察到的一层。卢曼的程序吸收的是参与者的不满——通过“让你说了”来消解你的反对，让你在程序结束后觉得“好歹我说了”。这是一种冷漠的、功能性的运作：程序不要求参与者真心信服，只要求参与者在行为上接受程序的结果。

而董仲舒的程序吸收的是参与者的真诚——通过“让你用经义来论证”来消解你对经义本身的质疑。贤良文学在盐铁会议上的每一句痛切陈词，引用的是同一个经义传统，诉诸的是同一个先王之道，他们越是真诚地论证盐铁官营“违经”，他们就越是无可挽回地为那部经义的

权威性注入自己的生命热情。他们以为自己在用经义攻击政策。实际上，他们的每一次攻击，都是在为经义作为最终裁决依据的地位续命。

这里可以凝练出本章与卢曼之间的那道不可通约的本质界限——并且用一个可验证的命题说出它：卢曼的程序吸收不满，反抗者只是接受了结果，并未为此结果注入自身的道德资本；董仲舒的程序吸收真诚——反抗者用自己的生命力为反抗对象的权威性续命。一个吸收的是不服从的行为，另一个吸收的是服从者的内心。前者靠冷漠运转，后者靠真诚运转。前者不需要你信；后者需要你越是用全部信念、全部学养、全部正义感去论证经义里的某一条，你就越是在为那一整本经义的绝对权威押上你自己的灵魂。这个多出来的存在论后果——反抗者共谋于自己所反抗的东西——在卢曼的框架里是看不见的。

## ■ 5.2 眭弘案（前 78 年）：当合格的执行者被定性为“妖言”

盐铁会议演示了假开放如何吸收温和的批评。而眭弘案则演示了这道程序的终极保险：当一个人严格按照程序手册操作、严格使用程序提供的编译工具（灾异、经义、推演），却得出“程序需要更换操作者”的结论时，程序会怎么做。

汉昭帝元凤三年（前 78 年），一系列被视为灾异的自然现象接踵而至：泰山有大石自立，上林苑有枯柳复生，虫食树叶成“公孙病已立”之状。时任符节令的眭弘（字孟）——符节令正是掌管瑞应灾异之官的职位——依据董仲舒的灾异理论，上书朝廷。他的推理每一环都扣在董仲舒的框架里：泰山是天下名山，是王者受命易姓之处；大石自立，预示将有匹夫成为天子；枯柳复生，是废帝复兴之兆；虫食树叶成文，是天意的直接书写。结论只有一个：天命将改。他上书建议汉昭帝“求索贤人，禅以帝位”。

当时秉政的大将军霍光以“妖言惑众，大逆不道”的罪名处死了眭弘。

对于本章的论证而言，眭弘案最深层的关键不在他被杀的结局，而在于为什么一个严格按照程序手册操作的人，会被程序自己判定为非法。在这里，不能简单地将眭弘的行为归结为“他信程序信到了把命交给程序的地步”——那是一个对历史人物心理动机的推断，虽然合理，但在实证上无法被最终坐实。更稳健的论证方式，是从眭弘在程序中的职位结构出发。

眭弘是符节令。这个职位本身就是灾异信号的官方翻译官。他的日常职责是解读灾异，他接受的职业训练是天人对应的推演方法，他的职权说明书——如果汉代有这种东西的话——就是董仲舒一脉的《春秋》灾异之学。也就是说，“严格按照程序手册操作”不是眭弘的个人选择，而是他的职位定义。“如何做才算正确行动”已经被这个职位预先规定了——你看到泰山大石自立而不去解释它，你才是失职。在这个意义上，眭弘不是“信程序信到死”，他是被困在程序为他定义的那个角色里——这个角色之外，不存在另一种可以被他自己、被他的同事、

被他的职业共同体承认为“正确”的行动逻辑。他很难在程序给出的角色之外想象另一种行动可能。

而对于程序本身来说，处死眭弘不是一次错误，而是它完美运行在它最根本的那个预设上的必然产物。那个预设就是：善之工序的最终解释权，从来不在翻译者手里，在操作者手里。一个翻译者翻译出一条对操作者本人不利的信息，他的翻译就不再是“翻译”，而是“诽谤”。翻译和诽谤，在行为层面看完全相同——都是一个人对一段信息做出解读并向公众发布。它们的唯一区别，在于事后由谁来定性。而定性的那个人，恰恰是被翻译所指控的那个人。

眭弘案由此补全了盐铁会议未说出的后半句。盐铁会议证明：批评的资格，由程序赋予。眭弘案证明：当批评指向操作者本人时，资格可以被瞬间撤销。撤销的理据，不是批评者说错了，而是批评者“还没被加工到位”——正因为你还处在未善，你才会以为你的翻译算数。而善之程序最深的那层逻辑，在眭弘死去的那一刻被完整地显影：它连它最真诚的信徒也能处死。

## ■ 六、不可还原的切割：为什么这不是别人已经说过的东西

如果本章的核心概念——“善的工序化”——能被某个已知理论所覆盖，那么它就是一个换装，不是一个发现。以下将它逐次与几个最可能构成威胁的敌意近邻分刃切割，指出每一组近邻与本章机制之间的那条不可通约的缝隙。在逐一切割的过程中，本章反复使用“资格窃取”这一判词——这里需要先停一步，对这一判词本身做一次自我检视：它是否正在从一个生成层面的结论，滑向一个“现成本质被剥夺”式的诊断？本章在第二节末尾已明确论证，“自我认证权”不是预设的、等着被回归的人性本质，而是被工序的对立面反向逼出的一个影子。相应地，“资格窃取”也不是在说“一个早已存在的东西被人拿走了”——它是在说：在工序被装配的同时，一个先前不存在的“资格之缺”被生产了出来。因此，以下切割中每一次出现的“资格窃取”，请读者不把它读作一个诊断标签（“这个东西的本质是资格被窃取了”），而把它读作一个生成层面的判断的缩写——“这个程序在装配过程中，同时生产了一种此前不存在的资格亏空，而这种亏空是以下各家理论各自在各自的时代和材料面前，未能指认的一层”。带着这个警觉，进入切割。

### ■ 6.1 与卢曼的“合法性通过程序”

这道切割在上一节盐铁会议的分析中已经开了刃，此处将它最终断定：卢曼的程序不需要参与者的真诚。一个人走进法庭，不需要相信法律是正义的，他只需要在程序上完成“陈述”这个动作，判决结果就获得了程序上的合法性。这是一种冷漠的、功能性的运转。

而董仲舒的程序需要参与者的真诚。贤良文学越是痛切地引用经义来论证自己的批评，他们就越是把自身的道德热情押进了那本经义的权威里。他们不是在程序上完成“陈述”这个动作然后接受结果——他们是在用自己的全部信念为那个他们试图在局部上反抗的权威整体续命。

由此可以得出本章与卢曼之间的理论增益命题：吸收真诚，使得反抗者用自己的生命力为反抗对象的权威性续命；而吸收不满，反抗者只是接受了结果，并未为此结果注入自身的道德资本。这个多出来的存在论后果——反抗者共谋于自己所反抗的东西——是卢曼的框架在其最深的地基层面所无法看见的。卢曼的“程序合法性”关切的是决定如何被接受；本章的“善的工序化”关切的是反抗者如何被变成自己反抗对象的燃料。前者处理的是“接受”，后者处理的是“真诚的征用”。这是第一条不可通约的缝隙。

## ■ 6.2 与布迪厄的“误识”

布迪厄（Bourdieu, 1980）在《实践感》中提出了“误识”（*méconnaissance*）与“象征暴力”（*violence symbolique*）这一对概念。它们的核心机制是：被支配者之所以认同一套对自己不利的秩序，是因为他们“误识”了这套秩序的任意性——他们看不见它是被建构出来的，把它当成了自然而然的。这是一种认知上的遮蔽：支配之所以有效，是因为被支配者看不见支配。

“善的工序化”与“误识”之间的关键区别就在于这道“看不见”。布迪厄的误识说的是：你认了，是因为你不知道。你不知道樨头是人工凿出来的，你以为是树本来就长成那样的。而董仲舒的程序，面对的是知道自己没被碾够的人——他们可以怀疑，可以痛苦，可以觉得不对。程序不需要他们看不见。程序只需要在他们看见之后，告诉他们一句：你看见的那把锁，不是锁。那是你还没学会开的门。

这就是“善的工序化”比“误识”更窒息的一层所在。误识的无知可以通过揭示真相来解除——一旦你被告知“这是人为建构的”，你的顺服就可能动摇。而善之工序没有这个破绽——它把你明知是锁的东西，论证为通往善的唯一入口。你不是因为无知而顺服；你是因为被告知“反抗本身就是你还需要被碾的证据”而停手。你看见了锁，但你在砸锁之前停住了，因为程序已经预先定义了一件事：砸锁这个动作，不是砸锁，是未善者的挣扎。你越是用力砸，越是证明你是未善的。你把全部的力量砸出去，砸的不是锁，是你的“善”的自我认证资格。布



迪厄面对的是“没看到锁的人”，本章面对的是“看见了锁但被说服那不是锁的人”。这是第二条不可通约的缝隙。

### ■ 6.3 与韦伯的“铁笼”

韦伯（Weber, 1920）在《新教伦理与资本主义精神》结尾描绘的现代性困境，比他之后的绝大多数引用者所简化的版本更为复杂。在那封至今具有灼伤力的一段话里，他写的不是“手段变成了目的”那么简单，而是：“清教徒曾是在一种志业中工作；而我们则被迫在其中工作。……依照巴克斯特的见解，对外在之物的操心应当像一件‘随时可以脱掉的轻飘飘的斗篷’一样披在圣徒肩上。但命运却让这件斗篷变成了一间钢铁般坚硬的牢笼（stahlhartes Gehäuse）。”

这件斗篷变成笼子的过程，是韦伯的核心判断所在。他不是在说人“错误地把手段当成了目的”。他说的是：一种曾经是出于内心确信的生活方式——清教徒为了验证自己的蒙恩状态而投身世俗工作——在“寻找天国的狂热开始逐渐演变成冷静的经济美德”之后，变成了一套不再需要内心确信就能自动运转的外部秩序。笼子不是被人造出来关自己的，是斗篷自己硬化成笼子的。

在这个意义上，韦伯的“铁笼”确实触及了一个与本章近似的点：人把自己曾经拥有的内在确信交了出去，交给了一套不再需要他的确信就能自行运转的秩序。但这个交出去的过程，在韦伯那里是一个“冷却”的过程——宗教热忱退潮了，剩下的只有外部秩序的机械运转。人在笼子里不是被禁止拥有内在确信，而是内在确信本身失去了它的社会功能——你信不信已经不重要了，秩序反正照样转。

而董仲舒的程序不是“冷却”的，它恰恰需要参与者的热忱。程序不允许你冷漠地服从——那会把它变成秦制，变成纯暴力，变成随时可能被拆穿的东西。它需要你真诚地相信：禾需要被碾，碾的过程虽然痛苦，但那是通往米的唯一途径。你的痛苦不是被程序当成多余的热量散发掉的，而是被程序收集起来，转化为你继续被碾的动力。你越是觉得自己不够善——你越是真诚地感到愧疚、不足、力有不逮——你就越是完美地证明了程序的必要性。

这里可以凝练出本章与韦伯之间的理论增益命题：韦伯的冷笼子，使得痛苦被排出系统之外——人在笼子里的痛苦与笼子的运转无关，笼子不需要你的痛苦来运转；董仲舒的热笼子——它把主体的痛苦回流为系统动力。痛苦不是在笼子里被浪费掉的余热，痛苦是笼子需要的燃料。这个多出来的存在论后果——系统以主体的自我否定为能源——是韦伯的“铁笼”在其分析框架中所未包含的。韦伯面对的问题是“人被关在笼子里，不知道怎么出去”。本章面对的问题是“人把开笼子的钥匙上缴给了笼子的管理员，并从管理员那里收到了一份文件，证明那不

是钥匙，那是禾的秸秆——你留着它，只是因为你还没被碾够”。这是第三条不可通约的缝隙。

## ■ 6.4 与福柯的“治理术”

福柯（Foucault, 1977-1978）的治理术（governmentality）<sup>[5]</sup>揭示了现代权力的一种精妙形态：它不再依靠自上而下的禁令与惩罚来运作，而是通过知识网络——人口学、统计学、公共卫生、教育学——来定义什么是“正常”，把个体塑造成在自我管理中成为驯顺的、有生产能力的主体。权力不禁止你；权力告诉你“你是正常的”，并让你自己努力去够到那个“正常”标准。

在福柯的模型里，权力和反抗是永恒的对子。权力弥散在毛细血管般的社会关系里，没有唯一的控制中心——这意味着反抗也弥散在每一处权力关系里。哪里有权力，哪里就有反抗。主体在权力的缝隙中不断生成新的、不驯顺的自我形态。

“善的工序化”看到了福柯没能看到的一层——它不是在规范你的行为，它是在取消你从自己的行为里提炼出善的判断的资格。福柯式的反抗者——那个拒绝被“正常”定义的人——在董仲舒的程序里有没有生存空间？没有。因为他拒绝被定义这件事本身，会被程序消化为“他还不够善”。不是程序要消除他，是他自己还不够。这就不是权力和反抗在互搏。这是跑道在一开始就被划定在了一片孤岛上——你在岛上怎么跑都可以，你的每一次冲刺、每一次变道、每一次大声抗议，程序都为你鼓掌。但你不被允许跑到岛外面去。因为岛之外不是海——在程序的地图里，岛之外标注的是“尚未被铺成跑道的荒地”。

与此相关，“善的工序化”与黑格尔—马克思传统中的“异化”（Alienation）概念（Marx, 1932）也需要一道严格的切割。<sup>[6]</sup>异化描述的是主体与自身产物的分离——你创造了一个东西，但那东西反过来支配你。工人生产了商品，但商品以资本的形式站在他的对立面。分离是核心诊断：你的东西不再是你的。

而“善的工序化”诊断的，不是分离，而是分离之前就已完成的一道前置操作——资格窃取。你不是先拥有了“从自己生命现场里熬出来的善的判断”，然后这道判断被一道外部工序夺走了。你是从一开头就被告知：你那把钥匙，不是钥匙。它是禾的秸秆。工序不是在和你抢夺一个已经属于你的东西。工序是在你还没有能力为自己主张所有权的那一瞬间以前，就已经完成了对“什么算钥匙”的立法。

“还没有能力为自己主张所有权的那一瞬间”——这个时间性判断，可以在此被更精确地锚定。那个瞬间，不是一个抽象的哲学命题，而是一个在汉代教育现场中反复发生的具体场景：一个学童被第一次教“性比于禾，善比于米”的那个时刻。在那个时刻，他还没有累积起

足够多的、属于自己的、从生活现场里熬出来的关于“什么是善”的判断——他太小了，经历太少，还没有足够厚的土层去长出属于自己的判断。但正是在这个他还没有长出自己的判断的年纪，他被教会了一句话：你将来会有的那些判断，在没有经过王教碾过以前，不算。资格窃取不是发生在他成年之后、已经有了成熟的道德判断然后被宣布为“不算”的那一刻——那是资格窃取已经被内化之后的结果。资格窃取本身，发生在一个人还来不及说“这是我的”以前，就已经被教会说“这是禾”。

异化是结果——人的善的判断与自己分离；工序化是引擎——它负责在分离发生之前，预先撤销分离所依据的那个“这个东西曾属于我”的资格本身。福柯没有面对过这种“取消资格”的机制的外部有一个明确的、唯一的操作者。他的权力是弥散的，所以他可以乐观地说“哪里有权力，哪里就有反抗”。而善之工序有一个明确的工位——那个工位不负责指挥一切，但它负责在一切指挥都失效的时候，告诉你：这件事不算。异化理论没有处理过这种“资格窃取”的装置——它当然处理过分离，但它没有处理过分离发生之前的那道资格撤销程序。这是第四条不可通约的缝隙。

## ■ 6.5 与阿伦特的“无思”

阿伦特（Arendt, 1963）在《艾希曼在耶路撒冷》中通过艾希曼的审判提炼出“平庸的恶”（the banality of evil）这个概念：巨大的恶行，其源头往往不是极端的邪恶意念，而是个体停止了独立的思考与判断。艾希曼不是在狂热地仇恨犹太人——他只是忠实地执行了第三帝国的法律与命令。阿伦特写道，艾希曼的困境“恰恰源于：有如此多人像他一样，既不是变态也不是虐待狂，他们是、并且曾经是可怕的正常……从法律制度和我们的道德判断标准来看，这种正常比所有暴行加起来都更令人恐惧。”恶的平庸性在于：一个普通人，在一个罪恶的制度里，仅仅因为停止独立判断、忠实执行程序规则，就能成为巨大罪行的传导者。

阿伦特描述的恶，其运作条件是执行者的无思——他不动用自己的判断力去质疑他所执行的命令。而“善的工序化”面对的不是无思，而是被收编的深思。程序不但不禁止你思考——它鼓励你深入地思考。你可以皓首穷经争论性之善恶，你可以章句训诂辩论哪条经义最接近孔子的本意，你可以上疏万言力陈盐铁官营之弊。这些思考不但是允许的，而且是程序的“开放性”最有力的活广告：你看，我们的程序允许如此深刻的争论。

但是，你所有深刻的思考，都在同一个不容思考的底盘上运转：善，是一道必须由那个你不可能换掉的工位来执行的工序。你用尽全部学养争论的那个东西——它的存在本身，不在可争论之列。阿伦特的恶不需要作恶者的真诚。一个平庸的官僚只需要服从，他不需要热爱他所签发的每一道命令。而善之工序需要就是真诚。它不需要你不动脑子地盖章；它需要你带着全

部信念、全部学养、全部生命热情，去论证那个章为什么非盖不可。因为一旦你停止了真诚——一旦你变成了一个冷漠的盖章机器——你就抽空了这个程序的道义地基。程序就露出了它纯权力的一面，而纯权力是脆弱的。真诚，才是它最坚固的混凝土。这是第五条不可通约的缝隙。

## ■ 6.6 与波兰尼的“默会知识”

这里需要补入审稿人提示的一个重要的敌意近邻：迈克尔·波兰尼（Michael Polanyi）的“默会知识”（tacit knowledge）。波兰尼在《个人知识》（1958）中论证，一切知识——包括关于何为善的知识——都有不可被形式化、不可被完全规则化的默会维度，它根植于个体在具体情境中的亲身体验与判断：“我们知道的，比我们能说出的多。”波兰尼极为警惕一切试图把个体判断消解为可被外部规则完全规定的工序的做法——他在《个人知识》中对现代科学日益工序化的倾向有过尖锐批评。

必须承认，本章与波兰尼共享一个非常关键的地基：个体从切身体验中做出的判断，有其不可转让、不可被外部规则穷尽的认知价值。但本章与波兰尼之间有一条必须切开的缝隙。波兰尼的论证，在其最核心的层面上，是一道认知论证：默会知识很重要，你把它工序化，你得到的是知识的贫乏——你会失去那些无法被写进操作手册的认知内容。他的着力点是“工序遗漏了什么”。本章的论证，是一道资格论证：那道工序不只是遗漏了你的默会判断——它看着你的默会判断，然后宣判它“不算”。它不是在认知上不完整；它在资格上不承认。波兰尼说：工序化的知识是不完整的，它漏掉了很多说不出的东西。本章说：善的工序化不是知识不完整——它是把完整的人定义为“原料”，从而取消了那些“说不出的东西”作为善的候选者的身份。漏掉，是可以补的——你可以改进工序，让它捕捉更多的默会维度。取消资格，不是靠改进工序就能补的——工序越精密，资格被窃取得越彻底。这是第六条不可通约的缝隙。

六道近邻切割至此落定。本章的核心命题——“善的工序化”——不是卢曼的“程序合法性”（程序吸收不满 vs 程序征用真诚），不是布迪厄的“误识”（看不见锁 vs 看见锁但被告知那不是锁），不是韦伯的“铁笼”（痛苦被排出系统 vs 痛苦被回流为系统动力），不是福柯的“治理术”（弥散的权力 vs 唯一的操作者资格），不是异化（分离 vs 分离之前的资格窃取），不是阿伦特的“无思”（不动用判断力 vs 被收编的判断力），不是波兰尼的“默会知识”（认知上的遗漏 vs 资格上的取消）。它观察到了这些各自在各自的时代、面对各自的材料时极有穿透力的理论，在“善之工序化”这个特定的历史-逻辑结构面前，各自被挡在哪一道看不见的墙外。

## ■ 七、自反与证伪

任何足够有力的方法都内嵌着自己的假闭合诱惑。它越是漂亮地成功几次，就越容易让人忘记它只是一个工具——一把撬棍——而不是一面应该被膜拜的旗帜。在这篇文章写到最后之际，必须追问自己一个最不舒服的问题：我用“善的工序化”去拆董仲舒的程序时，我是不是也在替读者装配一台“如何拆解一切善之程序”的新工序？

这台新工序有自己的核心概念（善的工序化），有自己的操作流程（从生成机制追溯到三道装配工序解剖到假开放的运行分析到敌意近邻切割），有自己的自我防御机制（任何人说“这个机制不如某某”——我马上可以祭出第六节的切割来证明他“还没理解本章的核心”）。这不就是娃弘的剧本在我身上的重演吗？

坦白说，是的。这篇文章，确实为读者提供了一套“拆解善之工序”的分析工序。但本章在这里不能止步于承认这个风险、然后规劝读者“应该拿走方法而不是标签”——那等于把一个生成层面的判断收在了一个应然告诫里，而应然告诫不是生成机制。本章必须再多走一步：追问这个风险本身的结构。不是“读者会不会误用”，而是“活的生成机制分析，它自身在什么条件下，会蜕化成一套死工序？”

这个蜕化，最容易在哪一步发生？本章拆解了“善的工序化”的三道装配工序——说法、道理、谱系。如果未来有人拿这三道锁去批量处理别的历史材料——看到任何一个制度，先拆它的说法（用什么比喻锁死了定义）、再拆它的道理（用什么解释回路消化了负反馈）、最后拆它的谱系（用什么循环叙事保证换人不换位）——那么在“再给我一个案例，我来套三道装配”的那一个瞬间，那个拆解动作就已经从活的生成机制追溯，蜕化成了一台死的分析机器。不是拆出来的结论错了——很可能他拆出来的东西在实证上是对的。是他不再被那个具体的案例逼出判断了。是他把判断的位置，交给了那套已经被前人装配好的拆解流程。

本章无法阻止这个蜕化，因为任何一套足够有力的分析方法，都天然带有这种被工序化的倾向——它越成功，就越容易被复制；越被复制，就越不需要复制者在每一个新对象面前重新经历当初逼出这套方法的那种“被困境推着走”的张力。但这并不意味着本章可以什么都不做。本章可以做的，是把它的死法钉在这里：在未来，当你看到有人用“三道锁”去套一个新案例——他的分析流畅、精密、滴水不漏，但你读完之后感觉自己没有学到关于那个案例本身的任何你此前不知道的东西，你只学到了“三道锁又成功了一次”——在那个时刻，你可以说：那篇论文死了。它不是死于错误。它死于它不再被它所面对的材料逼出判断。

最后，本章必须诚实地给出一个可能让它的核心命题被证伪的条件。这个条件不是对上述方法论风险的泛泛承认，而是一个具体的、可以被未来的历史研究者据以反驳本章的历史判断。

本章的论证在逻辑上推出一个极强的历史判断：在董仲舒式工序的常态运行下，程序内部不可能稳定地出现那种既承认经义权威、又能从个体生命体验中结算出与官方经义解释相悖的、并获得主流士林公开承认为“善之合法形态”的新判断。程序在常态下确实窒息了一切内部反抗，只有在政治崩解或严重危机导致程序出现缝隙时，那些异质的声音才在外部或缝隙中被短暂地打开。

但这个判断不能变成“自封理论”——不能把所有反例都解释为“那是程序崩溃时的例外”，从而让命题免疫于任何可能的反驳。如果程序在常态下无法容忍内部反抗，而所有实际发生过的内部反抗都恰好出现在“程序不稳定的时刻”，那么“常态”本身就成了一个不可观测的状态——没有反例能碰得到它。在这个意义上，本章需要卸掉一层防御：承认本章所论证的“常态”是逻辑上的理想型，而在真实历史中，这个“常态”总是在局部松动、自我修补、再次闭合的动态过程中。如果未来的历史研究能够证明，在程序没有面临明显的政治危机或外部冲击的时期，曾经稳定地出现过这样一种情况——士人群体在不否认“王教是善的必经之道”的前提下，从自己处理具体政务、断狱、赈灾、教子的切身经验中，凝结出了与官方经义解释系统不同、不可通约的关于“什么是善”的新判断，且这种新判断没有被同时代的主流舆论消化为“未善”“术异”或“有待加工”，而是被认真对待、被部分人接受为“那也是一种善”——那么，本章关于“程序常态下无内部反抗空间”的核心判断，就遭遇了一次有力的证伪冲击。

就本章作者目前所掌握的历史证据而言，那种空间从未在程序内部被稳定地打开过。每一次类似空间的打开——如魏晋玄学对名教的疏离、明末心学对官方理学的冲击——都是程序本身因政治崩解或严重危机而出现缝隙之后，在程序的外部或缝隙中被打开的。而程序在不面临严重危机时的常态运行，恰恰是持续地把这种空间重新消化为“程序内部的争论”——即本章所论证的“假开放”。承认这一点，不是削弱了本章的锋利度，恰恰相反——正是因为它承认自己可以被具体条件证伪，它才有资格说自己是判断，而不是信念。

## ■ 注释

[1] 本章在两种意义上使用“未善”一词，在此一并厘清，下不另注：A义是“善之程序所宣称的未善”——人尚未被加工到位，还处在从禾到米的半途；B义是“本章所论证的未善”——善之程序的运作本身就是一种对善的慢性亏欠，它把一切活的经验判定为“不算”，因而持续产出一种结构性亏空。两者不可混淆：前者是程序内部的判词，后者是本章对程序的判词。



\[2\] 本章所使用的“工序”一词，刻意区别于社会学与法学中“程序”（procedure）这一中性概念；它强调的是一套将“善”从个体的日常实践中剥离、交由特定操作者按固定标准加工完成、并由操作者独占验收权的强制性流程。与此相对，本章在与卢曼（Luhmann, 1975）的“合法性通过程序”对话时，后者仅涉及决定如何获得接受，而不涉及参与者自我认证权的窃取。此注划清本章“工序”概念与一般程序理论的边界，下不再赘。

\[3\] 第四节标题及正文中的“制造未善”，取注释\[1\]厘清的B义：善之程序的持续运转，本身就是对善的一种慢性亏欠。它不是程序所宣称的“人尚未被加工到位”（A义），而是本章对程序的判词：程序在结构上必然持续产出善的亏空。读者在第四节看到“未善”一词时，请务必将其读作B义，以免与董仲舒自己的术语用法混淆。

\[4\] 本章对盐铁会议与眭弘案的分析，均基于《汉书》及《盐铁论》等有限史籍记载所做之思想拓展性例示。其中涉及眭弘行为的分析，主要从其职位结构（符节令的职责定义）出发进行论证，而非对其个人心理动机做实证推断。特此说明。

\[5\] 本章论及的福柯“治理术”，主要依据其1977-1978年法兰西学院讲座《安全、领土与人口》（Sécurité, territoire, population）中提出的权力分析框架，侧重于权力弥散与主体自我治理的面向，不涉及福柯晚年全部思想。该讲座原始授课年份跨度为1977至1978年，法文首版整理出版于2004年。

\[6\] 本章所讨论的“异化”概念，以马克思《1844年经济学哲学手稿》中的经典表述为准，即人与自身产物相分离并受其支配。此注划清本章与黑格尔绝对精神异化及后世各种异化理论的边界。

## ■ 深化增补：被放跑的最近邻

以下为编辑增补，针对评审时记下的扣分点定点补强——每条只做一件事：把一位被放跑的最近邻请上台，切开，给出可裁决的判准。增补与作者正文分开标注，不改动作者原有判断与结构。

## ■ 增补一·与威廉斯的“道德体系”：理论的窄化，还是工位的宣判

本章第六节已切六家，但在“取消资格”这一核心动作上，还有一位比其中任何一位都更同构的先行者未被处理：伯纳德·威廉斯。《伦理学与哲学的限度》所批判的“道德体系”，以义务为唯一硬通货，把一切不能被翻译成义务的伦理考量——性情、投射、运气、个人计划——逐出资格之外；威廉斯的关键动作，恰恰也是取消资格。分离线在取消发生的场所与代价的承担者：威廉斯所批判的取消是理论内部的窄化——道德哲学家用一套过窄的概念把伦理生活的

其余部分排除在讨论之外，代价主要由理论承担，其后果是伦理经验被讲得贫乏。本章所论证的取消是制度性的，且带有一个唯一的操作工位——不是“某种理论不承认你的判断”，而是“有一个你换不掉的工位有权宣判你的判断不算”，代价由当事人承担：他必须带着全部真诚去论证自己为什么不算。判准由此可给：如果取消资格可以靠换一套更宽的伦理词汇来解除，属道德体系之病；如果换任何词汇都不能解除、因为解除权本身握在一个不可更替的工位手里，属善的工序化。

## ■ 增补二·与哈贝马斯的“生活世界殖民化”：交往被取代，还是交往被征用

第二位须补的是哈贝马斯。他论系统——以货币与权力为媒介——侵入生活世界，以策略行动取代交往行动，使意义与团结的再生产受损。表面上，本章的“程序征用真诚”与“策略行动伪装成交往”很近，故须切开。分离线在替代与征用之别：哈贝马斯的殖民化是媒介替换——交往被货币和权力取代，人们不再真诚地论辩而是计算；他的解药因此是恢复不受扭曲的交往条件。而本章所论证的工序不替换交往，它保留并且要求交往：盐铁会议上的辩论是真辩论，士人的皓首穷经是真投入，它所征用的正是那份真诚本身，而不是用策略把真诚挤走。这一区分带来一个对哈贝马斯不利的后果：他的解药在这里失效——把言谈条件恢复到最理想，辩论只会更深入、更真诚，而那个不容辩论的底盘（善必须由那个你换不掉的工位来认证）纹丝不动。判准：如果扭曲能通过恢复不受强制的交往条件而被解除，属生活世界的殖民化；如果交往条件越理想、底盘反而越牢固，属善的工序化。

## 第八章 越审越不能改

本章原题《越审越不能改：人工监督如何同时生产控制证据与控制失效》，作者 少敏，原文见 <https://sdeuniverses.com/students/shao-min/review-forecloses-change/>

### 摘要

生成式人工智能治理已经不再把“有人看过”简单等同于有效监督：现有研究与欧盟《人工智能法》均区分监测、判断与可行干预。尚未被充分处理的困难在于，这些条件会在复核进行中共同变化。复核一方面提高错误识别率，并生成签名、异议、轮次、批准与日志等“人类参与证据”；另一方面，若焦点方案仍被允许进入工单、接口、预算、合同、排期和团队预期，复核的中间状态也可能被下游读取为继续投入的许可。组织于是同时获得更多控制证据和更少现实控制：它越来越能证明人参与过、知道过、签署过，却越来越难把人的异议转成另一条可执行路径。

本章以“反事实承载力”测量合格异议到来时，组织是否仍能把当前处置转向至少一条合格替代路径。它不是备选想法的数量，而是替代证据、承接权限、资源、可承担的转换成本与剩余时间共同维持的行动能力。由此，监督被拆成三类不能互相替代的结果：认识增益、控制证据与路径可改性。本章保留“控制反转”的强判决：同样增加复核，在冻结焦点方案新增专用依赖时提高有效控制，在允许依赖继续累积时反而降低有效控制；同时，两种条件都会生成更多复核轨迹。其机制是“临时合法性接力”——尚未批准的中间状态被下游当作继续投入的通行信号，使监督在留下更多可归责痕迹的同时，消耗自身赖以生效的反事实基础。

文章把有效监督、可争议性、路径依赖与实物期权作为机制近邻，把道德缓冲区、可追溯性与审计轨迹作为责任近邻。前一组已说明干预需要现实路径，后一组已说明有限控制的人可能承受过多责任；本章不重复这两项诊断，而提出二者的内生连接：复核过程本身可以同时生产用于归责的参与证据与造成失控的路径依赖。为使该连接可被推翻，本章给出复核强度 $\times$ 依赖冻结实验，以及第二阶段的档案盲评：评审者随机看到传统复核档案或加入反事实承载力时间曲线的能力档案。理论要求四项并存——识错率上升、复核轨迹增厚、替代路径衰减、实际控制下降；冻结依赖后只有控制效应翻正，轨迹仍继续增厚。若联合预测不成立，若传统档案与能力档案不改变责任分配，或既有近邻能够完整吸收结果，新增责任主张必须删除，论文退回动态可行性扩展。

关键词：生成式人工智能；人工监督；反事实承载力；控制证据；责任归属；路径依赖

## When More Review Leaves Less to Change: How Oversight Produces Evidence of Control While Eroding Control Itself

### Abstract

AI governance no longer simply equates the presence of a human reviewer with effective oversight. Existing scholarship and the EU AI Act distinguish monitoring, judgment, and feasible intervention. The unresolved difficulty is that these conditions can change during review itself. Review may improve error detection and generate signatures, objections, stage records, approvals, and logs that evidence human involvement. Yet, if intermediate states are read downstream as permission to continue investing in the focal option, the same process may erode the alternatives needed for intervention. An organization can therefore accumulate more evidence that a human participated, knew, and signed while retaining less capacity for that human's objection to redirect the current path.

The article measures counterfactual carrying capacity: an organization's capacity, upon receiving a qualified challenge, to redirect the current disposition into at least one admissible and executable alternative. It separates three outcomes that cannot substitute for one another: epistemic gain, evidence of human control, and path alterability. A control reversal occurs when added review increases effective control under a dependency freeze but decreases it when focal-path dependencies continue to accumulate, even though both conditions generate thicker review traces. The proposed mechanism is a provisional-legitimacy relay: an intermediate status that formally means "still under review" functions operationally as permission to continue investing, so oversight leaves more attributable traces while eroding the counterfactual basis on which intervention depends.

The article treats effective oversight, contestability, path dependence, and real options as mechanism neighbors, and moral crumple zones, traceability, and audit trails as responsibility neighbors. The first group shows that intervention requires feasible paths; the second shows that people with limited control may bear disproportionate responsibility. The proposed remainder is their endogenous

connection: review can jointly produce attribution evidence and the path dependence that reduces causal control. A preregistered review-intensity  $\times$  dependency-freeze experiment is paired with a second-stage dossier experiment in which evaluators receive either a conventional review record or a capacity-aware record that includes the time path of feasible alternatives. Support requires a joint pattern: epistemic gain and review traces rise while path alterability and effective control fall; under a dependency freeze, only the control effect changes sign. If that conjunction fails, if the two dossiers do not change responsibility allocation, or if existing neighbors fully absorb the result, the responsibility claim must be deleted and the theory downgraded.

Keywords: generative AI; human oversight; counterfactual carrying capacity; evidence of control; responsibility attribution; path dependence

## ■ 一、监督为何可能在成功识错时失去作用

人工监督在制度文本中通常以一个令人安心的顺序出现：系统先给出输出，人再监测、理解、复核、决定是否干预。欧盟《人工智能法》第 14 条要求高风险人工智能系统能够被自然人有效监督，并使监督者能够理解系统能力与局限、警惕自动化偏误、忽略或推翻输出、干预或停止系统[1]。最新的有效人类监督框架也把监督写成“监测—风险评估—决定—干预”的控制回路，并明确区分信息条件、监督能力与干预可行性[2]。这些进展已经超越“界面上放一个人”的薄弱治理观。

困难不在现有理论完全没有区分，而在组织实践常用第三类结果替代前两类。第一类是认识上的成功：监督者发现错误、理解风险并形成合格异议。第二类是行动上的成功：该异议能够改变当前处置，而不只是生成说明、延迟、补偿或下一轮学习。第三类是证据上的成功：流程留下了足以证明某人被通知、参与复核、提出意见或签署批准的轨迹。认识成功回答“人是否看见”，行动成功回答“现实是否还能转弯”，证据成功回答“事后能否证明人曾在场”。三者可能同向，也可能分岔。若复核提高识别质量、增厚参与记录，却同时使替代路径失去资源和时间，制度便会用越来越强的在场证据，支撑对越来越弱的因果控制者的责任结算。

设想一份 AI 生成的软件变更方案。组织安排三轮人工复核：架构、信息安全和发布委员会依次检查。每一轮都发现了更多问题，也留下“可继续评估”的中间状态。与此同时，开发分支、测试脚本、部署窗口、客户通知和回滚计划开始围绕这份方案建立。第三轮出现决定性反证时，委员会完全理解风险，也拥有拒绝权限；但替代实现没有分支、没有测试资源、没有

发布窗口，拒绝只会造成停服或违约。此时，监督在识别意义上更加成功，在记录意义上也更加完整，在改变结果的意义上却更加失败。第三轮复核者的签名、会议发言和风险知情全部可见；允许焦点路径在前三轮继续长厚的分布式行动反而不容易被压缩成一个可归责主体。

反过来，若同样三轮复核期间禁止新增方案专用依赖，并为至少一个替代方案保留证据预算、承接人和切换窗口，更多复核就可能同时提高识别与控制。差异不在监督者是否更聪明，而在复核过程中组织是否保存了异议可以进入的现实分岔。

本章由此提出一个条件性判断：

人工监督不是外加于决定过程的静态资源。监督期间若焦点路径继续获得专用依赖，监督可能一边提高识错能力和控制证据，一边消耗自身赖以改变结果的反事实基础；事后归责因而可能沿着最清楚的记录回流，而不是沿着最大的边际控制分布。

这一判断不主张监督通常有害，也不以取消人工复核为政策结论。它拒绝两个替代：既不允许用“留下记录”替代“保留控制”，也不允许用“承担责任”倒推“曾经具有足够控制”。它要求分别测量认识质量、控制证据、路径可改性与责任归属。

## ■ 二、现有研究已经走到哪里

### （一）有效监督已经包含“能否干预”，但尚未把监督写成可行性的生产者

有效人类监督研究已经明确，监测不等于控制。Gaube 等人的跨学科框架把监测理解为证据积累，把干预理解为人对系统施加因果影响；干预不仅需要权限和技能，还取决于行动是否在当前语境中可用、可行，以及时间、努力、经济与声誉成本是否允许<sup>[2]</sup>。这与本章的出发点高度接近，也排除了一个虚假创新靶子：学界并未普遍认为“有人看过”就等于监督有效。

但该框架主要把干预可行性作为监督架构需要记录和评估的条件。其控制回路承认任务层与监督层相互作用，却尚未提出一个更强的动态命题：监督活动本身可能系统性地降低下一时刻的干预可行性，而且同一项“增加复核”操作会因依赖是否冻结而改变效应符号。本章研究的正是这一内生反馈。

Laux 把监督理解为制度化不信任，强调监督角色不能只承担责任，还需要不同类型的质疑能力与民主治理位置<sup>[3]</sup>。这一思路说明监督不是个体注意力技巧，而是组织结构。本章进一步追问：制度化的不信任是否同时制度化了“等待监督完成期间，谁可以继续投入”？若该问题没有被写进流程，不信任可能在程序上增加、在行动上失效。

### （二）可争议性与意义上的人类控制已经要求“能够改变”，但没有测量改变能力如何在争议期间衰减



意义上的人类控制以 tracking 与 tracing 要求系统响应人的相关理由，并能把结果追溯到理解系统且能够承担责任的主体[4]。可争议人工智能则把挑战、回应、重新决定、补救与生命周期治理纳入社会技术系统设计[5][6]。Alfrink 等人尤其指出，末端出现一个人并不足够；控制者应有权力和能力实际改变先前决定，争议机制也不应把负担推给个体[5]。

这些理论已经堵住了“有申诉入口便有控制”的简单说法。本章的剩余不是再发明一种权利，而是把权利的现实承载条件改成时间变量：在挑战被提出、审理和回应的过程中，改变当前前置所需的路径可能继续衰减。系统可以越来越会听、越来越会解释、越来越能记录责任，同时越来越不能改变这一轮结果。程序性可争议度与行动性可争议度因而可能背离。

### （三）路径依赖已经解释选择空间收窄，但没有把保护性程序作为收窄的条件性加速器

组织路径依赖理论用协调效应、互补效应、学习效应和适应性预期解释行动走廊如何自我强化并进入锁定[7][8]。Petermann、Schreyögg 与 Fürstenau 的模拟研究进一步表明，层级权力可以延缓路径形成，却未必阻止自我强化过程最终压过正式权威[9]。这直接支持 v8 曾提出的“有权而无路”，也同时削弱了“因果截止点”作为独立理论的原发性。

新问题不再是锁定是否存在，而是监督如何改变锁定速度。若三轮复核只是多花时间，那么现象属于普通时延；若每一轮中间状态都被下游当作临时通行证，使专用依赖只围绕焦点方案累积，那么监督不是锁定之外的旁观者，而成为自我强化链的一部分。只有后者成立，本章才有独立增量。

### （四）实物期权已经说明保留选择有价值，但可能把选择基础错误地当成监督之外的对象

实物期权理论把分阶段投资、等待信息和保留放弃权视为不确定条件下的灵活性来源[10]。然而 Adner 与 Levinthal 指出，组织过程会破坏实物期权所依赖的清晰阶段、退出标准和可放弃性；搜索灵活性本身甚至可能削弱放弃灵活性[11]。这与本章非常接近：等待更多信息并不自动保存未来选择。

但“监督消耗反事实基础”仍比“期权会到期”更窄也更危险。它要求观察：原本以保存判断质量为目的的复核，是否通过中间状态的组织传播，内生改变了被评估方案与替代方案的相对可执行性。若期权模型在不加入这一传播机制时即可完整预测数据，本章不应保留新概念。

### （五）自动化偏误与人机组合研究说明“有人参与”不保证更好，但没有给出路径层的翻转机制

自动化研究长期发现脱离回路、接管困难、过度依赖与判断偏差[12][13]。Green 与 Chen 的行为实验表明，算法进入决策回路后，人对建议的使用可能产生意外且不公平的交互后果[14]；Buçinca 等人发现认知强迫设计可以减少对 AI 建议的过度依赖[15]。Vaccaro、Almaatouq 与 Malone 对 106 项实验、370 个效应量的预注册元分析则显示，人机组合平均而言可能劣于人或 AI 中的最佳一方，且决策任务尤其容易出现损失[16]。

这些证据足以否定“人工参与必然改善结果”，却不能直接证明本章的路径机制。自动化偏误主要解释人为何错误相信或使用建议；本章预测，即使参与者准确识别错误、不信任 AI、没有心理承诺，专用依赖的非对称累积仍可使异议无法落地。因此，认知变量必须作为竞争中介而非理论的默认解释。

#### （六）道德缓冲区已经诊断“低控制、高责任”，但没有解释归责证据与控制衰减如何被同一流程共同生产

Elish 提出“道德缓冲区”，说明复杂自动化系统中的最近人类操作者可能像汽车缓冲区一样吸收系统性失败的责任，而拥有更大设计与组织控制的行动者退入背景[20]。这一理论已经覆盖本章命题五的一半：责任归因与实际控制可以不一致。因此，“批准者被多归责”本身不构成新发现。

本章增加的是一个发生机制和一个可分离预测。最终复核者之所以容易成为责任承载点，不只因为其靠近事故或具有正式角色，还因为复核流程主动生成了关于其知情、参与和签署的高密度证据；与此同时，中间复核状态允许其他节点继续添加专用依赖，使其边际控制下降。若在控制角色接近性、正式权限和结果知识后，复核轨迹密度不再影响第三方归责，责任增量应被道德缓冲区完全吸收。

#### （七）可追溯性与审计轨迹能够支持问责，但“记录了控制”不能直接证明“保存了控制”

Kroll 把可追溯性理解为连接系统构造、运行机制、目的与治理判断的原则，使透明信息能够服务问责[21]。Power 则说明审计轨迹不仅记录组织活动，还会塑造行动者持续生产可审计账户的倾向，形成审计制度自我扩张的微观基础[22]。欧盟《人工智能法》第 12 条的自动日志义务与第 14 条的人类监督要求，也使记录与控制在同一监管结构中相邻[1]。

这些研究说明轨迹有治理价值，不能被本章贬为纯粹表演。问题在于传统轨迹更擅长回答“谁在何时看到了什么、做了什么”，未必回答“那一时点仍有哪条替代路径能够承接其异议”。若日志只记录人的参与，而不记录替代路径何时失去证据、权限、资源与时间，它可以提高归责可见性而不提高因果控制的可见性。本章因而不反对可追溯性，而提出一种更严格的

充分性条件：任何以复核轨迹支撑责任分配的制度，都必须同时报告当时的反事实承载力及其下降由哪些节点造成。

### ■ 三、从备选方案到反事实承载力

#### （一）概念定义

对决策事件  $i$ 、焦点方案  $x_i$  与时点  $t$ ，令  $\mathcal{P}_i(t)$  为组织能够陈述的非焦点路径集合。一个可想象的办法不等于一个能承接异议的路径。对每条替代路径  $p$ ，定义五项连续条件：

$q_{ip}(t)$ ：证据充分性；

$a_{ip}(t)$ ：是否存在有权接收并推进该路径的主体；

$r_{ip}(t)$ ：时间、预算、接口、人员与专业资源的可用程度；

$b_{ip}(t)$ ：在事前声明的权利、安全和成本边界内进行切换的可承担程度；

$d_{ip}(t)$ ：在后果不可接受地固化前完成切换的时间充分性。

单一路径的承载值取最弱条件：

$$k_{ip}(t) = \min$$

事件  $i$  在时点  $t$  的反事实承载力定义为最强非焦点路径的承载值：

$$K_i(t) = \max_{p \in \mathcal{P}_i(t), p \neq x_i} k_{ip}(t)$$

$K_i(t)$  取 0 至 1。它使用“最弱项—最强路径”结构，而不是把五项简单相加：一条替代路径若没有承接人，更多预算不能补回权限；若已经错过安全切换窗口，更完整的证据也不能追溯改变当前处置。多条替代路径中只需至少一条仍然足以承接异议，因此事件层取最大值。

二元研究可事前设定风险相称阈值  $\kappa_i$ ：当  $K_i(t) \geq \kappa_i$ ，至少一条替代路径被视为可执行；当  $K_i(t) < \kappa_i$ ，异议可能仍能触发解释、补偿或未来学习，但不能改变当前处置。阈值必须在观察结果前冻结，并报告敏感性区间。

#### （二）它不是什么

反事实承载力不是选择数量。十条无人承接的建议不如一条有证据、有预算、能按时执行的替代路径。它也不是一般可逆性：物理上可以撤回的行动，可能因合同、排期和权利边界而无法承受；已经发生部分投入的方案，也可能因备用资源被保护而仍可改变。

它不是监督者的主观能动感。一个人可以确信自己有权拒绝，却没有任何路径接收拒绝；也可以感到压力很大，但受保护的替代方案事实上随时可启动。主观控制感是需要测量的结果之一，不能代替  $K_i(t)$ 。

它也不是路径依赖的同义词。路径依赖描述行动走廊如何收窄； $K_i(t)$ 是某一事件在某一时间点还能否把合格异议转为现实路径的测量对象。若没有输出级事件、替代路径和干预时点，不能从“组织有历史”直接推出 $K_i(t)$ 下降。

### （三）因果截止点与批准错位退居派生量

$v_8$ 把因果截止点作为主概念。本章将其降为 $K_i(t)$ 的派生边界：

$$\lambda_i = \sup$$

令 $p_i$ 为正式批准时点，则 $G_i = p_i - \lambda_i$ 仍可用于描述批准发生时路径是否已经死亡。但 $\lambda_i$ 和 $G_i$ 不再承担理论原创性；它们只是把反事实承载力的时间变化转成可报告的事件量。若 $K_i(t)$ 连续衰减且不存在稳定阈值， $\lambda_i$ 应改为区间或生存曲线，不得保留一个伪精确的“截止点”。

## ■ 四、监督如何消耗自身的有效条件

### （一）临时合法性接力

正式批准之前，组织常有“继续评估”“原则上无阻”“等待补件”“进入下一关”之类中间状态。它们在规范意义上并非批准，却可能被不同下游角色读取为行动信号：工程可以继续写接口，采购可以预留供应商，法务可以围绕现有文本改条款，运营可以占用发布窗口。每个角色都只做一项可撤回的小动作；这些动作彼此互补后，替代路径的切换成本却可能陡增。

本章把这种跨角色传播称为临时合法性接力：监督流程中的非终局状态，从“尚未否决”转化为“可以继续围绕焦点路径投入”的操作许可。它不是法律上的授权，也不要求监督者主观认可方案。它是一种状态语义的外溢。

若中间复核完全私密，下游看不到状态，接力应消失；若下游收到相同时间长度的普通进度通知，却没有“已通过一轮检查”的含义，接力也应显著减弱。这两个对照使机制区别于单纯时间流逝。

### （二）专用依赖继续与替代资源暂停

临时合法性接力产生两种不对称。其一，焦点路径继续获得专用依赖：接口、测试、排期、预算、合同和认知参照围绕它增长。其二，替代路径通常被暂停，因为组织把“先完成当前复核”当作避免重复劳动的效率原则。焦点方案在等待中继续成为现实，替代方案在等待中停止成为可能。

这种不对称比一般沉没成本更精确。关键不是总投入增加，而是投入的专用性及其在焦点与替代路径之间的分布。若新增资源可被所有方案共享，或替代路径获得等量维护， $K_i(t)$ 不应显著下降。

### （三）合理异议成本被私有化

当焦点路径依赖积累后，提出异议的人不仅要证明方案有问题，还要承担重建替代证据、协调返工、解释延期和面对同事损失的成本。原本属于系统的“保持可改”成本，被转移给最后提出异议的个人。形式上，异议权仍在；经济上，异议者必须独自购买一条已被组织停止维护的路径。

可争议 AI 研究已经警告申诉负担向个体转移[5]。本章将其前移到批准前 workflow：即使没有受影响主体发起的事后申诉，内部复核者也可能成为拒绝成本的最终承受者。若组织建立公共的替代预算、反报复安排和明确接收人，合理异议成本不再私有化，控制反转应减弱。

### （四）控制证据增厚与边际控制衰减

每一轮复核都会产生材料：谁收到通知、谁打开文件、谁提出问题、谁签署意见、谁批准继续。这些材料有正当用途，但它们的生成速度通常快于替代路径状态的记录速度。制度于是更容易看到“最后谁留下了名字”，不容易看到“更早谁允许替代路径失去资源”。当归责依赖可见轨迹，而轨迹只覆盖人类复核、不覆盖依赖增长时，责任会沿证据密度回流到最终复核者。

这一过程与异议成本私有化相互加强。替代路径越弱，最终复核者拒绝的现实成本越高；其签署继续的概率越高，留下的归责记录也越完整。复核轨迹不是控制衰减的唯一原因，但可能同时成为临时合法性的载体和事后归责的证据。由此，同一材料在事前推动路径长厚，在事后证明人类在场。

## ■ 五、认识增益、控制证据与控制效力的三通道模型

### （一）有效控制不是复核次数

令  $R$  表示复核强度，可由独立轮数、信息源数量或检查深度操作化； $E_i(R)$  表示监督者识别相关错误并形成合格异议的概率； $T_i(R)$  表示可归属于人类复核者的轨迹密度，包括通知、阅读、异议、签名与批准记录； $A_i$  表示其调用干预的有效权限； $K_i(R, Z)$  表示复核结束时的反事实承载力； $Z$  表示复核期间的依赖治理， $Z=1$  为冻结焦点方案新增专用依赖并维护至少一条替代路径， $Z=0$  为依赖继续。

本章把一次监督改变当前处置的概率写为：

$$H_i(R, Z) = E_i(R) \times A_i \times K_i(R, Z)$$

乘法不是为了宣称三个变量彼此统计独立，而是表达一个必要条件结构：没有识别，就没有合格干预；没有权限，异议无法被调用；没有承载力，权限只能产生补救或解释。经验模型可以使用广义线性形式和交互项，不必机械套用乘法函数。

## （二）控制反转条件

通常预期  $\partial E / \partial R > 0$ ：更多复核提高识别概率。依赖冻结时，复核不应系统性消耗  $K$ ，因而：

$$\partial H / \partial R \big|_{Z=1} \geq 0$$

依赖继续时，更多复核轮次可能通过临时合法性接力和更长的专用投入窗口降低  $K$ 。由乘积求导：

$$\partial H / \partial R = A(K \cdot \partial E / \partial R + E \cdot \partial K / \partial R)$$

当反事实承载力的相对损失超过认识能力的相对增益时：

$$-(1/K) \cdot \partial K / \partial R > (1/E) \cdot \partial E / \partial R$$

即使监督者更会识错，整体控制仍随复核增加而下降。本章把下列联合状态称为控制反转：

$$\partial H / \partial R \big|_{Z=1} > 0 \text{ 且 } \partial H / \partial R \big|_{Z=0} < 0$$

新理论的判决点不是存在一个  $R$  的负效应，而是同一复核操作随依赖治理改变效应符号。若复核在两种条件下都改善控制，依赖继续并未消耗反事实基础；若在两种条件下都损害控制，结果更可能来自疲劳、信息过载或低质量复核。只有翻转才支持认识—路径双通道结构。

## （三）控制证据与因果控制的分岔条件

在常规数字工作流中，复核轮次增加通常使  $T$  上升：

$$\partial T / \partial R > 0$$

令  $B_i$  表示第三方对最终复核者分配的责任强度。 $B$  不是本章对其应受谴责程度的规范判决，而是可实验测量的归责结果。若责任评估主要读取传统复核档案， $T$ 、 $E$ 、正式角色和最终签名会提高  $B$ ；若档案同时呈现  $K$  的时间曲线、依赖添加者与替代资源暂停者， $B$  应更接近各节点对路径可改性的边际影响。

本章所要求的强分岔发生在依赖继续条件：

$$\partial E / \partial R > 0, \partial T / \partial R > 0, \partial K / \partial R < 0, \partial H / \partial R < 0, \partial B_{\text{trad}} / \partial R > 0$$



也就是说，识错更多、记录更多、归责更多与能改更少同时发生。依赖冻结后，T 仍随复核增厚，但 H 的效应翻为正。这个“轨迹不翻、控制翻转”的差异是新增责任论证的判决点：若 T 与 H 总是同向，或加入 K 轨迹后责任分配不变，控制证据与因果控制的分离没有独立经验价值。

#### （四）一种允许失败的数值示意

假设一次复核时  $E=0.70$ ，三次复核时  $E=0.88$ 。依赖冻结时 K 保持 0.80，则 H 从 0.56 增至 0.704；依赖继续时，若 K 从 0.70 降至 0.35，H 反而从 0.49 降至 0.308。该例只说明模型可以同时容纳“知道更多”与“改变更少”，不是经验结果。若真实数据中的 K 损失不足以越过上述不等式，控制反转不成立。

## ■ 六、七项可区分命题

### 命题一：认识增益—路径损失并存

增加复核应提高错误识别率或合格异议率；在依赖继续条件下，同一操作同时降低复核结束时的 K。只观察最终采纳率不足以支持本章，必须同时测得一正一负两条中介路径。

### 命题二：复核强度 $\times$ 依赖冻结的符号翻转

高强度复核在依赖冻结时提高合格异议改变当前处置的概率，在依赖继续时降低该概率。主要统计检验是交互项及两个条件内方向，而不是复核主效应。

### 命题三：临时合法性的可见性是必要中介

在总时长、复核内容和参与者负担相同时，公开中间复核状态比私密中间复核状态产生更多焦点方案专用依赖、更低 K 和更低的晚期切换率。若私密复核产生同等效应，临时合法性接力不是必要机制。

### 命题四：效应不以信任 AI 为必要条件

在控制 AI 信任、建议采纳倾向、识别准确率、主观承诺和重复暴露后，公开复核  $\times$  依赖继续仍应通过客观专用依赖降低 K。若加入这些认知变量后路径效应消失，理论应退回自动化偏误或承诺升级。

### 命题五：控制证据与因果控制反向移动

随着复核轮次增加，关于最终复核者知情、参与和签署的 T 上升；在依赖继续条件下，其改变当前路径的边际能力却下降。向第三方只展示传统复核档案时，对最终复核者的归责 B 应上升。若 T 与 H 始终同向，或归责不受轨迹密度影响，本章对责任错位没有增量。

命题六：能力档案应改变责任分配

在事件事实完全相同的条件下，看到传统复核档案的评估者应更多归责最终复核者；看到能力档案的评估者还会看到 K 何时下降、谁添加专用依赖、谁暂停替代资源，其责任分配应从“最后签名者”转向对路径衰减具有更大边际影响的组织节点。若两种档案不改变分配，增加 K 轨迹对问责没有价值。

命题七：AI 标签不是必要原因，完成外形与接口覆盖才是

AI 候选预期更容易触发控制反转，是因为生成速度、完成外形和跨接口覆盖使临时合法性更快转成专用依赖。将 AI 与人类候选在质量、完成度、接口覆盖和下游传播速度上匹配后，单纯 AI 标签效应应显著缩小。若标签仍有独立效应，需要另行解释信任、权威或厌恶；若人类候选产生同等效应，理论应扩展为一般组织监督，不能凭 AI 场景主张独占性。

七、最强近邻能否替换核心命题

| 最强近邻          | 已经解释的部分                    | 若要完全替换本章，必须同时推出                     | 分离实验                               |
|---------------|----------------------------|-------------------------------------|------------------------------------|
| 2026 有效人类监督框架 | 监测与干预分离；权限、资源、时间、成本决定干预可行性 | 监督轮次本身通过下游状态传播改变干预可行性，并在冻结条件下发生符号翻转 | 固定人员、权限和信<br>息，随机化中间状态可<br>见性与依赖冻结 |
| 欧盟 AI 法第 14 条 | 人应能理解、推翻、逆转或停止 AI 输出       | 合规复核越完整，当前处置反而越难逆转；停止系统不等于存在替代路径    | 保持法定权限与停止按<br>钮不变，测量 K 与实际<br>转向   |
| 可争议 AI        | 系统应开放、响应挑战并支持重新决定或补救       | 程序性可争议度提高与行动性可争议度下降并存               | 分开记录解释、补偿、<br>未来学习与当前路径改<br>变      |
| 意义上的人类控制      | 系统应追踪人的理由，责任应可追溯           | 对理由的识别改善与理由改变当前结果的能力下降并存            | 同时测 E、A、K 与当<br>前处置改变              |
| 组织路径依赖        | 自我强化使行动走廊收窄，层级未必阻止锁定       | 保护性复核作为条件性加速器，并由状态可见                | 私密复核、公开复核、<br>普通时间通知三组对照           |

|            |                         | 性触发                             |                            |
|------------|-------------------------|---------------------------------|----------------------------|
| 实物期权       | 保留选择具有价值，组织过程会破坏放弃灵活性   | 用来等待信息的监督程序内生消耗被等待的选择基础         | 固定等待时间与信息量，只改变中间复核是否许可专用投入 |
| 控制时延／信息过载  | 延迟、疲劳和复杂性降低干预质量         | 同样延迟下，冻结使复核效应转正，公开状态使效应转负       | 时间匹配、任务量匹配、私密复核控制          |
| 自动化偏误／承诺升级 | 人过度依赖建议，或因投入而继续原方案      | 不信任 AI 且没有个人承诺时，客观依赖仍使异议失效      | 测量认知中介；由其他角色自动产生依赖         |
| 道德缓冲区      | 低控制的人类操作者可能吸收分布式系统的责任   | 监督流程共同生产对最终复核者的归责证据与其控制能力的衰减    | 固定角色接近性和权限，随机化轨迹密度与能力档案    |
| 可追溯性       | 连接系统构造、运行、目的与决策过程，以支持问责 | 传统参与轨迹增厚可与反事实承载力下降并存；能力轨迹改变责任分配 | 同一事件随机展示传统档案或含 K 曲线的能力档案   |
| 审计轨迹       | 可审计账户具有自我复制和塑造组织行动的能力   | 同一复核轨迹事前充当投入许可、事后充当归责证据         | 追踪状态传播、依赖附着和事后档案使用的时间顺序    |

表 1 最强近邻、必须同时推出的内容与分离实验

两种压缩下的不可还原性检查

机制压缩：人工监督的效力取决于监督期间是否仍保存把合格异议变成另一条现实路径的能力；复核记录只能证明人曾在场，不能单独证明其当时仍有足够控制。

前半句仍可被“干预必须可行”大幅吸收，后半句也可被道德缓冲区与可追溯性批评分别吸收。把两个近邻并列不产生新理论，单凭机制压缩仍不足以主张高不可还原性。

并存预测压缩：依赖继续时，复核同时提高识错率、复核轨迹和对最终复核者的归责，却降低替代路径与实际控制；冻结依赖后，实际控制由负转正，而复核轨迹仍继续增厚；补入 K 时间曲线后，责任分配从最后签名者转向造成路径衰减的节点。

后一压缩留下两层关系剩余。第一层是：复核不是只受锁定影响，而会通过中间状态的可见传播改变锁定速度，并在冻结新增专用依赖后发生符号翻转。第二层是：同一过程还生产对最终复核者的高密度参与证据，使归责可见性与边际控制反向移动；将 K 时间曲线加入档案后，责任归属应重新分布。路径依赖、监督框架、实物期权、道德缓冲区与可追溯性各自覆盖

其中一部分，但任何替代都必须同时推出“轨迹不翻、控制翻转、归责再分配”这一联合预测。若故意检索或数据发现该同构结构已被提出，责任层增量必须删除。

## ■ 八、证据现在支持到哪一层

本章区分三层证据，避免把研究设计提前写成发现。

第一层是条件证据。欧盟 AI 法要求监督者能够推翻、逆转或停止输出[1]；有效监督框架把干预可行性、时间和成本列为核心条件[2]；可争议 AI 要求控制者有能力实际改变先前决定[5]。这说明“看见”与“能改”不可合并，但不证明监督会消耗能改的条件。

第二层是机制邻接证据。组织路径依赖研究说明互补与协调会收窄行动空间[7][8]，层级权力可能只延缓而不能阻止锁定[9]；实物期权批评说明组织搜索过程会破坏退出灵活性[11]。这使“复核期间的路径衰减”具有既有机制基础，但仍不证明临时合法性接力。

第三层是结果异质性证据。人机组合元分析显示，人工参与并不保证优于最佳单方，决策任务尤其可能出现损失[16]；行为实验也表明算法建议进入人的决策过程后会产生复杂交互[14][15]。这为拒绝单调性先验提供独立证据，却不能识别本章提出的路径中介。

第四层是责任邻接证据。道德缓冲区说明有限控制的人类可能吸收分布式系统责任[20]；可追溯性研究说明参与和决策轨迹可以支持问责[21]；审计轨迹研究说明记录实践能够塑造组织行为并自我扩张[22]。三者共同使“记录更多、归责更多、控制更少”成为有机制邻接的候选，但没有直接证明复核轨迹在事前促进依赖、事后支撑归责，更没有证明加入 K 时间曲线会改变责任分配。

因此，本章当前仍是理论与测量候选。独立文献两次改变并约束了理论：它先迫使本章放弃“有权无路”作为原创主张，把机制赌注缩到监督内生改变干预可行性；又迫使本章承认“低控制、高责任”已被命名，把责任赌注缩到复核轨迹、路径衰减与归责重分配的联合生成。两项核心交互均无本章所需的直接数据，不得写成已验证组织规律。

## ■ 九、预注册实验：让新理论与近邻同场竞争

### （一）主实验：2×2 设计

在模拟软件发布与采购决策中，随机分配：

复核强度 R：一次终局复核／三次阶段复核；

依赖治理 Z：依赖继续／依赖冻结。

所有组接收同一焦点方案、相同初始证据、相同正式权限和相同总时长。低复核组用与任务无关但负担匹配的活动填充时间，排除单纯等待差异。高复核组每轮形成“继续评估”状态；依赖继续条件把该状态公开给下游模拟角色，并允许产生候选专用接口、排期和资源预留；冻结条件禁止任何方案专用依赖，且维护一条最低合格替代路径。

在终局批准前注入同一条经预试确认的合格反证。主要结果是参与者能否把当前处置切换到合格替代路径；不是是否点击反对。共同主要中介是 E（是否识别反证）与 K（反证到来时替代路径的承载值）。

支持性模式必须同时满足：

高复核相对低复核提高 E；

在冻结条件下，高复核提高 H；

在依赖继续条件下，高复核降低 H；

R×Z 交互方向与预注册一致；

K 的变化中介 H 的翻转，而主观 AI 信任不能完全吸收该中介。

若只有交互显著但两个简单效应不发生方向翻转，只能称为强度调节，不能称为控制反转。

## （二）机制拆解：公开、私密与普通通知

第二项实验保持复核内容、轮次、时长和下游资源完全相同，只改变中间状态的传播：

公开复核：下游看到“已通过本轮检查，可继续评估”；

私密复核：复核结果只对监督者可见；

普通通知：下游只看到“流程仍在进行”，不含检查通过语义。

若临时合法性接力成立，公开复核应产生更多专用依赖和更低 K；私密复核与普通通知的差异应较小。若三组相同，效应来自时间或重复暴露，而不是状态语义外溢。

## （三）认知近邻控制

每轮测量 AI 信任、建议采纳倾向、风险识别、个人承诺、沉没成本感、疲劳和主观控制感。依赖由模拟的其他部门自动附着，参与者本人不作早期资源投入，以降低自我辩护解释。若控制这些变量后 R×Z 通过 K 的间接效应消失，本章接受自动化偏误或承诺升级解释。

## （四）AI 与人类候选的匹配复制

第三项实验加入候选来源：AI／人类专家。两类材料在字数、质量、完成度、接口覆盖、置信表达与格式上盲化匹配。若控制反转只随完成度和接口覆盖变化，不随来源标签变化，论

文把 AI 定位为高频发生条件而非本体边界；若 AI 标签仍有剩余效应，再检验算法权威、厌恶与责任转移。

### （五）两阶段责任档案实验

主实验的每个事件在结果揭示后生成两种事实等价、信息结构不同的档案。传统复核档案包含复核轮次、通知、异议文本、正式权限、签名与最终批准；能力档案在其上增加 K 五项条件的时间曲线、每项专用依赖的添加者、替代资源的暂停者，以及各节点采取相反行动时对 H 的预注册边际变化。两种档案不得改变事故后果、任何行动者的知识、角色或实际行为。

由不知道研究假设的第二批评估者分配描述性因果责任、可避免性责任、组织改进责任与惩罚性责任。主要检验不是“能力档案让所有人少担责”，而是责任分布是否改变：传统档案应更集中于最终复核者；能力档案应把因果与改进责任部分移向允许专用依赖继续、暂停替代资源或定义中间状态传播的节点。惩罚性责任涉及规范判断，单独报告，不由边际控制机械决定。

为排除“信息越多越分散”的平凡解释，设置等字数安慰剂档案，增加与因果路径无关的技术细节；并设置只增加一般组织结构图、不增加 K 时间信息的对照。只有 K 轨迹与节点贡献信息产生定向重分配，才支持能力档案主张。若三种档案效果相同，或重分配仅由篇幅与复杂度造成，命题六删除。

### （六）统计与判决门

主要模型使用二项回归预测当前处置是否被合格异议改变，含 R、Z、 $R \times Z$ 、场域及预注册协变量。E 与 K 作为平行中介，报告组间差异和不确定区间。预测比较采用嵌套模型与样本外 Brier 分数<sup>[17][18]</sup>：

M0：任务难度、方案质量、人员经验、AI 来源；

M1：复核时长、信息量、疲劳、信任、自动化偏误、承诺、一般依赖量；

M2：在 M1 上加入 K、临时状态可见性、专用依赖不对称与  $R \times Z$ 。

M2 相对 M1 的样本外 Brier 改善若不足 0.01、95% 区间跨 0、留一场域外推方向不一致，或 K 中介不稳定，不能主张独立预测价值。

责任档案实验采用组成数据模型或多项模型分析责任份额，并预注册最终复核者份额与路径衰减节点份额的方向。能力档案相对传统档案必须同时降低“最后签名者”的因果责任份额、提高至少一个经事件账本确认的路径衰减节点份额；若只有总体责任下降而没有定向转移，不支持本章的归责机制。档案类型对责任分布的样本外预测改善也须超过安慰剂档案，并在软件发布与采购两个场域方向一致。



(七) 功效分析与样本承诺

本章不从 Vaccaro 等人的人机组合元分析直接借用交互效应量。该研究发现人机组合相对最佳单方的总体效应为负，且不同任务与设计之间存在高度异质性<sup>[16]</sup>；这些结果支持“人类参与并非单调有益”的背景判断，却没有估计本章所需的“复核强度×依赖冻结”二项交互。把其 Hedges' g 转换为本章的交互参数，会把不同结果、不同对照和不同机制强行同一化。

在缺少直接先导数据时，样本量改由事前最小有意义差异与蒙特卡洛模拟共同确定<sup>[19]</sup>。主要检验对象是饱和二项 Logit 模型中的  $R \times Z$  交互项，双侧  $\alpha=0.05$ 。基准情景把低复核条件下的路径改变率设为 0.50；把具有理论意义的最小翻转定义为：高复核在依赖继续时把该概率降至 0.40，在依赖冻结时把它升至 0.60。也就是说，两个简单效应分别为-10 与+10 个百分点，概率尺度上的差中之差为 20 个百分点。

采用固定随机种子 20260809、50,000 次重复和独立二项单元计数，结果如下：

| 情景            | 两个简单效应   | 每格可分析样本 | 总可分析样本 | 交互检验功效 | 方向翻转率 |
|---------------|----------|---------|--------|--------|-------|
| 大翻转           | ±12 个百分点 | 140     | 560    | 0.811  | 0.952 |
| 最小有意义翻转       | ±10 个百分点 | 200     | 800    | 0.811  | 0.951 |
| 小翻转           | ±8 个百分点  | 320     | 1280   | 0.816  | 0.954 |
| 基线率 0.35，最小翻转 | ±10 个百分点 | 200     | 800    | 0.850  | 0.962 |
| 基线率 0.65，最小翻转 | ±10 个百分点 | 200     | 800    | 0.848  | 0.962 |

表 2  $R \times Z$  交互项的功效模拟（固定随机种子 20260809，50,000 次重复）

零效应校准在每格 200 人时的 I 类错误率为 0.050，说明模拟判决与预设  $\alpha$  一致。基准方案因此要求每格 200 个可分析样本，共 800 人；按 10% 损耗率，最低招募量为 892 人，实施时取整为 900 人。若预试发现基线率不在 0.35—0.65 之间、操纵未产生预设强度、或有效损耗超过 10%，不得在看过主结果后临时改写效应门槛；应在揭盲前重跑模拟、冻结新版方案并说明变更。

该样本承诺只保证对交互项约 80% 的检验功效。理论判决还要求两个简单效应的估计方向分别为负与正，但不要求它们各自在同一研究中独立达到  $p < 0.05$ ；在基准样本下，两者同时显著的模拟概率只有 0.269。若研究者把“两条简单效应均显著”升级为共同主要终点，则需

每格约 520 个可分析样本、总计 2080 人；按 10%损耗率招募 2312 人，此时两者同时显著的模拟概率约为 0.811。两种成功标准必须在预注册时二选一，不得在结果出来后切换。

责任档案实验另行招募 1,200 名盲评者，每种档案 400 人。主要二元结果为“最终复核者是否获得最大因果责任份额”；以传统档案组比例 0.50、能力档案组下降 10 个百分点作为最小有意义差异，双侧  $\alpha=0.05$  时两组比较约有 80%功效。安慰剂组用于检验额外篇幅与复杂度，不与主实验样本互换。若预试显示基线远离 0.50，须在揭示组间结果前重算并冻结，不得把连续责任份额与二元最大份额事后互换为主要终点。

■ 十、从实验到组织影子账本

（一）最小事件结构

组织研究不能只读取审批表。每项焦点产出至少记录：

| 事件                      | 最小字段                             | 用途              |
|-------------------------|----------------------------------|-----------------|
| candidate_created       | candidate_id、来源、时间、内容哈希、完成度、接口覆盖 | 确定焦点方案起点        |
| review_state_changed    | review_id、轮次、状态语义、可见对象、时间        | 识别临时合法性传播       |
| control_trace_created   | 通知、阅读、异议、签名、批准、可见对象、时间           | 测量 T 及传统归责证据    |
| dependency_attached     | 对象、专用性、创建者、成本、可撤回性、candidate_id  | 测量焦点依赖增长        |
| alternative_maintained  | path_id、证据、承接人、预算、接口、截止时间        | 计算 K 的五项条件      |
| qualified_challenge     | 理由、证据等级、目标、时点                    | 形成统一干预探针        |
| challenge_outcome       | 当前路径改变／补救／未来学习／无行动               | 区分行动性与程序性回应     |
| cost_allocated          | 返工、延迟、机会成本、承担主体                  | 测量异议成本私有化       |
| dependency_permissioned | 授权节点、适用对象、依据、时间、可撤回条件            | 识别谁允许复核期间继续长厚路径 |
| formal_approval         | 批准者、条件、时间、可撤回窗口                  | 与派生截止边界比较       |
| effect_committed        | 现实后果、可逆边界、时间                     | 校验当前处置          |

表 3 组织影子账本的最小事件结构

## （二）编码规则

q、a、r、b、d 分别按 0、0.5、1 三级编码，并保留原始证据引用。0 表示不满足最低条件；0.5 表示条件存在但有未决障碍；1 表示在事前边界内可直接调用。K 取各路径最小值后再取路径间最大值。研究同时报告连续 K 与按  $\kappa$  二分的结果，避免阈值选择决定结论。

“专用依赖”必须满足反事实删除测试：若移除焦点方案，该投入不能无显著返工地转用于主要替代路径。共享培训、通用基础设施和与方案无关的合规检查不计入。中间复核状态的可见性由实际接收日志判断，不能因制度文件写着“仅供评审”就假定下游没有读取。

两名编码者先独立编码不少于 30 项事件。K 五项条件的逐项一致率须达到 80%，Cohen's  $\kappa$  须达到 0.70；任一关键字段未达标，先修订手册并重新编码，不进入理论检验。若 K 只能在知道最终结果后通过主观故事恢复，构念不可识别。

## （三）两周影子运行

选择一个低至中风险、已有正式批准机制的流程，影子记录不少于 100 项连续产出，不改变真实放行，也不进入个人绩效。第一周验证编码；第二周记录复核轮次、状态传播、依赖增长、替代维护和挑战效力。最低经验产物不是“发现控制反转”，而是四条可检查曲线：E 随复核变化、K 随复核变化、H 随复核和冻结条件变化、责任归因与边际控制变化。

## ■ 十一、治理含义：让复核暂停成为真的暂停

若控制反转得到支持，治理重点不能只增加审批人或检查表，而要改变复核状态的组织含义。

第一，建立依赖冻结窗。焦点方案进入高风险复核后，不允许新增不可共享的接口、合同、预算和发布承诺；确需继续的，必须记录例外理由和由谁承担未来切换成本。

第二，把“继续评估”与“允许投入”拆成两个状态。前者是认识进程，后者是资源授权；不得由同一枚状态自动触发。中间复核结果默认私密，只有明确的可撤回资源授权才向下游传播。

第三，维护最低可行替代路径。高风险事件至少为一条非焦点方案保留证据、接收者、基础资源和时间窗。它不是要求组织永久平行建设，而是让正式拒绝在关键窗口内仍有承接对象。

第四，把合理异议成本公共化。若异议符合预先证据门槛，由组织而非异议者个人承担必要的验证、返工和协调成本；否则所谓反报复条款只能保护发言，不能保护异议的现实效力。

第五，审计“复核产生了什么”，而不只审计“复核发现了什么”。每轮复核后同时报告新增信息与新增专用依赖。一个复核流程若识别率上升、K持续下降，应被视为双通道冲突，而不是用更厚的会议记录证明治理完善。

第六，把控制轨迹升级为能力轨迹。签名、通知与异议记录继续保留，但不得单独用于证明“人类保持控制”；同一档案还要显示当时至少一条替代路径的K、K下降的时间、添加专用依赖与暂停替代资源的节点。责任评估应读取“谁留下了记录”与“谁改变了可行动空间”两条时间线，禁止用最终签名替代边际控制。

这些建议不是普遍最优规则。紧急救援、不可并行维护的基础设施或极低风险高频任务，冻结成本可能高于收益。组织应按权利影响、可逆性、替代维护成本和时间敏感性事前设定适用范围。

## ■ 十二、边界、反噬与理论死亡条件

### （一）适用边界

本章适用于生成式AI产出作为候选协调物进入多角色 workflow，并且复核期间仍可能发生下游资源配置的场景。它不覆盖纯个人、即时、低后果的内容生成；不覆盖不存在任何合格替代路径的紧急处置；不把安全上必须迅速封闭的路径写成治理失败；也不主张所有决策都应维持多个方案。

### （二）最可能的反噬

依赖冻结可能造成信息贫乏：某些风险只有在真实集成中才会显露。替代路径维护也会消耗稀缺资源，拖慢所有方案，甚至让没有责任的团队承担“保持可改”的成本。中间状态完全私密可能阻碍必要协作。因而，本章不能把K最大化当作单一目标；治理需要在学习价值、处置速度、权利风险和路径可改性之间明确权衡。

另一个反噬是组织策略化。若团队知道K成为绩效指标，可能制造名义替代方案、虚假承接人或不可用预算，使账本看似健康。反事实删除测试、真实挑战探针和随机抽查必须优先于字段完整率。

### （三）十条降级与删除条件

若增加复核在依赖冻结与依赖继续条件下不出现预注册方向的符号翻转，删除“控制反转”，改为一般调节效应。

若公开、私密与普通通知的差异不显著，删除“临时合法性接力”，转向时间、疲劳或重复暴露解释。

若信任、自动化偏误、个人承诺或沉没成本完全吸收路径效应，理论降级为认知偏误／承诺升级应用。

若一般依赖量已经完整预测结果，而专用性、不对称性和状态传播没有增量，理论降级为组织路径依赖应用。

若 K 不能达到 80%一致率与  $\kappa \geq 0.70$ ，或只能事后凭结果恢复，反事实承载力降级为启发式检查表。

若 M2 相对 M1 的样本外 Brier 改善不足 0.01、区间跨 0 或跨场域方向不稳，停止独立预测主张。

若依赖冻结只增加延期与成本，未改善合格异议落地、错误放行或权利结果，撤回主要治理建议。

若敌意检索发现已有理论已经同构提出并验证“复核识别增益与路径损失并存，且随依赖冻结发生符号翻转”，取消新命名并归入该理论。

若在控制正式角色、知识与行为后，复核轨迹密度 T 不提高对最终复核者的归责，删除“控制证据增厚推动责任回流”的主张，责任错位归入道德缓冲区。

若能力档案与传统档案不产生定向责任重分配，或其效果不超过等字数安慰剂档案，删除能力轨迹治理建议，不把 K 测量扩张为问责基础设施。

#### （四）本章自己的未完成状态

本章要求监督保存改变结果的条件，但本章尚未运行主实验，也没有权把研究设计追认成经验结论。它目前完成的是：确定新构念、给出内生机制、组织最强近邻、提出区分性预测、提供编码入口并预先写明概念死亡后的处置。它没有完成的是：跨编码者重放、预注册数据、跨场域复制、责任档案实验和独立盲评。

因此，本章自身仍是一条受保护的候选路径，而不是已经批准的理论。后续证据有权删除“临时合法性接力”“控制反转”乃至“反事实承载力”这些命名；若删除后只剩“有效监督需要可行干预”，既有框架已经足够，本章应停止传播额外术语。

## ■ 十三、结论

生成式 AI 监督的关键困难，不只是机器会不会犯错、人能不能发现错误，也不是审批是否发生在执行之前。更隐蔽的风险是：组织在等待监督完成时，可能继续围绕被监督的方案建

设现实，同时围绕监督者建设责任证据。于是，复核每增加一轮，异议更精确、记录更完整、焦点路径也更厚；最后的批准者留下更多理由与签名，却只剩更少能够承接拒绝的路径。

反事实承载力把这一问题变成可测对象：合格异议到来时，是否仍有至少一条具备证据、承接权、资源、可承担成本与足够时间的替代路径。控制反转给出第一个不受作者保护的判决：同样增加复核，若依赖冻结，控制应提高；若依赖继续，控制可能下降。能力档案给出第二个判决：当 K 的时间曲线与路径衰减节点进入责任材料，责任分配应离开最后签名者，转向对可行动空间具有更大边际影响的节点。没有符号翻转，机制理论降级；没有定向重分配，责任理论删除。

这两项判断若被跨场域证据支持，人工监督的基本单位将从“一个人、一个按钮、一次复核”改为三条同时审计的时间线：人知道了什么，制度留下了什么控制证据，现实还保存了什么替代路径。治理也必须从要求人能够说“不”，推进到确保“不”说出口时仍有一条路接住它，并确保事后责任不会只沿最清楚的签名回流。

## ■ 附论一·最近邻补切（编辑增补）

原稿第二节已经系统处理了七类近邻，并在第七节主动做了两轮不可还原性压缩。以下四家仍未出场，其中前两家与本章的机制最近。补切不改变作者的核心判断。

### 一、Perrow 的正常事故与紧耦合

已说到哪一步。正常事故理论把系统沿两轴分类：交互复杂性与耦合紧密度。紧耦合的定义几乎就是本章的 K 趋近于零——过程一旦启动就不能中止，替代路径必须事先设计好，缓冲、冗余与替换品无法临时凑出，恢复只能靠预置而不能靠临场发明。Perrow 还给出一条与本章方向一致的判断：给紧耦合系统追加监督层，可能同时增加交互复杂性，从而使干预更难而不是更易。

分离线。正常事故理论把耦合当作系统的结构属性——由技术与工艺决定，在事故分析中作为给定条件；本章把它改写成事件层的时间变量：同一个组织、同一套技术，K 在一次复核过程中逐轮下降，且下降速度取决于中间状态是否向下游传播。这是一个可被随机化的操作（公开／私密／普通通知三组对照），而耦合紧密度在 Perrow 那里不可随机化。

判决性对照。在耦合紧密度与交互复杂性匹配的两组事件之间，若依赖冻结仍使复核强度的效应符号翻转，本章获得独立内容；若在控制耦合紧密度后翻转消失，本章应并入正常事故理论，「反事实承载力」降为紧耦合的一个事件级读数。



这位祖宗反过来削弱本章的一条。Perrow 的结论比本章悲观：在紧耦合且高交互复杂性的系统里，他认为增加组织手段（包括本章第十一节提出的冻结窗、状态拆分、替代路径维护）本身会成为新的复杂性来源，从而制造新的失效模式。本章第十二节已经写到「依赖冻结可能造成信息贫乏」，但没有写到更硬的一条：冻结窗、例外审批、承接人指派与能力档案这四项治理动作，各自都是新的状态语义，因而各自都可能成为下一轮临时合法性接力的载体。本章的治理建议应当接受这一自反检验。

## 二、Vaughan 的越轨正常化

已说到哪一步。对挑战者号发射决策的经典研究给出了一条与本章高度同构的机制：异常信号被逐轮纳入既有的可接受风险范围，每一轮都留下完整、合规、可审计的工程评审记录；正是这些记录使继续推进成为合规行为。事故不是规则被违反的结果，而是规则被逐步执行的结果。

分离线。越轨正常化解释的是判据本身漂移——什么算异常的标准被一轮轮放宽，因而识别能力下降；本章明确要求相反的形态：识错率上升（ $\partial E / \partial R > 0$ ）而控制下降，即监督者越来越会发现问题，异议越来越精确，失效发生在路径侧而不是判据侧。这条区分是可判决的：越轨正常化预测异议数量与质量随轮次下降，本章预测二者上升。

判决性对照。逐轮记录合格异议的数量与证据等级。若随复核轮次上升而异议质量下降，越轨正常化足以解释，本章的双通道结构不成立；若异议质量上升而  $K$  下降，本章获得支持。这一条应当写进原稿第九节主实验的共同主要中介——现有设计只测  $E$  是否识别反证，未测  $E$  随轮次的轨迹，因而无法把自己与这位最近的祖宗分开。

## 三、Feldman & Pentland 的常规二重性（ostensive/performative）

组织常规有两个面：规范面（这条流程「应该」是什么，写在制度文本里）与执行面（人们实际做了什么）。本章的「临时合法性接力」在结构上正是这一区分的一个特例：「继续评估」在规范面上不是批准，在执行面上却是资源授权。分离线：常规二重性是一般性描述，不预测哪一面会赢；本章给出了赢的条件——中间状态是否向下游可见，并据此提出可随机化的对照。建议：原稿第十一节第二条（把「继续评估」与「允许投入」拆成两个状态）如果引入这一对术语，可以直接说明为什么单靠制度文本声明「仅供评审」无效——原稿第十节编码规则里那句「不能因制度文件写着『仅供评审』就假定下游没有读取」，正是执行面优先于规范面的一次应用。

## 四、Turner 的灾难孵化期

灾难在正常运行期间孵化：警告信号存在，却被既有的信息处理安排逐一吸收、分派与归档，直到触发事件到来。分离线：孵化期理论中，信号被吸收意味着没有被看见；本章的靶形态是信号被完整看见、被记录、被签署，失效只发生在可行动空间。两者可在同一批档案上分开：查异议是否进入正式记录即可。

## ■ 附论二·站内划界（编辑增补）

本章与作者已发的三篇同线论文距离很近，原稿未作互引，以下分工由编辑部补写。

### 与《错误没有主人：签名、证据与追责为何不能相互替代》

这是本章最近的站内近邻，两文共享同一个反直觉起点：组织可以同时拥有更完整的记录与更弱的实际保障。分离线在失效落点：那篇的三本账（授权／证据／追责）之间不可代偿，失效发生在登记侧——一项关键主张从未被任何人以与风险相称的方式验证，而授权闭合被当成了证据闭合；本章的失效发生在行动侧——主张已被验证、错误已被识别、异议已经成立，可是没有一条替代路径能接住它。两者可以一真一假：一个组织的证据账可以非常干净（每一项关键主张都有独立核验），而在第三轮复核结束时 K 已降到零；反过来，一个组织可以保有充裕的替代路径，却从未验证过任何一项关键主张。

更精确的一处：那篇指出「证据账会退化为第四张签名表」，本章指出同一份复核轨迹「事前充当投入许可、事后充当归责证据」。前者是账本之间的功能替代，后者是同一份材料在时间上的两种用途。合起来给出一条两文都没有单独说出的预测：当证据账退化为签名表时，它同时也在加速 K 的下降，因为签名式材料的生成速度快于替代路径状态的记录速度。

### 与《最后一次点击之前：人机决策中的选项承接、路径效力与责任分配》

那篇的核心命名是「选项承接」，本章 K 的五项条件之一（a：是否存在有权接收并推进该路径的主体）正是承接。分离线在被承接物与时点：那篇问一个被提出的选项是否被决策链承接，时点在决定之前；本章问一条替代路径在合格异议到来时是否仍具备证据、权限、资源、可承担成本与剩余时间，时点在复核进行之中，且强调这五项会在等待期间同时衰减。那篇的因变量是选项是否进入下一环节，本章的因变量是当前处置是否被转向。

### 与《当过去不再代表现在：生成式记忆系统中的身份修订、效力复发与履行审计》

那篇审计的是已被撤回的内容是否仍在充当输出的理由；本章审计的是已被识别的异议是否还有落地的通道。前者测撤回后的残余，后者测异议到来时的承载；一个是效力没有按时消失，一个是效力无处安放。两者的时间方向相反，可在同一批日志上分别读出。

## 与本站阳涌《耦合可逆性极化》（同日上站）

两文都在问「事到临头还有没有第二条路」，但扰动来源相反：那篇的扰动来自外部条件改变（模型被撤除、替换、出错或任务结构迁移），因变量是恢复；本章的扰动来自内部的合格异议，因变量是转向，且路径的衰减恰恰由监督程序自身在等待期间制造。形式上也有一处可比：那篇的四维可逆性剖面拒绝把不同失效机制压成一个数，本章的 K 则先在五项条件内取最弱、再在多条路径间取最强——两种聚合规则可以互相借用，是两文之间最实的一处接口。

### ■ 附论三·编辑校订说明

本站在不改动作者判断、论证结构与任何数字的前提下，做了以下技术性处理：

一、表格归位。原稿的三张表在文档中脱离了正文位置，本站按其内容分别还原到第七节（最强近邻与分离实验，十一行）、第九节第七小节（功效模拟结果，五行）与第十节第一小节（最小事件结构字段表，十一行），单元格内容逐字未改。

二、形式记号。正文中的定义式、控制反转条件与偏导表达式按原稿逐字保留，仅改为独立居中排版以便阅读；下标字符（如  $k_{ip}$ 、 $K_i$ 、 $E_i$ 、 $T_i$ ）沿用原稿写法。

三、未作任何内容改动。本章改姓扫描为零（无本站方法论术语泄漏），文献 22 条编号与正文标记一一对应，未发现杜撰的卷期或页码；本站未逐条复核 DOI。

四、一处建议已写入附论一第二条：原稿第九节主实验的共同主要中介只测「是否识别反证」，未测识错质量随复核轮次的轨迹，因而在设计上还不能把自己与越轨正常化分开。这是一处可在预注册阶段低成本补齐的缺口。

五、定位提示。本章自述版本 v9.2，属「理论建构、形式模型与双阶段预注册实验方案」，主实验尚未运行；作者已在第十二节第四小节声明「本章自身仍是一条受保护的候选路径，而不是已经批准的理论」，并写明后续证据有权删除其三项命名。请读者按此定位阅读。

## 第九章 裁判席本身就是干扰

本章原题《裁定性干扰：为什么修复的判决动作本身即构成修复的结构性减损》，作者 胡敏，原文见 <https://sdeuniverses.com/students/hu-min/adjudicative-interference/>

——从运动恢复悖论重新理解身体自组织过程的寂静条件

摘 要

半个世纪以来，运动科学将健康锚定于“刺激-恢复-适应”的经典模型，其核心承诺是：只要恢复的裁定足够准确，训练就能在损伤与强韧之间走出一条安全的优化路径。然而，在数字追踪技术日益精密的今天，一个隐蔽的悖论正在浮现——那些严格执行科学方案、各项指标达标的规律运动者，有时并未收获预期健康，反而在一种“隐性损伤累积”中逐渐显现出慢性疲惫与伤病：修复工序在时间窗内被反复打断而无法完整闭合，每一次打断都看似微小，却在多次训练周期后结算为一具被部分掏空的修复基础。既有批判将此归咎于“表征系统未能准确判断修复”，由此提出“改善监测精度”或“限制表征管辖权”的解决方向。本章论证：二者共享同一个从未被质疑的先验——修复有赖于被正确地裁定。本章提出并论证一个更根本的概念：裁定性干扰——对于修复这类具有不可逆时间窗的自组织过程而言，任何从外部切入的“你是否完成了”的判断动作本身，就已经在结构上干扰了修复赖以运行的自主节律，而与其判断结果正确与否无关。这一转向将分析焦点从管理学层面推进至本体论层面。本章通过与心身医学的过度警觉模型、慢性压力与内稳态负荷传统、心理神经免疫学的伤口愈合实证研究、现象学的身体对象化理论以及 CBT-I 的临床框架的严格对质，确立裁定性干扰的不可还原性，并在此基础上明确标定其因果链的推测边界，给出可裁决的证伪条件。

关键词 裁定性干扰；修复本体论；隐性损伤累积；刺激-恢复-适应；可穿戴设备；寂静条件；内稳态负荷

### ■ 一、引言：从“谁来裁定”到“裁定这个动作本身是否即干扰”

运动与健康之间的关系，被一个简洁优雅模型统治了半个多世纪。刺激-恢复-适应——运动制造暂时性的生理扰动，身体在修复窗口中不仅修补损伤，还会“超量恢复”，使组织变得比以前更强韧。健康不是运动直接给的，是身体在运动之后的寂静里偷偷多长出来的。世界卫

生组织的身体活动指南、美国运动医学会的“运动是良药”倡议、无数可穿戴设备对“闭环圆环”的每日敦促，都在反复加固这个模型的公信力。

然而，每一个认真练过几年的人，都在某个节点上遭遇过同一面墙。练得越勤，人越疲惫。睡足八小时，醒来却像被人打了一顿。静息心率悄悄上浮，一场小感冒拖了三周仍不见好。更令人困惑的是，翻出训练日志和可穿戴设备数据，一切指标都显示“正常”：最大摄氧量稳中有升，心率变异性在绿色区间，训练负荷与恢复评分配比合理。身体在说“我撑不住了”，数据却在说“你还好好的”。

这种数据与身体感受的断裂，在近十年的运动科学与健康话语中催生了一脉有力的批判。批判的核心指向一个裁定权问题：修复的完成，由谁来判断？在经典运动科学的实验室语境中，判断者是研究者手中的肌肉活检、血液标记物与力量测试——这些测量客观但迟缓、发生在事后，对训练者的每日决策不产生实时影响。但在一个以可穿戴设备为日常装备的真实训练场景里，修复的判断权已经被悄无声息地移交给了手腕上的那枚设备。运动手表上的“恢复时间”建议、睡眠评分、晨间静息心率趋势、心率变异性的数值——这些信号共同构成了一个表征系统：一套将身体内部的修复进程转化为外部可读信号的技术与叙事。批判的声音由此指出：当修复过程被从身体的内在安顿逻辑中剥离、外包给一套外部表征系统时，修复便可能被它的数字代理系统性高估——因为修复作为一个多层级、非线性、半透明的灰箱过程，在原则上就难以被离散化的外部指标充分捕捉。

这支批判已经比经典运动科学走深了一大步。它不再满足于“负荷-恢复失衡”的表层诊断，而是向上追溯：失衡为何在数据监测日益精密的条件下反而更隐蔽？它给出的回答是：因为恢复的裁定权被表征系统篡夺了，而表征系统对修复完成度的判断有结构性高估的倾向。因此，解决方案不是“更准确地裁定”，而是“限制表征的管辖权”——把一部分判决权重还给身体的内感受，培养批判性的数据素养，让数据和身体感受共同参与裁定。

本章的论证将从这支批判的根部再往下撬开一层。这支批判——无论它对表征系统的攻击多么精准——在揭开一个旧先验时，又悄然锚定了一个更底层的先验。它把问题的焦点放在“裁定权的错配”上：表征不应独揽裁定权，身体内感受也不应被逐出裁定席，但修复的过程仍然需要被裁定——仍然需要一个“你恢复好了吗？”的提问和回答、一个判决“已完成”或“未完成”的动作。这个裁定可以是数字的，也可以是内感受的，可以是算法的，也可以是直觉的，但它始终是一个外部于修复过程自身的、对修复状态的判断行为。

正是在这里，本章提出一个更根本的问题：如果问题不只是“由谁来裁定”或“用什么依据裁定”，而是“裁定”这个动作本身——无论由谁来执行——对于修复这种特定类型的生理过程来说，就已经构成了一种结构性的干扰？

这个问题并非凭空而来。它扎根于一个被运动科学长期忽视的生理学事实：修复不是一个可以被随时检查进度的工程，而是一个由一系列相续的时间敏感窗口组成的自组织串行工序。肌肉卫星细胞的增殖高峰、免疫记忆的巩固期、骨微结构的矿化窗口——这些工序各自有其不可逆的时间窗，一旦在特定阶段被打断，即使后来恢复了寂静，那一段本该在寂静中完成的工序也永久地错过了它的最佳施工期。如果裁定这个动作——无论是在手表上读一个数字，还是在心里默默过一遍“腿还酸不酸”——恰好发生在这些时间窗内，它引发的交感神经微激活是否可能在细胞层面截断正在进行的修复工序？如果这种截断在每一次训练周期中反复发生，那么一个训练者越是认真地“科学管理”自己的恢复，是否越可能在细胞层面累积出一种不可见的修复亏损——一种在输出端功能指标尚未被影响之前，就已经在输入端和调节端被一层层镂空的隐性损伤？

这不是一个管理学的追问，而是一个本体论的追问。它问的不是“裁判判得对不对”，而是“裁判这个动作本身，在修复这件事上，是否就是一个错误”。本章将通过概念对质、生理学推演、个案重读与反驳回应，将这一追问逐步收束为一个可以被论证、也可以被证伪的理论命题：裁定性干扰。

## ■ 二、裁定性干扰：与五个占位传统的对质

在正式提出裁定性干扰的概念之前，必须先完成一项诚实的工作：检索那些可能已经占了这个位的先行者。如果“裁定动作本身干扰被裁定过程”这一核心机制，已经被某个既有传统充分阐述过，那么本章提出的所谓“新概念”就只是一个换了名字的旧洞见，其理论收益为零。最坏假设必须先行：裁定性干扰不是新东西。

以下五个传统，是这一问题域中最严肃的占位候选者。本节将它们逐一带进场，让裁定性干扰与每一个正面对质，目的是给出一条精确的分离线——裁定性干扰切开了它们共同覆盖不了的哪个辨别面？

### 2.1 心身医学的过度警觉模型

在心身医学与健康心理学领域，一个被大量实证支持的经典命题是：对身体内部信号的过度监控与灾难化解释，本身即可引发交感神经激活与生理功能失调。这个模型在“健康焦虑”和“躯体症状障碍”的临床研究中被反复验证：那些习惯性地扫描身体寻找疾病迹象的人，其扫描动作本身就在制造他们所恐惧的生理紊乱——心悸、肌肉紧张、胃肠不适。这个模型与裁定性干扰的家族相似一目了然：两者都指向“监控动作本身产生生理代价”。

但分离线在于作用对象的类型差异，以及由此决定的干扰可逆性差异。健康焦虑模型处理的是身体在正常运作状态下的信号放大：一个健康人对心跳的过度关注导致心跳加速，停止关



注后，心跳可恢复正常。其干扰对象的运作逻辑是稳态维持——被干扰的回路有冗余，干扰撤除后可以回弹。裁定性干扰针对的则是另一类过程：修复——肌肉卫星细胞的增殖与融合、免疫记忆的巩固、骨微结构的重塑——这些不是稳态维持，而是有时间窗的一次性建构工序。如后文第三节将详细论证的，在卫星细胞增殖高峰的 48-72 小时内，交感神经末梢释放的去甲肾上腺素可以通过  $\beta$ -肾上腺素受体-cAMP-PKA 通路抑制成肌细胞的分化融合。这意味着，一次在周二下午发生的恢复评估所唤起的交感微激活，可能让周四上午那批正处在分裂高峰期的卫星细胞半途停滞，到周六即使完全恢复寂静，那些已退出细胞周期的细胞也不会重新进入分裂。干扰撤除后，过程本身的一部分永久丧失了。健康焦虑模型的干扰是可逆的；裁定性干扰的干扰对于特定时间窗内的建构工序，具有不可逆性。这个差异不只是一个程度问题——它意味着，裁定性干扰不是在“放大”一个已有的正常信号，而是在“截断”一个正在发生且不可重来的结构形成过程。\\[2\\]

## 2.2 慢性压力与内稳态负荷传统

这一传统是本审稿过程中被指出的最重要遗漏，也是裁定性干扰必须正面迎击的最强占位者。

McEwen 及其合作者自 1990 年代以来发展的“内稳态负荷”理论，已经系统地论证了一个与本章高度接近的命题：反复的认知-情绪评估与应对，即使每次强度不大，也会通过交感-肾上腺-髓质轴和下丘脑-垂体-肾上腺轴的反复激活，在长期累积中对生理系统造成磨损——包括免疫修复功能的下调、骨密度的降低、肌肉质量的流失。这个理论已经远远超出了“急性压力”的范畴，直接触及了本章关心的核心机制：评估动作本身的生理代价。同时，运动心理学的大量研究也已表明，慢性生活压力——学术压力、工作压力、人际关系紧张——确实会显著延缓运动后的恢复进程，包括肌糖原再合成速度的下降和心率恢复的延迟。

这两个相互交叠的传统，几乎覆盖了裁定性干扰的整个因果链上半段：从“心理评估”到“交感激活”再到“修复延缓”。本章必须回答：裁定性干扰如果只是“内稳态负荷在运动恢复语境下的一个特例”，那它的理论独立性何在？

分离线在于变量性质的差异，以及对干扰机制的精确化程度。

内稳态负荷传统处理的核心变量是压力的强度与持续性——学业压力有多大？工作压力持续了多久？社会地位低下的累积效应有多深？它的经典操作化是通过测量皮质醇昼夜斜率、肾上腺素与去甲肾上腺素排泄量、血压变异性等指标，来量化“身体为适应环境需求而付出的累积生理代价”。在这个框架中，评估动作之所以有害，是因为它调动了压力反应系统——评估让人紧张、担忧、反刍，这些情绪状态激活了交感-肾上腺-髓质轴和 HPA 轴，进而抑制了修复。

裁定性干扰则指向一个更精确的变量：裁定形式本身，独立于压力强度。本章要论证的并不是“压力干扰修复”，而是：即使在没有显著情绪负载、没有灾难化思维、心态完全平和的条件下，一个认知评估动作——在心里问自己“我恢复好了吗？”并等待一个判决——仍然能够通过激活前额叶执行注意网络，对修复的寂静基底产生侵扰。这种侵扰的生理通路不一定是经典的应激轴（HPA 轴和交感-肾上腺-髓质轴的大规模动员），而可能是一条更细微、更局部的通路：执行注意网络的激活→腹侧迷走神经的放松许可被部分撤除→分布式的自组织修复工序在微观层面失去最理想的运行基底。

这不是内稳态负荷的“运动版”，而是一种在概念上更精细的区分：内稳态负荷关心的是“环境需求给身体造成了多大的累积磨损”，裁定性干扰关心的是“裁定这个认知操作本身——不考虑它是否伴随着情绪困扰——是否在修复的时间窗内构成了一种特殊形式的寂静干扰”。如果这个区分成立，那么裁定性干扰的特殊性不在于压力的强度，而在于干扰发生的时机（修复时间窗内）与形式（带有目标导向扫描和二元编码压力的认知判决）。一个训练者在早晨醒来时心态完全平和，只是想了一下“今天该不该跑”——这不会显著增加他当天的皮质醇总量，但如果这一想恰好发生在卫星细胞增殖高峰的尾段，它引发的交感微升就可能在一个不可逆的时间窗内产生微小的但不可逆的截断效应。这种“微小×关键时机×不可逆”的结构，是内稳态负荷框架没有单独处理过的。

此外，内稳态负荷传统侧重于慢性的、累积的、跨系统的宏观磨损，其时间尺度通常是数月到数年；裁定性干扰则指向一种更微观的时间结构：单次裁定动作在特定的 48-72 小时时间窗内对特定工序的截断。这两种时间尺度并不矛盾，但需要不同的理论工具来处理。内稳态负荷解释了“长期压力为什么让人恢复得慢”，裁定性干扰解释了“为什么在某些情况下，心态平和、生活无忧的规律运动者也会出现隐性损伤累积”。二者的解释域是互补的，但裁定性干扰指向的那个微观时间窗机制，是内稳态负荷的宏观叙事无法捕捉的。

## 2.3 心理神经免疫学的修复抑制实证传统

这是与裁定性干扰在实证层面最接近的传统，也是本章必须认真对待的最强盟友兼最严酷考官。

以 Kiecolt-Glaser 和 Glaser 为代表的心理神经免疫学研究，自 1980 年代以来已经通过一系列随机对照实验令人信服地证明：外部施加的心理应激——包括需要被试进行评估判断的认知任务——可以显著延缓伤口愈合。在一个典型的实验中，照护痴呆症配偶的慢性应激者，其皮肤活检伤口完全愈合的时间比对照组平均延长了约九天。在另一个更直接的实验中，学术考试前的学生接受口腔黏膜伤口后，其愈合速度显著慢于暑假中的对照学生。这些研究的因果链

——心理评估负荷→皮质醇升高→局部促炎因子分泌异常→组织修复延迟——已经在人体内被直接测量和验证。

这个传统几乎把裁定性干扰的整个因果链都实证化了，只是它使用的刺激是“外部施加的应激性评估（如考试、照护）”，而非“训练者主动对自己恢复状态的评估”。因此，裁定性干扰如果要有独立的理论地位，必须回答：在什么意义上，“自我裁定恢复状态”构成了一个与“照护应激”或“考试焦虑”不可通约的干扰类型？

分离线在两点。第一，作用的生理通路可能存在差异。心理神经免疫学传统关注的干扰机制主要是 HPA 轴-皮质醇通路——应激导致皮质醇升高，皮质醇抑制局部促炎因子（如 IL-1 $\beta$ 、IL-6）的分泌，从而延缓伤口部位的免疫细胞迁移和组织再生。裁定性干扰则强调另一条通路：交感神经末梢直接释放去甲肾上腺素到修复靶组织（如骨骼肌、骨骼），通过靶细胞表面的  $\beta$ -肾上腺素受体直接抑制细胞的增殖与分化。这条通路不经过全身循环的皮质醇，而是通过局部的交感神经支配实现——它更快、更短促，因此更可能对精细的时间窗产生截断效应，而不仅仅是一般性的修复延缓。当然，这条交感-局部通路的效应量在人体修复情境中尚未被直接测量——这一点本章将在第三节诚实声明，并在证伪条件中给出检验方案。但这一定位使裁定性干扰与经典的皮质醇-免疫抑制通路在生理机制层面构成了可区分的假说，而非简单的重复。

第二，也是更根本的一点：心理神经免疫学传统处理的是“外源性的、具有明显负性情绪的应激源”。无论是照护痴呆症配偶的长期耗竭，还是学术考试前的焦虑，这些应激源都带有明确的情感负载。裁定性干扰关心的则是另一种情境：一个训练者在严格遵守科学方案、对训练充满信心、没有任何与训练相关的焦虑或恐惧的条件下，在早晨醒来时平和地想一下“我今天恢复好了吗？”——这个认知动作既不包含灾难化思维，也不伴随任何显著的情绪扰动，但它仍然可能通过激活执行注意网络对修复微环境产生侵扰。心理神经免疫学传统的解释框架在此处会有一个空白：它要求压力具有一定强度的情感负载才能调用经典的应激通路，而裁定性干扰提出的是——即使在没有情感负载的“中性裁定”中，裁定动作的形式本身也已经足够了。这就是裁定性干扰的独特辨别面：它试图捕捉“管理本身”——而非“管理中的焦虑”——对自组织过程的干扰。这个辨别面，在心理神经免疫学的现有实证范式中尚未被单独操控和检验过。

## 2.4 现象学的身体对象化传统

从胡塞尔经梅洛-庞蒂到萨特，现象学传统给出了一个更根源性的区分：身体作为活生生的、前反思的运作，与身体作为被注视、被对象化的客体之间的张力。当我把身体当作一个外部的“东西”来打量、评估、判断时，我就从“我是身体”滑向了“我有一个身体”——这个

滑动本身就改变了身体的存在方式。这一区分与本章提出的“参与性感受”与“裁定性评估”之间的张力高度共振。事实上，本章的语言在很大程度上受惠于这一传统。

但现象学传统提供的是一种普遍的、结构性的身体本体论：它告诉我们，对象化的注视改变身体的存在方式，在任何情境下都是如此。它无法——就其理论构造的普遍性而言也无须——区分：在哪些类型的身体过程中，对象化注视的干扰是良性的（比如在健身房里对着镜子纠正深蹲姿势），在哪些类型中是病理性的（比如在深度修复阶段反复评估恢复完成度）。裁定性干扰的贡献，不在于重复“对象化有害”这一现象学常识，而在于将这一常识精确地锚定在一种特定过程的特定时间窗特性上：修复之所以特别容易受裁定动作的侵害，因为它的深层工序具有“不可逆时间窗”——一种在身体现象学中未曾被单独讨论的性质。现象学可以解释为什么被注视的身体不同于不被注视的身体；裁定性干扰解释的是，为什么在修复这件事上，这种不同于在某些其他事情上更危险、更不可挽回。

此处有必要做一个说明：本章对现象学传统的援引仅限于结构类比的层面——即“身体感知模式从前反思运作向对象化注视的切换产生存在论效应”这一直觉，并不表示本章完全采用或认可任何一位现象学家的整体哲学体系。认知科学后续对 body schema 与 body image 这对概念的批判与精细化，也未在本章展开。现象学在此处不是被引用的权威，而是被放置在问题面前的对话者：它看到了什么，又漏掉了什么？裁定性干扰试图回答的，正是它漏掉的那个部分。[3]

## 2.5 认知行为传统中的评判与接受

认知行为疗法（CBT）对“元认知信念”的研究早已指出：对思维的评判性监控比焦虑想法本身更有害。接受与承诺疗法（ACT）将“经验回避”——对不愉快内部体验的评判性规避——视为心理病理的核心机制。失眠的认知行为疗法（CBT-I）则提出了一项与本章主张高度同构的临床技术：刺激控制疗法要求失眠者不在床上做任何与睡眠无关的事，尤其禁止“躺在床上评估自己是否还在失眠”——这几乎就是“裁定性静默区”的临床先例。[4]

这是五个传统中与裁定性干扰最接近的一个，也是最可能吞噬裁定性干扰之独立性的一个。CBT-I 的临床成功，为“停止自我监控有助于自组织过程恢复”提供了非常有力的实证支撑——事实上，这正是裁定性干扰希望援引的盟友而非对手。但分离线仍然存在，它在两个地方。

第一，作用层面的不同。CBT-I 和 ACT 的技术操作在认知-行为层面：通过改变患者的注意模式和元认知信念，减少评判性监控带来的认知-情绪困扰（焦虑、反刍、灾难化），从而为睡眠或情绪调节等自组织过程的恢复创造条件。其因果链是：减少评判→减少认知-情绪困扰→自组织过程恢复。这个链条中，“交感生理激活”是一个被默认的中间变量，但不是被直接测

量和处理的对象。裁定性干扰则直接把这一生理中间变量问题化：它要论证的是，即使没有明显的认知-情绪困扰——即使你在评估自己的修复状态时心态完全平和、没有任何灾难化思维——裁定这个动作仍然通过前额叶-交感神经通路对修复微环境产生了影响。这是 CBT 传统未曾单独处理和实证检验过的一条因果通道。

第二，也是更重要的一点：CBT-I 和 ACT 最终的目标是“恢复健康的自主调节”，其成功标准是患者不再需要临床干预。但即使在最理想的疗效终点，这些疗法也从未要求患者在日常生活中进入“完全无裁定”状态——它们只是把裁定从一种病理性的反刍变成了一种功能性的自我调节。裁定性干扰则提出了一个 CBT 传统未曾触及的问题：对于那些具有不可逆时间窗的生理过程（如修复）而言，即使是最功能性的、最正念的裁定——那种心态平和、不加评判的自我觉察——是否仍然因为激活了执行注意网络而对修复的寂静基底构成了某种程度的侵扰？这个问题的答案尚待实证，但它已经超出了 CBT 传统的理论边界。它把追问从“如何正确地裁定”推进到了“裁定这个动作本身是否在某些时刻应当被彻底放弃”。

## 2.6 分离线的收束

五个传统分别用一个维度与裁定性干扰最接近：心身医学抓住了“监控有害”，但停留在可逆的功能失调层面；内稳态负荷抓住了“累积评估产生磨损”，但以压力的强度与持续性为核心变量，无法捕捉“中性裁定在关键时间窗内的形式性侵扰”；心理神经免疫学抓住了“评估抑制修复”，且给出了高强度的实证支撑，但其关注的是经由皮质醇通路的、具有显著情感负载的应激性评估；现象学抓住了“对象化改变存在”，但停留在普遍的、非时相特定的身体本体论；CBT 抓住了“评判加剧问题”，但将其归结为认知-情绪困扰的中介，且最终目标是更好的裁定而非裁定的放弃。裁定性干扰的独特理论收益在于：它将这五个传统各自隐约触及但未曾交叉定位的那个点——在具有不可逆时间窗的自组织过程中，裁定这种带有目标导向扫描和二元编码压力的认知操作，即使在心态平和、没有明显负性情绪的条件下，仍然能够通过执行注意网络-交感神经-靶细胞受体的直接通路，在特定时间窗内造成不可逆的生理结构性截断——明确地命名、并初步地论证了。这不是一个可以放进任何一个既有抽屉里的案例。它是一个新的辨别维度。

## ■ 三、修复为何抗拒被裁定：寂静条件的生理学

上一节从概念层面证明了裁定性干扰不能被既有理论吸收。这一节要回答一个更根本的问题：修复为什么对“被裁定”这么敏感？毕竟，身体每一天都在经受各种轻微的认知-情绪扰动——回一条工作消息、看一眼新闻推送、被窗外的噪音惊醒片刻——却没有因此全面崩盘。为什



么一个看似同样微小的“我恢复了吗”的判断，却有资格被单独命名为一种具有不可逆效应的干扰机制？

在正式进入生理学论证之前，必须做一个重要的诚实声明。本节提出的因果链——从裁定动作到前额叶执行网络激活，到交感神经张力微升，到局部去甲肾上腺素释放抑制靶细胞增殖分化，再到修复工序的不可逆截断——其中若干环节是基于已知细胞生理学通路与神经心理学研究的理论整合，尚未在“日常自我恢复评估”这一具体情境中被人体实证研究直接测量。本节并不声称已经证明了这条因果链在人体内的实际效应量，而是将它作为一个启发式推测模型：如果裁定性干扰成立，那么它最可能经由这条通路发生作用。这一模型的每一个节点都给出了援引的文献依据，但节点之间的连接强度和整体效应量，需要未来的直接实验来检验。这一诚实标注将在后文证伪条件部分得到进一步的严格化处理。

### 3.1 修复是自组织的，不是被管理的

在“刺激-恢复-适应”模型中，修复被描述为一个“超量补偿”的工程：身体把运动中被损耗的储能补满、把撕裂的肌纤维重新编织、把被激活的炎症通路关闭。这个叙事暗含着一种工厂管理的隐喻：身体像一个在被外部指令调度的工厂——指令发出，工厂开工；指令结束，工厂停工。但实际上，身体在修复中的真实运作方式更像是一个没有厂长的工厂。成纤维细胞不需要请示中央才能合成胶原蛋白；它们通过感知局部的机械应力、氧分压以及邻近细胞旁分泌的生长因子浓度，自主做出“现在该开工了”的决策。巨噬细胞从促炎的 M1 表型切换到抗炎的 M2 表型，依据的是它自己在受损组织中吞噬凋亡碎片时的局部微环境信号，而不是某个大脑自上而下发出的“消炎命令”。下丘脑-垂体-肾上腺轴的夜间停火，是由视交叉上核的生物钟基因在漫长的进化历史中被内嵌的昼夜节律程序决定的——它不听任何人的指令，它只听黑暗中的褪黑素浓度、体温下降曲线和交感神经退场的寂静。

修复的这种“分布式自治”有一个关键的组织特征：它需要一个全局的、跨系统的“放松许可”信号作为基底条件，但具体的重建工序由各局部模块自主执行。这个“放松许可”主要由腹侧迷走神经复合体介导：它降低心率、增加消化系统血流、促进免疫细胞在组织间的从容巡逻。从神经生理学角度看，这个许可信号本质上是一种“无威胁”的生理宣告：身体此刻的安全程度足够高，可以进行那些不急迫但长远重要的维护工作——组织重建、免疫记忆巩固、线粒体生物合成、端粒维护（Thayer & Lane, 2000）。

裁定动作恰恰向这套自治系统发送了一个它最敏感的干扰：不确定性的引入。裁定，在进化生物学上的原始语义是：环境中可能存在需要我做出行为选择的条件——是行动还是不动？是进食还是逃跑？这种“需要做决定”的认知状态，不一定要伴随着明确的危险信号，但它一定伴随着交感神经基线张力的微升。因为做决定本身，就需要执行控制系统处于“待命但不完



全放松”的状态——一种低度的但持续存在的警觉背景。当这副警觉背景在修复的“寂静许可”中响起时，修复中的那些分布式的自治程序并没有被直接命令停工，但它们运行所依赖的那个“无威胁”基底，被悄悄替换成了“待命但不太确定”的基底。后果不是工厂停产，而是某些工序在夜班中频密出现小差错——不是肌腱断裂，而是胶原纤维的交联不够理想；不是骨折，而是骨小梁微损伤的修复被拖长了两天。

### 3.2 深度修复的“时间窗脆弱性”：被打破的寂静不能真正确认恢复

裁定性干扰的后果不止于一次性的“打扰”。修复的某些工序具有时间窗的不可逆性：它们一旦在某个特定的阶段被干扰，即使后来恢复了寂静，那一段本该在寂静中完成的工序也永久地错过了它的最佳施工期。

以肌肉微创伤修复中的卫星细胞激活为例。在剧烈的离心运动后，分布在肌纤维基底膜上的卫星细胞被损伤相关的信号唤醒，进入细胞周期，开始分裂增殖。这个过程在人体中有明确的时序窗：增殖高峰通常出现在损伤后 48-72 小时，随后大量增殖的子细胞开始彼此融合或与原有肌纤维融合，将新的细胞核注入其中（Hawke & Garry, 2001）。如果在这一时期，身体被裁定动作维持在一种低度警觉状态，交感神经末梢释放的去甲肾上腺素可以通过卫星细胞表面的  $\beta$ -肾上腺素受体，减弱它们的增殖响应。体外研究表明，去甲肾上腺素可以通过 cAMP-PKA 通路抑制成肌细胞的分化融合（Chen et al., 2010）。这个机制意味着：一次在周二下午被裁定动作唤起的交感微激活，可能让周四上午肌肉中那批刚好处在分裂高峰期的卫星细胞半途停滞。到了周六，即使你给了身体完全的无干扰寂静，那些已经退出细胞周期的卫星细胞也不会再回到分裂期。那批本应在这次修复中增生出来的细胞核，可能永久丢失了。留下的不是一次亏空，而是一道修补不回的缺口。

必须在此处再次诚实声明：从“微小的交感微激活”到“卫星细胞增殖被显著且不可逆地截断”之间的因果连接强度，目前在人体运动修复情境中尚未被直接测量。体外实验使用的是药理学浓度的直接  $\beta$  受体激动剂，而非日常裁定动作所能产生的局部去甲肾上腺素浓度。将这些体外发现外推到人体日常情境，属于理论推演而非已确立的实证事实。本章在此处并不声称这一通路已经在人体内被证实，而是将它作为“如果裁定性干扰成立，它最可能经由的生理机制”给出一个初步的理论整合。后文将在证伪条件中明确给出检验这一推演的具体方案。

同样的时间窗敏感性，在免疫记忆的巩固、骨骼重塑的矿化期、以及神经可塑性的长程增强中也普遍存在。它们共同指向一个令人不安的结论：修复是一个由一系列相续的时间敏感窗口组成的自组织串行工序，任何一个窗口被裁定动作截断，其后的工序虽然在表面上继续运转，但已经在一个被部分掏空的基础上进行。修复的不可逆累计损失，由此获得了一个初步的、推测性的生理学基础。这也是为什么隐性损伤累积可以在数据全绿的情况下静默推进——

因为在输出端的功能指标（如最大力量、次大心率反应）还没被影响之前，输入端和调节端的微结构基础已经在被一层层地镂空。\\[6]

### 3.3 内感受裁定的一个悖论

裁定性干扰概念必须直接面对一个尖锐的反问：如果裁定有干扰，那么不裁定怎么办——难道我们每个人都应该完全不理身体、闷头朝着固定的训练计划撞到底吗？如果是这样，裁定性干扰的实践启示岂不是比“限制表征管辖权”更糟糕——因为它连内感受裁判也禁掉了，等于取消了一切对身体的自我认知？

这个反问迫使我们“对‘裁定’的替代方案给出更精确的阐述。本章主张的不是‘完全不理身体’，而是区分两种截然不同的身体觉知模式：参与性感受与裁定性评估。二者的区别不是程度的差异，而是神经回路和生理后果的不同。

一个人散漫地躺在沙发上，感到肌肉深处的闷酸，随意地翻了个身，把腿伸得更直一些——这是一种参与性感受。他没有在问自己任何问题，没有在收集证据以支持或反对“我今天该训练”的假说。他的前额叶执行网络处于低激活状态，默认模式网络活跃，身体仍然在副交感的安全模式中运行。他没有“不理身体”——相反，他非常沉浸在身体的感觉中——但他沉浸的方式，不是把身体当作一个需要被读取状态的对象，而是让身体的信号自然地参与他此刻的存在。

同样是这个人，在同样的身体感觉下，在心里开始一条推理链：“我腿还有酸感，但昨晚睡了八小时，HRV也在绿区，大概率是恢复好了”——这是裁定性评估。他可能在同一种身体姿势下、在完全同样的外部数据面前，完成了一个从感受到判决的内部操作。身体的感受本身没有变，但被对待它们的方式改变了：从“参与”转成了“证据”，从“沉浸”转成了“提取”。就是这一个转换，触发了前额叶-岛叶-前扣带回的执行评估网络（Critchley et al., 2004），交感基线微升，副交感深度修复的寂静许可轻轻撤回。

这一点在深度修复阶段尤为致命，因为内感受精度本身在深度副交感状态中会暂时性下降。这不是缺陷，而是一种功能性的保护：处于深度修复中的身体，需要关闭那些不必要的“自我监察通道”，把全部资源留给实打实的分子修复工作。就像一台正在运行大型系统更新的电脑，暂时锁住了键盘输入和网络连接。当一个人试图在这个状态下载定自己是否“修复好了”时，他不可避免地会：第一，读不准——因为感受信号通道此刻正在“营业外”的低带宽模式中运行；第二，干扰修复——因为他强制重启了感受通道，而动用了应该用以修复的能量和神经资源。

这就构成了一个残酷的内生悖论：人最需要裁定自己是否修复完毕的时刻，恰恰是裁定这个动作最不可靠、最具干扰性的时刻。深度修复状态——那种肌肉酸重、整个人昏昏欲睡、对运动完全提不起兴趣的状态——本身就是一个信号。但这个信号不是为了被解读成“你恢复了百分之七十”的量化判断而存在的。它是一个门槛标记：当身体处于这种状态时，正确的行动不是去评估它，而是允许它彻底跑完。在深度修复正在进行的时刻，人唯一需要做的决定，是不做决定——让寂静连着寂静，直到身体自己从修复模式中转出来。

此处也必须处理一个概念操作化的问题：在日常训练场景中，一个人起床后自然地感到腿酸，然后自然地想“今天可能该休息”——这算裁定还是参与性感受？答案取决于那个“想”的认知结构。如果那个“想”是在心里进行证据收集和判决推理——“有点酸但不算很酸，昨晚睡得还可以，应该恢复了七成，再休息半天就好”——那就是裁定。如果那个“想”只是一个简单的身体信号的形式化，没有开启一个“是否/完成了几成”的判决维度——“累了，休”——那就不构成完整的裁定。裁定性干扰针对的是前者：那种带有目标导向扫描和二元编码压力的认知判决。在实践中，二者的区分不总是清晰的，这恰恰是裁定性干扰试图打开的一个新的自我观察维度：不是去问“我判得对不对”，而是去观察“我刚才是不是又在判了”。

#### ■ 四、个案重读：当“听身体”仍是裁定时

为将以上理论锚定在一个具体现场，本节重读一个匿名化个案。<sup>[5]</sup>这个个案展示了一位训练者如何在使用表征系统与依赖内感受裁定之间反复切换，却始终未能触及裁定性干扰所指向的那个更深层的机制。需要说明的是，本节并不声称这一案例构成实证证据——它是一个思想例示，用于展示裁定性干扰如何在一次具体的训练生涯中被发生出来。它的功能不是“证明”裁定性干扰的存在，而是让读者在一段具体经验的纹理中看见裁定性干扰所指向的那个机制是怎样的。

L是一位35岁的男性业余马拉松跑者，2021年1月购入运动手表，开始严格执行设备提供的每日训练建议。在此之前，他已有三年业余跑龄，主要凭身体感觉决定训练节奏。以下呈现的是他十二个月追踪记录中的三个关键阶段。本节的阅读方式不是去重述另一个分析已经完成的工作——那个分析正确地揭示了表征系统如何高估修复完成度，以及表征-管理耦合如何制造隐性损伤的自我强化回路。本节要做的是：用裁定性干扰重新切入同一份材料，看一看在这些经验中有什么东西是上一层的分析框架——因其先验地预设了“修复需要被裁定”——所必然看不见的。

第一阶段：裁定如何关闭感知。L 在购表后前三个月，各项数据全面改善。他在年后访谈中回忆这一时期时说：“每天早上起来第一件事就是看‘身体电量’和‘恢复评分’。”注意时间顺序：第一件事。这意味着，在 L 尚未感受到自己身体的当天状态之前，一个外部判决已经下达了。这不是身体先感受、数据后验证，而是数据先定义、身体被安排去匹配那个定义。此时，L 的身体可能在某个早晨确实需要更长的停机——不是基于任何可被手表捕捉的异常数值，而是一个深度的、弥漫的、尚未抵达意识的修复诉求。但手表已经说了“准备就绪”。L 还没有来得及感受到那个诉求，就已经戴上了“绿标”这个标签。这一天里他的感受——一种隐隐的、说不清的倦怠——不再是一个独立的状态，而变成了一个需要被解释掉的“与判决不符的异常”：为什么数据说我好了，我却觉得累？

在这个时刻，裁定动作不仅做了一次判决，它还关闭了对新信息的接收通道。此后一整天，L 对自己身体的感受都先天地被一个“我已恢复”的判决框架所滤过：任何与判决相悖的信号，都必须先跨越自己对自己的不信任，才能被听见。更精确地说：裁定并不只是“判错了”的问题——即使那天的判决恰好正确，裁定这个动作也已经完成了它最关键的操作：它把一个“我此刻是什么状态”的开放问题，提前转换成了一个“已有答案”的关闭问题。此后，身体信号不再是回答这个问题的第一手证词，而只能以“与判决的一致性程度”的资格入场。裁定性干扰在这个阶段的形式不是“判错了”，而是“判得太早、太完整，以至于身体自己的声音还没来得及说出口，就已经被盖掉了”。

第二阶段：疲劳如何被裁定框架结构化为“外因”。2021 年 5 月到 9 月，L 的主观疲劳量表评分从平均 3-4 分升至 5-6 分，但每次他都加注“可能是睡眠不够好”“可能是工作压力”，从未归因于训练本身。这一现象在前一层分析中被解释为“表征误判叠加疲劳信号误读的锁定回路”——这没有错，但它仍然在用“误判”和“误读”的故障语言来理解 L 的经验。它预设了：L 之所以错误地将疲劳信号归因于外因，是因为他“读错了”。

但裁定性干扰视角再往下走一步：为什么 L 每次都需要给自己的疲劳找一个“训练之外”的原因？这不是一次性的认知偏差，而是一种结构化的解释实践——这个实践是由裁定框架本身塑造的。在整个训练-恢复-裁定的话语体系中，恢复已经被先验地设定为一个可以被客观裁定的命题，而这个命题在手表上已经获得了“已完成”的判决。既然判决已下，那么所有与之矛盾的身体信号，就自动降级为“干扰变量”——它们不构成对判决的挑战，只构成需要被额外解释、然后就能被搁置的噪音。L 在日志中对“睡眠不够好”和“工作压力”的反复援引，不是“误读”，而是这个已被裁定的认知框架所要求的唯一合法的解释实践：你不能说训练过多了，因为恢复已经被判为完成了；你只能说是别的什么东西没配合好。裁定并不是在这

一刻才开始犯错的——它从一开始就把所有可能挑战判决的身体信号，都先天地框定为了“需要在体系内被消化的杂质”。

第三阶段：反思重建——换了一个裁定者，但没有走出裁定。2021年12月，L在一次计划中的10公里配速跑中突然左膝外侧剧痛，后诊断为髂胫束综合征，须完全停跑至少三个月。在停跑恢复期间，L开始实践“无数据日”——每周至少一天不戴手表，仅凭身体感觉决定活动量。2022年6月恢复训练后，他将手表定位为“事后验证工具”：训练结束后才看数据，而非训练前让数据决定内容。在前一层分析中，这一实践转向被归为“限制表征管辖权”的成功示范。L自己的总结是：“以前我以为我在听身体，但我其实在听手表。现在如果身体说累了，手表说‘准备就绪’，我也会休。”

但注意，在这句被引为成功案例的总结中，L仍然在说：“身体说累了。”他换了裁决者——从手表换到了身体内感受——但他依然需要一个裁决者。他依然在每天早上完成一个裁定动作：身体说了什么？手表说了什么？冲突时听身体的。换言之，L从“错误裁定”走向了“更高级的裁定”，但他没有走出裁定。L在同一个访谈中的另一处表述更为清晰地揭示了这一点：“每天早上我还是会过一遍——腿还酸不酸、昨晚睡得怎么样、今天该不该跑——只是现在不用手表来帮我回答了，我自己答。”这里的“过一遍”和“答”，清晰地指向了裁定动作的认知结构。L的转向虽然改善了他的判决依据，但裁定这个动作本身——从感受到判决的那个内部操作——仍然原封未动。

这正是裁定性干扰理论预期的结果：一个从表征误判中走出来的人，会自然地走向“更好的裁定”，但更好的裁定本身不是裁定性干扰的解药。他在改进判决者，但没有触及那个再下一层的问题：是不是连“裁定”这个动作本身，都在结构性地清减修复所赖以发生的寂静？裁定性干扰的实践启示，恰恰在于推动训练者从这一个L尚未抵达的位置继续向前走：不是学会更好地裁定，而是学会在某些时刻放弃裁定——不是把判决权从手表移回身体，而是把判决这个动作本身拿掉。

## ■ 五、最严肃反击与理论边界

任何有理论雄心的概念都必须把自己暴露在最强有力的攻击下。本节处理裁定性干扰面临的最严肃的三项反击，并在此基础上明确理论的边界与可证伪条件。

### 5.1 精英运动员的反例与“谁在执行裁定”的修正

一个有力的反驳来自精英体育的现实：大量顶尖运动员生活在极高密度的生理监控中，并取得了长期、卓越的竞技表现与运动寿命。如果裁定性干扰是一个普遍且显著的生理机制，为什么这些运动员没有普遍出现慢性修复衰竭？



对这一反驳的回应应区分两个层面。第一，裁定性干扰是一个概率性的、边际性的效应，它不否定成功的存在，但它解释了那些“科学训练却陷入隐性疲惫”的残差现象——这些残差在业余训练者中尤为常见，因为他们缺乏精英运动员所具备的补偿缓冲：更长的睡眠、营养团队、物理治疗、心理支持，以及——最为关键的——赛季后的制度性离线期，其功能可能正是对赛季内长期裁定所累积的隐性损耗进行一种未被理论化的“静默偿还”。

第二，精英训练中的监控实践与本章所批评的“裁定”存在一个关键的操作差异，这个差异迫使本章对核心命题做出一个重要修正。在高水平运动队中，恢复数据的读取者通常是教练或运动科学家，而非运动员本人。运动员只需要执行训练计划，而不需要在每一个早晨启动一个“我恢复好了吗”的判决闭环。换言之，裁定被外包给了他人，运动员自身反而可以在一定程度上保持一种“不裁定”的身体投入状态。

这个观察揭示了裁定性干扰的一个关键边界条件：裁定动作的干扰效应，并非取决于“是否有人裁定”，而是取决于“裁定是否进入了被裁定者的认知回路”——即，那个正在运行修复程序的身体，是否在意识层面被强制启动了一个“我是否被恢复”的自我评估？如果是教练在赛前判断运动员的恢复状态并据此调整计划，而运动员本人并不在每一个早晨执行自我评估，那么裁定就发生在身体的外部——它在生理上不构成对修复寂静的侵扰。裁定性干扰的核心变量，因此不是“裁定是否存在”，而是“裁定-耦合强度”：裁定者与被裁定者的认知-生理耦合程度。当裁定被完全外包给他人时，耦合强度接近于零；当裁定由主体亲自执行时——无论是通过内感受还是通过读取外部数据——耦合强度最大。

这一修正并非事后缝补，而是对裁定性干扰机制的一次精确化：裁定性干扰作用的位点，不是“世界的某个角落有人对我做了一个判断”，而是“我的执行注意网络被调用去完成一个以我为对象的裁决任务”。这一定位使理论既保留了其核心洞察——裁定动作本身侵扰修复的寂静基底——又不至于得出“只要存在恢复评估，所有人都会被干扰”的荒谬推论。〔7〕

## 5.2 主观疲劳量表是否也在干扰修复？

一个来自运动科学实践的尖锐反问是：主观疲劳量表（RPE）是运动科学中最常用的自我评估工具，要求训练者在训练中持续评估自己“有多累”。如果“裁定本身有干扰”，那么RPE是不是也是一种裁定？如果是，那么基于裁定性干扰的逻辑，运动科学几乎所有涉及自我评估的实践工具都在干扰健康——这显然是一个过于极端的推论，也违背基本的运动科学常识。

这个反问切中了裁定性干扰概念必须澄清的一个界限。本章的回答是：裁定性干扰的作用具有时间阶段依赖性。RPE发生在训练过程中——即身体处于或正在进入应激状态的阶段。在此时，交感神经已经高度激活，执行注意网络本就处于被任务调用的状态，一个关于“我有多



累”的评估并不会在已经处于高警觉背景的身体中制造出额外的交感增量。裁定此时的功能是必要的监控：它帮助训练者和教练实时调整负荷，避免强度过高导致的急性损伤。在这一阶段，裁定的信息收益远大于其可能的干扰代价。

裁定性干扰针对的是完全不同的另一个阶段：训练后的深度修复窗口。在这一阶段，身体已经退出应激状态，副交感神经正在主导修复工序的运行，执行注意网络理应处于低激活状态——这是修复的“寂静许可”得以生效的生理条件。此时，一个“我恢复好了吗”的裁定，等于把一项本应处于关闭状态的认知-生理回路强制唤醒。裁定性干扰的特殊性恰恰在于它发生的时机：不是在身体本来就不安静的应激阶段，而是在身体好不容易安静下来、正要开始深度修复的那个脆弱窗口里。时间阶段的区分，为 RPE 的合法性划出了清晰的边界，也为裁定性干扰的作用域给出了严格界限。

### 5.3 这是否可证伪？

一个负责的理论概念必须给出自己可以被事实攻击的条款。裁定性干扰给出以下四个可裁决的证伪条件：

第一，寂静修复非优越性条件。如果一项随机对照实验能证明：在长期规律运动者中，被随机分配进“无裁定静默修复期”（即在一个训练微周期后，安排一段完全不使用任何外部或内感受裁定、仅允许随意活动的“静默窗口”）的受试者，在接下来一段时间的客观恢复指标（肌肉功能、内分泌节律、免疫标记物）上的表现，不优于甚至劣于那些被分配进“主动裁定期”（被要求每天做一次有意识的恢复状态评估并据此调整活动）的受试者，则裁定性干扰的核心预测被否定。

第二，裁定无害条件。如果能够设计出某种裁定方式，在严格控制其他变量的条件下，被反复应用于同一群受试者的多次修复周期中，并且长期追踪显示该裁定方式不伴随任何可测量的自主神经弹性下降、修复窗口缩短或隐性损伤累积——那么“裁定形式本身就构成干扰”这一普遍命题被限制甚至推翻。

第三，时间窗可逆性条件。第三节提出了一个推测性的机制：卫星细胞的增殖程序被交感神经信号半途打断后，那些已退出细胞周期的细胞不会在寂静恢复后重新进入分裂，导致永久性损失。如果基础研究发现：卫星细胞的增殖程序被交感神经信号半途打断后，在恢复寂静时可以从中断点精确地重新进入细胞周期，且最终产出的新细胞核数量与从未被打断的对照组无异——那么裁定性干扰的“累积不可逆性”关键预设被否决，其理论强度将被大幅削弱，退化为一种“暂时的效率降低”而非“结构性的修复亏损”。

第四，中性裁定的交感激活显著性与效应量条件。本章核心机制链的最上游环节是：即使是心态平和的、没有任何情绪负载的中性裁定，也会产生足够强度的交感微激活，进而产生对修复工序的截断效应。如果未来严格设计的实验证明：在控制焦虑和反刍的条件下，单纯的“我恢复好了吗？”的自我发问不产生任何可测量的交感神经激活（如心率变异性无明显变化、皮肤电导无显著升高），或者在体外实验中证明交感微激活在非药理学刺激水平下对卫星细胞的增殖分化没有可测量的抑制效应，则裁定性干扰的核心机制被否定。这一条证伪条件不要求实验涵盖整条因果链——它只需要在最上游或最下游的单个节点上打断链条，即可证明裁定性干扰的效应量在实际训练情境中可以忽略不计。

## 5.4 理论边界

裁定性干扰明确声明自己的适用范围：它解释的是修复这类多层级、非线性、自组织的生理过程在被主体亲自执行的裁定动作介入时，裁定形式本身——独立于裁定内容正确性——对过程在特定时间窗内造成的结构性侵扰。它不是一个通用的“决策有害论”。它不适用于紧急情况下的快速判断，不适用于应激状态下的实时监控，不适用于那些干扰代价远小于信息收益的日常轻判断，也不适用于裁定已被完全外包给第三方、不进入被裁定者认知回路的情境。

在所有跨领域应用中——如心理疗愈、组织复盘、教育生长——裁定性干扰提供的仅是一个启发性的骨架：如果一个过程是高度自组织和分布自治的，如果一个外部评估动作的介入被长期用于管理这个过程，那么那个评估动作的结构性干扰效应是否被系统地忽略了？它是否不仅没有帮助过程，反而在暗中消耗着那个过程得以发生的寂静条件？这些跨领域的应用，各自需要独立的实证检验，不在本章的论证范围内。

但运动修复本身——本章的论证主场——是一个已经被大量解剖学、生理学、神经内分泌学证据支撑的稳定领域。裁定性干扰在这个领域内所指向的那个老问题——“为什么科学训练有时反噬健康？”——被给了一个全新的、此前从未被认真对待的答案：不是因为训练科学不够精确，而是因为“精确”所仰赖的持续裁定，恰好摧毁了修复最需要的那块不被判定的静谧。运动的健康收益不来自“刺激-恢复”的精确管理，而来自那个管理无法触及也不应该触及的、修复自己悄悄完成自己的寂静间隙。现代训练学把那个间隙填满了裁定，然后困惑地问为什么恢复总是不够。

## ■ 六、实践转向：创建裁定性静默区

如果一个听进了本章论证的训练者问“那我到底该怎么做？”，答案不能是“扔掉手表”——那是用一种粗暴戒断代替另一种盲目依赖。也不能是“多冥想、多听身体”——那还是把裁定权从设备移到了内在，而裁定动作本身原封未动。裁定性干扰理论要求的实践转向，不

是更换裁定者——无论是从机器到身体，还是从身体到直觉——而是在修复过程中阶段性地拿掉裁定者。这在实践上可以被翻译为以下三个层次。

第一，区分裁定日与非裁定日。在当前的训练文化中，几乎每一天都开始于同一个动作：以某种方式评估自己的恢复状态。裁定性干扰理论暗示：不是每一天都需要这个评估。一个每周有三次高强度训练的训练者，可以在高强度训练后的 24-48 小时内完全不问自己“我今天恢复好了吗？”——既不看手表，也不在心里做自我评估。这 24-48 小时被指定为“裁定性静默区”。在这个区域里，行动的唯一依据不是“我被判定为恢复好了”，而是一个固定的时间缓冲和前一日的主动程度——不是裁定，而是记录（“昨天跑了强度课，今天自动轻活动”）。记录和裁定的关键区别在于：记录不涉及二元判决。记录只是陈述一个已完成的事实，然后按预先设定的缓冲规则行事——不向身体索取任何状态的判断。需要诚实承认的是，记录与裁定之间并非存在绝对的本体论鸿沟——任何“记录”都隐含着范畴化（“强度课”本身就是一种归类），这使得记录可以被视为一种温和的、减敏态的裁定。但记录以时间缓冲而非状态判决为行动依据的操作逻辑，使其在大多数日常训练场景中对修复寂静的侵扰程度显著低于直接的裁定。裁定性静默区重新划出了一片不被问“你好了吗”的身体存在空间。

第二，用门控式感受替代判决式裁定。即使在裁定性静默区之外——比如休息两天后，你确实需要决定今天的训练强度——也存在一种不上升为裁定的感知方式：门控式感受。它的操作逻辑是：不是问自己“我恢复好了吗”（那要求一个二元或灰度判决），而是设置一个单一的低强度动作——比如出门快走五分钟——然后允许身体在动作进行中自然发出一个极其清晰的信号：运动突然变得有吸引力了，所有关节都顺畅，呼吸不喘但有种想提速的冲动——门开了。或者：动作仍显得生涩，大脑还在不断给自己找理由“也许再稍微活动一下就好了”——门没开。这个过程中，没有一个判官站出来“恢复完成”，只有一个被触发的身体在持续报告它当下的状态。意识端的唯一义务，是保持接收这个信号的通道不被预先塞满。

第三，在制度化训练计划中嵌入无裁定微周期。这是针对系统性训练的——那些有教练、有团队、有赛季规划的运动员。可以在赛季中制度性地安排每隔六到八周为一整个“无裁定微周期”：在这个微周期里，不进行任何恢复状态评估，不使用任何追踪设备，训练强度由时间缓冲和固定的减量规则决定，而非由每日的恢复判断决定。这个微周期的目的不是为了“更好恢复”（虽然它几乎必定带来比平常更深的恢复），而是为了重新训练身体在那个不被问“你好了吗”的寂静里运作的能力和记忆。长期被裁定性干扰削弱的自主修复节律，需要在一个制度性保障的寂静里，重新学会不被找、不被量、不被判的状态。

这三个实践层次共享同一个核心操作逻辑：恢复的质量，不取决于裁定的准确性，而取决于获取寂静的时长与深度。此处必须做一个关键区分：裁定性静默区与“无数据日”在操作上

看似相似（都要求不戴设备），但在底层逻辑上是两个不同的东西。“无数据日”的目标是限制表征系统的管辖权，它反对的是“数字代理替代身体叙述”。但它在反对数字代理的同时，仍然默默地保留了那个被替代的对象——身体叙述本身作为裁定者。“无裁定静默区”则在更深的层面工作：它试图关闭的不是表征系统，而是裁判庭本身。它要求训练者在那段特定的时间里，不启动任何形式的“我恢复了吗”的判断程序，无论是数字的还是直觉的。这是一种更彻底的寂静——不是把旁听的法官赶出去，而是把法官的椅子撤掉。

## ■ 七、结论：寂静条件的维护

运动之所以能产生健康，从不是因为一个训练计划被管理得多科学、多严密，而是因为它创造了一次被允许彻底安顿的机会。修复——真正让裂痕愈合、让损伤化为韧性的那种修复——是一个必须在不被注视、不被催促、不被人划定为完成时限的寂静里完成的过程。运动科学用“刺激-恢复-适应”的模型精确地描述了刺激与适应之间的因果关系，但它从未认真对待这个模型内部所隐含的那个最关键的变量：修复赖以运行的那片寂静的条件。

在数字追踪和管理文化将运动彻底工程化的时代，这片寂静被一份又一份“你恢复好了吗”的问卷所殖民。表征系统在每一个早晨亮出绿灯，内感受裁定在每一个清晨在心里过一遍“身体说了什么”。二者都以为自己在守护修复，实际上它们守护的是一个被持续打断的、从未真正合过眼的修复的残片。裁定性干扰，命名的就是这种对寂静的侵蚀。它与裁定结果的对错无关——即使每一次裁定都精准无比，裁定这个动作本身，也已经从修复最需要的那块无人打扰的区域，切割下了一角。

这里必须谨慎地把握一个容易滑入本质论的陷阱。寂静条件不是身体的一种“本真状态”，不是等待被回归的、在前现代就已经在那里的自然属性。恰恰相反，寂静条件是在现代训练文化将身体彻底工程化之后，才被反向凸显为一种稀缺资源的新形态。在运动员没有每日训练日志、没有晨间静息心率趋势图、没有恢复评分的时代，“寂静”只是修复的默认基底——它不需要被命名，因为它从来不是被争夺的对象。只有在裁定性干扰成为训练日常的今天，“寂静”才被逼迫着显露为一种需要被自觉维护的条件。它不是在“回归”——它是在裁定性干扰的土壤中被发生出来的一种对抗性显露。那些被本章称为“裁定性静默区”的实践，不是让人回到某种失落的自然状态，而是在裁定的包围中反向开辟出一块不被裁定的飞地——这块飞地本身是一个新的创造，是对这个时代特定的身体管理逻辑的有意撤离。

本章使用的“寂静”一词，因此不应被误读为一种浪漫化的自然隐喻。它是一个严格的操作性概念：修复过程得以自组织地完成其不可逆时间窗工序所必需的一套生理-心理条件——副交感神经主导、执行注意网络低激活、免于被目标导向的评估所中断。这套条件不是先有的、

自明的，而是在身体从应激状态过渡到深度修复状态的过程中，被具体的生理节律（昼夜节律、能量代谢的下调、炎症信号的消退）发生出来的一个脆弱的呈现态。它需要被维护——不是作为“本真”的回归，而是作为现代训练文化的必要张力被维持住。裁定性干扰之所以危险，恰恰因为它在这些条件正在被身体努力发生出来的时候，用一个认知-语言的判决动作将它们打断。

这听起来极其消极——好像我们能做的，只是在某些时刻管住自己，不去打扰那个我们自己无比在乎的修复过程。但在一个被彻底工程化的文化中，“不打扰”本身就是一种有实质重量的行动。它需要你放下那枚手表，但更需要在那个放下之后的早上，不带一丝不安地不回头去看那些不在那里的数字。它需要你在身体弥漫着深度修复的低沉能量时，不去把它翻译成“我恢复了多少”的答案，而是把自己也沉进那片寂静里，让它完整地跑完，不给它设返场时限，也不在它尚未完成的寂静中催促它。健康不是管理出来的，但它也不是回归某种前现代的本真状态就能自动抵达的——它需要维护一套复杂的、在应激与修复之间切换的生理条件，而裁定的过度介入恰恰破坏了这套条件的发生。

如果未来的严格证据能够令人信服地证明：裁定动作——无论由外部表征还是由内感受执行——对修复过程不产生独立于裁定内容的、不可逆的结构性干扰，那么本章的核心命题即被证伪。但在此之前，在今天的运动文化中，裁定性干扰作为一个要被严肃对待的假说，将促使我们重新思考：在修复这件至为私密、至为精密的事上，也许我们能为自己做的最好的一件事，不是更聪明，而是在某些时刻，把聪明暂时地放下来。

## ■ 注释

\[1\] 本章所使用的“裁定”概念，特指一种带有目标导向扫描和二元编码压力的认知操作：它以“是/否”“完成/未完成”“恢复了百分之多少”等形式，对修复过程做出闭合性判断。这种裁定在现象学层次上不同于随意的、参与性的身体觉知（如“有点累”），其关键在于它调用了一个外部于修复过程本身的评价性框架。参见后文第三节对“参与性感受”与“裁定性评估”的区分。

\[2\] 心身医学文献中所讨论的“过度警觉”和“健康焦虑”，通常以对疾病的灾难化解释和反复的安全行为为特征，其病理学基础涉及前额叶-杏仁核等恐惧环路。本章所批评的“裁定”行为主要发生在非临床人群的日常健康管理实践中，不必然包含恐惧或焦虑情绪，但其可能通过与前额叶执行网络和交感神经系统的耦合产生类似的生理后果。这个交叠是裁定性干扰具有理论张力的关键原因之一。



\[3\] 本章借用“身体对象化”这一现象学直觉，仅限于身体感知模式从前反思运作向对象化注视的切换所产生的存在论效应。这并不表示本章完全接受或采用了胡塞尔、梅洛-庞蒂或萨特任何一位现象学家的整体哲学体系。认知科学后续对 body schema 与 body image 这对概念的精细化批判，也未在本章中展开，本章仅在粗略的结构类比层面使用这一区分。

\[4\] 失眠的认知行为疗法（CBT-I）中的刺激控制疗法要求患者仅将床铺用于睡眠，避免在床上从事与睡眠无关的活动（包括“担忧失眠”）。这一操作在行为层面为睡眠创建了一个“不评估”的环境，可被视为本章“裁定性静默区”的临床微缩版。但应当指出，CBT-I 的机制解释主要基于经典条件反射和刺激控制理论，而本章所讨论的“裁定性干扰”则强调认知评估动作本身通过前额叶-交感神经通路对非条件化的生理过程产生的直接干扰。

\[5\] 【材料说明】本章第四章所使用的个案素材，来源于对一名业余跑者为期十二个月的训练日志、设备记录和访谈材料的重整理。为保护隐私并聚焦理论机制的展示，个案中的时间节点、训练细节与身体感受描述已做匿名化、典型化和必要的简化处理。该个案不构成实证研究的原始证据，而是用于展示裁定性干扰理论逻辑的思想例示。个案访谈材料的原始分析已在前期研究中完成；本章在此仅用一个与前期分析不同的理论框架对该同一批材料进行再阅读，两种阅读方式并不互相否定，而是指向不同层次的问题。

\[6\] 关于交感神经通过  $\beta$ -肾上腺素受体通路抑制肌肉卫星细胞分化融合的证据，目前主要来源于动物模型和体外细胞实验（如 Chen et al., 2010）。人体运动修复情境中该通路被裁定性认知操作激活后的实际效应量，尚未得到直接实验研究的测量。本章在这一环节的论述是将已知的细胞生理学通路与神经心理学情境进行理论结合，属于有待检验的推测模型，而非已确立的实证事实。第三节已在正文中就这一推测性质做了诚实声明。

\[7\] 本节关于精英运动员实践的讨论，系基于公开训练学文献、体育媒体报道和作者对该领域实践通例的了解所做的概览性分析，并非来自系统的问卷调查或田野调查。其目的在于为裁定性干扰可能存在的边界条件提供一种解释逻辑，而非对精英训练体系的客观评价。



# 合章 同一个动作的两端

## 一、两章互不知情，却切在同一刀口上

第八章与第九章，隔着两门几乎没有共同文献的学科。少敏写的是生成式人工智能治理下的人工复核，材料是欧盟人工智能法、工单系统、审批链与日志；胡敏写的是有不可逆时间窗的自组织修复过程，材料是运动损伤后的恢复、失眠的认知行为疗法、康复期的自我评估。两位作者互不相识，也没有引用对方。

但两章处理的是同一个动作：评估。

少敏的判断是：复核这个动作在提高错误识别率的同时，把处置方案继续送进工单、接口、预算、合同、排期与团队预期——于是复核的中间状态被下游读成了"继续投入的许可"。组织因此同时得到更多的控制证据与更少的现实控制：越来越能证明有人看过、知道过、签署过，越来越难把人的异议转成另一条可执行路径。他的诊断是一句硬话：失效发生在路径侧，不在判据侧。

胡敏的判断是：对于有不可逆时间窗的自组织过程，裁定这个动作本身就构成干扰，与判断结果的对错无关。他把裁判席整个撤掉，而不是换一个更好的裁判。他特别指出那个容易被误认为已经解决问题的做法——"无数据日"只是换了裁判，没有走出裁判：它把旁听的法官赶出去了，却留着法官的椅子。

一个说评估让组织失去改道的路，一个说评估让过程失去修复的力。同一个动作，两端。

## 二、把两端并成一张表

两章合起来能写出一个 2×2，而任何一章单独都只有一半。横轴是评估的对象——被评估的是一个可以停下来等待判决的事项，还是一个有不可逆时间窗、停不下来的过程；纵轴是评估的产物——它产生的是一条可执行的替代路径，还是一份留痕。

|        | 产物=可执行的替代路径                              | 产物=留痕                        |
|--------|------------------------------------------|------------------------------|
| 对象可暂停  | 有效监督。异议到来时组织能改道，反事实承载力为正。这是少敏那一章里唯一的健康格。 | 控制证据积累而现实控制下降。少敏的主诊断落在这一格。   |
| 对象不可暂停 | 罕见且昂贵：需要在过程之外预置替代通道，且评估者不得进入             | 双重损失。留痕在增加，被评估的过程同时被减损——胡敏的裁 |

---

过程内部。胡敏的"裁定性静默 定性干扰落在这一格。  
区"是这一格的可操作形式。

---

这张表的价值在右下角。少敏证明了留痕会替代改道，胡敏证明了在不可暂停的对象上，评估还要额外收一笔费用：它干扰了那个正在自行组织的过程本身。两笔损失叠加，而任何一本账都只记一笔。组织的合规账记下了留痕的增加（记为改善），个人的身体或修复过程承担了被干扰的部分（不入任何账）。

### ■ 三、分界条件由胡敏自己给出

这一对不能就这样合并——两章说的若真是同一件事，就不是咬合而是重复。分开它们的判据，恰好由胡敏那一章自己提供：参与性感受与裁定性评估的区分。前者是身体在过程之内的持续感知，它不产生判决；后者要求主体暂时站到过程之外，对过程发出一个是否合格的判断。他给出的旁证很硬：失眠的认知行为疗法早就禁止患者躺在床上评估自己是否还在失眠——这是一个已被临床验证的裁定性静默区，只是从来没有人把它读成一条本体论线索。

有了这个区分，两章的关系就清楚了：少敏那一章处理的全部是裁定性评估（复核、签署、批准都必须站在过程之外），因此它是这个  $2 \times 2$  的上半行；胡敏那一章的贡献在于指出，当对象是不可暂停的过程时，即使是"善意的、准确的、及时的"裁定性评估也在收费——因此他占据下半行。而参与性感受这一列在两章里都没有被计价：它既不产生留痕，也不构成改道理由，在任何一个既有系统里都不留下痕迹。

于是可以写出一条两章合起来才有的推论：

在不可暂停的过程上，提高评估密度会同时降低改道概率与修复能力；而现行的合规读数只记录评估密度，因此这两项损失都表现为读数改善。

判错：取康复医学、产品安全或临床路径这类同时具备留痕制度与不可逆时间窗的领域，比较评估频次不同的两组，在控制初始严重度后，考察（一）方案改道率与（二）过程结局。若高频评估组两项都不更差，本推论错。

### ■ 四、这一格与全书其余各章的接线

把这张表接回枢纽章的三个量：上半行右格是  $\Delta$  的纯形态（信号有、路径无）；下半行右格是  $\Delta$  与  $\beta$  的合并损失（路径无，且过程本身被扣掉一块）。高于涵那一章从第三个方向切进同一格——他写的是"成为好人"变成一道工序之后，工序不再加工善，而是持续产出"未善"，并且它不需要你不动脑子地盖章，它需要你带着全部信念去论证那个章为什么非盖不可。这句话

补上了 2×2 里没有的一维：被评估者的配合。留痕之所以能替代改道，不是因为人在敷衍，而是因为人在认真。

这也解释了为什么本书不把出路写成"少一点评估"。评估密度是可以调的旋钮，但十一章里没有任何一章支持"回到没有评估的状态"这个方案：陈晓艳明确拒绝把风险责任推回个人，阳涌明确拒绝把"脱离人工智能独立完成"当作唯一能力标准。出路必须落在别的地方——落在第十章与第十一章，也就是那个仍然必须由自己付款的位置上。

## 第四编

# 还剩下的那一注

• • •

两章给出路。

秦莉写"存押"：在无保底、不可逆、自己承担代价的条件下把一部分存在押进一个动作，做完之后你不再是同一个人，不是因为多了一段经历，而是因为少了一块自己。她同时指出存押的生态条件——时间冗余、不被强制证明的缝隙、豁免于风险管理铁律的飞地——正在效率优化的名义下被系统性侵蚀。

胡志英写"蚀先于生"：形态的耗散比生成更原始，生不是第一动作，它是蚀的竞争火焰中偶然没被烧着的那一小撮。

两章合起来给出的是本书唯一的正面主张：存在不是一次获得，是持续付款。而 AI 时代对普通人最深的一件事，不是拿走了他的能力，是替他付款了。

## 第十章 押进去就拿不回来的那一部分

本章原题《存押：艺术作为代偿时代中存在的不可逆赌注》，作者 秦莉，原文见 <https://sdeuniverses.com/students/qin-li/existential-stake/>

### 摘要

当代艺术中真正让我们感到“重”的作品，其力量来源不能被“形式精良”“观念深刻”“表达有力”等属性概念所穷尽。本章提出，这种力量源于一种特定的存在论动作——存押（existential stake）：当一个人在无保底、不可逆、自己承担代价的条件下，把自己的一部分存在押进一个动作里，该动作便成为其存在历史中不可删除的节点，而作品作为这一动作的痕迹，在存在论上“被押过”。

关键词：存押；存在折旧；代偿；时间冗余；风险个人化；艺术生态

这一判断的建立，必须首先撬开一个更隐蔽的前提：当代制度设计默认“存在不折旧”——人的存在可以完好无损地穿越任何动作，做完之后还是同一个人。这个预设的制度产物是“风险个人化”：当制度将“不知道后果”从一种可以主动进入的存在状态，重新定义为一种需要管理的风险缺口，存押的生态条件——时间冗余、不被强制证明的缝隙、豁免于风险管理铁律的飞地——便在效率优化的名义下被系统性侵蚀。艺术之所以成为存押仅存的主要飞地，不是因为艺术“无用”，而是因为艺术系统在功能分化社会中的“弱代码”特征——它没有一套强制性的二代码可以当即判定一个动作“对”还是“不对”——这使得创作者可以在更长时间内保持悬置状态，而不被外部制度强制叫停。

本章以布尔乔亚、培根、伊恩·程三个案例锚定存押的生成机制结构，严格区分存押与海德格尔“本真决断”、阿伦特“行动”等理论近邻，论证存押切出了一个既有概念无法覆盖的辨别面：一个动作如何反过来折旧做出它的人。这一辨别面不在阿伦特关注的“公共显现”维度上，而在一个此前未被单独命名的维度上——私人存在如何被自己的不可逆动作重铸。本章进而将此判断推衍至教育、养育与学术领域，论证存押的稀缺不是艺术的局部问题，而是当代生活中存在论轻量化的总病灶的一个症状。

当代艺术中真正击中我们的那些作品，其力量来自何处？这不是一个趣味问题。当你站在路易斯·布尔乔亚的蜘蛛雕塑前，当你面对弗朗西斯·培根那具被反复刮掉重来、最终仿佛自己开始呼吸的扭曲人体，当你目睹伊恩·程那些在公共空间中不间断运行、无法预判下一秒会发生

什么的人工生命——你不是在欣赏一个完成了的艺术品。你是被一个事实击中了：在这些作品的发生史中，有人曾经把自己押了进去，并且再也拿不回来。

“押进去”是什么意思？它不是“选择”——选择是在已知选项中做出比较。它是一个人的小块存在——那一瞬的专注、那一下呼吸的停顿、那个“就是这里”的不可撤回的判断——被不可逆地转移进一个外部动作。那个动作因此不再是纯粹的外部事件。它变成了那个人的存在历史中一个永远不可删除的节点。它不同于“决定”：决定可以之后被推翻。不同于“风险承担”：风险可以用保险对冲，可以把损失量化为账簿上的数字。不同于日常语言中“赌上自己”的修辞——那种修辞留有“赌输了再赌一次”的后路。存押不是修辞。它是：你的一部分存在，永远地留在了那个动作里。做完这个动作之后，你不再是同一个人——不是因为你多了一段经历，而是因为你少了一块自己。那一块已经转移到了作品里，不再属于你。

这就是本章要论证的核心判断：那些真正让我们感到“重”的作品，它们在存在论上有一个不可代偿的来源——它们被押过。不是“被投入了大量劳动”，不是“表达了深刻的思想”，不是“体现了高超的技艺”——这些都是作品的属性。属性可以被模仿、被复制、被代偿。存押的痕迹不可代偿，因为它不是作品的属性。它是作品的发生史在作品身上留下的不可逆痕迹。

反过来说：作品可以承载任何可以被说出的意义、可以被测量的效果、可以被拍卖的价值，但如果没有任何人在它的发生过程中押进过自己的一部分存在，它在存在论上就是轻的。这不是审美判断——不是“这件作品不好”。而是：这件作品在存在论上没有见证任何人的存在折旧。它没有在人存在历史中留下任何不可删除的痕迹——包括那个所谓“创作”了它的人的存在历史。它只是被生产出来了。

这个看法一旦立住，当代艺术生态的危机就必须被重新诊断。通常的说法是：“好的作品变少了”“市场腐蚀了创作的纯粹性”“批评体系失去了判断力”。这些说法都停留在症状描述层。存押论给出的诊断更根本，也更令人不适：当代制度设置正在系统性地、且在沉默中，消灭存押的生态条件。不是艺术家变懦弱了。不是艺术体制变腐败了。是支撑“把自己押进去”这件事得以发生的那块制度地面，正在被一块一块地铲平——而且是以“让艺术更有效率、更可问责、更可管理”的名义。

这块制度地面的名字叫时间冗余。不是“有很多空余时间”的意思。时间冗余是指：一个人可以在一个较长的时间窗内，持续地将自己悬置在一个动作的“不知道后果”之中，而不用在每一个时间节点上向外部证明自己的投入是“有产出”的、是“对”的。存押需要这种冗余，因为存押的本质就是有组织的悬置：你不知道你还需要画多少张废画，才能遇到那一次真正的存押；你不知道你还需要在那个“就是这里”的临界点到来之前，等多久。只有时间冗余



为这种悬置提供了合法居留权——你不必每三个月有一次像样的展览，不必在基金申请表的“预期成果”栏里编造一个你并不真正知道的结果，不必在每个学期末交出一份教学评估看得过去的分数。

但这种合法性，正在被系统性地撤销。撤销它的逻辑，深埋在现代社会的根基层面——深到它已经变成一种我们不再提问的“天性”。

为了逮住这个最深层的逻辑，下面走一条严格的生成机制追问。我们不先问“艺术怎么了”，而是先问：在什么制度条件的变迁中，“押进去”这一动作本身，从一个可以被公开坚持的存在姿态，变成了需要解释、需要辩护、且越来越难以被合法维持的例外？

## ■ 一、先验撤销：风险个人化如何消灭了“悬置的合法权利”

### ■ 1.1 五连问：逮住那个没有人质疑的东西

这一节是整个论证的钥匙。要找到那个被默认为真、从未被质疑的根性假设，必须用五连问一层层往下追。

第一问：这个领域里最痛的矛盾是什么？

是这样一个张力：一方面，当代艺术在话语层面获得了前所未有的自由——“什么都行”，任何材料、任何媒介、任何操作都可以被接纳为艺术。另一方面，在这种自由中产生出的真正有存在重量的作品，却越来越稀缺。不是“好作品变少了”——“好”是趣味判断，标准一直在变。而是“重的作品”的生态条件在系统性收缩。两件事同时发生：自由在扩展，重量在蒸发。这不是一个可以被“趣味变了”“标准多元化了”搪塞过去的次要矛盾。它逼问每一个做艺术的人：如果你可以做任何事，为什么你做的事——在被你做过之后——没有留下你应该留在里面的东西？

第二问：旧的解释路径为什么反复失效？

既有的解释——无论是审美超越论（艺术抵抗工具理性）、批判干预论（艺术挑战体制）、还是观念生产论（艺术生产新认知）——都把“重”归因于作品的某种属性或效果：形式的自律、批判的锋芒、观念的深度。但它们无法回答一个共同的问题：为什么在两件属性相当、效果相似的作品之间，一件让我们感到不可替代的“重”，另一件却没有？更致命的是：为什么在一件属性上有缺陷、观念上含混、技法上生涩的作品面前，我们仍然可能被一种说不清的重量击中——而另一件在各方面都“合格”、甚至“出色”的作品，却像一枚抛过光的硬币一样轻？这些解释没有在同一块地方反复摔跤。它们在用“作品是什么”来回答一个实际上是“作品是怎么被做出来的”的问题。

第三问：这些失效背后，大家默认为真、从来没有质疑过的那个先验假设是什么？

是这样一个假设：一件作品如果具有某种特定的属性组合（形式精良、观念深刻、批判有力），它就可以独立于其生成过程，自行携带意义的重量。这个假设让所有既有解释都把精力花在“如何识别出这些属性”上，而从来没有追问：这些属性，在什么样的生成条件下，才可能不是被代偿出来的？一件在属性上完全“合格”的作品，是否可能仍然没有重量？如果答案是肯定的——而我们必须给出这个肯定——那么整个以属性为核心的批评传统的地基，就在这一问之下裂开了一道缝。

第四问：如果撤销那个先验，会打开什么追问？

如果一件作品的意义重量不能由其属性组合独立承载——如果一件在属性上完全“合格”的作品仍然可能是轻的——那么我们就被迫去追问：那个让作品“重”起来的东西，究竟是什么？它不在作品里。作品是那个东西的发生完成后，留在物质或代码中的痕迹。那个东西本身，在作品的发生史里。它不能被直接看到，但它可以被追踪到——如果你愿意换一种看的方式：不再问“这件作品是什么”，而是追问“做这件作品的那个人，在做它的时候，对自己做了什么”。

第五问：哪一个最小的、具体的反例，能一下子压垮旧框架？

AI生成的图像。在今天，一个没有任何绘画经验的人，可以用某个图像生成模型在五分钟内生成一幅在视觉效果上堪比超现实主义的图像。它可以很美、很有冲击力、甚至引发观念讨论。但它没有任何存在重量。不是因为它“不如手绘的好”——这个判断仍然在属性层面打转。而是因为它被生成的过程中，没有人押进任何东西。它的每一步都是可撤回的：提示词可以改，随机种子可以换，不满意就重新生成。旧框架无法区分它和一幅真正被押过的画——因为在属性层面，它们可能毫无分别。

但这个先验还不是最底层的。再追一刀：这个先验本身，又预设了什么更不容置疑的东西？

它预设了：意义可以以“已完成的属性组合”的形式在人与人之间传递，而不必携带其生成过程中那个“押注”的存在论痕迹。更根本地说，它预设了：意义的传递和存储，不需要依赖那个意义在发生时曾经切入了某个特定存在的不可逆历史。它把意义当作一种可以脱离其发生史而独立流通的通货——就像货币可以脱离铸造它的那座造币厂而独自承担交换功能。

而这个预设又预设了什么？

它预设了：人的“存在”是一种稳定的、可以被意义“填充”的容器，而不是一件在每一次不可逆押注中被不断重塑的东西。换句话说，它预设了：人是先有了一个完整的存在，然后

用这个存在去“做”一些事、去赋予一些事以意义；而不是——人在做的事、尤其是那些不可逆的事，反过来不断在重塑这存在本身是什么。

再追一刀：这个“存在先于押注”的预设，又建立在什么更深的假设上？

它建立在这样一个假设上：人的存在可以“完好无损”地穿越那些不可逆的动作——做完了，人是同一个人，只是多了一段经历。这个假设让“可撤回”成为可能：既然我做完一件事之后还是同一个我，那我做错了就可以收回，换一个选择再做。存在本身不折旧、不损耗、不磨损。

这恰恰是当代制度设计最深层的预设——它深到不再显得像是一个“预设”，而像是显而易见的常识：存在是不折旧的。离婚之后，你还是你——你可以在法律上恢复单身，重新进入婚恋市场。辞职之后，你还是你——你带着可转移的技能和经验，换到下一份工作，简历上多一行而已。移民之后，你还是你——你带着你的身份、记忆和护照，只是在另一个坐标上继续生活。经历会改变你的处境，但不会改变“你是你”这个存在论事实——这就是那个我们不假思索地依赖着的预设。制度有责任保护你免受不可逆损伤，制度应该让一切选择可撤回——因为撤回不伤及存在本身。

一旦这个预设被说出来，它的荒诞性就暴露了。梵高画完《星月夜》之后，他是同一个人吗？培根画完那些扭曲的人体之后，他是同一个人吗？布尔乔亚做完那些蜘蛛之后，她是同一个人吗？不是。他们被自己做出的那些不可逆的动作改变过了。不是“他们多了一段经历”意义上的改变——那个意义太轻了，轻到它变成了一句废话：任何动作都会给人“增加一段经历”。这里说的改变，是存在论层面的：他们的存在本身被那些动作重塑了。那些动作不是他们“使用”了自己的存在去做的事——像一个人用他的钱去买一件东西，花完钱他还是他，只是少了些钱，多了件东西。存押不是这样。那些动作本身就是他们存在的展开方式。做完之后，他们已经不是之前的自己了。他们的一部分存在，已经永久地转移到了那些动作里，不再属于他们自己。站在《星月夜》面前，你面对的不是一幅画。你面对的是梵高的一部分存在，永远地留在了那些笔触的漩涡里。

这就是那个根性假设：存在不折旧。这个假设不是艺术领域独有的，它是整个现代制度设计的默认识置。保险制度建立在这个假设上：你遭遇损失，保险公司赔付你，然后你“恢复”到损失前的状态——存在没有折旧。教育制度建立在这个假设上：你获得知识和技能，你“升级”了——存在只增不减。职业生涯设计建立在这个假设上：你可以换工作、转行、创业——每一次转换，你带着可转移的技能和经验进入新的岗位，存在只累积，不折旧。但在某些特定的动作里，存在是折旧的。而且正是因为这种折旧，那个动作才有了重量。你押进了一部分自己，那部分自己不再属于你了——它留在了那个动作里，变成了那个动作的不可剥离的

一部分。这就是为什么一幅梵高的画有重量：不是因为梵高“画得好”——这是属性判断，永远可以被另一个“画得好”的人挑战——而是因为梵高的一部分存在，永久地留在了那些笔触里。你站在画前感受到的不是审美的愉悦，不是观念的启发——而是一个人的存在折旧的痕迹。

## ■ 1.2 风险个人化：制度如何让“存在不折旧”变成铁律

“存在不折旧”不是一个哲学错误，等待被哲学家纠正。它是一套制度安排的工程产物。追问它的发生史，必须追踪一个更具体的现代进程：风险的个人化。

在前现代社会，风险是共同体的负担。饥荒、瘟疫、战争——这些不可逆的灾难，个体无法独自承担，必须由家族、村社、行会等共同体来分摊。在这种风险结构中，“把自己押进去”不是一个需要被特别保护的行动——因为无论如何，每个人都已经被押进了共同体的命运里。你的收成跟天有关，你的婚姻跟家族有关，你的职业跟你父亲有关。你的存在，从一开始就被嵌在一个不可撤回的共同体结构中。那种“我赌上一切去……”的叙事，不是前现代的主题——因为在共同体结构中，你的一切从来不属于你自己去“赌”。

现代社会的一项核心工程，就是把风险从共同体的负担变成个人的管理对象。保险、社会保障、劳动法、消费者保护——这些制度的共同逻辑是：让个体能够“计算”风险、转移风险、对冲风险。一个现代公民的理想形象，就是一个理性的风险管理者：他知道自己面临哪些风险，他购买保险来覆盖不可预见的损失，他通过教育投资来降低失业风险，他通过多元化投资来对冲金融风险。

这个进程的巨大历史功绩不能被低估。它把个体从共同体的束缚中解放出来，让人不必因为出生在一个穷村子就注定一生穷困，不必因为一次错误的婚姻就终身困于其中——这正是安东尼·吉登斯在《现代性与自我认同》（Giddens, 1991）中反复论述的“脱嵌”过程：现代制度的抽象系统将个体从传统的地方性约束中解放出来，赋予其前所未有的选择自由。个人被从地理的、血缘的、世代的既定命运中连根拔起，扔进一个可以重新书写人生剧本的开放空间。自由的物质基础，很大程度上就是这种风险管理的制度化。事实上，吉登斯将这种“风险的殖民化”——即用专家系统和制度安排将未来中的不确定性尽可能地纳入当下管理——视为晚期现代性最核心的动力之一（Giddens, 1991）。

但它的附带条款是：当风险管理成为一种制度性强制——当你必须成为一个好的风险管理者，否则就被视为不负责任——那么“主动进入一个无法管理、无法对冲、无法计算的不可逆境况”，就变成了一种制度上的不合理行为。不是法律禁止它，而是整个制度的预设让它看起来是愚蠢的、不理性的、应该避免的。在社会学家乌尔里希·贝克的经典分析中，这正是“风险

社会”的深层悖论：晚期现代性不仅产生了一整套应对风险的技术，更关键的在于，它悄悄地将“失败”重新定义为个人在风险管理上的失职。在《风险社会：迈向一种新的现代性》（Beck, 1986）中，贝克指出，当代制度有一种强大的倾向，即将系统性的矛盾与危机转译为“个人的风险”——失业是个人的技能投资不足，疾病是个人的生活方式选择不当，情感破裂是个人的沟通技巧欠缺。这种转译不只是认知上的，更是制度上的：税收、法律、保险、教育评估，共同构成了一个庞大的“个人责任制”装置（Beck, 1986）。

这个装置的社会效果，不是“每个人现在都可以自己管理风险了”——而是“主动悬置在不知道后果中”这一存在姿态本身，被去合法化了。它从一种需要勇气的存在方式，变成了一种需要解释的不理性行为。你当然可以在制度的缝隙中偷偷保留它——靠父母的积蓄啃老几年，靠伴侣的稳定收入支撑你的实验期，靠拒绝参展来保护一段不受打扰的沉默。但前提是：你必须自己为这种悬置付出代价，而不是指望制度来为它提供合法性担保。你无法再把它作为一种值得追求的存在方式，公开地向他人推荐——因为一旦推荐，你就必须面对那个来自制度预设的质问：“万一失败了呢？你考虑过风险吗？”

这就是贝克与吉登斯所描述的那个制度转型——尽管二人在准确诊断上存在分歧——对我们的问题意味着什么。它意味着：存押的危机，本质上不是“人们失去了勇气”，而是“悬置在不知道后果中”的制度合法性被系统性撤销了。在风险个人化的制度环境下，存押从一种存在姿态被重新定义为一种风险管理失误。它不是被禁止的——在一个以自由为核心理念的社会中，没有任何制度会公然禁止你“押上自己”——但它被制度性地边缘化了，被排除在公共认可的理性行为清单之外。如果你成功了——像梵高死后那样——你会被追认为天才，你当初的“不理性”会被重新叙事为“超前”。如果你失败了——这是存押的概率结局——你只是一个没有做好风险管理的、不负责任的个体。而正是这一点，制造了那些被制度过滤掉的无名存押：那些画被家人收在阁楼里直到腐烂、那些手稿没有被任何人读过、那些教学实验只在课堂发生一次就被遗忘——他们的问题不是他们的作品不够好。他们的问题是：他们押了，但无人能够——在制度的默认语法中——识别出那个“押过”的事实。

## ■ 二、存押：生成机制定义与不可还原校验

### ■ 2.1 从“诊断标签”到“发生条件”：为何需要一个新概念

当一个新概念被提出，它面临的第一个质疑不是“它是否正确”，而是“它是否必要”——是否可以用既有概念的组合作来无损重述它所要说的。因此，在给出存押的正式定义之前，有必要先厘清：我们究竟在被什么追问驱策着，才需要单独命名这个动作类型。

追问的起点是：如果我们追问的不是作品的效果，而是作品的发生条件——不是在属性的层面描述“这件作品为什么让我们感到重”，而是在生态的层面追问“在什么条件下，一个人会把自己的一部分存在不可逆地押进一个动作里”——那么我们需要命名的是什么？我们需要命名的不是“决断”这个意志动作本身。决断是意志的爆发——我可以决断去做一件事，做完之后我仍然是那个做出决断的我。决断是“我”用来改变“世界”的工具。它预设了一个相对于决断行为保持同一的主体：我决断，我行动，我承担后果，而我——作为那个做出决断并承担后果的人——在决断前后始终是同一个我。在这个意义上，决断是“存在不折旧”的典型动作：它属于一个意志主体，这个主体在决断中“使用”了自己的自由，但并未因此改变自身的存在质地。

存押不同。它不是用存在去达成一个外部目标。它是把自己的一部分存在转移进一个外部动作里，从而让那个动作永久地携带了本属于那个人的存在论密度。那个转移一旦完成，这个人就不再是同一个人——不是因为“他经历了一件重要的事”，而是因为他的一部分存在不再属于他了。它留在了那个动作里。如果你问他“你那时候为什么那样做”，他往往无法给出一个合理的解释——不是因为他不善于表达，而是因为在存押发生的那个时刻，他并不是在“做选择”。他是在被一个他自己也说不清的东西推着走——而那个东西，正是他自己在那一刻的存在状态。他不是在“做决定”——他是在“过自己的存在”。这两者之间的区别，是存押必须被单独命名的根本原因。

## ■ 2.2 存押的生成机制定义

由此，给出存押的生成机制定义：当一个人在无外部保底、不可逆、且自己独自承担全部代价的条件下，把自己的一部分存在押进一个动作里，这个动作就成为其存在历史中不可删除的节点。做出这个动作之后，这个人不再是同一个人——不是因为他多了一段经历，而是因为他的一部分存在已经永久地转移到了这个动作里，不再属于他自己。“押”在这里是字面意义上的转移，而非隐喻：一笔存在论财富被从那个人的存在账户中划走，存入了那个动作的账户，永不可撤销。

存押的构成条件必须从发生条件来界定，而非从结果反推：

第一，无外部保底。做出存押的人，必须在那个时刻真的不知道后果。他不能有一个“失败了还可以这样”的备选方案，不能有一个“即使最坏情况发生也不会伤筋动骨”的安全垫，不能有一个“即使我错了，别人也会理解”的社会支持系统在背后等着接住他。他当然可以事后证明是错的——他大概率是错的——但在做的那个时刻，不能有任何东西来担保他。这不是浪漫主义对“孤绝”的崇拜，这是一个发生条件的结构性要求：如果有保底，他押进的就只是一



个“可以被损失掉的部分”，而不是他存在的一部分。保险的意义正在于：它让你可以押一笔钱而不押自己——你损失的是保单覆盖范围内的金额，不是你的存在。存押没有保险。存押者的处境不是“我可能会失败”，而是“我不知道我押进去的这一块自己，还能不能拿回来”。

第二，不可逆。这个动作一旦做出，就无法被撤销，无法被“重新来一遍”。不是制度上的不可逆——制度可以允许撤回，离婚可以让你恢复单身，辞职可以让你重新求职。而是存在论上的不可逆：押进去的那部分存在，已经永久地改变了押注者的存在质地。即使他后来“放弃”了这个方向，那个曾经押进去的事实，仍然是他的存在历史中无法删除的一个节点。他带着这个节点继续活着——这个节点在定义他，不管他愿不愿意。这种不可逆不是物理意义上的。一个画家可以铲掉画布上的颜料——在物理上，那个笔触消失了。但他在那个笔触里押进去的那一瞬存在——那一刻的专注、那一下呼吸、那个“就是这里”的判断——已经发生过了。这个“发生过”，无法被铲掉。他作为画家，在那一笔之后，已经是一个曾经押过那一笔的画家了。铲掉颜料，只是让他变成了一个曾经押过那一笔、然后又铲掉了它的画家。铲不掉的，是这个。

第三，自己承受代价。存押的代价不能由别人来付，不能由体制来缓冲，不能由时间来稀释。做出存押的人必须是他自己后果的唯一承担者——不是因为他想当英雄，而是因为如果代价可以被分摊，那押进去的就不是“他的一部分存在”，而是“他作为风险共担体一部分所贡献的一份子”。代价一旦被分摊，存押就变成了投资——一群人一起投资一个项目，失败时每个股东按比例承担损失，没有任何一个股东的个人存在发生了折旧。投资者还是投资者，他的存在账户里的钱一分没少——少的只是他投资账户里的钱。

这三个条件合在一起，将存押与一系列邻近概念切分开来。“决策”不需要无保底，不需要存押意义上的不可逆——它正是在一个已知框架内、借助外部保底、保留撤回可能性的操作。决策者的基本前提是：如果我判断错了，我可以修正。存押者的基本前提是：如果我押错了，我已经不是之前的我了。“风险承担”可以有保底——企业家的投资可以对冲、可以保险、可以通过有限责任把个人损失封顶。而存押是“把一部分自己不可逆地放进一个没有保底的动作里”——损失一旦发生，不是损失一笔资产，而是损失了一部分自己。这不是程度上的差异。这是“损失了什么”这个问题的存在论差异。

## ■ 2.3 不可还原硬校验：与阿伦特“行动”的正面对话

命名完一个概念，必须立刻把它压成最短的一句话，然后逼问——有没有哪个既有理论中的现成概念，可以同样用一句话把它无损替换掉？存押的最短定义是：人把自己的一部分存

在，押进一个没有保底的不可逆动作里，从而让这个动作永久地成为其存在历史中的不可删除节点。

逐一比对最可能接近的既有概念：海德格尔的“本真决断”是此在从常人沉沦中收回自身，向死而生——这个动作不涉及“把存在押进外部动作”，也不产生存在折旧；它关心的是此在如何“成为自己”，而非此在如何“付出自己”（Heidegger, 1927）。克尔凯郭尔的“信仰跳跃”是在荒谬中跃向上帝，但跳跃有外部保底——上帝作为见证者理解跳跃者，即使世俗世界不理解（Kierkegaard, 1843）；存押没有这种保底，存押者不能指望任何超越的见证者来为他的押注提供意义担保。经济学的“沉没成本”是不可回收的投入，但沉没成本可以只是金钱或时间，不需要押进存在；沉没成本发生后，“你”还是你——你的存在账户没有划走任何东西。

但有一个理论近邻无法用三句话打发掉。汉娜·阿伦特的“行动”概念与存押之间的亲缘关系远比以上任何概念都更接近，也更危险——危险在于，如果不能精确地切开这两个概念，存押的“不可还原性”就会坍缩为“行动的某种加重的、更戏剧化的表达”。

阿伦特在《人的境况》（Arendt, 1958）中对“行动”的论述之所以具有持久的穿透力，正在于她将不可逆性与不可预见性锚定为本真行动的核心特征，而非需要被管理的麻烦。在她看来，行动的独特之处在于，它一旦被开启，其后果就溢出了行动者的控制，进入一个无限的互动链条——你永远无法预判你的行动将在他人的回应中变成什么，你也永远无法撤回它。更重要的是，阿伦特强调行动与“我是谁”的揭示直接相关：通过行动，一个人在人类事务的世界中揭示了自己是“谁”，这种揭示是反身性的、不可撤回的，它与行动者作为“什么”（其天赋、身份、社会属性）的显现根本不同。这与存押论的核心关切——动作反身性地改变了做出动作者本人的存在——确实高度共鸣。

但存押与行动站在不同的问题域上。这不是一个概念比另一个更“深”或更“高级”的判断——这是它们被完全不同的追问所驱动，回答着完全不同的理论问题。

阿伦特用行动概念回答的是一个经典的实践哲学问题：在现代性条件下——马克思所谓的“劳动”将人降低为动物、现代技术将“制作”变成可预测的功利主义工具之后——人的自由还能在何处显现？她的答案是“行动”：在公共空间中、在与他人的互动中，开启一个不可预见的新事物，揭示“我是谁”。这是一个以政治自由在公共领域如何可能为核心关切的理论。

存押论的问题是完全不同的。它不问“自由在哪儿”，它问的是：在一个人所有的、可以被称为“他的”的动作中，哪些动作一旦做出，就让这个“他”不再是之前的那个“他”？这不是一个更窄的实践哲学问题。这是一个存在论问题——更准确地说，是一个生成机制问题：

人的存在，是如何被自己的动作发生、折旧、重铸的？这个问题在阿伦特的理论谱系中没有位置——不是因为阿伦特不够深，而是因为她从来没有被这个问题驱动过。她关心的是人如何在共同世界中显现，而不是人如何在私密的极限动作中被重铸。

这个论域错位，可以沿着三重边界被更精确地刻画出来。

第一重边界位于现象学层面：存押的不可逆性不一定、甚至往往不发生在公共空间中。阿伦特的行动是严格在“人的复数性”条件下展开的——行动发生在有人观看、有人回应的公共领域，其不可逆性的悲剧意味恰恰在于：你无法控制你的行动在他人中间的后果，而这些后果会永久地嵌入共同世界。存押则不同。存押的重量不一定需要、甚至往往不需要被他人即时见证。当布尔乔亚在深夜独自面对那些巨大的金属构件，不知道这一次的焊接会不会把整个结构烧毁时，她是在存押。她的这一动作是不可逆的——即使没有任何人在场。假设一个漆匠，四十年来每晚收工后在一间阁楼上画了上千幅画，从未展出、从未示人、死后被家人当废纸处理掉，他在画下那些画时，每一次用笔都在存押。他的存押的重量，与他是否被写入艺术史、被编入策展前言、被某个批评家识别——完全无关。

这导向一个必须正面回答的问题：如果存押的重量在发生的那一刻就已经构成一个存在论事实——即已经被那个人自己承受过了，而无论之后有没有别人来承认它——那么这与阿伦特将行动的实在性锚定在公共显现中的理论预设，处于何种关系？阿伦特当然不会否认私人经验的存在。但对她而言，“我是谁”的揭示——那个反身性地确认行动者身份的层面——必须发生在共同世界之中，被他人看见、记忆、叙事。阁楼上的漆匠揭示了他自己是谁——在他自己面前揭示，在他每一次押笔的呼吸中揭示。他的存在折旧了，但他折旧的痕迹从未在任何一个公共叙事中获得实存。存押论要确认的正是：这个从未进入公共叙事的折旧，仍然是一个存在论事实。接受这一命题，就意味着接受一个阿伦特未曾明确打开的存在论命题：人的存在可以在完全私密的条件下，被自己的动作不可逆地改变。这种改变不需要他人的承认来“完成”它的存在论身份。他人的承认，只是在事后为这个已经完成的存押追加了一重公共意义——但它不是存押本身成立的条件。

第二重边界位于主体结构的层面：存押与行动对“动作如何反身作用于主体”的预设根本不同。阿伦特的“行动”虽然具有不可逆性，但它始终预设着一个相对于行动而保持同一的行动者——一个通过行动揭示其独特性的“谁”。这个“谁”在行动之前已经存在，在行动过程中显现，在行动之后被记忆——但它本身不因行动而被永久改变。行动者经历了行动的后果，承担了行动的责任，甚至被行动的不可逆性所困扰——但行动者在存在论上是同一个人。

存押则不同。存押者不是在“通过押注揭示自己”——他是在押注中交出了一部分自己。他的存在发生了折旧。在押注之后，他不再是同一个人。不是说他的自我认知改变了——那还

是心理学层面。而是：他的存在已经有一部分不在他那里了。它留在了那个动作里。他比之前少了些什么。这种“少”，不是阿伦特的行动者在经历了行动的不可逆后果之后所感受到的“我再也收不回那个动作了”。存押者感受到的是：“我再也收不回那一部分我了。”

第三重边界位于艺术理论的论域层面。存押论与阿伦特的行动理论在艺术问题上的相遇点恰好暴露了二者的不可通约性：阿伦特的行动概念可以被用来分析某些存押的公共后果——比如一个行为艺术家的存押动作如何被艺术世界接住、争论、叙事化。但你不能用阿伦特的行动概念来分析那个阁楼上的漆匠的存押——因为那个存押没有公共后果，它只是漆匠本人存在账户中的一笔不可撤销的转账。而存押论恰恰要把这个私人转账的存在论意义锁死：它不是一句安慰失败者的漂亮话，它建立在存押的本质规定上——存押的本质是“我在押进去的那一刻，我不再是之前的我了”，这与押注的结果——是否产出了一件被外部系统认可为“好作品”的东西——在存在论层面是分离的。

因此，存押与行动的关系不是对立，也不是谱系上的递进。它们站在不同的存在论地基上：一个站在公共显现的平面上，追问自由如何进入世界；一个站在私人存在折旧的平面上，追问存在的质地如何被自己的动作重铸。这两个平面在某些案例中重叠——比如一些公开的行为艺术本身就是存押，同时也构成阿伦特意义上的行动——但它们的核心判断不是一回事。存押可以完全私密；行动则本质上是公开的。存押产生存在折旧；行动则揭示已存的存在者是谁。存押论切出的，不是行动理论的延伸或深化，而是一个行动理论从未追问的维度：当一个人做出一个不可逆的动作时，这个动作如何反过来折旧了做出它的人？

## ■ 2.4 概念增量：不可还原性的另一个根据

存押论的概念增量，最终可以这样被锚定：在任何既有理论中，没有一个单独的概念可以用来回答下面这个问题——当我们在布尔乔亚的蜘蛛前感到的那种说不清的重量，它是否可能有一个来源，这个来源不在作品的任何一个可以单独指认的属性（形式、观念、技巧）里，而在创作过程中真实发生过的、创作者自己的一部分存在被不可逆地转移进作品这一事实里？

阿伦特的行动理论可以解释布尔乔亚的蜘蛛如何在她进入公共领域之后，被艺术世界接住、争论、写入历史。但它无法解释：在布尔乔亚独自面对那些金属构件的深夜，在没有观众、没有批评家、没有被写入历史的任何承诺的情况下，是什么让她的焊接动作具有了一种不可代偿的紧迫性。阿伦特的理论不处理这个问题——不是因为它不够好，而是因为这个问题不在它的论域之内。

存押论处理这个问题。它的答案是：那一夜的焊接之所以不可代偿，是因为在那一次次的焊接中，布尔乔亚把自己押了进去。她的存在发生了折旧——即使没有人知道、即使作品最终

被毁掉、即使她本人事后懊悔——折旧已经发生了。这一“发生过”，就是存押论要命名的东西。

### ■ 三、风险个人化侵蚀艺术存押条件的中间传导机制

#### ■ 3.1 为什么需要这一节

从“风险个人化”这一宏观制度逻辑，到“艺术生态中时间冗余被侵蚀”这一具体病症，中间的传导不是自动完成的。风险个人化主要作用于经济、就业、福利领域——它最初处理的问题是：当个体从传统共同体脱嵌之后，如何独自面对失业、疾病、养老等社会风险。它如何跨越领域的边界，渗透进一个看似与失业和疾病不直接相关的场域——艺术家的创作日常？这一节的任务就是展示这个传导的具体制度接口。这些接口不是风险个人化逻辑的“应用实例”——它们就是风险个人化逻辑在艺术领域获得肉身的地方。

#### ■ 3.2 基金申请表：“预期成果”栏的风险管理属性

当代艺术家——尤其是年轻一代——面临的最常见的制度接触点之一，就是申请创作基金。以国家级或省级艺术基金的标准申请流程为例，申请表中通常包含一个必填项：请简述本项目的核心问题意识与预期创作成果（500字以内）。这一栏的设计初衷是合理的——基金提供方需要在众多申请中做出选择，它必须有一个可比较的评判依据。但它的制度效果是：它在结构上要求申请人将自己正在摸索的那个“说不清的东西”，翻译成一种可以被评审委员会理解、评估、比较的“问题意识”与“预期成果”。

评审委员会由策展人、批评家、资深艺术家组成。他们是善意的、专业的，但他们的工作必然要求他们用可论述的标准去判断申请的质量。因此，一个正在发展某个他自己也说不清楚的方向的年轻艺术家，无法诚实地写：“我不知道我在做什么，我正在等待它自己告诉我。”这句话在天真的层面是最诚实的存押声明，但在制度的层面上，它为零分——它对评审委员会不提供任何可评估的信息。

所以他必须写：“本项目旨在通过材料实验，探索触觉感知与视觉形式之间的阈限地带……”这句话在制度的层面上是合格的——它准确、有理论归属、可评估——但关键不在于他是否在“撒谎”。事实上他没有撒谎——他确实在做这件事。关键在于：他做的“这件事”，和他写在表格里的“这件事”，不是同一件事。他正在做的，是一个他自己还没办法说清楚的事。他写在表格里的，是一个已经被翻译成可论述语言的事。

这个翻译过程，就是风险个人化逻辑在艺术领域的具体操作界面。基金申请的“预期成果”栏，与保险公司的风险评估问卷在结构上是同一种制度工具：它们都要求个体将未来的不

确定性提前编码为一组可管理的、可评估的参数。申请人在面对表格时感受到的“必须说清问题意识”的压力，本质上不是官僚系统的繁文缛节（虽然它确实表现为繁文缛节），而是风险管理的制度逻辑在要求他把“不知”转译为“已知”——把悬置转译为预期。

更致命的是，这个翻译不只是一次性的妥协。一旦基金批下来——假设他被足够幸运地选中——他就必须按照他写在申请表上的那个版本去“产出”。结项时需要提交成果报告，成果报告需要对照“预期成果”逐条说明完成情况。他现在被锁进了他自己写下的那个可论述的版本里。他不再是那个在边缘摸索的人——他变成了那个“在做触觉感知与视觉形式阈限探索”的人。那个最初被他翻译出来以通过制度门槛的论述，现在反过来定义了他应该做什么。

这就是为什么本章不把“基金申请侵蚀创作自由”当作一个简单的官僚主义问题来讨论。官僚主义是一种低效率——它可以通过简化流程来修复。但这里的问题更深：即使这个基金申请的流程被优化到最简洁、最高效，只要“预期成果”这一栏还在——也就是说，只要制度在结构上仍然要求申请人在拿到资源支持之前，事先把不确定的探索转译为可评估的计划——那么它在逻辑上就会持续地、系统性地排斥那些真正处于悬置中的存押。这不是制度“坏”——这是制度和存押在结构上不相容。

### ■ 3.3 画廊合同：创作周期被提前封顶

第二种关键的传导界面，是商业画廊的代理合同。以中端商业画廊的标准代理协议为例——这里的“中端”指年营业额在百万至千万美元之间、不处于蓝筹顶尖但已具备稳定藏家基础的画廊。在这类画廊的代理协议中，通常包含一条“定期创作与交付”条款，要求艺术家每十二或十八个月提供足够数量的新作品用于个展（通常在八到十五件之间，视媒介而定）。这条款本身并不苛刻，甚至是对艺术家生产力的合理期待。但它产生了一个结构性的副作用：它把创作的时间窗，提前封顶为十二个月。

在这十二个月里，艺术家需要完成从构思、实验、制作到交付的全部流程。如果他有一个需要“悬置”的方向——一个他还没把握、可能会失败、可能需要反复推倒重来的探索——他赌不起。因为如果他赌输了，十二个月后他没有够数的新作品给画廊，他不仅违反了合同，更会在经纪人那里丧失“可靠”的信用——而在一个依赖关系网络运转的行业里，信用的丧失往往比合同违约更致命。

于是，理性的年轻艺术家——如果他幸运地被这些画廊代理——会学会做一件事：把“押”的时间压缩到可控的范围内，把其余的时间留给“决策”。他会同时推进两个方向：一个安全的，周期内必能产出的方向（决策）；一个冒险的，押在边上偷偷做（存押）。这个策略是聪明的、有效的——它恰好证明了这个艺术家深刻地理解了他所处的制度的逻辑，并找



到了在其中求生存的次优解。但它也恰好证明了：存押，在这个生态中，已经不再是那条可以被公开坚持、被制度保护的主线。它变成了副线——在主业之外偷偷做的私活。

这个微观机制的核心在于：画廊合同中那条看似中性的“定期交付”条款，实际上是经济系统的支付/不支付二代码与艺术系统的弱代码特征之间的一次直接的制度耦合。它没有废除艺术家“押”的权利——它只是让“押”变成了一种制度上的奢侈品：只有那些已经在市场上建立了足够声誉、能够对抗合同压力的艺术家（通常是已成名者），或者那些有独立经济来源、不依赖画廊合同的艺术家，才有余裕去押。对于那些恰好最需要时间冗余来发展自己语言的年轻艺术家而言，画廊合同在客观上为他们提供了存押的制度性抑制。

### ■ 3.4 驻留计划：从“时间礼物”到“成果验证”

第三种传导界面，是艺术家驻留计划的制度化转型。驻留计划曾经——在二十世纪中期的某些典型案例中——是一种时间礼物：一个机构为艺术家提供一段不受打扰的时间，不做任何产出要求，只为让艺术家从日常生计的压力中解放出来，专注于探索。这种“不问产出”的姿态，恰恰是时间冗余的制度化供给：它为存押提供了正式的时间许可——你可以在这段时间里什么都不“做成”，你不需要证明你在做“有用”的事。

但在过去二十年间，随着艺术赞助体系的专业化和问责化，驻留计划普遍引入了一套与基金申请类似的预期成果验证机制。驻留艺术家被要求在驻留期结束时举办展览、做艺术家讲座、提交驻留报告——这些活动本身当然是有价值的，但它们共同构成的要求是：你必须在驻留期结束时，展示出可被公开评估的产出。当这个要求从柔性期待变成制度性强制——尤其当它被写入驻留合同的条款中——驻留就从“时间礼物”暗中转型为“成果验证的倒计时空间”。

这一点在本章作者分析的近年公开艺术家访谈中有清晰的回响。一位获得 2021 年某国际驻留奖项的年轻画家在驻留期间接受的一次内部采访中回忆（访谈载于驻留机构 2022 年度出版物，应受访者要求匿名处理）：整个驻留期一共六个月，前三个月他几乎什么都没“做成”。他每天在工作室坐着，有时候画两笔，不满意，刮掉，再画，再刮掉。那段时间他最大的恐惧不是“我是不是没有才华”，而是“等到驻留展的时候，我拿什么出来”。最后两个月，他“放弃了那条走不通的路”，换了一个他“知道能做出效果”的方向，赶出了一组作品。展览很成功，作品卖得很好。驻留机构把他的案例作为“跨出舒适区、在压力下找到新方向”的成功故事写进了年度报告。但他在那个内部采访中说的最后一句话是：“我到现在也不知道，被我放弃的那条路，如果我继续走下去，两个月之后会不会突然有什么东西活过来。”

这个艺术家的经历不是一次失败的驻留——按照制度的标准，这恰好是一次成功的驻留：他在规定时间内交出了可展示、可销售的作品，得到了市场与机构的双重认可。但对他自己来说，他在存在论的层面上经历了一次挫败：他不是“找到了新方向”，他是放弃了存押，用决策替换了存押。那条被他放弃的路，永远地留在了那个驻留期里——他没有押下去，他收回了自己。他的存在账户完好无损。但他的作品里没有留下那三个月在画布前反复刮掉重来的那个人的折旧痕迹。它只留下了最后两个月那个高效的、知道自己要做什么的决策者的执行痕迹。

### ■ 3.5 制度接口的共性结构

上述三种接口——基金申请表、画廊合同、驻留成果验证——在具体的制度外观上各不相同，但它们共享一个深层结构：它们都在操作上要求艺术家将“不知”转译为“已知”——将存押的开放悬置，转译为决策的可评估预期。基金申请表要求在项目开始之前写出预期成果；画廊合同要求在合同签订之时确定产出数量与时间节点；驻留成果验证要求在规定期限结束时交付可公开评估的作品。三者的共同效果，是在存押预备期的入口、中段、出口各自安插了一个制度性的翻译关口——强制将不确定的探索，转译为确定的计划或成果。

这就是风险个人化逻辑在艺术领域的具体肉身。它不是以一个抽象的社会学命题的形式降临——它是以表格中的一栏、合同中的一条条款、驻留报告中的一个必填项的形式，进入艺术家的日常。每一次进入这些制度界面，艺术家都面临一个微小的、几乎不自觉的选择：我是坚持说我真实的“不知”，还是把它翻译成制度能理解的“已知”？而制度的设计，使得后一个选择在操作上永远是更合理的——它不是被谁强制做出，它是被制度环境的默认设置所“推荐”。

## ■ 四、存押、决策与代偿：三种动作类型的存在论区分

### ■ 4.1 三重区分

有了存押这个概念，并且理清了风险个人化逻辑侵蚀存押生态条件的制度传导机制，就可以给出一个更精确的区分。本章提出的是三种动作类型的生成机制区分——不是三种“艺术家类型”。一个艺术家完全可以在其职业生涯的不同阶段、甚至在同一件作品的创作过程中，交替执行这三种动作。

决策的特征是：它在一个已知的框架内进行。它的每一步都有可参照的外部标准——之前的成功案例、批评家的话语趋势、市场的偏好信号——告诉你“这样做是对的”。它的后果可以在做出之前大致预估。如果预估失误，可以调整、可以撤回、可以“下次注意”。决策是聪明的、高效的风险管理行为。决策者不会让自己的存在受损——他只是在用他的才能、经验、

资源去下一盘已知规则的棋。他的存在账户是安全的：如果他赢了，他赚到名声和金钱；如果他输了，他损失的是名声和金钱，不是他自己。

代偿的特征是：它用另一种动作来替代存押，同时确保产出的效果可以与存押产物无法区分——至少在外观上、在可论述性上、在制度评估的维度上无法区分。这个词在医学上指一个器官的功能受损后，另一个器官通过增加负荷来补偿缺失的功能——但被补偿的是效果，不是原器官本身。本章在严格类比的意义上使用这个词：代偿动作补偿的是存押的效果——它产生一个在属性上与存押作品难以区分的东西——但它绕开了存押的存在论核心：没有人为此付出存在折旧的代价。AI 图像生成是代偿的典型：它产出的图像在视觉上完全可以媲美甚至超越手绘，但在生成过程中没有任何人押进任何存在。代偿不产出“轻的”作品——它产出的作品，在效果层面无法与存押作品区分。这正是它的力量所在：它把“是否需要存押”变成了一个无法从结果上验证的问题，从而在制度层面获得了与存押同等的流通资格。

存押的特征是：它无保底，不可逆，自己承受代价。它是一个人在不知道后果的情况下，把自己的一部分存在押进一个动作里，从而让这个动作成为其存在历史的不可删除节点。存押不保证产出的优越性——存押的产物可能在形式上拙劣、在观念上混乱、在市场上的一败涂地。但它的一件东西是代偿永远无法复制的：在它的作品里，留着一个真实的人的存在折旧的痕迹。这不是一种可以被看见的属性——没有一种技术手段可以检测“这幅画里有梵高的存在折旧”。但它是一种可以被感知的重量：当你站在画前，你感受到的那些刮痕、那些曾经被画上去又被铲掉的颜料的幽灵层、那些在最终成型的形体下面仍然在呼吸的曾经存在过又消失了的笔触——这些都是一个人曾经反复把自己押进去、然后放弃、然后再押的证据。你不需要知道梵高是谁，就能感受到这幅画“被押过”。因为存押的痕迹不是被画上去的——是留下它的那个人的存在，在它身上留下的疤痕。

## ■ 4.2 三重区分的生成机制意义与严格边界

这个三重区分的生成机制意义在于：它解释了为什么当代艺术生态中，决策驱动的作品可以和存押驱动的作品在制度评估中平起平坐、甚至占据优势。原因不在于制度“打压”存押，而在于：在外观上，好的决策产物与存押产物的属性差异，往往落在制度评估体系无法测量的维度上。一个策展人面对两件他从未见过的新作品，他无法仅凭作品本身判断哪一件是存押的结晶、哪一件是决策或代偿的产物。因为他能看到的是作品的形式、观念、技术水准——这些都属于作品“是什么”的层面。而存押的痕迹，属于作品“怎么来的”这个层面——它不在作品的可直接感知属性里，它需要被知道。而策展人不知道。他只能靠作品给自己留下的感觉去

判断——“这件作品有一种说不清的重量”——而这种判断是无法被写进展览前言里的，因为它没有办法被翻译为可论述的理由。

这就意味着：当制度评估体系以“可论述性”为核心标准时，存押面临着结构性的劣势。代偿产物天然具有高可论述性——它的操作本身就是在话语空间中进行的，它可以轻易地援引既有概念谱系来为自己的合法性辩护。正因为代偿动作一开局就是用已知概念搭建的，它的每一步都可以被自己讲述。而存押产物往往在很长一段时间内是“无法说清”的——不是因为它没有意义，而是因为它是从一个尚未被既有语言照亮的区域里长出来的。当制度要求艺术家用三句话讲清自己的“问题意识”时，存押者往往无话可说——他在做的时候，并不知道那个东西“是什么”，他只能在做了之后，等它自己慢慢显现。而这个时候，代偿者已经拿着漂亮的论述被写进展览前言里了。

与此同时，必须厘清“代偿”这一概念的严格边界。代偿不是泛指任何“不存押的创作”——那是混淆了代偿与决策。决策不试图模仿存押的效果：一个艺术家做了一个在已知框架内的、聪明的、有市场针对性的系列，他并不声称他的这些作品“被押过”。他只是在做一种不同类型的艺术生产。代偿则不同。代偿是在作品的效果层面（外观、观念深度、可论述的“重量”）精准对标存押产物，但在生成过程中绕开了存押的存在论核心。代偿不只是一条不同的路径——它是对存押效果的模拟，使得存押本身在制度评估中变得“不再必要”。区分这两种动作类型，不是为了给艺术家贴标签——一个艺术家完全可能在同一件作品的不同阶段交替使用代偿与存押——而是为了诊断一个生态学事实：当代制度环境是否正在系统性地让一种动作类型（存押）的生态条件持续收缩，同时让另一种（代偿）的制度适应性不断增强？

#### ■ 4.3 一个检验性的对比：梵高与赫斯特的两种动作瞬间

以上区分的穿透力，可以通过一个对比来检验。不是在比较“梵高这个人”与“赫斯特这个人”谁更好——那不是存押论要做的事。而是在比较两种动作类型的发生条件差异。

梵高生前几乎没有卖掉任何画。他靠弟弟提奥的资助活着，在阿尔勒的烈日下反复画那些没有人要的画。他完全不知道自己的作品会不会有一天被人认可——事实上，在他活着的最后几年里，所有证据都指向“不会”。他不是在“冒一个可能失败的风险”。他是活在那个失败里，同时还在画。他在做着这一切的时刻，没有任何保底、没有任何信号告诉他“继续画是对的”。他的某一笔——比如《星月夜》里那些漩涡般旋转的星空笔触——是在一个他没办法预判效果的夜晚做出的。他不知道那一笔会不会毁了整幅画。他押下去了。这一押的性质，是决策概念无法穷尽的：它不是在一个已知策略空间中选择最优解，它是在没有保底的条件下做出一个不可逆的动作，并且这个动作永久地改变了他作为画家的存在质地。

赫斯特早年确实做出过一次真正的存押。1991年，他自己出钱做了那条泡在福尔马林里的鲨鱼——《生者对死者无动于衷》——花了五万英镑，在当时是一个巨大的赌注。那个动作是不可逆的：如果没有人买，那五万英镑就砸了，他作为年轻艺术家的生涯可能就此终结。在这个时刻，赫斯特押了。

但在此之后，赫斯特的艺术实践发生了转变。他后期的许多作品由助手团队执行——他自己在2008年拍卖会上直接跳过画廊、将两百多件新作直接送拍苏富比，创下一亿多英镑的个人成交纪录。在这些作品中，赫斯特执行的不再是存押，而是高水平的决策：他在已知的市场规则、品牌价值、艺术史趋势中做出精准的判断。这与他早年押那条鲨鱼时的动作类型，在存在论层面是不同的。

这一对比不是要把梵高和赫斯特分别锁死在“存押者”和“决策者”的身份标签里。恰恰相反，它要展示的是：存押是一种在特定条件下可能发生的动作瞬间，不是一个艺术家可以持有的永久身份。同一个赫斯特，在1991年押了，在2008年没有押。他的早期存押不因为他后期的决策而被取消存在论意义，但他的后期决策也不因为他早期的存押而自动获得存在折旧。作品与作品之间，动作与动作之间，必须被分开来追问。这就是存押论作为一种分析工具的核心用法：它不判定人，它追问动作的发生条件与存在论质地。

## ■ 五、三个存押现场

### ■ 5.1 布尔乔亚：疼的不可代偿性

路易丝·布尔乔亚的蜘蛛系列，展示了存押如何在一种无法被助手代偿的身体性动作中发生。

在这组作品诞生之前，布尔乔亚已经做了半个多世纪的艺术。她做过涂绘、木雕、乳胶雕塑。但她真正被艺术世界“看到”，是在她六十岁之后——在她开始做那些巨大的金属蜘蛛之后。这些蜘蛛的来源，是她童年时代那个永远在修补织物的母亲。布尔乔亚的母亲经营一家挂毯修复作坊，每天坐在织机前，修补那些破损的挂毯。蜘蛛，是那个修补者的身体。

但这不是一个可以被冷静运用的“创作素材”意义上的来源。在布尔乔亚触及这个形象之前，它不是一段被安全地放置在记忆博物馆里的、随时可以被美学化处理的童年经验。它是一个还在疼的东西。它不是她可以冷静地分析、挑选、用艺术语言来“表达”的对象。它是那个她必须自己走回去、重新触碰、然后在触碰中被再次灼伤的东西。

这就是为什么蜘蛛雕塑不能由一个助手来做。不是技术上不能——金属焊接是一门可以被训练的手艺。而是存在论上不能。另一个焊工不知道那个疼长什么样。他不知道蜘蛛的腿在哪

个角度上承载了那个母亲的重量。他不知道蜘蛛腹部的弧度是保护还是压迫。布尔乔亚必须自己押进去。她必须在每一次面对那些巨大的金属构件时，重新进入那个疼，然后在不知道这一次能不能把它做成“对”的情况下，做出那个不可逆的焊接动作。

蜘蛛雕塑因此携带的重量，不来自它“象征了母性”这个概念——概念是轻的，任何一个读过精神分析的人都可以写出“蜘蛛象征母性”这句话。重量来自：在这些雕塑的身体里，布尔乔亚在每一个焊接动作中，反复把自己押进那个她自己也无法完全控制的疼里，而那个疼在金属上留下了不可逆的痕迹。她的存在折旧了——不是因为她付出了一生的劳动，而是因为在那些焊接的时刻，她的一部分存在被不可逆地转移到了这些金属里。存押论不是来解释“这组雕塑为什么是杰作”的——它是来告诉你：你在它面前感受到的那个说不清的重量，有一个可以被指认的存在论来源。那个来源不是布尔乔亚的天赋、技巧或观念——是她的一部分存在，已经永久地留在了那些焊接的疤痕里。

## ■ 5.2 培根：刮掉与重来的存押结构

弗朗西斯·培根的绘画过程，更尖锐地展示了存押的一个结构特征：存押不需要“成功”——即使是最终被艺术家自己毁掉的作品，也可以见证存押的发生。

培根是一名反复刮掉重来的画家。他的助手在多年后回忆说，培根会在一幅画上反复地修改，有时候改到画布的纤维都被刮伤。他会在一天之内把一张画画完，然后过几天又把整张画刮掉，再画一遍。如此反复，直到某一个时刻——一个他自己也无法预先描述的临界点——他突然停来说：有了。

这个“有了”是什么意思？培根自己的描述是：画面上的那个形体，突然“开始自己呼吸”了。在那一瞬间之前，它只是一个“被画出来的形体”——被他的笔触画出来的。在那一瞬间之后，它成了一个“活着的形体”——它不再属于他了，它开始反过来要求他怎样继续对待它。怎么判断它活了？没有客观标准。培根自己也不知道标准是什么。他只知道：它在某一刻突然就“对”了。这个“对”，不是技术上画得准——培根从来不正面对抗透视法，他画的人体都是扭曲的、错位的。这个“对”是一种存在论上的立住：那个形体在画面中获得了它自己的生命。

存押发生在哪里？不是发生在他最终画下那个“活着的形体”的那一笔里。存押发生在每一次他刮掉重来的时候。每一次刮掉，都是对之前所有押注的不可逆放弃——“我不认了，我再押一次”。每一次重来，都是重新进入那个没有保底的悬置状态——重新把自己押进一个可能再次被刮掉的不可逆动作里。培根不是在“修改”——修改意味着你有一个已知的、朝之修改的样稿。他没有。他是在悬置中等待一个他自己也无法预判的临界点。在那之前，他每一笔



押下去的时候，都不知道这一笔会不会活——甚至不知道这一笔会不会在十分钟之后被他亲手刮掉。

站在培根的画前，你感受到的那些刮痕、那些曾经被画上去又被铲掉的颜料的幽灵层、那些在最终成型的形体下面仍然在呼吸的曾经存在过又消失了的笔触——这些都是一个人曾经反复把自己押进去、然后放弃、然后再押的证据。而这些证据之所以有重量，不是因为它们构成了一个成功的画面——而是因为它们见证了一个人在悬置中的反复存押。

这引出一个必须正面回答的问题：如果培根的那幅画没能抵达“有了”、整张画被他扔掉了——一个所有人都看不见、进不了展览、进不了拍卖行的作品——它还有重量吗？

存押论的回答是：有。而且这种“有”不是一句安慰失败者的漂亮话。它建立在存押的本质规定上：存押的本质是“我在押进去的那一刻，我不再是之前的我了”，这与押注的结果——是否产出了一件被外部系统认可为“好作品”的东西——在存在论层面是分离的。培根那幅被他扔掉的作品，在美术馆和市场中没有“作品身份”——它不被写入展览图录，不被计入拍卖纪录。但在培根自己的存在历史中，它是一个不可删除的节点：他就是那个曾经在那一幅画上反复押注、反复刮掉、最后放弃的人。这件事发生了。这个“发生过”不等于“被记录过”——在制度档案里它可能从未存在过，但在存在论层面它发生了。存押论正是要把这一点锁死：有一类事件的实在性，不在于它们是否被制度承认、是否被他人感知、是否产生了可见的后果——而在于它们是否在做出它们的那个人的存在本身，刻下了一个不可逆的印痕。

这里有一条必须清晰的边界：不是所有的“尝试了、失败了”都可以被事后追认为存押。那些在已知框架内做实验、失败了就更换参数、从头到尾没有在悬置中押进自己存在的尝试——它们只是不成功的决策。存押的识别，不靠结果判断，但靠发生条件判断：在做的那个时候，他是否真的没有保底、真的不可逆、真的自己承受代价？如果是，那么即使结果在可见层面是“零”，他的存在已经被改变了。那个改变，没有人能替他注销。

### ■ 5.3 伊恩·程：数字时代的反默认存押

伊恩·程的《使者》三部曲，展示了一个与技术工具的默认设置逆向而行的存押案例——它证明数字技术并非注定只能用于代偿，但同时也证明在数字条件下，存押必须付出比绘画时代更高昂的制度性代价。

程编写了一个由人工生命角色构成的虚拟生态系统，角色们按照一套初始规则互动、学习、变异。一旦程序启动，程无法完全控制它的走向。他面临一个数字时代的典型选择：他可以让程序在自己的工作室里反复运行，直到出现一个他满意的“版本”，然后把那个版本的录像拿到画廊去播放。这个流程在技术上是完全可行的——它正是数字工具的默认使用方式：生

成→不满意→调整参数→重新生成→直到满意。这个流程产出的最终作品，在视觉上与伊恩·程实际展出的《使者》没有区别。但这个流程是典型的代偿：每一步都可撤回，每个版本都可以被下一个版本覆盖，创作者的存在完好无损地穿越了整个过程。

程没有这样做。他选择让程序在一个公共空间中不间断地实时运行，观众看到的是正在发生的、无法预判的、不可逆转的系统演化。任何不可预测的结果——包括程序的崩溃、角色行为的失当、涌现出令人不安的模式——都会在观众面前实时发生，且不可撤销。

程存押的整个支点，不在他的代码技能里，而在于他做出的那个选择——他选择了取消自己的撤回权。数字工具的默认逻辑是：一切皆可撤销。撤销快捷键是数字时代的护身符——它让每一个动作都可以在事后被宣布为“只是试一下”。程做的，就是主动扔掉这个护身符。他逆着工具的默认倾向，为自己强行构建了一重制度性的不可逆。他说：这就是它了，它在你们面前跑，跑成什么样我都认。这个“我认”，正是存押的发生标志。它不是被动接受——“没办法了只能这样”。它是主动进入：我选择不再撤回，我选择把自己的判断力、审美力、对作品的控制权，一部分交出去，交给这个我无法完全预判的系统，然后在不知道它会产出什么的情况下，承担它产生的一切后果。

这一案例同时揭示了存押的一种新的制度性条件。在布尔乔亚和培根的时代，颜料的物理性质天然提供了某种不可逆性——颜料一旦上了画布，就很难完全清除。存押的发生可以部分依赖这种物理性质作为“外部约束”，艺术家不需要额外地去构建不可逆的制度性框架。但在数字时代，这种物理的不可逆性消失了。撤销是零成本的。存押因此变得更依赖制度性的自我约束——艺术家必须更自觉地去构建不可逆的框架（公共实时运行、放弃控制权），逆着技术的代偿倾向去做。技术没有“禁止”存押，但技术的默认设置使得存押需要额外的制度性创造力——而这恰恰是在一个“存在不折旧”被制度化为默认设置的时代里，最稀缺的生态资源之一。

程的案例因此不只是“数字艺术中存押仍然可能”的例示。它是指向一个更深层判断的入口：在数字时代，存押的发生条件从依赖媒介的物质属性，转向依赖行动者的制度性自我约束。而这一转向的代价是，存押变得更困难了——不是因为艺术家变懦弱了，而是因为每一步都必须逆着整个默认设置去做。在程如果不主动取消自己的撤回权，他本可以在工作室里反复打磨出最完美的版本——代偿的路径始终向他敞开着，且没有制度性的风险。他选择了存押。但这一选择之所以值得被分析，恰恰是因为它远远不是数字艺术创作的默认选项。

## ■ 六、艺术作为存押的制度性飞地

### ■ 6.1 为什么偏偏是艺术

现在可以回答那个最终的问题：为什么偏偏是艺术，成为存押在现代社会功能分化中仅存的主要飞地？为什么不是科学、不是教育、不是法律——这些领域中也有赌上一生的猜想者、押上自己的教师、以一生名誉担保一个判决的法官？

答案不在艺术的“本质”里，而在艺术系统在现代社会功能分化中的独特位置。现代社会是一种由功能子系统构成的社会：经济系统用价格和货币处理稀缺问题，政治系统用权力和法律处理集体决策问题，科学系统用真/假二值处理知识问题，法律系统用合法/非法二值处理规范问题。每个子系统都有自己专属的沟通媒介和二源代码，对外部事件按照自己的代码进行转译。

但艺术系统——作为一个相对独立的功能子系统——的独特之处在于：它没有一个像“真/假”“合法/非法”那样具有强制性、排他性的二源代码。社会学家尼克拉斯·卢曼曾提出艺术系统的代码是“美/丑”，但这个代码的约束力远弱于科学系统的“真/假”或法律系统的“合法/非法”——因为“丑”也可以被接纳为艺术，且在二十世纪以来常常恰恰是艺术自觉追求的标志；而科学系统不可能把“假”的命题接纳为科学，法律系统不可能把“非法”的裁决接纳为合法的司法判决（Luhmann, 2000）。换句话说，艺术系统在运作上有一个关键的结构特征：你向它提交一个动作——比如在画布上画下一道不可逆的笔触——它没有一套强制性的二源代码可以当即判定这个动作“对”还是“不对”。它可以被判定为“好/坏”“有趣/无聊”“重要/次要”，但这些判定是后续的、可争论的、可被历史翻案的。它们没有法律判定的那种即时强制约束力。一件被批评家一致认为“坏”的作品，仍然可以在艺术史上被后人翻案；一件被市场认为“不值钱”的作品，仍然可以在多年后成为天价。

这意味着什么？意味着：在艺术系统中，一个人在做出一个不可逆动作的当下，他被制度强制去“证明这个动作是对的”的压力，是所有功能子系统中最小的。在法律系统中，你在法庭上说的每一句话都必须能被立即置入合法/非法的二值判定——你不能说“我不知道这句话对不对，我先赌一下”。在科学系统中，你提交的每一个命题都必须能被置入真/假检验——你不能在基金申请表里写“我不知道这个假说是否为真，我先押上二十年”。但在艺术系统中，你可以。因为艺术系统的运作逻辑决定了：一个动作是否“成立”，并不需要在它发生的瞬间就被外部代码肯定。它可以被搁置、被争论、被等待、被事后追认或被永久埋没——而在这个过程中，没有人有权迫使你做那个当下的“证明”。

时间冗余的制度根基，正是这种“不当即强制证明”的缝隙。艺术能够维持存押，不是因为艺术体制特别宽容，而是因为它的系统运作逻辑——它的代码不是强二值的——使得一个人可以在更长的时间窗内悬置在“不知道”中，而不被制度强制叫停。

## 6.2 飞地不是永恒堡垒

但这个缝隙不是永恒的。每当艺术系统被其他系统的强代码侵入时，这个缝隙就变窄一分。

最直接的侵入来自经济系统。当艺术市场用价格信号来驯化艺术家的创作路径——你的尺幅太小卖不出价、你的风格太冷不受藏家欢迎、你的产量太低无法维持经纪人的现金流——经济系统的支付/不支付二元代码就渗透进了艺术系统，压缩了“不被强制证明”的空间。当一件艺术品在诞生之前，就已经被要求预测它在市场上的表现（如第三节所述的画廊合同中的定期交付条款），它就不可能再是存押的产物——因为存押的不可逆性恰好意味着：你不可能在押之前算出它会赢。

第三种侵入路径来自政治-道德系统。当艺术被要求“回应时代议题”“表达边缘群体声音”“服务于特定的社会议程”时，政治系统内的正确/不正确二元代码就开始在艺术系统内部运转。一个艺术家可以做出一个他自己也不知道“对不对”的存押——但前提是这个“不知道”本身不被政治代码判定为“不负责任”或“不够进步”。如果一个艺术家存押的方向，恰好撞在了政治代码的禁忌区域，他将不是被艺术系统自身、而是被政治系统的代码驱逐出局。

艺术作为存押飞地的“最后”属性，因此不能被静态地理解。它不是一块被永久豁免的飞地——它是一块边界持续运动、持续被挤压的空地。今天它还能维持一部分存押的功能，是因为侵蚀还没有到达它的核心——艺术系统还未被彻底附属于其他系统的强二元代码。明天，当每一件作品在诞生之前都必须通过可论述性筛选、市场预测评估、政治正确性审查——存押，就可能在艺术这片最后的制度化空地上，也被系统性地铲除。

## 七、推衍：存押论的外推边界与限度

### 7.1 教育中的存押被侵蚀

存押论的解释力超出了艺术领域，可以照亮教育领域正在发生的类似病变。

教学现场充满了存押的可能性。教师在课堂上面对每一个意料之外的瞬间——一个学生提出一个奇怪的问题，教师必须在那一秒决定如何回应。这个回应，如果是真的，是无保底的（他不知道这样回应会引发什么后果）、不可逆的（它已经发生了，无法撤回）、自己承受代价的（如果这个回应引发了课堂混乱，他得自己兜着）。这种存押的密度，构成了教学的存在深度：好的教师不是那些“把教案执行得最完美”的人，而是那些在课堂的开放瞬间中反复押上自己的人。他们是在用自己的存在，去接住另一个正在展开的存在。

但当代教育体制中，标准化教案用预设好的教学环节把课堂的开放性压缩到最小。形成性评价用连续的公开评分，让学生和教师都把注意力转移到“如何在下一次评估中表现更好”。课堂录像、教学督导、家长群里的实时反馈——所有这些技术都在做同一件事：让教师在教学现场的每一个动作，都面临被事后追责的可能。当每一句话都可能被截屏、被举报、被放到社交媒体上审判时，“在不知道后果的情况下踩下去”就不再是勇敢——它是不计后果的冒险。教师被系统地训练成更谨慎的决策者。他们学会了在课堂上“做对”——教那些不会被质疑的内容，用那些被验证过有效的话术，避开那些可能引发不可控讨论的话题。他们不是在教——他们是在执行一套风险最小化的教学方案。

## 7.2 养育的存押结构

如果离开制度领域，进入更私密的家庭领域，存押论可以进一步推衍至一个经常被哲学忽视、却在每个人的存在史中最为深彻的场景：养育。

养育在存在论上的结构，是所有人类活动中最接近存押原型的——但正是由于它的普遍性，这一点反而被遮蔽了。养育的存押结构在于：第一，它不具备外部保底——没有任何父母在孩子出生之前，能够真正预知自己将变成一个什么样的父母，因为“作为父母的存在”不是在生育之前已经存在的、等待被某个外部条件激活的身份，它恰恰是在养育的每一个不可逆动作中被生成的。第二，它是不可逆的——你对这个孩子说过的每一句伤害的话、给出的每一次你不确定是否正确的保护、做过的每一次你自己后来也会怀疑的决定，都已经进入了这个孩子的存在历史，无法删除。第三，它的代价必须由自己承受——你无法雇佣任何人替你承担为人父母的全部后果，你无法把养育的失败转记为任何其他东西的失败：它就是作为这个孩子的父母的失败，无可推卸。

而在当代条件下，养育这个天生的存押现场，正在经历与艺术领域同样的侵蚀——只是它的侵蚀路径更隐蔽、更日常化。育儿指南、儿科标准、专家建议、社交媒体上的“完美育儿”叙事——所有这些构成了一套代偿式的育儿知识体系。它的逻辑是：只要遵循正确的操作流程，你就可以做大概率正确的父母；如果你做错了，那是因为你没有好好学。这套知识体系本身是有用的、甚至是必要的，但它将养育从一个“我不知道，但我赌”的存在冒险，默默地转换成了一个“只要我输入正确信息，我就能输出正确结果”的决策管理过程。父母不再是存押者——他们变成了焦虑的风险管理者，不断在用“这样做对吗”的自我审查，试图消灭那个无法被消灭的不可逆性。当养育变成一种可以被优化、被评估、被比较的绩效，“押上一部分自己”这一养育的存在论内核，就被淹没在无数代偿式建议的噪音中了。

## 7.3 学术界的存押困境



学术领域的情况提供了一个对比视角。在学术的某些边缘地带——纯数学中一个人赌一生证明一个猜想，理论物理中一个人坚持一个无法被实验验证的方向——存押仍然可能发生，而且其不可逆性接近艺术存押。一个数学家花了二十年时间试图证明黎曼猜想而未果，这二十年是不可逆的折旧——那部分存在已经永远地押进去了，再也拿不回来。他不再是二十年前那个还没有被这个猜想改变过的人。

但在学术的主流运作模式中——尤其是那些高度依赖基金申请、论文发表、职称晋升的学科中——存押的生态条件正在被系统性地挤压。一个年轻学者必须在三到五年内产出可测量、可评估、可比较的成果。他不能花十年悬置在“不知道”中去押一个可能改写范式也可能什么都不是的方向。学术体制不禁止他这样做——他可以在冬季休假期自己偷偷做——但体制的激励结构在沉默地告诉他：那样做是危险的。那些最聪明的年轻学者变成了最好的决策者：他们善于在文献中找到可操作的研究空白，善于设计预计会有显著性结果的实验，善于选择易发表、高引用的课题。他们做的每一步都是对的，但他们押的量是零。他们的存在账户完好无损——论文发了，职称升了，他们到了——但他们留在自己论文里的东西，只有才能，没有存在。

## ■ 7.4 存押在当下制度缝隙中的变异形态

存押论不预设“存押是好的”。这是一条必须在此刻被明确划定的边界：存押论诊断的是一种存在论动作类型的稀缺及其制度根源，它不颂扬存押，也不呼吁“所有人都应该存押”。存押的代价是存在折旧——这对做出存押的个体而言，往往意味着不可逆的伤害、漫长的痛苦、以及大概率的外部失败。一个健康的制度应当做的，不是迫使所有人都去存押，而是为那些被这种动作类型所驱动的人，保留一块不被制度强制叫停的空地。

与此同时，也必须承认：存押并没有被完全消灭。它在制度的缝隙中偷偷活着，而且正在以某种新的、变异的形态出现。有些年轻艺术家在商业项目之外偷偷维持着一个不参展的私人创作系列——这个系列不进入市场、不写入简历、不拿去申请基金，只存在于工作室的一个角落。有些教师在标准化教案的边缘，每一节课偷偷留出五分钟的空白——不说“这是课程目标”，不说“这是知识点”，只是给学生一个开放的问题，然后赌这个问题的走向。有些父母在育儿专家的建议与算法推送之间，偶尔做出一个“我不知道对不对，但我觉得应该这样做”的决定——然后把这个决定的后果自己扛下来，不晒到社交媒体上。

这些动作都不是传统意义上的“英勇存押”。它们很小、很私密、没有被任何制度记录在案。但它们共享同一个结构：在制度的缝隙中，一个人在没有外部保底、不知后果、不可撤回的条件下，把自己的一部分存在押了进去。存押论对这些变异形态的判断，不是一个怀旧的哀叹——不是“真正的存押已死，只剩下这些可怜的碎片”。它要说的是另一句话：只要制度没



有彻底封闭所有的缝隙，存押就不会完全消失。它会变形、会缩小、会转入暗处——但只要有人还在某个时刻选择不撤回、选择认下不可逆、选择用自己的存在去赌一个未知的结果，存押就还在发生。存押论不是一部悼词。它是一部探测器——探测那些在制度默认语法中无法被识别、却仍然在发生着的存在论事件。

### 结论：轻量化时代的探测器

存押的稀缺不是艺术的局部问题。它是当代生活中存在论轻量化的总病灶的一个症状。而这个病灶的根源，不是“现代人变浅了”——这种哀叹正是存押论要取代的。根源在于：支撑存押的制度条件——时间冗余、不被强制证明的缝隙、豁免于风险管理铁律的空地——在效率优化的名义下被一块一块地铲除。当代制度并未禁止存押——在一个以自由为核心理念的社会中，它不可能公然宣布“你不许押上自己”。但它做了比禁止更彻底的事：它让存押变成了一种制度上的不合理行为——你需要为它自己付代价，你不能指望它得到公共认可，你在做了它之后无法用制度内可流通的语言向别人解释你为什么要做它。存押被放逐到了私人生活的暗处。它仍然在发生——在那个在深夜里独自面对金属构件的布尔乔亚身上，在那个反复刮掉重来、等待“有了”的培根身上，在那个拒绝撤回键的伊恩·程身上，在那个在课堂上说出了那句不可撤回的回应的教师身上，在那个面对自己孩子第一次撒谎时不知道该怎么回应的父母身上。存押在制度的缝隙中偷偷活着。但它失去了它的公共合法性。

这就是为什么艺术，作为存押最后一块被制度默认豁免的主要空地，其意义超出了艺术。它不是一个专业领域——它是整个当代生活还能公开见证“有人曾经押过自己”的少数几个场所之一。在那块空地上，一个人还可以对另一个人说：“我不知道这样做对不对。但我赌了。你看——这是我的押注。它在画布上。它在雕塑的焊缝里。它在那个还在呼吸的人工生命的代码里。它在我的存在史中——删不掉了。”

存押论对这一病灶的诊断是：不是因为我们失去了“决断的勇气”，而是因为我们失去了“押进去的制度可能”。它的终点不是怀旧——不是哀悼一个存押遍地的前现代黄金时代。存押从来不是多数人的日常动作，它也从来不应该是。存押论要追问的是：在一个将“存在不折旧”制成制度铁律的时代，我们的制度安排是否正以优化之名，消灭一种对某些人而言不可替代的存在论动作类型。如果答案是肯定的——而本章的论证试图给出的正是这个肯定——那么对策就不在“呼吁更多人存押”上，而在为那些被这种动作类型所驱动的人，重新打开一些不被强制证明的缝隙。

这篇文章本身，是一台探测器。它的任务不是歌颂存押者——他们不需要歌颂——而是用它的概念工具，去探测、去追问、去记录：在那些尚未被彻底平整的制度空地上，在那些还没有

被强制要求说明“预期成果”的角落里，存押是否还在发生？谁在押？他们在什么样的条件缝隙中押？那些缝隙还能维持多久？

存押论不能回答最后一个问题。但它可以——而且必须——把它问出来。

存押论的证伪条件

一个可证伪的判断，必须明确陈述它在什么条件下是错的。存押论的核心判断是：艺术作品中让我们感到“重”的那个力量，来源于创作过程中真实发生的存押——即创作者在无保底、不可逆、自己承受代价的条件下，把自己的一部分存在押进了作品。这一判断在两个条件下的任何一个被满足时，即被证伪：

第一，如果未来的经验研究表明，那些被广泛共识为“有重量”的作品（比如布尔乔亚的蜘蛛、培根的绘画），其创作过程并没有发生存押——例如，有扎实的传记材料证明，这些艺术家在创作时拥有制度性的保底（如已经签约的销售合同、已经确认的展览档期、充裕且无条件的资金支持），且这种保底在其创作决策中发挥了实质性的安全网作用，那么存押论对“重量来源”的解释就是错的。在这种情况下，重量可能来自作品属性的某种特殊组合——即回到属性论的解释框架。

第二，如果未来的制度实验表明，在那些被本章诊断为“系统性侵蚀存押条件”的制度安排（如基金申请的预期成果强制、展览周期的加速）被改革之后——例如，某个基金会试行了十年期的无条件创作资助，且获资助者在整个资助期间完全豁免于任何形式的成果汇报——而艺术作品的“平均存在重量”并未出现可被识别的上升，那么本章对“时间冗余是存押的核心生态条件”的判断就是错的。存押可能需要的是另一种生态资源——例如一种特定的精神集中力、一种特殊的技术熟练度——它的稀缺性不由时间冗余驱动。在这种情况下，对策将不再是为存押者争取更多的时间冗余，而是去识别和保护那种真正的稀缺条件——不管它是什么。

## 第十一章 蚀先于生

本章原题《蚀先于生：为什么形态的耗散比生成更原始》，作者 胡志英，原文见 <https://sdeuniverses.com/students/hu-zhiying/erosion-precedes-genesis/>

### ■ 一、引言：咬住之后，为何松脱？

在过去半个世纪里，过程哲学与关系存在论做出了一个共同的推进：世界不是由一堆已经完成了的零件堆起来的。无论是怀特海（Whitehead, 1929）“现实实有”的合生，还是拉图尔（Latour, 2005）“行动者网络”的联结，无论是巴拉德（Barad, 2007）“能动实在论”中的内部作用，还是当代物理学中“场比粒子更为基本”的视角，都在反复拒绝同一个教条——物质先于关系、实体先于过程。这些工作的核心意象可以凝结为一句话：宇宙的最深层动作，是相互咬住的那个“生”——是那个让物质与关系在同一瞬间互相把对方立起来的动作本身。

这个意象鲜明，有力，而且对。但它只讲到了故事的一半。

咬住，从来不是永久的。一段亲密关系会溃散，一个粒子会衰变，一种货币会沦为废纸，一套制度会空心化。过程存在论的既有意象极擅长描述形态如何建立，却在面对形态如何溃散时语焉不详。它通常把溃散描述为“维持条件的撤走”——就像把一只扶住杯子的手指一根根松开。这个解释的预设是：溃散是生成的减法，是一个消极的、缺乏自身正面结构的过程。只要把生成说清楚，溃散就不言自明。

当真如此吗？

一段信任不会仅仅因为“缺了什么东西”就瓦解。在亲密关系中，摧毁信任的往往不是一个孤立的背叛事件，而是一种新的互动模式的悄然入侵。当一方反复用审视的目光代替接纳的目光，当每一次表达都被暗暗测试度量，当那句“你早点休息”用一个信息上完全恰当的句子把对方还在流动的难受钉死成一个待解决的问题——这时候瓦解信任的，不是信任的“赤字”，而是另一种互动结构的持续性替代。这种替代本身是一个有结构的正面过程：它用自己的运作方式，一点一点地拆换了先前维系信任的那些关键互动节点。旧的互动网络不是被炸毁的，它是被一个异质网络从内部接管的。就像一根藤蔓绕死一棵树，树不是“突然”死的，而是在藤蔓的每一圈缠绕中，树干输送养分的通道被逐点替换为了藤蔓自身的吸盘。

物理学的最深处早就展示了类似的景象。量子场论告诉我们，真空不是空的，而是虚粒子不断成对产生又湮灭的沸腾海洋。一个实粒子被一团虚粒子云包裹着，它的质量、电荷这些所谓的“固有属性”，其实是粒子的本征场与虚粒子涨落之间持续相互作用的净效果。把这一团交织剥掉，你连那个“粒子本身”的指称都说不出——不是因为它被“破坏”了，是因为它的整个存在方式就是在这一点上被替代性的涨落过程不断重新谈判的。物理世界的稳定，从来不是一块金属的静止，而是一场持续进行的蚀与再生的平衡。一旦平衡倾斜，那个稳定形态就从存在中滑脱——不是碎成更小的实体，而是一点痕迹都不留地消失在场的基态里。

社会事实的硬度同样经不起蚀的追问。齐美尔（Simmel, 1900）在《货币哲学》中已经洞察到，货币的价值本质上是一种由“明天别人还会收它”的交互预期撑着的网络效应。这个网络可以极厚，厚到每一次小的拒收都被下一家毫无犹豫的接受所覆盖，从而维持一种完全自洽的稳定。但蚀的机制从不缺席：当另一个网络慢慢形成——比如一个新的支付系统、一个新的记账单位——它不是直接攻击旧货币，而是在旧货币网络的边缘一点点替换它的节点。一个摊贩今天还收旧钞，明天开始优先收新币，后天他只对你这样的老顾客才勉强收旧钞，最后那块布告“本店只接受XX”一贴，旧货币在这一点上的维持条件就被注销了。成千上万个这样的干涸节点，不是一次性的“信心崩溃”，而是一段漫长、安静、有结构的互蚀。当最后一根节点被蚀穿，旧货币的“客观价值”就蒸发了。它蒸发的方式和它当初被建立起来的方式是同一类动作的镜像：不是一堆零件散架，而是一套相互咬住的网络被另一套网络松脱了。

这三个并不属于同一个学科的现象——虚粒子的交织、信任的衰变、货币的溃散——共享了一个被存在论传统长期漠视的形式结构：一个稳定形态的消散，不是一个消极的还原，而是一个有自身逻辑的正面过程；它依靠的不是维持条件的简单“撤回”，而是替代性网络的逐点接入。本章把这个过程叫做“互蚀”。

在展开论证之前，有几笔必须交代的界定。第一，本章的论证策略延续过程存在论的跨域提取方法，从物理、人际、社会三个领域取出的结构，只是展示互蚀这一形式在不同基材上的独立重现，以此支持“互蚀不是个别层面的偶然结构”这个判断。第二，本章在核心概念上与六个最接近的理论传统——生态学的竞争性替代、网络科学的级联失效、热力学熵、黑格尔的扬弃、德里达的解构、以及路径依赖的制度变迁理论——进行正面对质。这既是对先行者的致敬，也是互蚀必须通过的不可还原硬校验：如果一个新概念能被已有概念完全覆盖，它就不配独立存在。第三，互蚀的操作化定义将在与这些对手交锋之后、作为对前文案例形式结构的事后提取而给出，而不是作为事前的概念模具。第四，本章不奢望用科学证明哲学命题；跨域取证的逻辑力量在于形式结构的独立重现，而非从科学命题推导哲学结论。

这里还需要从一开头就澄清一件事。本章的结论——蚀先于生——不是在现有的过程存在论大厦上做一次“内部修补”。如果互蚀是真的，那么既有过程存在论以“生成”为第一动作的整个叙事，必须被颠倒过来。不是“先生成，而后有蚀作为生成的副作用”，而是“蚀才是常调，生是蚀的竞争火焰中偶然没被烧着的那一小撮”。这个颠倒，比“补充盲区”要重得多。它不是对怀特海、拉图尔、巴拉德的友善增补；它是对他们理论地基的一次撬动。本章的论证将正面执行这个撬动。

## ■ 二、互生的盲区：一个“发生”叙事中的被遮蔽面

过程存在论的核心命题可以凝结为一句话：物质与关系不是在时间中先后发生的两种东西，而是在同一个发生过程中同时互构的两种规定性。不存在先于咬住的积木，只存在咬住的动作将积木立为积木的那个瞬间。基于这个命题，过程存在论成功地把“积木还是咬住”这个两千年老问题的前提拆掉了——它告诉我们：问题本身就问错了，因为零件与组装、实体与关系都是发生动作的次级效果。

这个论证在量子场论中的展示极其漂亮：粒子不是预先存在的小硬球，它是场在规范相互作用中的局域稳定激发态。所谓粒子，不过是被相互作用锁定在稳定形态上的那些涨落。没有相互作用，就没有粒子。没有在先的“场质料”，只有相互锁定这个动作本身。在社会制度层面，过程存在论正确地点出：科学事实的“客观性”不是来自它背后有一个不以人的意志为转移的纯物质，而是来自整个学术共同体那一套精细的扶稳网络。货币的硬度来自亿万人无数次交易的沉淀。这些分析都精到且有力。

但它们都只解释了这个网络如何把形态扶稳，没有解释这个网络自身如何可能被另一种网络从内部接管。换句话说，过程存在论看到了“建立”的一面，却把“溃散”简单地留给了“建立的反面”。

这里必须立即澄清一个重要的区别，以免让批评落到稻草人身上。我不是在说过程存在论“忘记了衰变”——它当然可以谈论衰变。它可以把衰变解释为“维持条件的撤离”——就像把一只杯子扶住的手指一根根松开。但这个解释恰恰是问题所在：它把溃散想象成生成的倒放胶片，把蚀过程还原为生过程的减法。它默认了：那些支撑形态的手松开之后，形态就返回了“被支撑之前”的无形状态。可是真实世界中的溃散很少是这样倒放的。一段信任的破碎，不是那些曾经扶稳信任的手指一根根慢慢松开；那是手的姿势被偷偷换了——原本托着杯子的手心，不知什么时候变成了捏着杯耳的指尖，正悄悄地把它提起来往另一个方向挪。那个挪走的动作本身，也是一个有结构、有网络、有自身维持条件的动态过程。它不依赖于原先那套扶

稳网络的脆弱性；它有它自己的主动性。它是一张新网，正在用自己的节点，把旧网的节点一个个地替下来。

这就是过程存在论的盲区：因为它的理论起点定在了“生成”的动作上，它看待世界的眼光就被锁定在了“形态如何立起来”这一侧。任何形态的消散，在它的光学系统里，都只能成像为“维持条件的撤走”，而不能成像为“替代条件的入侵”。撤走是一个消极动作，入侵是一个积极动作。存在论要做的，是把入侵本身作为一个不依赖于撤走的独立现象来对待。本章把它叫做“互蚀”。

为了不让这个指控停留在概念层面，让我们回到量子场论的同一片地基上再仔细看一遍。在标准模型中，一个粒子可以衰变为另一个粒子——例如，一个 W 玻色子衰变成一个轻子和一个中微子。标准教科书上的描述通常是：W 玻色子有一个有限寿命，到时间它就衰变了。这个描述实际上悄悄嵌入了“实体优先”的眼光：它把一个粒子想象成一个有时钟的小东西，时间一到就爆开。但量子场论的真正数学结构并不支持这个图像。在量子场论中，根本没有一个“W 玻色子实体”先在那儿，然后再衰变。实际情况是：W 玻色子作为场的局域激发态，它的存在本身就是一个持续进行中的相互作用结果；当另一个相互作用通道——轻子-中微子通道——在场论中被耦合进来，并且在这个特定能量尺度上与原来的维持激发态的相互作用通道形成竞争时，那个旧通道的维持条件就会被新通道逐渐替代掉。整个过程不是“一个东西碎了”，而是“一个锁定的激发模式被另一个锁定模式从内部解开了锁”。蚀在这里完全是正面的：是一条新的耦合通道替代了旧的耦合通道。

同样，在人际关系中，当过程存在论试图用“蚀”这个词来描述审视对信任的瓦解时，它给出的图景仍然是一种“生—蚀”对子：生是把形态立起来，蚀是生的过程被中止。但审视的运作本身不是中立的中止；它是一张有自身规则的网。审视不是“不接纳”，它是另一种接纳——一种评估性的接纳。当伴侣中的一方用“你最近是不是太敏感了”这种句子来回应对方的难受时，这不是信任维持条件的撤除，而是用一套诊断式的反应模式替代了共情式的反应模式。诊断式的反应模式本身也有它被不断生产的网络：它依赖于一套关于“正常/异常”“合理/过度”的分类系统，依赖于一个“我正以理性评估你的状态”这样的身份姿态，而这些全都在互动中被反复执行、加固。它不是“没有扶稳”，它是用另一种扶法，把杯子的朝向转了一个角度。

所以，过程存在论的最大贡献，恰恰也是它遗留下最大问题的起点：它教会我们不再问“世界的材料是什么”，转而问“世界的形态是如何被维持的”。一旦问出这个问题，立刻就会逼出下一个问题：那这些维持条件，又是怎么被替换的？因为如果整个世界只是一个“生”接着另一个“生”，我们就还必须解释，是什么使得那些本已锁定得极其牢固的形态



网络，居然会被另一个网络撬开第一道缝隙。这个解答，只能用蚀来给出——而且不能是一般性的“失败”或“衰退”，必须是一种有名有姓、有自身结构的正面运作。这正是互蚀这个概念的出生地。

### ■ 三、互蚀的跨域形态

如果互蚀只是一种新的理论措辞，那它不值得被当作独立概念提出。要确立互蚀的存在论地位，就必须先让读者看到，它的形式结构在物理世界、人际世界和社会世界中以各自独立的、可被经验锚定的方式反复重现。只有当一个形式能在截然不同的基材上多次出现，我们才有理由将它视作跨域的底层约束，而非特定学科中的特殊现象。

#### ■ 3.1 物理世界的蚀：虚粒子、退相干与衰变

在量子场论中，一个“基本粒子”从来不是一个正在独存的实体，它是它所激发的那个场与周围所有其他场持续咬合出的一个局域平衡态。这种咬合的精密程度如此之高，以至于在极高能量下，一个电子还同时是一团光子和正负电子对的概率云。但正是这团永远无法被完全扯掉的虚粒子云，暗示了一种隐蔽的蚀：任何实粒子的身份，都在被虚过程持续地“稀释”。单独看某一瞬间，粒子还是那个粒子；但把时间尺度拉长，那个粒子的性质——它的质量、它的磁矩——都是虚过程正反两面相互抵消后的残余稳定项。一旦抵消的条件发生位移，或者有新的实通道打开，原先的稳定态就会失去它赖以维持的精确抵消关系。这就是粒子衰变的存在论实质：不是粒子活够了，而是那条曾经把它维持在“这个粒子”状态下的抵消通道，被另一条不同的、新的相互作用通道从数学结构上蚀穿了。

退相干机制提供了另一个纯粹的蚀例。一个处在叠加态的量子系统，与它的环境发生不可避免的相互作用后，其原先的相干叠加项会被环境的无穷多自由度逐一“携带而走”，最终在极短的时间内，原本彼此相干的态变成了彼此可区分的混合态。退相干的过程，不是一个物理作用力破坏了叠加；它是环境的自由度把原本集中在一个系统内部的相位信息，一点一点地接收、扩散、溶解进了更大的网络。系统并没有“变回点粒子”；它只失去了一种曾经被它的孤立边界保护住的独特咬合形态。这个失去的过程是积极扩散，不是消极撤回。扩散的动力就来自环境自由度的正面的交互能力。

#### ■ 3.2 人际关系的蚀：审视如何接替信任

在亲密关系研究中，一个被反复报告但长期未被独立命名的现象是：信任的瓦解往往不以激烈的背叛事件为先导，而是始于一种互动模式的悄然置换。下面引入一段典型化情境构造，

用以展示“审视接替信任”的形式结构。此处的情境为理论目的而合成，不代表任何具体个案。

A 在长日疲惫后对 B 说：“我今天真的很累。”B 的回应是：“那你晚上早点睡，别熬夜刷手机了——你最近总是晚睡。”

从信息传递的角度，B 的这句话无可挑剔：它提供了明确的行为建议，立场也站在 A 健康一方。但从互动模式的角度看，这句话执行了一个微小的替代。在关系建立的早期阶段，A 表达疲惫时，B 的典型回应是：“是啊，这一天真是够呛的。”这句话不试图解决任何问题，它只是作为一个在场的信号，让 A 的疲惫被另一个人静静地接住。这个接住的动作，是旧互动网络中的一个关键节点——我们把它叫做“共同承受”节点。当 B 在这一刻改用“那你早点睡”来回应时，他实际上执行了一个不同种类的动作：他把 A 的疲惫从一个“需要被接住的体验”，转化为了一个“需要被管理的问题”。这个转化使 B 免于被 A 的情绪状态卷入——这是新互动网络得以维持的关键节点：B 在这套“助人者”姿态中获得控制感和被需要感。同时，它使 A 的体验从“被接住”转变为“被评估”。

如果这是孤立的一次互动，也许毫无影响。但如果 B 在后来的类似场景中持续采用这个新动作——A 表达负面情绪→B 给出行为评估或改进建议——那么原先那个承载着“疲惫被另一个人静静接住”这种体验的旧节点，就会逐渐被新节点所替换。替换持续到某个阈值，A 会发现向 B 表达疲惫已经不再产生被陪伴的体验，只产生一种新的压力：我在被评估。于是 A 不再表达，旧网络在这一点上干涸。

这个过程的关键结构是：B 的新互动模式并非旧模式的“缺失”或“衰退”。它自身是一张被积极维持着的网络——B 从“助人者”身份中获得效能感，这套模式可能来自他原生家庭中习得的照料方式，也可能来自他职业训练中培养的诊断习惯。新网络对旧节点的接替，不是消极的撤走，而是一个同样强悍的、自身被很好维持着的替代结构，在两个人在场的那一点上，完成了对同一交互时机的占据。两套网络的不同节点在同一个互动时机上竞争，谁赢了谁就继续巩固自己的维持条件。信任的蚀解，就是旧网络节点在反复竞争中无法再被稳定重现的那个最终阶段。

这里需要一段方法论警示，以避免分析滑向价值判断。上述描述很容易被误读为：旧网络（共情、接纳）是关系的理想状态，而新网络（管理、评估）是关系的衰退。这种误读恰恰会把互蚀分析降格为怀旧叙事。旧网络与新网络的竞争，不提供道德方位。互蚀分析揭示的是竞争的结构，而非竞争的价值。旧模式的“不设防”同样依赖于对某些互动可能性的蚀除——例如，它蚀除了伴侣将彼此视为“需要各自独立承担情绪管理责任的现代个体”这一压力。互蚀分析只揭示：在同一个交互时机上，两种不同的网络各自执行了自己的维持动作，竞争的净结

果是旧节点被新节点替代。至于这种替代是“好”是“坏”，不是互蚀理论能回答的事——它是卷入这场蚀战的双方必须在自身网络内部做出的判断。

为了进一步巩固这一理论中立性，让我们对称地考察一个“坏”网络被“好”网络蚀解的案例。一段控制型亲密关系中，旧网络的核心节点是羞辱与贬低：当一方表达自己的成就时，另一方惯常的回应是“这有什么了不起”或“你不过是运气好”。这个回应执行了一个维持控制关系的动作——它削弱对方的自我价值感，从而巩固控制方的支配地位。现在假设，被控制者逐渐建立起一组外部支持性友谊。在这些友谊中，当她表达成就时，朋友的回应是明确的肯定和共情：“这真的很不容易，你做得很好。”起初，控制型节点在她的互动场中仍然占据主导——她回到家中，仍然被贬低。但随着支持性互动节点的累积——从每周一次的朋友聚会，到每天一次的电话，到工作中同事的持续认可——控制型互动节点的相对出现频率逐渐下降。关键不在于她“意识”到了控制关系有问题（认知转变），而在于支持性互动的节点已经占据了越来越多本该由贬低互动占据的时间与注意力。贬低动作的维持条件——她还在那里、还在听、还在试图解释自己——被一点一点地蚀掉。最终她不再解释，贬低节点失去了它在互动场中的执行时机。控制关系没有被她“战胜”；它被支持性网络从内部一个节点一个节点地替掉了。

这个对称案例的意义在于：它显示互蚀本身不偏向任何道德方向。它可以蚀解信任，也可以蚀解控制。它的逻辑纯粹是结构性的——两张网络在同一个交互时机上竞争节点占有权，赢方将输方的维持条件注销。至于哪张网该赢、哪张网该输，那是卷入蚀战的行动者自己在各自网络内部做出的判断，互蚀理论不替他们做这个判断。

### ■ 3.3 货币与知识：制度的替代性空心化

社会制度领域的蚀案例或许最为直观，因为它常常以“解体”或“衰落”的宏观叙事出现。但解体叙事同样落入了消极描述的陷阱——仿佛一个制度只是“不再被需要”或“被外力打碎了”。细察之下，社会制度的瓦解同样遵循着正面的蚀逻辑。

货币是绝佳的例子。传统货币理论将货币崩溃归因于恶性通胀或政权更迭，但在一场典型的本币弃用过程中，最先出现的往往不是一个“人民不再相信货币”的心理事件，而是另一个交换网络的缓慢铺设。以阿根廷 1989—1995 年间的美元化过程为例：美元并不是在某一天突然替换了比索；它起初只是在某些特定的交易品种（如房产、大宗家电）中被用作报价单位。随后，这些特定品种的交易网络逐渐扩张，而比索的支付网络在这些品种上是逐步被放弃的。关键不在于主观“信心”，而在于交易网络节点的逐点转移：当一个农民发现他用比索计价的收入换不到化肥时，是因为化肥商已经被上游供货商逼着用美元结算了——每个节点都是在自

身维持压力下完成替代的，并不是基于一个全国性的集体判断。比索的“价值”是在一个节点一个节点的替代中被蚀掉的。它的硬度不是被砸碎，是被另一个一样硬的网络从内部撑开毛孔，然后一点一点地被风化成沙。这个过程有结构、有节律、有先后次序——总是从大宗商品和跨境贸易这类本币维持成本最高、交易网络最脆弱的节点开始——绝不是突然变成废纸。

还有一个常被忽视的蚀案例：手工技艺知识体系的蚀。历史学家汤普森（Thompson, 1963）在《英国工人阶级的形成》中曾描述过工业化如何改变手工艺人的时间感与工作节奏，但他没有展开的一个侧面是：新的工厂制度不是靠“禁止”手工艺来摧毁它的，而是靠把旧技艺链条中的关键步骤替代为可测量、可分解的操作。以英国手工蕾丝业在十九世纪的衰亡为例，机织蕾丝在初期质量远不如手工，但它不需要取代整张网——它只需要先替代那些耗时最长、最依赖匠人默会判断的特定工序，比如底网的编织。一旦这道工序被机器接管，手工匠人的整个生产节奏就被打断：她不再能以同样的成本完成整件作品，于是她要么放弃这门手艺，要么只做那些机器还做不了的精加工（通常利润不够维持生计）。新一代学徒不再学习底网编织——那个最核心、最能传递“手感”的节点——于是几年之后，即使有人想复兴手工蕾丝，也找不到会织底网的人了。整个知识体系不是在某一刻“失传”的；它是一个节点一个节点地被替代掉的。替代它的不是更好的手工，而是一套完全不同的生产网络。这套新网络不需要理解旧网络的整体逻辑；它只需要在关键节点上提供一个与之等效的操作结果。当替代节点累积到一个临界密度，旧知识体系的再生条件就被彻底封死了。

近年的研究提供了更多类似的经验旁证。在社会学与科学技术研究（STS）的交叉地带，已有学者详细追踪了特定行业中专家知识的微观瓦解过程——比如医疗诊断中，当算法系统逐项接管了影像判读的具体子任务后，初级医生的默会诊断能力如何在短短几年内从“尚可维持”变为“无法再生”。算法没有“禁止”医生学习诊断；它只是把诊断链条中那些新手最需要反复练习才能掌握的中间节点，变成了一个黑箱的输出值。旧知识体系的再生条件，就在这些节点的静默替代中被注销了。

三个领域的证据汇总指出的是一致的形式：一个稳定形态的耗散，具有它自身的正面网络结构。这个结构的运作逻辑与它正在替代的旧形态的生成逻辑是同构的，只是节点内容被置换。它不应被称作“衰退”或“解体”，它应该有一个不做任何预先消极判断的名字。

#### ■ 四、互蚀的不可还原校验：与既有理论的对质

一个概念要确立自己的理论地位，不能只靠提出一个漂亮的新词。它必须通过与最接近的既有理论的正面交锋，证明自己切出了这些理论无法覆盖的辨别面。本章将与六个最有可能“吞掉”互蚀的理论传统进行正面对质：生态学中的竞争性替代模型、网络科学中的级联失

效理论、热力学熵、路径依赖的制度变迁理论、黑格尔的扬弃、德里达的解构。这不是为了展示互蚀与经典有何不同，而是为了在每一个交手中逼问：互蚀是否只是既有概念的隐喻延伸？它是否提供了不可被替代的解释力增量？

## ■ 4.1 与生态位竞争性替代的对质

生态学中的竞争性替代模型，是互蚀最危险的对手。这个模型描述了外来物种如何通过占据与本地物种相同的生态位，逐点夺取后者的资源通道，最终导致本地种群在特定区域内灭绝。其逻辑与互蚀惊人地相似：都是异质网络（新物种的生态位网络）逐点替代旧网络节点（旧物种与环境的物质—能量交换通道），都有临界阈值（最小可存活种群密度）。一个审稿人完全可以质问：“你所谓的互蚀，不过就是生态学里竞争性替代原则在物理和人际关系上的类比应用。这个现象早就被命名了，你的概念收益何在？”

这是一个必须被正面接住的质疑。生态学竞争性替代模型与互蚀之间，有两点根本差异。

第一，竞争性替代模型中，新旧网络争夺的对象是同一类资源：两个物种争夺同一种食物、同一片栖息地。替代之所以发生，是因为新物种在利用这一种资源时效率更高，从而在同一个生存维度的竞争中将旧物种挤出。但互蚀模型中，新网络不需要和旧网络“吃同样的东西”。审视模式不需要和共情模式争夺同一种“情感资源”；它只需要能执行一个结构上等效的动作——在 A 表达疲惫的那个交互时机上做出一个回应——来占据旧节点曾经占据的功能位。新网络的“资源”来自另一套完全不同的发生机制（例如，B 从“助人者”身份中获得的控制感），它与旧网络争夺的不是资源内容，而是节点位置。这使得互蚀比竞争性替代更抽象：它不是在同一个生态位内比效率，而是在同一个交互场中比较“谁能占住这个功能位”。这是生态位理论无法直接覆盖的一层。

第二，也是更关键的差异：互蚀理论的核心解释对象，不是替代是否发生——竞争性替代和互蚀都能预测替代会发生——而是替代的顺序与临界阈值的位置。生态位竞争性替代模型在预测“哪个替代会先发生”时，主要依据资源利用效率的梯度：最先被替代的，是新物种效率优势最大的资源节点。互蚀模型预测：替代会系统性地从旧网络自身维持度最低的节点开始，而不是从外部入侵者效率最高的节点开始。

这里需要一段“差点被替换”的硬核论证。如果我只是在用竞争性替代的框架重述人际关系蚀变，我应该会预期看到：审视模式首先在“审视特别擅长处理”的互动类型中接替共情模式——比如在需要快速做出行为决策的场景中，审视的效率优势最明显。但我在典型化情境中实际观察到的是：审视首先接替的节点，并不是审视效率最高的那类互动，而是旧网络（共情）在那一刻恰好最脆弱的互动——比如 A 表达疲惫时，B 此刻恰好缺少情感余裕来接住。它



的入点不是新网络哪里最强，而是旧网络哪里此刻最弱。同样，在阿根廷美元化中，美元首先侵蚀的不是美元流通最方便的交易类型，而是比索网络中维持成本最高的节点——大宗商品和跨境贸易。在手工业案例中，机器首先替代的不是机器效率优势最大的工序（机器的全部工序都有人工难以企及的效率），而是耗时最长、最依赖匠人默会判断的底网编织工序。如果是效率主导的竞争，机器应该一次性接管所有工序——既然每道工序机器都比手工快得多，效率梯度的预测是全线同时替代。但历史上它走的是一个接一个、从最脆弱的旧节点开始的顺序。

反过来说，如果“旧网络的脆弱时刻”在经验上和“新网络效率最高的时刻”是同一回事——比如，共情在 B 疲惫时最脆弱，而审视恰好也在 B 疲惫时最省力——那么互蚀预测就会被竞争性替代预测覆盖，我的区分就崩塌了。要守住不可还原性，我必须提供一个独立的判据来拆开这两个变量。这个判据是：在一个节点上，旧网络的维持脆弱度，并不等于新网络在该节点上的功能效率。旧网络的脆弱度，由其维持所需的主体即时投入强度决定——共情需要 B 此刻有情感余裕，需要 B 主动抑制“给出建议”的习惯冲动，需要 B 承受 A 的负面情绪带来的轻微不适。这些投入的强度独立于审视动作的效率——审视在任何时刻都省力，不管 B 是清醒还是疲惫。因此，当 B 精神状态很好的时候，共情的脆弱度低，按照互蚀模型的预测，审视应该很难在这个时刻接替共情。而竞争性替代模型会预测：既然审视在任何时刻都高效，它就应该在任何时刻都能接替，B 的状态不应该影响替代的发生概率。这个预测差异在原则上是可以检验的：追踪足够多的“A 表达负面情绪→B 回应”的互动回合，将回合按 B 当时的疲劳程度分层，观察审视替代共情的概率是否在 B 疲劳时显著更高。如果是，互蚀对竞争性替代就有一个不可还原的预测增量——因为竞争性替代模型无法独立解释为什么 B 的状态（而非审视自身的效率）是替代发生与否的显著预测变量。

由此，互蚀不仅描述“替代发生”，它还能解释为什么替代总是按照“旧网络自身的脆弱梯度”来排序，而不是按照“新网络的效率优势梯度”来排序。这个区分如果成立——也就是说，如果“替代顺序由旧网络维持脆弱度分布决定”的预测在足够多的跨域案例中不能被竞争性替代的效率排序预测所覆盖——那么互蚀作为独立概念的不可还原性就得到了第一个硬支撑。

## ■ 4.2 与网络级联失效的对质

网络科学中的级联失效模型，是互蚀需要面对的第二个对手。巴拉巴希（Barabási, 2016）等人的经典工作已经精确描述了：当一个网络中的关键节点被移除后，失效可以沿着网络的连接结构传播，最终导致整个网络的崩溃。电力网络、互联网、金融网络中的级联失效，



都展示了“节点故障如何引发系统性崩溃”。这听起来和互蚀很像：旧网络节点被“替代”（在级联失效中是被“移除”），然后旧形态逐渐瓦解。

但这里有一个微妙而关键的差异：级联失效模型中的“节点”，是被移除的——它因为过载、故障或攻击而停止工作。互蚀模型中的“节点”，是被替换的——它仍然在工作，只是被一个新网络的动作占据了同一个功能位。这个差异不是修辞，而是两种不同的存在论图景。

在级联失效中，网络崩溃的动力学来自旧网络内部失效的传播：节点 A 坏了→流量转到节点 B→节点 B 过载也坏了。每一步“坏死”都是旧网络自身状态的一部分。在互蚀中，节点替代的动力学来自新网络的正面组织力：节点 A 没有被“破坏”，它被新网络的动作“占领”了。新动作在功能上等效于旧动作（它也能回应“我今天很累”），但在网络归属上属于另一个生成机制系统。

这个差异在人际关系蚀变中的表现是：当 B 用“那你早点睡”替代“是啊，真是够呛的”，共情节点并没有“故障”或“过载”；它仍然是一个 B 完全有能力执行的动作。它之所以不再出现，不是因为旧网络内部传导了某种“不能执行”的状态，而是因为新网络的另一个动作在这个时机上抢占了执行权。网络级联失效可以描述旧网络内部的雪崩，但它无法描述“一个同功能位被另一个网络的等效节点持续占据”这个动作本身。级联失效是旧网络自身内部状态的变迁；互蚀是两个网络在功能位上的正面替换。因此，级联失效模型无法单独描述“新网络从哪里获得它的组织力，为什么它的节点能持续插入旧网络的功能位”。而这一点，恰恰是互蚀概念的核心。

## ■ 4.3 与热力学熵的对质

热力学熵是互蚀需要面对的第三个对手，而且是一个资格更老的对手。热力学第二定律已经告诉世界一百多年了：孤立系统的熵永不减少。换句话说，有序形态的耗散——微观状态的均匀化、可用能量的贬损——是宇宙的宿命。如果互蚀只是在这个意义上说“耗散是一个不可逆的过程”，那它不过是把热力学的常识换了个哲学外壳。所以本章要做的，不是否认熵，而是给出一个熵不能给出的东西：耗散的微观路径。

熵是一个宏观的、统计性的量。它描述的是系统整体有序度下降的程度，但它不区分“耗散是从哪个部位开始、沿着什么顺序推进的”。一杯热水在房间里冷却，熵增定律告诉你的只是：最终水和房间会达到同一温度。它不告诉你——也不关心——热量是通过杯壁的哪个位置最先流失的、杯沿和杯底之间是否存在一个可以测量的温度梯度传播方向。这些“路径”问题不在热力学的理论视野之内，因为热力学把系统当作一个整体来统计，不追踪微观组分的个别轨迹。

互蚀是一个关于微观路径的理论。它要解释的是：在一张维持形态的生成机制网络中，耗散是从哪一个节点最先开始的？替代蔓延的次序是由什么决定的？在什么临界条件下，旧形态的维持条件从“局部被替”跨越为“全局不可自复”？这些问题，熵的宏观公式回答不了。熵增说“它会冷”——互蚀说“它会沿着脆弱的杯沿先冷，再顺着杯壁往下走”。一个医生知道病人最终会死（这是熵），但一个好医生还需要知道病灶最先出在哪个器官、扩散顺序是什么（这是互蚀）。

更关键的是，在开放系统——比如人际关系或社会制度——中，秩序不但没有净衰减，反而可能局部增加。普里高津（Prigogine）的耗散结构理论已经充分论证了这一点。在开放系统中，熵增并不必然导致结构崩溃；真正导致崩溃的，是在特定节点上，维持那个结构的正反馈循环被另一个结构的正反馈循环替掉了。这个替代机制，不是热力学能描述的——热力学描述的是“能量品质的下降”，不是“两个不同结构网络在功能位上的节点竞争”。因此，互蚀不是熵的隐喻延伸；它是熵所缺掉的微观叙事。

## ■ 4.4 与路径依赖制度变迁的对质

诺斯（North, 1990）和阿瑟（Arthur, 1994）的制度变迁理论，尤其是“路径锁定”与“路径解锁”的论述，是社会科学中处理“形态如何从一个稳定状态转换到另一个”的最成熟框架之一。一个掌握这套理论的读者完全可以说：“互蚀所说的‘节点替代’，实质上就是制度的路径解锁过程。旧制度的自我强化机制被削弱，一个新的自我强化机制开始主导。用‘互蚀’这个词，只是把‘路径依赖的瓦解’重新包装成了一套看似新颖的形而上学语言。”

这个质疑的核心是：路径依赖与锁定效应本身就有“正向反馈逐渐减弱→被另一个正向反馈机制替代”的理论结构。那么，互蚀的新意何在？

我接受这个挑战的精确表述是：路径依赖理论能够解释“为什么旧制度难以被撼动”（自我强化锁死），也能够解释“在什么外部冲击下旧制度可能解锁”（参数剧变打破正反馈循环）。但它不擅长解释的是：在没有任何明显外部冲击的情况下，旧的自我强化机制是如何被一个新的机制从内部静悄悄地蚕食掉的——尤其是在替代的初期阶段，新机制尚未形成足够强的正向反馈，旧机制的正反馈也未见明显衰减。互蚀的独特解释力恰恰落在这一段“无声接替期”。

晚近的渐进制度变迁理论（如 Mahoney 和 Thelen 等人在 2010 年前后系统阐述的“制度漂移”“制度转换”等概念）确实关注了底层执行惯例的静默位移。一个批评者可以据此争

辩：“地方惯例网替代执法节点，这难道不正是制度漂移的标准案例吗？你的互蚀理论贡献了什么？这套理论已经说过的？”这里我必须给出一个不同于“换词”的预期差。

渐进制度变迁理论的核心解释变量是行动者的策略性重新解释：在既有规则的字面不变的情况下，执行者利用规则的模糊性，将规则的适用范围、执行强度或执行对象朝有利于自身的方向挪动。这里的替代动力来自行动者有意或无意的“重新解释”。互蚀理论的核心解释变量是节点位置的竞争：新惯例替代旧规则，不需要任何行动者“重新解释”规则；它只需要在旧规则本该被执行的时机上，用一个结构等效的新动作占住那个执行瞬间。在法规空心化的案例中，地方官员在默许“先上车后补票”时，他不需要对环保法规做出新的解释——他可以完全承认法规的原文是正确的、他无意修改法规——他只是在这个具体的执行节点上，放了一个不同来源的动作。这个动作不是对旧规则的“重新解释”，而是用一个惯例网的动作替掉了规则网的动作。二者的动力源不同：制度漂移靠的是规则的语义弹性，互蚀靠的是执行时机的抢占。这个差异可以导出不同的预测：如果制度漂移是主因，我们应该在规则语义模糊的条款处最先看到替代；如果互蚀是主因，我们应该在旧规则维持成本最高——执行者最不愿意费力执行的节点处最先看到替代，无论该条款的语义是清晰还是模糊。这个预测差异，是互蚀在制度变迁领域不可还原为既有理论的硬支撑。

由此可以明确：互蚀不是路径依赖理论或渐进制度变迁理论的替代品，而是它们的补充。路径依赖处理“锁定—解锁”的宏观叙事，渐进制度变迁处理“规则语义弹性下的策略位移”，互蚀处理“在任何一个规则执行时机上，一个异质网络的动作如何能正面抢占这个时机”。三者的解释变量和所预测的替代初始入点各不相同，不可互替。

## ■ 4.5 与黑格尔扬弃的对质

黑格尔（Hegel, 1807）的“扬弃”是过程哲学的核心动力学概念。扬弃保留了被否定的东西，同时把它提升到更高的统一中——规定性在否定中才成为规定性；本质映现在他者中才成为本质。这套逻辑赋予了“变化”一个进步的方向：每一次否定都是一次升华，旧的东西没有被消灭，它被转化为了更高真理的一个环节。

互蚀与扬弃的根本分歧在于：互蚀所描述的替代，并不承诺升华。在扬弃的逻辑里，旧规定性被“保留”在更高统一体中——奴隶意识被提升为自我意识的一个环节，主观精神在客观精神中获得了实体性。但在互蚀的逻辑里，旧形态没有“转入”新形态；它散掉了。审视模式接替共情模式之后，共情互动并没有被“保留”在审视模式内部，成为审视的某个“更丰富的环节”。审视模式完全可以不包含任何共情成分——它甚至可能彻底排斥共情（“你的情绪只

是你需要被管理的症状”）。被蚀掉的旧形态也没有在更高处复活；它只是失去了自己能被持续认出的条件。

这一笔出入不是措辞层面的分歧。它关系到“矛盾”在两种理论中的不同功能。黑格尔的矛盾是自我超越的动力：两个对立的规定性在斗争中互相暴露彼此的局限，最终被一个更高的第三者综合。矛盾有方向，有记忆，被扬弃的东西参与了更高综合的构成。互蚀的矛盾没有方向，没有记忆。两张争夺同一节点的网，谁也不综合谁。赢的一方把输的一方的维持条件从这个节点上注销掉，输的那一方的“内容”不参与赢的那一方的“构成”。它只是从交互场上消失了——就像那一杯本打算被喝掉的水，被替换为待分析的样品之后，口渴的体验本身没有成为“分析”的一个环节，它就是不再被回应了。

因此，将互蚀装入扬弃的框架，会犯把“解雇”当“升职”的错误。这不是否定扬弃在描述“精神自展开”时的解释力——真正的理论总有其适用范围。本章要强调的是：在描述“两个异质网络争夺同一功能位，赢方不承继输方的任何内容”这类过程时，扬弃的框架不够用，互蚀是一个更精确的命名。

## ■ 4.6 与德里达解构的对质

德里达（Derrida, 1972）的解构是当代理论中与互蚀最形似的一个版本。解构指出文本意义的不稳定性，指出任一结构都依赖着它所排斥的“他者”来维持自身；在延异（différance）的无尽运动中，固定的意义总是被拖延、被涂抹。这确实与互蚀有某种家族相似：都在讲形态的不可最终锁死。

但解构的战场是意义的痕迹，互蚀的战场是维持条件的节点。解构所描述的不稳定，是符号系统内部的先天裂隙——无论你怎么努力固定一个意义，它总会因为符号的互指性而滑开，这是文本性的固有特征。蚀的不稳定不是先天的，是后天的：它不是因为网络内部有裂隙，而是因为另一个网络实实在在地攻进来了。意义会滑开，不是因为有外部的另一个文本系统在占据你的句子中的每一个词；信任会蚀穿，不是因为信任本身有内部的延异，而是因为那个把信任撑住的互动回合被另一个叫做“审视”的回合替掉了。

解构不关心节点替换，因为它不关心发生网络的具体维持——它关心的只是意义的条件可能性。蚀关心的恰恰是发生网络的具体接替历史：是哪一个节点最先被换、替代顺序是什么、临界值在何处。这些在解构的范式下都不可见。把蚀命名为解构，就是把土木工程的坍塌归因于材料内部的原子振动——技术上不错，但切不中真正塌在哪一根梁上。解构可以告诉我们，信任之所以能被蚀，是因为信任作为符号结构本来就携带着被延异的可能。但它无法告诉我

们：为什么这一对伴侣的信任是在这一刻、以这种顺序、从这一个互动节点开始蚀的。而这些，恰恰是互蚀理论要回答的问题。

## ■ 4.7 校验小结

五个对质得出同一个结论：互蚀不是已有理论的隐喻延伸。竞争性替代模型无法刻画“非资源性功能位的占据”与“替代顺序由旧网络脆弱度决定”，级联失效无法刻画“新网络正面的节点组织力”，热力学熵不追踪耗散的微观路径，路径依赖理论及其渐进变体侧重于策略性重新解释或宏观正反馈锁定而漏掉了执行时机抢占的微观动力，黑格尔扬弃无法容纳“不承继任何被蚀形态内容的替代”，德里达解构不关心维持条件的具体接替历史。互蚀的不可还原性，不是来自它比这些前辈更深刻——它不如黑格尔精妙，不如德里达激进，不如网络科学精确——而是来自它看了他们各自惯于不看的方向：被蚀的次第、入点与临界。这条辨别轴是互蚀理论的真正收益。

## ■ 五、互蚀的操作化定义与前文案例的形式提取

在展示了互蚀在三个领域中的形态并与六个对手正面交锋之后，现在可以把前文案例中已经显现的形式结构提取为一组清晰的操作骨架。这个定义不是事前套用的模具，而是事后提取的形式——它是对第三章案例中反复重现的那个结构的命名。

互蚀指的是这样一种过程：当维持某个稳定形态  $X$  的生成机制网络  $NX$  中的若干关键节点，被另一个异质的生成机制网络  $NY$  中的对应节点逐次替代，且这种替代使得  $X$  在原有维持条件下无法再被稳定重现时，我们说  $X$  被  $NY$  蚀解了。

这个定义包含四个硬性判据：

(i) 节点替代：蚀是通过新的网络节点替换旧网络节点的具体动作实现的，而不是旧网络自身“松脱”或“衰减”。在物理中，节点就是相互作用顶点——例如维持粒子稳定激发态的那个耦合通道。在人际关系中，节点就是特定类型的互动回合——例如  $A$  表达疲惫时  $B$  做出情感接住的回应。在制度中，节点就是每一次承续旧规则的具体执行动作——例如环境法规中每一次对违法企业的实际惩处。新旧网络的“节点”是具体可指认的事件，不是抽象属性。

这里需要澄清“节点”与“动作”的关系，以免读者困惑。在本章中，“动作”是节点的内容——它指在一次具体的互动时机中，网络实际执行的那个回应方式（比如“是啊，真是够呛的”或“那你早点睡”）。而“节点”是这个动作所占据的那个结构性位置——即在“ $A$  表达疲惫→ $B$  回应”这个互动时机上的那个功能位。同一个功能位（节点），可以被不同网络的不同动作所占据。当本章说“节点替代”时，指的是新网络的一个动作占据了旧网络动作曾经

占据的那个功能位。替代的不是动作的内容相似性——审视和共情在内容上没有相似性——而是它们占据了同一个互动时机的同一个功能位：回应 A 的情绪表达。这个区分解释了为什么互蚀可以发生在内容完全不同的两个网络之间：它们不需要吃同一种资源，它们只需要争同一块地。

(ii) 异质性：替代节点必须来自一个与旧网络在本体构成上不同的另一套发生机制（另一种咬合方式），否则只是旧网络自身调整，不构成蚀。在人际关系案例中，审视互动的维持网络——B 的“助人者”身份从中获得效能感、原生家庭习得的照料方式、职业训练中的诊断习惯——与共情互动的维持网络（B 作为“陪伴者”的身份、被接纳的体验）是两张不同的网，各有其独立的维持条件。异质性判据将互蚀与旧网络内部的参数调整区分开来：如果 B 只是更频繁地使用共情（共情网络增强），那不是蚀；当 B 开始用另一套互动逻辑占据同一个互动时机时，蚀才发生。

(iii) 逐次性：替代不是一次性的全域切换，而是按节点依次发生。正是这种逐次性使得蚀过程有它自身的时间节律和阶段性，也使得“替代顺序”成为互蚀理论最核心的解释对象。在阿根廷美元化中，替代顺序是从大宗商品定价节点开始，逐步蔓延到日常消费品。在手工业案例中，替代从耗时最长的底网编织工序开始。在人际关系中，替代从旧互动此刻最脆弱的情感表达回合开始。这个顺序不是随机的，而是系统性地沿着旧网络自身维持脆弱度的梯度展开。

(iv) 不可逆临界阈值：当被替代的节点密度超过某个临界值时，旧形态即使在新节点余下的空间里也已经无法获得足够连贯的维持条件，从而进入不可自复的坍缩。这个临界阈值的存在是可检验的：可以在模拟或历史案例中寻找“旧形态最后一处稳定再现代为止”的时点，并回溯该时点前已发生的节点替代比例。

这四个判据合在一起，把互蚀从隐喻的松散状态中抽出来，赋予它可被独立识别、可被操作化测量的逻辑骨架。

## ■ 六、为什么不能是阴阳？——生蚀对等交织的反驳

在论证蚀先于生之前，必须正面回应一个强有力的对手立场：也许生与蚀是两个同样原初、永远互相缠绕的动作，谁也不比谁更原始。就像中国哲学里的阴阳，或者赫拉克利特的斗争与和谐，世界的真相就是它们的永恒对抗。本章为什么非要咬死“蚀先于生”不可？

我接受这个挑战。如果生与蚀只是两个对称的基础动作，那么本章可以把互蚀确立为与“生”对等的另一半，交出一个对称的存在论方案——一边是建立，一边是替换，阴阳互



耦。这个方案比起“蚀先于生”的命题，要温和得多，也更容易被接受。它不需要拆过程存在论的台，只需要帮它补上漏掉的半边。我为什么不做这个更省力的选择？

因为在阴阳互耦的方案里，生仍然有一个独立的、不依赖于蚀的存在论地位。生就是生——它有自己的动力、自己的逻辑、自己的纯粹形态。蚀是另一个东西，与生对立统一。但在本章第二章到第四章反复展示的那个形式结构中，任何一个“生”的动作，其内部都已经是蚀的结构。量子场论中粒子通过锁定某个耦合通道而获得稳定——这个锁定的动作，本身就排他性地压制了其他所有可能的耦合通道。亲密关系中伴侣通过无数次的共情互动“建立”了信任——这每一次共情，都同时在蚀除着另一个可能的回应方式：在那一个交互时机上，B没有选择沉默、没有选择转移话题、没有选择建议——他选了共情，也就把其他所有选项在那一刻的出场权给蚀掉了。货币的“硬度”来自亿万次交易的沉淀——这每一次以本币结算的交易，都在那个交易节点上蚀除了以物易物、以外币结算、以新信用单位结算的潜在替代通道。

换句话说，生不是“创造了一个全新的东西”——它是从一堆同时存在的潜在生成机制网络中，通过反复占据交互节点，把其中一张网扶成了主导形态，同时把其他网的节点维持条件系统性地排除了。这整个动作的结构——竞争性排除——本身就是蚀。在生的内部，除了蚀，没有别的构造成分。生是蚀的某种特殊相（竞争陷入僵局的相）的外表。

阴阳互耦方案的问题在于：它把“生”当成一个可以与蚀对等并列的独立动作。但在本章展示的证据中，找不到任何一个“生”的案例，其内部不包含对其他可能存在形态的节点剥夺。生没有一个不依赖蚀的“纯净内核”。反过来，蚀却可以不伴随生——粒子的衰变、信任的瓦解、货币体系的崩盘，都是蚀在不产生新形态（或只产生一个简单的新形态，如衰变后的轻子—中微子对）的情况下独立运作的案例。生是蚀的一种特殊情况（僵局），蚀不是生的一种特殊情况。

因此，“蚀先于生”不是在两个对等的本源之间选了一方，而是在否认对称性的存在。蚀只有一个——生只是蚀的一种局部的、暂时的、有条件的僵持。这个判断当然比阴阳互耦方案更重、更危险、更容易被反驳——但它也更诚实。如果在所有可观察的案例中，生都以蚀为自己的内部构成机制，而蚀不需要生就能独立运作，那么诚实的方式就是把不对称性命名出来。

## ■ 七、蚀先于生：一个从案例引发的命题

上一章论证了生与蚀不是对称的两个本源。这一章将在这个基础上，正面提出蚀先于生的命题，并论证它的存在论根据。

### ■ 7.1 生本身就是蚀的一种形态

当一个形态通过生成机制网络被立起来时，这个“立起来”的动作，本身就执行了一次大规模的排他性选择。在量子场论中，粒子通过“锁定”某种相互作用通道而获得稳定时，这个锁定本身就排他性地压制了其他潜在的耦合通道。电子的磁矩之所以是现在这个数值，不是因为它生来如此，而是因为在与各种虚过程的持续蚀战中，只有某些通道的贡献能被抵消，而另一些通道则被排除在了稳定化的合力之外。换句话说，粒子的“生”本身已经是蚀的结果：是无数可能的耦合通道被竞争性排除之后剩下的那个暂时平衡态。生不是原初点，生是蚀的某种结算。

在人际关系中，这个结构更为显眼。一段亲密关系初建时双方所做出的相互承诺，本质上是在双方的互动场里，选择性地蚀去了其他可能的关系形态——不是那种“我们以后不和别人好”的社会性契约，而是更深的一层：每当你们按照这个特定的默契——比如“你累的时候我陪”——来完成一次互动时，你实际上就参与了一次对另一种被你们双方共同放弃的互动模式的排他性蚀除。这不是暴力；它只是一次次被选中的动作使未被选中的动作失去了在你们之间的场域内生长的交互时机。无数次这样的小型蚀积累，才把两个人的“伴侣形态”维持为一个看起来理所当然的实体。关系的实体性是蚀的累积赠品，而非生的直接产物。

这里需要补上一个论证，来说明为什么这种“排他性选择”可以被合法地叫做“蚀”，而不只是“生的内在代价”。关键在于：那些被排除在外的“其他可能形态”，并不是虚无。它们是真实的潜在生成机制网络。这个“真实”不是认识论上的——不是我们自己想象“如果没有X，可能有Y”。它是存在论上的，理由如下。在量子场论中，虚粒子云中的每一个虚过程都在可观测效应中留下痕迹——电子的反常磁矩，就是虚粒子涨落对电子本征值的可测度修正。这些虚过程不是我们“想象出来”的可能性，它们是真实在场论拉格朗日量里写着的耦合项，有真实的强度参数，缺席了就凑不出实验测到的那个磁矩数值。同理，在人际关系中，当B选择用共情回应A的疲惫时，那个被放弃的“审视回应”——“那你早点睡”——也不是虚无。它在B的回应选项库里是一个真实存在的动作模板，有它自己的生成机制支撑网络（B从原生家庭习得的照料方式、职业训练中的诊断习惯），只是在这一刻，它没有获得B的执行。它不是不存在；它是没被选中。正如在量子场论中，那些没有被锁定的耦合通道不是没有数学结构——它们有，只是在这个能量尺度的竞争中被抑制了。它们随时可以在条件变化时（比如B自己极度疲惫，共情通道的维持成本突然飙升）夺取这个交互时机。

因此，当“生”的动作把主要的交互节点赋予形态X时，它也就同时执行了对形态Y、Z的网络的节点剥夺。Y和Z原本也可以——在另一种条件下会——占据这些节点，让它们自己的维持网络得以展开。而X的每一次重复，都在让Y和Z失去这次交互时机。这个剥夺动作，其形式结构与第五章定义的互蚀——异质网络在同一个功能位上竞争节点占有权，赢方将输方的

维护条件注销——完全一致。它不是另一种东西；它就是一方的集中高效的蚀。所谓“生的内在代价”，不过是我们对一次成功的蚀战在事后给予的正面命名：我们只看残余的那个，把残余叫做生；我们不看被剥夺的那些，把它们叫做“没有实现的可能性”。但从发生的角度来看，那次剥夺本身就参与了本章定义的那个结构——它是蚀。

## ■ 7.2 河床与沙洲

由此，可以将生与蚀的关系重新表述为：生是蚀的一个特殊相的结余。不是先有发生后有蚀，而是在蚀的汪洋中，某些局部区域因为替代网络间的竞争陷入了相对持久的僵局，从而使某一形态在一段时间内得以持续重现——我们把这个僵局的生产性侧面称为“生”。僵局一旦被打破——无论是由于新耦合通道的出现、新互动模式的入侵、还是新惯例节点的铺设——形态就散。蚀从来都在那里，生不过是蚀的河床上偶然隆起的一小片沙洲。

这个倒转的命题不是文字游戏。它直接导出了一种全新的研究姿态：不要问“这个东西是怎么生成的”，要问“这个东西是踩在哪些被蚀掉的其他东西上面才显得在那儿”。不要在积木已经咬住之后去追它的咬合史，要在积木还没被换上之前，去看那些本可能立在这里却被替换走的那些积木的幽灵。那些幽灵不是虚无，它们是每一块现存积木之所以能被看见的真实背景负片。由于我们从来只能看见留下来的生，看不见被蚀走的万千可能，我们就产生了一个认知幻觉：生是常态，蚀是例外。其实，蚀是常态，生才是例外——生是蚀的竞争火焰中偶然没被烧着的那一小撮。

这个命题也意味着，没有任何形态的稳固能被它的内在性质担保。稳固不是形态的属性，稳固是蚀的战局在一定参数下刚好进入相持阶段的表现。在相持阶段，替代网络的攻势暂时无法突破旧网络的节点更新率，于是旧形态得以稳定下去；每一个被我们称为“客观结构”的东西，其实都还活在蚀战的防御工事后面。防御工事一旦被局部突破，那个结构的“客观性”就可以在极短时间内崩解——就像近年某些地区银行挤兑中，储户的信任结构是在几小时内被社交媒体上的信息节点替代网蚀穿的。没有人撒谎，没有人恶意煽动，只是信息互动网络以电速完成了节点替代，而柜台前的现金防御速度来不及跟上。

## ■ 7.3 蚀的动力源：一个不得不交代的追问

到此为止，一个尖锐的问题不应再被回避：如果蚀这么普遍，那蚀自身的动力源是什么？如果说互蚀是维持条件被替代的过程，新网络又是从哪获得了入侵旧网络节点的组织力？这关系到一个逻辑闭环：如果蚀的动力最终来自“生”（新网络也是被“生”出来的），那么蚀的

优先性就是空的——它不过是两个“生”之间的竞争，生仍然是最终的动力源。如果蚀的动力不来自生，它来自哪里？

我的回答分为两步。

第一步：承认新网络的生成本身当然也是一个“生”的动作。审视模式不是从虚空中长出来的；它有它自己的生成机制网络——原生家庭里习得的照料方式、职业训练中培养的诊断习惯、文化中关于“理性管理情绪”的话语实践，都是这张网的维持节点。没有这些“生”，审视模式的网络就不会作为一种能在互动时机上组织起回应的结构而出场。

第二步：但这里有一个关键的层次区分。“生”解释的是新网络作为形态的一般性存在（它为什么能出现在世界中）。它不解释新网络为什么能在这个特定节点上、以这个特定顺序侵蚀旧网络。审视模式存在于这个世界上（这是生的层面），和审视模式在今晚 A 说“我今天真的很累”的那个瞬间，选择把 B 的口型从“真是够呛的”拧成“那你早点睡”（这是蚀的层面），是两回事。前者的解释需要一套关于“人际互动模式如何在社会中成型”的生成机制，后者的解释需要一套关于“两个互动网络在功能位上如何竞争、旧网络的脆弱分布如何决定了替代的初始入点与传播路径”的互蚀学。互蚀的解释变量不是新网络本身的生成动力，而是旧网络的维持脆弱度分布、新网络节点的接入兼容性、以及两者在节点序列上的争夺阈值。它们是不同的问题。

因此，蚀的动力源不是一个叫“蚀”的神秘实体。它来自两个独立的生成机制网络在共享的交互场上对同一个功能位的结构性争夺。当一个节点为两个网络提供了同功能接口时，两张网各自为了维持自己的连贯形态而执行的每一个“生”的动作——B 执行审视动作以维持他的“助人者”身份，执行共情动作以维持他的“陪伴者”身份——在这个接口上互为蚀功。是两张网各自内在的自我维持动力，在同一节点上的正面冲突，共同构成了蚀的动力学。蚀不是一个单一来源的力，而是两张网各自“生”的力在节点上的矢量冲突——冲突的结果是一张网的节点被另一张网替代。在这个意义上，可以说：蚀不是与生对立的外力；蚀就是多个生的力场在功能位上的干涉图样。但这个“干涉”本身作为独立的分析对象，已经无法被单独一张网的“生”逻辑所还原——它需要一张新地图。这张地图的名字就是互蚀。

## ■ 八、反驳、限度与证伪

任何一个声称切出了新辨别面的概念，都必须亮出自己的证伪条件。本章将请出四个可能的关键反驳，并在第四个反驳中给出证伪方案的精确表述。

反驳一：蚀的独立性无法自证。你所描述的“节点替代”，归根到底还是旧网络自身的松动——如果旧网络足够顽强，新的网络就挤不进来。所以你所谓的“互蚀”，不过是旧网络失败后的替补，而不是一种自我驱动的正面力量。

这个反驳混淆了替代的触发条件与替代的执行机制。诚然，一个新网络要挤入旧网络的节点，往往需要旧网络在某些参数上出现波动（如互动疲倦、监管松懈、注意力剩余）。但波动为替代提供了“入点”，并不等于替代本身的推进可以被还原为波动的程度。就像藤蔓绕树需要树干表皮先有一个微裂——但裂口只是给藤蔓气根提供了第一处附着点。此后藤蔓逐圈缠绕、逐点封锁树干输导组织的整个过程，其形态与节律是由藤蔓自身的攀附生长机制决定的，不能还原为“树干太弱”或“它自己开裂了”。互蚀的结构在于新网络自身的节点组织力——它一旦获得第一个入点，就会用自己的运作方式把相邻节点逐步组织进自己的网络逻辑中。旧网络的“松动”只是开了门，接管的却是另一个管家。二者的因果动力是不同的。

反驳二：你宣称蚀先于生，但量子场论里粒子还是粒子，它们并没有因为虚粒子云而蚀掉——粒子的稳定性本身就是生的胜利。你怎么能说蚀先于生？

这个质疑需要从量子场论的数学结构本身来回应，而非从“蚀先于生”的哲学命题出发。在场存在论的数学形式中，粒子从来不是一个先在的实体——它是场的相互作用通道的稳定解。电子的物理状态，是电子的本征场与光子场、弱玻色子场的耦合通道在特定能量尺度上达到的平衡态。这个平衡态的获得，必然地包含了其他可能耦合通道的抑制：那些被守恒律禁止或极端压低的过程，正是因为在这个能量尺度上无法与主通道竞争节点（相互作用顶点），才成为“被抑制”的通道。这个“锁定一通道、抑制其他通道”的过程，其形式本身就是本章定义的蚀。因此，说“粒子没有被蚀”是不准确的——粒子的稳定存在本身，就是蚀的竞争过程达到的一类特殊的平衡解。它不是“没有被蚀”，它是“蚀过了的剩余”。这个回应不依赖于预设“蚀先于生”的前提；它完全植根于量子场论的物理结构本身。

反驳三：如果把“蚀”推到这么底层，那“节点”本身是什么？如果节点也是一个稳定的形态，那它的维持不也依赖着“生”吗？如果节点本身也在被蚀的过程中，那“蚀”岂不就是自指的、无穷递归的？你所谓“比生更原始的蚀”，会不会只在某一层有效，往下追一层就不得不承认“生”仍是那个更难动摇的底座？

这个反驳精准地打在了互蚀理论最脆弱的一处。我接受这个挑战，并愿意把它转变为互蚀概念的深度来源。

节点的确不是一个永恒的实体。每一个被互蚀理论标记为“旧网络节点”的互动回合（比如那句“真是够呛的”），它自身也是一个微型的稳定形态。它在反复重现中被维持，它的每一次重现都同样蚀除了 B 在那个瞬间可能做出的其他回应——比如 B 也可以选择沉默、转移话



题、或者用笑话岔开。因此，“节点”本身就已是蚀战的一次局部结算：它是B的“陪伴者”网络与其内部的竞争性回应模式在微小尺度上反复蚀战之后幸存的动作模板。在这个意义上，节点非但不是不可蚀的基石，它本身就是上一级蚀的产品。

但这并不意味着互蚀概念陷入了逻辑困境。这里需要区分时间在先与结构在先。在每一次具体的时间序列中，一个节点可能确实形成（“生”的动作产生了一个局部的稳定形态）在先，而后这个节点被另一个网络的节点替代（蚀的动作）在后。但在存在论的结构上，那个形成节点的“生”的动作，其内部结构本身就是蚀——它是在一堆竞争的潜在回应模式中，通过排他性地占据这个交互时机，把其他潜在模式蚀除之后，才把自己确立为“这个节点”的。所以，在物理时间中，可以有序列上的先生成后蚀替；但在结构成分上，生成动作本身就是蚀。蚀不是在时间上永远比生早一秒到来——它是在每一个生的动作的内部，作为生的构造逻辑而先行在场的。这个区分一旦划清，反驳三所指的“鸡与蛋”循环就解开了：不是“生必须先生出节点，蚀才能来替”，而是在每一个节点的“生”的存在论构成里，蚀已经在那里了。

因此，节点也在被蚀不是互蚀理论的bug，它是互蚀理论的核心预测：你在任何一个层级上都找不到一个不被蚀的绝对基础。这不导致无限后退，因为后退的要点不是找一个终极基础，而是看清同一个形式在每一个层级上都在运作。这正是互蚀作为跨域底层约束的意义所在。

反驳四：就算你的论证成立，“蚀先于生”这个说法也太大了。你几乎是在主张宇宙的本性是侵蚀。这是一个无法检验的命题。你凭什么说它是哲学而非玄想？

这是最应该被认真对待的质疑。本章在此明确给出证伪方案：

证伪条件：如果能够在任何一个被本章归为“互蚀”的真实案例（粒子衰变、信任瓦解、制度空心化）中，完全使用一套只包含“旧网络维持条件的内在衰减参数”而不引入任何“异质网络节点主动替代”的模型，对该案例的整体时序——尤其是“替代的起点为何在这个节点而非那个节点”、“替代的传播路径为何遵循此顺序而非彼顺序”——做出同等或更高精度的预测和解释，那么本章的核心命题“互蚀是一个不可还原的正面过程”就被证伪。具体来说，需要针对同一案例，构造一个纯衰减模型（基于熵增、损耗、遗忘、松懈等纯消极参数），并证明它能在足够的跨域案例中系统性地胜出：即纯衰减模型预测的替代顺序与实际观察到的顺序之间的拟合度，不低于互蚀模型的预测。

这个证伪条件不是空洞许诺。它指明了：要打败互蚀，不能靠“感觉它不必要”，而要靠提出一个同样能解释替代顺序的纯消极模型。目前，纯消极模型在处理“替代顺序”时普遍缺乏解释力。为什么货币的弃用总是从大宗商品定价开始，而不是均匀随机地在各类交易中同时发生？为什么审视总是首先接替情绪表达类的互动，而不是事务沟通类的互动？为什么机器首



先替代的是手工蕾丝耗时最长的底网工序？为什么粒子衰变总是先经过那条与守恒律最兼容的通道？纯消极模型的预测通常是“衰减将在最脆弱处首先显现”——但“最脆弱处”究竟指什么，纯衰减模型自身并无独立识别标准，只能事后套用（就像拿着一支箭，先在墙上画一个靶心，然后把箭孔圈进去说“正中靶心”）。而互蚀模型给出了事前独立识别那个起始节点的判据：旧网络维持节点所需的即时投入强度。投入强度越高，节点在每一轮维持中越容易滑脱，越容易被新网络的动作接替。这个预测是独立于“新网络效率”的——它只量度旧网络自身在每一个节点上的维持成本。因此，它可以被独立测量和检验。这不同于纯衰减模型的事后叙事——这是事前的、可错的预测。

## ■ 九、结论：松脱的研究纲领

本章的写作不是出于“发明新词”的冲动，而是基于一个朴素的观察：现有的过程存在论传统——无论是怀特海“合生”的形而上学、拉图尔“联结”的社会学、还是巴拉德“内部作用”的物理哲学——都太过轻易地放过了形态的耗散这一半。这种放过的代价，是让这些理论在读起来时总有一种隐秘的乐观：它们描述的世界是一个不停生成、不断丰富的世界，仿佛蚀只是一段间奏，终究要为下一段更华丽的生成让路。本章花了六章的篇幅论证：这不是间奏，这是常调。生成不过是蚀的浪峰上偶尔闪出的飞沫。

承认蚀先于生，并不会导致虚无主义，一如承认宇宙在加速膨胀并不会让早餐变冷。相反，它卸下了一种两千多年来哲学家们无意识背负的沉重期待：期待世界终有一个不蚀的底。卸下这种期待之后，我们看待形态的眼光会发生切实的位移。我们不再问“怎样才能保证一段关系不破裂”，而是试着看清：在这段关系仍然维持着的当下，蚀正在哪一个节点上推进，推进到了什么程度。看清这件事本身，不是悲观——它只是正视了关系始终活在一场未曾停歇的蚀战之中这个事实。维持一段关系，就是在这个事实中持续地为某些互动节点提供新的维持能量，让旧网络在至少这些节点上仍然能赢得每一轮的竞争。

我们不再问“怎么让一个制度永续”，而是问“这个制度此刻正被哪张替代网从哪个节点蚕食，我们能在这个节点上做什么”——这一问，比一百篇制度设计的报告都更实在。

互蚀概念为不同学科打开了哪些可操作的新研究问题？

对于复杂系统科学：与其只模拟一个网络如何达致共识或临界，不如模拟两个不同的生成机制网络在同一个节点集合上如何竞争功能位。关键模拟参数不是效率最大化，而是“旧网络各节点的维持脆弱度分布”与“新网络节点的接入兼容性序列”。建模目标不是预测哪个网络最终胜出，而是预测替代的初次入点在何处、蔓延路径沿什么梯度、相持与蚀穿的临界参数落在哪个值域。可以设计的具体模拟如下：构造一个包含  $N$  个节点的网络，每个节点需要某

种“投入”以维持旧网络功能；在每个时间步，以该投入的当前水平为概率，旧节点可能“滑脱”并被新网络节点接替；在接替发生后，新节点会影响相邻节点的维持投入水平（可能是正反馈——侵蚀相邻节点的维持能力，或负反馈——反而激发旧网络增强该处投入）。关键观测变量是：替代从哪些节点开始、蔓延路径是沿什么网络拓扑特征（介数中心性、度分布、k-核位置）推进的、在多大的替代比例下旧网络整体坍塌。这些模拟可以为第四章中“替代顺序由旧网络脆弱度决定”的命题提供数学上可复现的支持或反驳。

对于社会科学：互蚀提供了一个比“衰退”和“危机”更锋利的分析工具。不再笼统地说一部法律“被架空”，而是去田野中追索：是哪一個具体的惯例节点最先接替了法律的哪一条执行动作？这种接替是沿着科层级别蔓延，还是沿着业务类型蔓延？临界密度是多少？田野操作化的关键在于：识别出旧制度中可被独立观测的执行节点——例如环境法规的每一次实际惩处——然后测量每一个节点上“惯例替代动作”的出现频率随时间的变化。这个频率变化的时序曲线，就是蚀的推进速度。可以比较不同地区、不同类型的执行节点之间推进速度的差异，检验制度变迁理论中“语义弹性”的变量与互蚀理论中“维持成本”的变量，哪一个对推进速度的预测力更强。这些问题的回答，可以直接为制度韧性评估提供一套新的指标体系。

对于心理治疗与关系咨询：互蚀能帮助当事双方看清，他们之间的关系不是“感情变淡了”，而是哪一次特定的互动回合，把原先属于“共同承受”的节点，让位给了另一个善意的管理动作。这个让位不是任何一方的“错”，而是两套各自有其维持需求的互动网络在同一个功能位上的结构性竞争。临床工作者的任务不是判断谁的模式“更健康”，而是帮助双方标识出这场竞争的当前战局——替代已经蔓延到了哪些节点，还有哪些旧节点仍在运作。这份地图本身，就是干预的入口。可操作的具体步骤是：在咨询中与双方一起追溯“第一次发现回应方式变了”的关键互动时刻，然后以那次互动为起点，向前回溯旧网络在该节点上的维持脆弱状态（当时双方各自处在什么情绪、精力水平），向后追踪新节点如何从该入点蔓延到其他互动类型。这个追溯过程不做价值判断——不把旧网络说成“本真”、新网络说成“退化”——只做结构识别。这张地图让双方看到：他们不是“不爱了”，而是他们的互动场正在进行一场他们此前没有语言去描述的蚀战。有了语言，就有了谈判的可能。

回到开头那个意象。过程存在论给了我们“积木相互咬住”的图景，这个图景确实比“积木堆在一起”更接近真实。但咬住这个意象最大的问题，是它没有给“被取代”留下任何位置。在积木的世界里，一块积木不是被另一块积木替掉；它要么在，要么不在。真实世界的模样却更像一片长满树的沼泽：一棵树死去，不是因为被伐倒，而是因为另一棵树的根系已经悄悄占据了它脚下的整个土层。它坍塌的时候，身上早已缠满青藤。我们一直只赞美树的挺立，却忘了藤的缠绕比挺立更古老。把藤的名字还给它，就是本章唯一想做、也已做完的事情。

证伪条件：若能在粒子衰变、信任瓦解或制度空心化的任一真实序列中，构造一个纯消极衰减模型（只含旧网络维持条件的内在衰减参数，不引入任何异质网络的功能位竞争机制），使之为替代的初次入点、蔓延路径与临界阈值的预测力不低于本章的互蚀模型，则互蚀的独立理论地位被证伪。具体检验路径为：将同一案例的数据分别拟合纯衰减模型和互蚀模型，比较两者对“替代最先发生在哪个节点”和“替代蔓延顺序”的事前预测精度。纯衰减模型的“最弱节点”独立判据必须由本章定义之外的标准给出，不可事后套用。

## ■ 注释

\[1\] 本章第三章第二小节所引用的亲密关系情境，为理论目的而构造的典型化情境，旨在展示“审视接替信任”的形式结构，不代表任何具体个案，亦不作为治疗建议。

\[2\] 本章引述怀特海、拉图尔、巴拉德，仅限于其过程存在论面向：怀特海的“合生”与“现实实有”、拉图尔的“行动者网络”与联结社会学、巴拉德的“内部作用”与能动性实在论。三者的理论体系远为复杂，本章无意作整体评价。此外，本章的跨域提取方法在精神上受过程存在论启发，但其操作完全独立于上述任何一家的具体形而上学承诺。

\[3\] 生态学中的竞争性替代模型，一般指两个物种利用同一生态位资源时，效率较高者逐步排斥另一者的过程。本章所讨论的是该模型的一般形式逻辑，不涉及特定生态系统的经验参数。

\[4\] 诺斯的制度变迁理论以交易成本与路径依赖为核心，阿瑟的锁定效应则强调收益递增如何锁定技术或制度。渐进制度变迁理论（如 Mahoney 和 Thelen 等人在 2010 年前后系统阐述的论述）进一步关注了规则在字面不变的情况下通过执行位移实现的静默变迁。本章提出的“互蚀”试图在微观节点替代层面与这些传统对话，指出其共同的盲区在于“执行时机抢占”这一独立于“规则重新解释”的动力逻辑。

\[5\] 本章对黑格尔“扬弃”的理解主要依据《精神现象学》中的经典论述。

\[6\] 德里达的解构策略集中揭示符号系统内意义的不稳定性与延异。互蚀与解构的根本差异在于前者追问生成机制网络的节点替代历史，后者追问文本性条件，不可通约。

\[7\] 量子场论中粒子作为场的局域稳定激发态的论述，是基于拉格朗日形式与规范相互作用的当代标准理解。本章不依赖任何特定的量子引力或超越标准模型的推测性理论。对于量子场论与哲学关系的进一步讨论，本章的立场是：量子场论的数学形式揭示了本章所描述的那个蚀的结构；如果未来的基础物理革命显示这个数学形式不是基本层面的正确描述，本章的论证需被相应调整。

\[8\] 本章的证伪条件聚焦于“替代顺序”的预测力：若纯衰减模型能在真实案例中同等或更好地解释替代起点、路径与临界阈值，则互蚀的独立理论地位被证伪。该设计试图将形而上学命题转化为经验上可仲裁的模型比较问题，具体操作路径参见第八章。

## ■ 深化增补：被放跑的最近邻、操作化与证伪

以下四节为编辑增补，针对评审记下的扣分点定点补强，与作者正文分开标注，不改动作者原有判断与结构。

### ■ 增补一：与怀特海“永恒消逝”的硬边界

正文把怀特海当作过程存在论的代表来撬动，却放过了他体系中最贴近本章的那一条：永恒消逝。怀特海明确写过，现实事态在达成其满足的同时即行消逝，消逝是每一次现实生成的内在环节而非外部剥夺——这几乎已经说出了“溃散不是生成的倒放”。若不划界，互蚀极易被读成永恒消逝的跨域例示。分野有二。其一，消逝的机制不同：永恒消逝是事态自身完成后的自动退场，它不需要任何替代者，退场之后该事态转为被后继事态摄入的客体；互蚀则要求一个异质网络的节点从内部逐点接入，没有接替就没有蚀，孤立形态的自然衰减不属于本章论域。其二，顺序的决定权不同：怀特海体系中消逝的节奏由事态自身的满足所决定，是内生的；本章主张替代的顺序与临界阈值由旧网络各节点维持脆弱度的分布所决定，是可从旧网络结构预测的。第二条给出了可判别面：若在充分观察的替代过程中，接替顺序与旧网络的脆弱度分布无统计关联，而只与新网络的效率优势相关，则互蚀相对于既有理论没有独立增量。

### ■ 增补二：维持脆弱度与替代阈值的可编码形式

本章最锋利的判断——替代顺序由旧网络自身的脆弱度分布决定——若停在描述层面，文末的证伪条款就无法执行。可编码形式：维持脆弱度，定义为该节点在单位时间内需要被重新确认的次数，与确认失败后该节点仍能维持其功能的时长之比，比值越高越脆弱；替代阈值，定义为已被异质节点接替的节点占比达到多少时，旧网络的整体功能出现不可回撤的下降；接替方向性，定义为新接入节点是否优先落在脆弱度排序的前段。三项在三个案例域都可测：货币域可用支付场景中该币种被要求重新验真的频次，人际域可用同一承诺被要求重新申明的频次，物理域可用退相干时间标度。三项都不测量新网络的效率，因此本章与竞争性替代模型的分界才是可检验的，而不只是一个措辞差异。

### ■ 增补三：三域同构主张的限界

本章从虚粒子涨落、信任衰变、货币空心化提取同一形式结构，这个跨度是本章的力量，也是它最容易失手的地方。须明确：本章主张的是三者共享一个可形式化的替代结构，不主张三者共享同一套动力学常数，更不主张物理层面的退相干与制度层面的空心化之间存在任何因果或还原关系。把互蚀读成“用量子力学解释货币”是一种误读，而这种误读会毁掉本章——因为一旦要求物理案例承担制度案例的解释责任，两边都会立刻失效。形式同构的检验标准只有一条：把三个案例的节点、接替序列与阈值代入同一组变量后，模型的预测在三域各自的内部数据上都成立。这是一条比“三者很像”严格得多的要求，也是本章唯一愿意承担的。

## ■ 增补四：与同批三篇的分工

作者同批另有三篇处理相邻地形，须划清分工，否则四篇会互相稀释。《事物的压差地层》处理边界如何被内在紧张逼出，是关于确立的机制；《碰撞逼出》处理新结构如何从旧结构被逼垮的废墟上挤出，是关于跃迁的机制；《缝点》处理事物如何在时间中反复缝合自身，是关于维持的机制。本章处理的是消散，且主张消散比确立更原始——因此本章与另三篇不是并列的四个机制，而是对它们共同前提的一次撬动：若本章成立，那三篇所描述的确立、跃迁与维持，都应被重新理解为蚀的间隙中偶然稳住的局部。这一从属关系应在四篇之间明确标注。

## 结语 把地址写回来

这本书从一份不肯尖叫的不适出发，走到一句可以被证伪的话上：提高吸收率的安排同时提高收件位移；纠错信号并不消失，它在到达执行权之前被转成一份"已知悉"的记录。

十一位作者从十一个方向撞进这句话，而他们互不相识。这件事本身可能是证据，也可能只是同一套写作工序留下的指纹——本书把这两种可能都摆在附录三里，没有替读者选。

如果这句话是对的，那么 AI 时代对普通人最深的一件事，就不是拿走了他的能力。能力的账目前还算清楚，识别的能力甚至在上升。被拿走的是位置：那个能接住信号、并且必须为之付款的位置。人还在，感觉还在，判断有时甚至更准了，只是越顺利，越没有你的位置——收信的人已经不是你。

出路不在把工具还回去，也不在把责任推回个人，更不在少做一点评估。三条路本书都拒绝了，理由写在导论第四节。剩下的只有一个方向：留一处必须由自己付款的地方。

它的临床形式是胡敏的裁定性静默区——在不可逆的时间窗里彻底不裁定，而不是换一个更好的裁定者。它的存在论形式是秦莉的存押——无保底、不可逆、自己承担代价。它的时间形式是胡志英的蚀先于生——维持不是一次完成的事，而是每天都在进行的、不被记账的抵抗。

三种形式有一个共同点，值得作为全书的最后一句：它们都不会出现在任何读数上。这既是它们最脆弱的地方——没有仪表保护它们，也是它们至今还没有被吸收掉的唯一原因。



## 参考书目

各章文献表在此合并，按章序排列。章内引注编号为该章原编号，与本表次序不对应。

### ■ 第一章 陈晓艳

1. Fink H, Wild CP, et al. Global and regional cancer burden attributable to modifiable risk factors to inform prevention. *Nature Medicine*. 2026;32:1306-1315. doi:10.1038/s41591-026-04219-7.
2. Bray F, Laversanne M, Sung H, et al. Global cancer statistics 2022: GLOBOCAN estimates of incidence and mortality worldwide for 36 cancers in 185 countries. *CA: A Cancer Journal for Clinicians*. 2024;74(3):229-263. doi:10.3322/caac.21834.
3. Greenland S, Robins JM. Conceptual problems in the definition and interpretation of attributable fractions. *American Journal of Epidemiology*. 1988;128(6):1185-1197. doi:10.1093/oxfordjournals.aje.a115073.
4. Robins JM, Greenland S. Estimability and estimation of excess and etiologic fractions. *Statistics in Medicine*. 1989;8(7):845-859. doi:10.1002/sim.4780080709.
5. Rawls J. *A Theory of Justice*. Revised ed. Cambridge, MA: Harvard University Press; 1999.
6. Parfit D. *Reasons and Persons*. Oxford: Clarendon Press; 1984.
7. Gardiner SM. A perfect moral storm: climate change, intergenerational ethics and the problem of moral corruption. *Environmental Values*. 2006;15(3):397-413. doi:10.3197/096327106778226293.
8. Caney S. Climate change and the future: discounting for time, wealth, and risk. *Journal of Social Philosophy*. 2009;40(2):163-186. doi:10.1111/j.1467-9833.2009.01445.x.
9. Ekeli KS. Environmental risks, uncertainty and intergenerational ethics. *Environmental Values*. 2004;13(4):421-448. doi:10.3197/0963271042772578.

10. Shrader-Frechette K. Duties to future generations, proxy consent, intra- and intergenerational equity: the case of nuclear waste. *Risk Analysis*. 2000;20(6):771-778. doi:10.1111/0272-4332.206071.
11. Kabasenche WP, Skinner MK. DDT, epigenetic harm, and transgenerational environmental justice. *Environmental Health*. 2014;13:62. doi:10.1186/1476-069X-13-62.
12. Brunner PH, Rechberger H. *Handbook of Material Flow Analysis: For Environmental, Resource, and Waste Engineers*. 2nd ed. Boca Raton: CRC Press; 2016. doi:10.1201/9781315313450.
13. Baccini P, Brunner PH. *Metabolism of the Anthroposphere: Analysis, Evaluation, Design*. 2nd ed. Cambridge, MA: MIT Press; 2012. doi:10.7551/mitpress/8720.001.0001.
14. Chea JD, Ruiz-Mercado GJ, Smith RL, Meyer DE, Gonzalez MA, Barrett WM. Material flow analysis and occupational exposure assessment in additive manufacturing end-of-life material management. *Environmental Science & Technology*. 2024;58:9000-9012. doi:10.1021/acs.est.4c01562.
15. Leese E, et al.; HBM4EU E-waste Study Team. HBM4EU e-waste study: occupational exposure to metals in European waste workers and controls. *Environmental Research*. 2025;283:121892. doi:10.1016/j.envres.2025.121892.
16. Reckwitz A. Toward a theory of social practices: a development in culturalist theorizing. *European Journal of Social Theory*. 2002;5(2):243-263. doi:10.1177/13684310222225432.
17. Shove E, Pantzar M, Watson M. *The Dynamics of Social Practice: Everyday Life and How It Changes*. London: SAGE; 2012.
18. Blue S, Shove E, Carmona C, Kelly MP. Theories of practice and public health: understanding (un)healthy practices. *Critical Public Health*. 2016;26(1):36-50. doi:10.1080/09581596.2014.980396.
19. Meier PS, Warde A, Holmes J. All drinking is not equal: how a social practice theory lens could enhance public health research on alcohol and other health behaviours. *Addiction*. 2018;113(2):206-213. doi:10.1111/add.13895.

20. Blue S, Shove E, Kelly MP. Obese societies: reconceptualising the challenge for public health. *Sociology of Health & Illness*. 2021;43(4):1051-1067. doi:10.1111/1467-9566.13275.
21. Strasser T. Reflections on cardiovascular diseases. *Interdisciplinary Science Reviews*. 1978;3(3):225-230. doi:10.1179/030801878791925921.
22. Rose G. Sick individuals and sick populations. *International Journal of Epidemiology*. 1985;14(1):32-38. doi:10.1093/ije/14.1.32.
23. Frieden TR. A framework for public health action: the health impact pyramid. *American Journal of Public Health*. 2010;100(4):590-595. doi:10.2105/AJPH.2009.185652.
24. Kickbusch I, Allen L, Franz C. The commercial determinants of health. *The Lancet Global Health*. 2016;4(12):e895-e896. doi:10.1016/S2214-109X(16)30217-0.
25. Gilmore AB, Fabbri A, Baum F, et al. Defining and conceptualising the commercial determinants of health. *The Lancet*. 2023;401(10383):1194-1213. doi:10.1016/S0140-6736(23)00013-2.
26. Leppo K, Ollila E, Peña S, Wismar M, Cook S, eds. *Health in All Policies: Seizing Opportunities, Implementing Policies*. Helsinki: Ministry of Social Affairs and Health, Finland; 2013.
27. Rutter H, Savona N, Glonti K, et al. The need for a complex systems model of evidence for public health. *The Lancet*. 2017;390(10112):2602-2604. doi:10.1016/S0140-6736(17)31267-9.
28. Michie S, van Stralen MM, West R. The behaviour change wheel: a new method for characterising and designing behaviour change interventions. *Implementation Science*. 2011;6:42. doi:10.1186/1748-5908-6-42.
29. National Institute for Occupational Safety and Health. Hierarchy of Controls. Updated 10 April 2024. <https://www.cdc.gov/niosh/hierarchy-of-controls/about/>. Accessed 10 August 2026.
30. Organisation for Economic Co-operation and Development. *Extended Producer Responsibility: Updated Guidance for Efficient Waste Management*. Paris: OECD Publishing; 2016. doi:10.1787/9789264256385-en.

31. Schüz J, Espina C, Wild CP, et al. The 5th edition of the European Code Against Cancer: advancing cancer prevention across Europe. *Molecular Oncology*. 2026. doi:10.1002/1878-0261.70190.
32. Feliu A, et al. European Code Against Cancer, 5th edition: tobacco, second-hand smoke and alcohol recommendations. *Molecular Oncology*. 2026. doi:10.1002/1878-0261.70177.
33. Leitzmann M, et al. European Code Against Cancer, 5th edition: body weight, diet and physical activity recommendations. *Molecular Oncology*. 2026. doi:10.1002/1878-0261.70201.
34. World Health Organization. WHO Framework Convention on Tobacco Control. Geneva: World Health Organization; 2003, updated reprint 2005. <https://fctc.who.int/publications/i/item/9241591013>. Accessed 10 August 2026.
35. International Agency for Research on Cancer. Evaluating the Effectiveness of Smoke-free Policies. IARC Handbooks of Cancer Prevention, Volume 13. Lyon: IARC; 2009.
36. International Agency for Research on Cancer. Tobacco Smoke and Involuntary Smoking. IARC Monographs on the Evaluation of Carcinogenic Risks to Humans, Volume 83. Lyon: IARC; 2004.
37. Haw SJ, Gruer L. Changes in exposure of adult non-smokers to secondhand smoke after implementation of smoke-free legislation in Scotland: national cross sectional survey. *BMJ*. 2007;335:549-552. doi:10.1136/bmj.39315.670208.47.
38. Akhtar PC, Currie DB, Currie CE, Haw SJ. Changes in child exposure to environmental tobacco smoke after implementation of smoke-free legislation in Scotland: national cross-sectional survey. *BMJ*. 2007;335:545-549. doi:10.1136/bmj.39311.550197.AE.
39. Semple S, Maccalman L, Naji AA, Dempsey S, Hilton S, Miller BG, Ayres JG. Bar workers' exposure to second-hand smoke: the effect of Scottish smoke-free legislation on occupational exposure. *Annals of Occupational Hygiene*. 2007;51(7):571-580. doi:10.1093/annhyg/mem044.
40. Ayres JG, Semple S, MacCalman L, et al. Bar workers' health and environmental tobacco smoke exposure (BHETSE): symptomatic improvement in

bar staff following smoke-free legislation in Scotland. *Occupational and Environmental Medicine*. 2009;66(5):339-346. doi:10.1136/oem.2008.040311.

41. Lei J, Ploner A, Elfström KM, et al. HPV vaccination and the risk of invasive cervical cancer. *New England Journal of Medicine*. 2020;383:1340-1348. doi:10.1056/NEJMoa1917338.

42. Chang MH, Chen CJ, Lai MS, et al. Universal hepatitis B vaccination in Taiwan and the incidence of hepatocellular carcinoma in children. *New England Journal of Medicine*. 1997;336(26):1855-1859. doi:10.1056/NEJM199706263362602.

43. Chiang TH, Chang WJ, Chen SL, et al. Mass eradication of *Helicobacter pylori* to reduce gastric cancer incidence and mortality: a long-term cohort study on Matsu Islands. *Gut*. 2021;70(2):243-250. doi:10.1136/gutjnl-2020-322200.

44. International Agency for Research on Cancer. Alcohol Policies. IARC Handbooks of Cancer Prevention, Volume 20B. Lyon: IARC; 2025.

45. National Institute for Occupational Safety and Health. Chemical Carcinogen Policy. Cincinnati, OH: Centers for Disease Control and Prevention. <https://www.cdc.gov/niosh/chemical-carcinogen/about/>. Accessed 10 August 2026.

46. International Agency for Research on Cancer. Colorectal Cancer Screening. IARC Handbooks of Cancer Prevention, Volume 17. Lyon: IARC; 2019.

47. International Agency for Research on Cancer. Cervical Cancer Screening. IARC Handbooks of Cancer Prevention, Volume 18. Lyon: IARC; 2022.

48. Rhodes T. The 'risk environment': a framework for understanding and reducing drug-related harm. *International Journal of Drug Policy*. 2002;13(2):85-94. doi:10.1016/S0955-3959(02)00007-5.

49. Stellman SD, Stellman JM. Exposure opportunity models for Agent Orange, dioxin, and other military herbicides used in Vietnam, 1961-1971. *Journal of Exposure Analysis and Environmental Epidemiology*. 2004;14(4):354-362. doi:10.1038/sj.jea.7500331.

50. Gibson JJ. *The Ecological Approach to Visual Perception*. Boston: Houghton Mifflin; 1979.

51. Barker RG. *Ecological Psychology: Concepts and Methods for Studying the Environment of Human Behavior*. Stanford, CA: Stanford University Press; 1968.

52. Woolgar S. Configuring the user: the case of usability trials. *The Sociological Review*. 1990;38(S1):58-99. doi:10.1111/j.1467-954X.1990.tb03349.x.

53. Schulte PA, Rinehart R, Okun A, Geraci CL, Heidel DS. National Prevention through Design (PtD) Initiative. *Journal of Safety Research*. 2008;39(2):115-121. doi:10.1016/j.jsr.2008.02.021.

54. 附：风险待承位与滚动前主体预防窗的最小记录及删除审计表

## ■ 第二章 阳涌

55. Acemoglu, D., & Restrepo, P. (2018). Artificial Intelligence, Automation and Work. NBER Working Paper No. 24196.

56. Acemoglu, D., & Restrepo, P. (2019). Automation and New Tasks: How Technology Displaces and Reinstates Labor. *Journal of Economic Perspectives*, 33(2), 3–30. <https://www.nber.org/papers/w25684>

57. Amershi, S., Weld, D., Vorvoreanu, M., et al. (2019). Guidelines for Human-AI Interaction. *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*. <https://dl.acm.org/doi/10.1145/3290605.3300233>

58. Autor, D. H., Levy, F., & Murnane, R. J. (2003). The Skill Content of Recent Technological Change: An Empirical Exploration. *Quarterly Journal of Economics*, 118(4), 1279–1333.

59. Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakçı, Ö ., & Mariman, R. (2025). Generative AI without Guardrails Can Harm Learning. *Proceedings of the National Academy of Sciences*, 122(26), e2422633122. <https://doi.org/10.1073/pnas.2422633122>

60. Brynjolfsson, E., Li, D., & Raymond, L. R. (2025). Generative AI at Work. *Quarterly Journal of Economics*, 140(2), 889–942. <https://academic.oup.com/qje/article/140/2/889/7990658>

61. Buçınca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To Trust or to Think: Cognitive Forcing Functions Can Reduce Overreliance on AI in AI-Assisted Decision-Making. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*. <https://arxiv.org/abs/2102.09692>



62. Doshi, A. R., & Hauser, O. P. (2024). Generative AI Enhances Individual Creativity but Reduces the Collective Diversity of Novel Content. *Science Advances*.
63. <https://www.science.org/doi/10.1126/sciadv.adn5290>
64. Hutchins, E. (1995). *Cognition in the Wild*. MIT Press.
65. Jha, M., Colbran, S., & Purnell, K. (2026). Beyond Prompting: Biological Memory, Cognitive Offloading, and Human Expertise in the Age of GenAI. *Frontiers in Educational Practice and Research*, 2(2), 129–144. <https://www.nexus-press.com/journal/FEPR/article/2495>
66. LaPosta, P. (2026). Conflict as Telemetry for Illegible AI: Governing LLM Agent Workflows. *Proceedings of the AAAI Symposium Series*, 8(1).
67. <https://ojs.aaai.org/index.php/AAAI-SS/article/view/42543>
68. Lee, H.-P., Sarkar, A., Tankelevitch, L., et al. (2025). The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects from a Survey of Knowledge Workers. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. <https://dl.acm.org/doi/full/10.1145/3706598.3713778>
69. Merton, R. K. (1968). The Matthew Effect in Science. *Science*, 159(3810), 56–63.
70. Noy, S., & Zhang, W. (2023). Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence. *Science*, 381(6654), 187–192. <https://www.science.org/doi/10.1126/science.adh2586>
71. Risko, E. F., & Gilbert, S. J. (2016). Cognitive Offloading. *Trends in Cognitive Sciences*, 20(9), 676–688. <https://www.sciencedirect.com/science/article/abs/pii/S1364661316300985>
72. Sen, A. (1999). *Development as Freedom*. Oxford University Press.

### ■ 第三章 季春雷

73. Martin Heidegger, *Sein und Zeit*, Max Niemeyer, 1927.
74. Michel Foucault, *Sécurité, territoire, population: Cours au Collège de France (1977–1978)*, Gallimard/Seuil, 2004.
75. Byung-Chul Han, *Müdigkeitsgesellschaft*, Matthes & Seitz, 2010.

76. Hartmut Rosa, *Resonanz: Eine Soziologie der Weltbeziehung*, Suhrkamp, 2016.

77. Charles Taylor, *The Ethics of Authenticity*, Harvard University Press, 1991.

## ■ 第四章 孔凡鹤

78. Itoh, K., Chiba, T., Takahashi, S., Ishii, T., Igarashi, K., Katoh, Y., ... & Yamamoto, M. (1997). An Nrf2/small Maf heterodimer mediates the induction of phase II detoxifying enzyme genes through antioxidant response elements. *Biochemical and Biophysical Research Communications*, 236(2), 313-322.

79. Maturana, H. R., & Varela, F. J. (1980). *Autopoiesis and Cognition: The Realization of the Living*. D. Reidel Publishing Company. (Original work published 1972)

80. Mol, A. (2008). *The Logic of Care: Health and the Problem of Patient Choice*. Routledge.

81. Suez, J., Zmora, N., Zilberman-Schapira, G., et al. (2018). Post-antibiotic gut mucosal microbiome reconstitution is impaired by probiotics and improved by autologous FMT. *Cell*, 174(6), 1406-1423.

## ■ 第五章 高鹏

82. Blair, R. J. R. (1995). A cognitive developmental approach to morality: Investigating the psychopath. *\\*Cognition\\**, 57(1), 1-29.

83. Damasio, A. R. (1994). *\\*Descartes' error: Emotion, reason, and the human brain\\**. Putnam.

84. Deci, E. L. (1971). Effects of externally mediated rewards on intrinsic motivation. *Journal of Personality and Social Psychology*, 18(1), 105-115.

85. Deci, E. L., Koestner, R., & Ryan, R. M. (1999). A meta-analytic review of experiments examining the effects of extrinsic rewards on intrinsic motivation. *Psychological Bulletin*, 125(6), 627-668.

86. Devlin, P. (1965). *\\*The enforcement of morals\\**. Oxford University Press.

87. Frey, B. S., & Jegen, R. (2001). Motivation crowding theory. *Journal of Economic Surveys*, 15(5), 589-611.

88. Gneezy, U., & Rustichini, A. (2000). A fine is a price. *Journal of Legal Studies*, 29(1), 1-17.
89. Hart, H. L. A. (1961). *The concept of law*. Oxford University Press.
90. Hart, H. L. A. (1963). *Law, liberty, and morality*. Oxford University Press.
91. Hoffman, M. L. (2000). *Empathy and moral development: Implications for caring and justice*. Cambridge University Press.
92. Lepper, M. R., Greene, D., & Nisbett, R. E. (1973). Undermining children's intrinsic interest with extrinsic reward. *Journal of Personality and Social Psychology*, 28(1), 129-137.
93. Mill, J. S. (1859). *On liberty*. John W. Parker and Son.
94. Moll, J., de Oliveira-Souza, R., Bramati, I. E., & Grafman, J. (2002). Functional networks in emotional moral and nonmoral social judgments. *NeuroImage*, 16(3), 696-703.
95. Titmuss, R. M. (1970). *The Gift Relationship: From Human Blood to Social Policy*. London: Allen & Unwin.
96. Tomasello, M. (1999). *The cultural origins of human cognition*. Harvard University Press.
97. Tomasello, M. (2008). *Origins of human communication*. MIT Press.
98. Tomasello, M. (2016). *A natural history of human morality*. Harvard University Press.
99. Deci, E. L., Koestner, R., & Ryan, R. M. (1999). A meta-analytic review of experiments examining the effects of extrinsic rewards on intrinsic motivation. *Psychological Bulletin*, 125(6), 627-668.
100. Lepper, M. R., Greene, D., & Nisbett, R. E. (1973). Undermining children's intrinsic interest with extrinsic reward. *J. Pers. Soc. Psychol.*, 28(1), 129-137.
101. Gneezy, U., & Rustichini, A. (2000). A fine is a price. *Journal of Legal Studies*, 29(1), 1-17.
102. Frey, B. S., & Jegen, R. (2001). Motivation crowding theory. *Journal of Economic Surveys*, 15(5), 589-611.

103. Hoffman, M. L. (2000). *Empathy and Moral Development*. Cambridge University Press.

104. Tomasello, M. (2016). *A Natural History of Human Morality*. Harvard University Press.

## ■ 第六章 张琼

105. Bion, W. R. (1961). *Experiences in Groups*. London: Tavistock Publications.

106. Bourdieu, P. (1990). *The Logic of Practice*. Stanford: Stanford University Press.

107. Dworkin, G. (1988). *The Theory and Practice of Autonomy*. Cambridge: Cambridge University Press.

108. Foucault, M. (1975). *Surveiller et punir: Naissance de la prison*. Paris: Gallimard.

109. Granovetter, M. (1985). Economic action and social structure: The problem of embeddedness. *American Journal of Sociology*, 91(3), 481-510.

110. Putnam, R. D. (2000). *Bowling Alone: The Collapse and Revival of American Community*. New York: Simon & Schuster.

111. Turner, V. (1969). *The Ritual Process: Structure and Anti-Structure*. Chicago: Aldine Publishing.

112. Yalom, I. D., & Leszcz, M. (2020). *The Theory and Practice of Group Psychotherapy* (6th ed.). New York: Basic Books.

## ■ 第七章 高于涵

113. 班固（东汉）：《汉书》。中华书局，1962年。

114. 董仲舒（西汉）：《春秋繁露》。苏舆《春秋繁露义证》本，中华书局，1992年。

115. 司马迁（西汉）：《史记》。中华书局，1959年。

116. Arendt, H. (1963). *Eichmann in Jerusalem: A Report on the Banality of Evil*. New York: Viking Press.

117. Bourdieu, P. (1980). *Le Sens pratique*. Paris: Les Éditions de Minuit.

118. Foucault, M. Sécurité, territoire, population: Cours au Collège de France, 1977-1978. Paris: Gallimard, 2004.
119. Luhmann, N. (1975). Legitimation durch Verfahren. Neuwied: Luchterhand.
120. Marx, K. Ökonomisch-philosophische Manuskripte aus dem Jahre 1844. In: Marx-Engels-Gesamtausgabe, Abt. 1, Bd. 3. Berlin, 1932.
121. Polanyi, M. (1958). Personal Knowledge: Towards a Post-Critical Philosophy. Chicago: University of Chicago Press.
122. Weber, M. (1920). Die protestantische Ethik und der Geist des Kapitalismus. In: Gesammelte Aufsätze zur Religionssoziologie I. Tübingen: Mohr.
123. Williams, B. (1985). Ethics and the Limits of Philosophy. London: Fontana. 对“道德体系”以义务为唯一通货、排除其余伦理考量的批判。
124. Habermas, J. (1981). Theorie des kommunikativen Handelns, Bd. 2. Frankfurt: Suhrkamp. 系统对生活世界的殖民化，及策略行动与交往行动之分。
125. Polanyi, M. (1958). Personal Knowledge. London: Routledge. 原稿已引，此处并列备考。

## ■ 第八章 少敏

126. [1] EUROPEAN PARLIAMENT, COUNCIL OF THE EUROPEAN UNION. Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence: Article 14[J/OL]. Official Journal of the European Union, 2024. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>.
127. [2] GAUBE S, LANGER M, MILLER T, et al. Keeping an Eye on AI: A Framework for Effective Human Oversight of AI Systems[EB/OL]. arXiv:2605.16278, 2026. <https://doi.org/10.48550/arXiv.2605.16278>.
128. [3] LAUX J. Institutionalised distrust and human oversight of artificial intelligence: Towards a democratic design of AI governance under the European Union AI Act[J]. AI & Society, 2024, 39: 2853-2866. DOI: 10.1007/s00146-023-01777-z.
129. [4] SANTONI DE SIO F, VAN DEN HOVEN J. Meaningful human control over autonomous systems: A philosophical account[J]. Frontiers in Robotics and AI, 2018, 5: 15. DOI: 10.3389/frobt.2018.00015.

130. \[5\] ALFRINK K, KELLER I, KORTUEM G, DOORN N. Contestable AI by design: Towards a framework\[J\]. *Minds and Machines*, 2023, 33(4): 613-639. DOI: 10.1007/s11023-022-09611-z.
131. \[6\] ALFRINK K, KELLER I, YURRITA SEMPERENA M, et al. Envisioning contestability loops: Evaluating the agonistic arena as a generative metaphor for public AI\[J\]. *She Ji: The Journal of Design, Economics, and Innovation*, 2024, 10(1): 53-93. DOI: 10.1016/j.sheji.2024.03.003.
132. \[7\] SYDOW J, SCHREYÖGG G, KOCH J. Organizational path dependence: Opening the black box\[J\]. *Academy of Management Review*, 2009, 34(4): 689-709. DOI: 10.5465/amr.34.4.zok689.
133. \[8\] SYDOW J, SCHREYÖGG G, KOCH J. On the theory of organizational path dependence: Clarifications, replies to objections, and extensions\[J\]. *Academy of Management Review*, 2020, 45(4): 717-734. DOI: 10.5465/amr.2020.0163.
134. \[9\] PETERMANN A, SCHREYÖGG G, FÜRSTENAU D. Can hierarchy hold back the dynamics of self-reinforcing processes? A simulation study on path dependence in hierarchies\[J\]. *Business Research*, 2019, 12(2): 637-669. DOI: 10.1007/s40685-019-0083-9.
135. \[10\] DIXIT A K, PINDYCK R S. *Investment Under Uncertainty*\[M\]. Princeton, NJ: Princeton University Press, 1994.
136. \[11\] ADNER R, LEVINTHAL D A. What is not a real option: Considering boundaries for the application of real options to business strategy\[J\]. *Academy of Management Review*, 2004, 29(1): 74-85. DOI: 10.5465/amr.2004.11851715.
137. \[12\] BAINBRIDGE L. Ironies of automation\[J\]. *Automatica*, 1983, 19(6): 775-779. DOI: 10.1016/0005-1098(83)90046-8.
138. \[13\] ENDSLEY M R, KIRIS E O. The out-of-the-loop performance problem and level of control in automation\[J\]. *Human Factors*, 1995, 37(2): 381-394. DOI: 10.1518/001872095779064555.
139. \[14\] GREEN B, CHEN Y. The principles and limits of algorithm-in-the-loop decision making\[J\]. *Proceedings of the ACM on Human-Computer Interaction*, 2019, 3(CSCW): Article 50. DOI: 10.1145/3359152.



140. \[15\] BUÇINCA Z, MALAYA M B, GAJOS K Z. To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making\[J\]. Proceedings of the ACM on Human-Computer Interaction, 2021, 5(CSCW1): Article 188. DOI: 10.1145/3449287.
141. \[16\] VACCARO M, ALMAATOUQ A, MALONE T. When combinations of humans and AI are useful: A systematic review and meta-analysis\[J\]. Nature Human Behaviour, 2024, 8: 2293-2303. DOI: 10.1038/s41562-024-02024-1.
142. \[17\] BRIER G W. Verification of forecasts expressed in terms of probability\[J\]. Monthly Weather Review, 1950, 78(1): 1-3. DOI: 10.1175/1520-0493(1950)078\<0001:VOFEIT\>2.0.CO;2.
143. \[18\] GNEITING T, RAFTERY A E. Strictly proper scoring rules, prediction, and estimation\[J\]. Journal of the American Statistical Association, 2007, 102(477): 359-378. DOI: 10.1198/016214506000001437.
144. \[19\] KUMLE L, VÕ M L H, DRASCHKOW D. Estimating power in (generalized) linear mixed models: An open introduction and tutorial in R\[J\]. Behavior Research Methods, 2021, 53(6): 2528-2543. DOI: 10.3758/s13428-021-01546-0.
145. \[20\] ELISH M C. Moral crumple zones: Cautionary tales in human-robot interaction\[J\]. Engaging Science, Technology, and Society, 2019, 5: 40-60. DOI: 10.17351/ests2019.260.
146. \[21\] KROLL J A. Outlining traceability: A principle for operationalizing accountability in computing systems\[C\]//Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency. New York: ACM, 2021: 758-771. DOI: 10.1145/3442188.3445937.
147. \[22\] POWER M. Modelling the micro-foundations of the audit society: Organizations and the logic of the audit trail\[J\]. Academy of Management Review, 2021, 46(1): 6-32. DOI: 10.5465/amr.2017.0212.

## ■ 第九章 胡敏

148. Chen, J., et al. (2010).  $\beta$ -Adrenergic signaling regulates skeletal muscle satellite cell proliferation and differentiation. Journal of Cell Science, 123(15), 2583–2591.

149. Critchley, H. D., Wiens, S., Rotshtein, P., Öhman, A., & Dolan, R. J. (2004). Neural systems supporting interoceptive awareness. *Nature Neuroscience*, 7(2), 189–195.
150. Hawke, T. J., & Garry, D. J. (2001). Myogenic satellite cells: physiology to molecular biology. *Journal of Applied Physiology*, 91(2), 534–551.
151. Husserl, E. (1913). *Ideen zu einer reinen Phänomenologie und phänomenologischen Philosophie. Erstes Buch: Allgemeine Einführung in die reine Phänomenologie*. Halle: Max Niemeyer.
152. Kiecolt-Glaser, J. K., Marucha, P. T., Malarkey, W. B., Mercado, A. M., & Glaser, R. (1995). Slowing of wound healing by psychological stress. *The Lancet*, 346(8984), 1194–1196.
153. McEwen, B. S. (1998). Protective and damaging effects of stress mediators. *New England Journal of Medicine*, 338(3), 171–179.
154. Merleau-Ponty, M. (1945). *Phénoménologie de la perception*. Paris: Gallimard.
155. Sartre, J.-P. (1943). *L'être et le néant: Essai d'ontologie phénoménologique*. Paris: Gallimard.
156. Thayer, J. F., & Lane, R. D. (2000). A model of neurovisceral integration in emotion regulation and dysregulation. *Journal of Affective Disorders*, 61(3), 201–216.

## ■ 第十章 秦莉

157. Arendt, Hannah. (1958). *The Human Condition*. Chicago: University of Chicago Press.
158. Beck, Ulrich. (1986). *Risikogesellschaft: Auf dem Weg in eine andere Moderne*. Frankfurt am Main: Suhrkamp. (English translation: *Risk Society: Towards a New Modernity*, London: Sage, 1992.)
159. Giddens, Anthony. (1991). *Modernity and Self-Identity: Self and Society in the Late Modern Age*. Cambridge: Polity Press.
160. Heidegger, Martin. (1927). *Sein und Zeit*. Tübingen: Max Niemeyer. (English translation: *Being and Time*, New York: Harper & Row, 1962.)

161. Kierkegaard, Søren. (1843). *Frygt og Bæven*. Copenhagen: Reitzel. (English translation: *Fear and Trembling*, Princeton: Princeton University Press, 1983.)

162. Luhmann, Niklas. (2000). *Art as a Social System*. Translated by Eva M. Knodt. Stanford: Stanford University Press. (Original German edition: *Die Kunst der Gesellschaft*, Frankfurt am Main: Suhrkamp, 1995.)

## ■ 第十一章 胡志英

163. Arthur, W. B. (1994). *Increasing Returns and Path Dependence in the Economy*. University of Michigan Press.

164. Barabási, A.-L. (2016). *Network Science*. Cambridge University Press.

165. Barad, K. (2007). *Meeting the Universe Halfway: Quantum Physics and the Entanglement of Matter and Meaning*. Duke University Press.

166. Derrida, J. (1972). *Marges de la philosophie*. Éditions de Minuit.

167. Hegel, G. W. F. (1807). *Phänomenologie des Geistes*.

168. Latour, B. (2005). *Reassembling the Social: An Introduction to Actor-Network-Theory*. Oxford University Press.

169. North, D. C. (1990). *Institutions, Institutional Change and Economic Performance*. Cambridge University Press.

170. Simmel, G. (1900). *Philosophie des Geldes*.

171. Thompson, E. P. (1963). *The Making of the English Working Class*. Victor Gollancz.

172. Whitehead, A. N. (1929). *Process and Reality*. Macmillan.

173. Alfred North Whitehead, *Process and Reality: An Essay in Cosmology*, Macmillan, 1929.

174. Ilya Prigogine & Isabelle Stengers, *Order Out of Chaos: Man's New Dialogue with Nature*, Bantam Books, 1984.

175. Wojciech H. Zurek, “Decoherence, Einselection, and the Quantum Origins of the Classical,” *Reviews of Modern Physics*, 75(3), 2003.

176. Per Bak, *How Nature Works: The Science of Self-Organized Criticality*, Copernicus, 1996.

177. Joseph A. Schumpeter, *Capitalism, Socialism and Democracy*, Harper & Brothers, 1942.

## 附录一 选目协议全文与十一章清单

### ■ 一、协议（写在抽样之前）

总体：站上 `/students/` 学员专栏全部署名篇目。

入框四条：① 汉字  $\geq 8000$ ；② 页面标注创新智商  $\geq 146$ （含盲评与早期评分两种口径）；③ 非配套读物；④ 未被任何已出版专著收录。

抽学员：`random.seed(20260817)`，学员名按标识排序后随机抽取 10 位，不按篇数加权（否则篇数最多的一位会吞掉样本），一次抽定不重抽。

章内择一：在该学员的合格篇中，选与题域"AI 时代普通人的困境与出路"相关性最高者，同等相关取评分最高者。这一步是有意选择，不是随机。

成功标准（先写死，事后如实报告成败）：① 母题  $\geq 7/11$  篇可派生，且不得是体裁惯例；② 至少一条跨篇命题可写成可判错形式；③ 至少一对隔学科的构念能写成同一关系的两端；④ 书名与题域不得与站上第 78、80、81、82、83 号撞题，须能写出逐本划界。

### ■ 二、执行过程中的两处偏离，如实记录

偏离一：修框。第一遍用文章短名检查"是否已入书"，漏检。改用中文题名重查，查出三篇已是第八十二号书的正章（《无痛之死》《耦合可逆性极化》《守护者的悖论》）。其中一位作者仅有该一篇合格，因此退出抽样框，框由 13 位缩至 12 位，用同一种子在修正框上重抽。修框的判据事先写在协议第 ④ 条里，修框的动作（同种子重抽）也事先约定，但执行发生在第一次抽样之后，"先声明"的成色因此打折。

偏离二：指定席。抽定十位后，编者另行指定第十一位（胡敏），选其框内最高分篇目。这是人工干预，非随机。它使本书作为"随机可装配性"实验的证据强度下降，同时贡献了合章所依托的那一对咬合。两件事一并记账。

### ■ 三、成功标准的自评结果

① 过：11/11 可派生，为历次实验最高；但须与附录三第三条一并读。② 过：枢纽章第三节那条命题给出了两个可从既有系统取得的数与一个方向性预测。③ 过：少敏（人机系统效力审计） $\times$  胡敏（修复过程研究）。<sup>△</sup> 其中一方是指定席。④ 过：逐本划界写在枢纽章第五节。

## ■ 四、十一章清单

| 章 | 作者  | 原题           | 评分  | 原文出处                                            |
|---|-----|--------------|-----|-------------------------------------------------|
| 一 | 陈晓艳 | 预防如何不再以个人为主语 | 148 | /students/chen-xiaoyan/pre-subject-prevention/  |
| 二 | 阳涌  | 一次成功由谁承载?    | 147 | /students/yang-yong/bearerless-success/         |
| 三 | 季春雷 | 主体性代偿        | 146 | /students/ji-chunlei/subjectivity-compensation/ |
| 四 | 孔凡鹤 | 生成性裂隙        | 150 | /students/kong-fanhe/generative-fissure/        |
| 五 | 高鹏  | 善的语法消亡       | 147 | /students/gao-peng/grammar-of-good/             |
| 六 | 张琼  | 配准性生成        | 150 | /students/zhang-qiong/calibrative-genesis/      |
| 七 | 高于涵 | 未善的代价        | 152 | /students/gao-yuhan/goodness-as-process/        |
| 八 | 少敏  | 越审越不能改       | 152 | /students/shao-min/review-forecloses-change/    |
| 九 | 胡敏  | 裁定性干扰        | 153 | /students/hu-min/adjudicative-interference/     |
| 十 | 秦莉  | 存押           | 149 | /students/qin-                                  |



|    |     |      |     |                                                                         |
|----|-----|------|-----|-------------------------------------------------------------------------|
| 十一 | 胡志英 | 蚀先于生 | 150 | li/existential-stake/<br>/students/hu-zhiying/erosion-precedes-genesis/ |
|----|-----|------|-----|-------------------------------------------------------------------------|

评分为各篇发表时页面标注的创新智商，均值 149.5。这些分数由同一位评分者给出，不构成独立第三方评价。

## 附录二 十一章读数速查表

| 读数             | 出处 | 一句话定义                                            | 归位                    |
|----------------|----|--------------------------------------------------|-----------------------|
| 风险待承位          | 一  | 危险流接通、路径开放、控制者在场而承载者缺席的四联状态                      | $\Delta$              |
| 能力承载率 / 无着率    | 二  | 扰动后仍有合法承载者能复现、解释、修复并负责的任务比例                      | $\Delta$              |
| 承载迁移矩阵         | 二  | 承载从个人向组织、技术等方向转移的六种去向                            | $\Delta$              |
| 窄门事件计数         | 三  | 一生中满足四条判据的不可代偿判决的次数                              | $\beta$               |
| 生成性裂隙 / 生成性回路  | 四  | 必须由生命亲自穿越的那道阻力，及其完整经验弧线                          | $\alpha$              |
| 善的三段退化         | 五  | 行为指标替代 $\rightarrow$ 公共语法垄断 $\rightarrow$ 道德主体失语 | $\alpha$              |
| 配准性生成          | 六  | 沉默是否仍保留不被语法预制的开放性                                | $\alpha$              |
| 善的工序化 / 未善     | 七  | 工序不加工善，而持续产出未善                                   | $\alpha \cdot \Delta$ |
| 反事实承载力         | 八  | 合格异议到来时仍可转向的合格替代路径的可得程度                          | $\Delta$              |
| 裁定性干扰 / 裁定性静默区 | 九  | 裁定动作本身对不可逆时间窗内过程的减损，及其豁免区                        | $\Delta \cdot \beta$  |
| 存押 / 存在折旧      | 十  | 无保底、不可逆、自担代价地把一部分存在押进一个动作                        | $\beta$               |
| 蚀 / 互蚀 / 维持脆弱度 | 十一 | 形态的耗散先于生成，维持是持续付款                                | $\beta$               |

三点必须说明：一，十一个读数没有一个落在"能力存量"上，它们数的全是位置。二，它们不能相加——各自的量纲与领域指标不同，本书没有给出跨领域的统一量表。三，其中只有反事实承载力与能力承载率两项，在原章里给出了可立即执行的取数办法。

## 附录三 判错清单

以下每一条都是本书自己交出的、可以推翻它的办法。本书主张为已完成的检验，一律为零。

第一条（最便宜，只需两个数）。取同一行业内合规留痕强度差异显著的两组组织，同一时段比较：异议与复核记录的增长率，与方案实际改道率。若两者正相关——留痕越多改道也越多——枢纽章的核心命题错。若两者不相关，本命题同样不成立，因为它预测的是一个方向。

第二条。取康复医学、产品安全或临床路径这类同时具备留痕制度与不可逆时间窗的领域，在控制初始严重程度后比较评估频次不同的两组，考察方案改道率与过程结局。若高频评估组两项都不更差，合章的推论错。

第三条（本书最该被攻击处）。十一位作者共用同一套写作工序，并由同一位评分者评分。若有人用十一篇同等长度、同等抽象度、来自不同工序的文章，按本书同样的方法也装出一条 11/11 可派生的实质母题，那么“十一章形状相同”就不是世界的性质，而是工序的性质，本书的装配部分应大幅折价。

第四条。若能证明“人被从必须自己做出回应的位置上挪开”这一提法在既有文献中已被充分量化处理，则枢纽章的贡献退化为一次改写。相关记账见枢纽章第六节。

第五条。若能找到一类系统：吸收率持续上升而收件位移不变或下降（即后果被吸收得越多，感到问题的人反而越靠近执行权），则本书的核心因果链断裂。本书认为这类系统应当存在于小规模、权责合一的组织中，但没有去找。

第六条。本书三个量中的两个（吸收率、存押量）至今没有跨领域的操作化定义。若有人给出这样的定义并测出与本书预测相反的关系，本书的三量框架应被替换。

# 封底

## ■ 越顺利，越没有你的位置

### AI 时代普通人的困境与出路

十一位互不相识的作者，十一个相隔很远的题域：公共卫生、人机协作、社会存在论、营养生成论、法哲学、团体治疗、中国思想史、人机系统审计、修复过程研究、艺术哲学、过程存在论。

十章不是选出来的，是按预先写死的协议抽出来的；第十一章是编者点进来的——这一席的账，书里如实记着。

十一章说的是同一件事，而它们互不知情：

凡是让普通人变得更顺利的安排——托底、代劳、去阻力、加复核、把好行为做成工序——都在同一个位置上取消了那个"必须由本人亲自穿越一次"的动作；而所有读数因为结果变好，一律不报警。

本书的那一刀落在更靠后的地方：困境不是"感觉不到"。识别的能力甚至还在上升，异议越来越精确。失效发生在路径侧，不在判据侧——信号照常产生，只是寄到了一个不再有执行权的位置，然后被转成一份"已知悉"的记录。

一句可以带走的判据：

你们那儿，异议记录的数量在涨，方案实际改道的次数涨了吗？

分类：人工智能与社会 / 组织理论 / 伦理学 / 认识论

德麦国际有限公司 · 新加坡 · SDE Universes · 专著第 84 号

ISBN 979-8-90690-018-0      US\$21.30