

## 第五章 善失去了自己的语法

《越顺利，越没有你的位置》· 王德生 + Claude 编 · 本章在全书中为第 143-169 页 · ISBN 979-8-90690-018-0

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，都会以更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——判据类型的分歧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東西。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处——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，道德认知不是一个密闭系统，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测的偏离。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。这一步的关键在于将心理学的“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——它有动机追溯硬要求，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岔：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——而不是任何其他人或任何规则——是这个善举的自发来源。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眼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——第一代人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”的语词。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，因为它们在加工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，需另行验证。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的道德积分制度（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。同时设置对照组：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。

---

本分册取自《越顺利，越没有你的位置》（德麦国际专著第 84 号）。全书 343 页，含前言、导读、导论、四编十一章、枢纽章、合章、结语、参考书目与三则附录；本章的判据与划界须与枢纽章、合章一并读。全书网页版：[sdeuniverses.com/books/m/84/text/](http://sdeuniverses.com/books/m/84/text/)