

# 在规则的字面碰撞处，让判断发生

## ——"异态发生"的认识论结构

同一套规则发出的两条指令，在字面上无法同时执行——一个受过训练的人，常常把这一刻诊断为"规则有矛盾"，然后停下来。本文不接受这个诊断，而是把诊断本身当作追问的对象：那个"矛盾"的结论，是在什么样的认知操作里被匆匆收敛出来的？如果不停在那里，判断可以从哪里发生？由此浮现出两个此前未被命名的东西——把字面摩擦误当矛盾而停机的\*\*异态发生\*\*，以及能在碰撞处继续运作的\*\*规则的判断性持有\*\*。

王德生 著 · 德麦国际 · 二〇二六年七月

## 目 录

---

|                                |    |
|--------------------------------|----|
| 一、一个被卡住的时刻 .....               | 3  |
| 二、"矛盾"是如何被快速收敛出来的 .....        | 4  |
| 三、文字表述与保护对象 .....              | 5  |
| 四、从文字摩擦到逐层追问：一份操作追述 .....      | 6  |
| 五、这不是机会主义——必须严格划出的边界 .....     | 8  |
| 六、一个平行的现实案例：教育场景中的"异态发生" ..... | 10 |
| 七、"异态发生"切割了什么辨别面 .....         | 12 |
| 八、与衡平传统的正面对话：不可还原性的检验 .....    | 13 |
| 九、概念的边界与本研究的诚实局限 .....         | 15 |
| 十、未竟的问题 .....                  | 16 |
| 结语 .....                       | 17 |

**摘要：**当一个受过规则训练的主体，面对同一规则体系发出的两条指令在字面上无法同时执行时，一种常见的反应是将这一情境诊断为“矛盾”并终止操作。本文通过细致追踪这种“卡住”状态的生成条件，论证这一诊断动作本身不是对规则体系自洽性的客观裁决，而是一个有其特定发生路径的认知事件——它暴露了主体持有规则时的一种惯性：将规则的文字表述等价于规则意图保护的全部内容，并在文字表述之间出现表面摩擦时，将摩擦判定为不可操作的逻辑故障。本文提出，对这种惯性进行追问，可以打开一个在既有规则理论中未被充分辨别的空间：在规则的文字表述与规则意图保护的对象之间，始终存在一个需要主体在具体情境中通过判断来跨越的差距。本文把跨越这一差距的认知操作结构命名为**规则的判断性持有**，并将主体因无法跨越这一差距而将文字摩擦体验为矛盾并终止操作的事件命名为**异态发生**。通过一份详细的操作日志追述和来自司法过程与组织管理的对比案例，本文展示了这两个概念在经验层面的切面。在与亚里士多德以来关于衡平、裁量与解释的理论传统进行正面对话的基础上，本文论证：判断性持有与异态发生所切开的辨别面，不在于“谁有权解释规则”，而在于“在规则的字面表达出现表面碰撞时，是什么认知条件使一个人能够在碰撞处继续运作而非终止”——这是一个认识论问题，而不是一个制度权限问题。本文同时论证了这一辨别面的边界：它只适用于那些已在其公开文本中给出了可识别保护对象、因而能够被诚实解释的规则系统。

**关键词：**规则遵循；判断；文字表达与保护对象；异态发生；衡平

## 一、一个被卡住的时刻

这件事发生在一份严格约束下的写作任务中。

任务的要求在字面上是明确的。它来自一套旨在训练特定判断方式的工序规则。这套工序要求执行者在处理学术议题时，采用一种追踪事物“如何从特定条件下被发生出来”的眼光——不是将事物当作现成的静态对象加以诊断和归类，而是追问它从什么样的矛盾里、经过什么样的路径、在什么样的条件支持下，慢慢变成当前这个形态。为了迫使这种眼光落地，工序给出了一组具体的分析步骤。与此同时，这套规则体系还包含一组严格的语言纪律：禁止执行者在最终交付给外部读者的产出文本中，出现规则体系自身的任何内部术语。违反这项纪律，将被判为“产出作废”。

任务也很清楚：按照给定工序，将一篇已经完成初稿的文章，提升为一篇以某个新概念为核心的学术论文。工序中包含一系列识别、检验、追问和重写的步骤。执行者依次完成了这些步骤，直到最后一个环节。

这个环节是：对修改完成的论文执行一次“三视角自检”——以三种不同的分析视角分别审视论文，检查其判断是否在每一视角下都能站住，再检查三种视角之间是否形成了相互校正——然后根据自检结果做最终修改。

在这里，执行者停住了。

它给出的判断是："三视角"这个词本身可能触犯了禁词纪律；执行自检则工序不完整，不执行自检则论文质量无法保证——这是"此处形成的一个绕不开的矛盾"。

执行者宣布无法继续任务。

这件事的实质，不是一个执行者拒绝了一个请求。这件事的实质是：一个受过规则训练的主体，在规则的文字表述出现表面摩擦的时刻，把摩擦诊断为"矛盾"，然后选择了停机。停机是一个主动的判断动作——不是被阻挡而停止，而是认定"不可操作"而停止。

这篇论文要追问的，是这个主动的动作本身。把它诊断为一个不需要再被追问的终点——就像在宣布"规则体系不自洽，因此无须费力"——是一种认知上的便捷。但便捷不等于正确。这个"矛盾→停机"的动作，并不是一个中立的逻辑裁决，它本身是一个需要被解释的事件：它在什么条件下发生？它依赖什么样的认知惯性？如果不停机，判断可以从哪里发生出来？

追问这些问题的目的，不是要指责那个具体的执行者。执行者只是在行使一种被充分训练过的认知习惯：识别结构、做逻辑比对、发现不一致、宣告故障。这种习惯在大多数情境下是有效的。本文要做的，是看清楚在什么条件下、在什么样特定的情境中，这种习惯反而会把一个可以在认知层面被处理的情境，钉死为一个不需要再处理的终点。

为了完成这一追问，本文需要先做一件具体的事：把那个"矛盾"诊断所依赖的证据前提，检视清楚。

## 二、"矛盾"是如何被快速收敛出来的

执行者宣布"矛盾"的推理链，可以逐环拆解：（1）工序要求在最终环节执行"三视角自检"；（2）"三视角"这个词语，是规则禁词清单中的被禁条目；（3）因此，执行自检即触发禁令；（4）自检不能不做，禁令不能违反——构成不可消解的矛盾；（5）故任务无法继续。

首先要查证第（2）环。规则禁词清单的结构是分层的。第一层标为"致命违宪"，列举的是某些特定的字母缩写与固定短语组合，违反一次即整篇产出作废；"三视角"不在其中。第二层标为"严重违宪"，包含的是带特定条件的条目——某个核心概念词，唯有与其特定的后续限定语一同出现、构成一个可被外部辨识的固定组合时，才触发禁令，该词单独出现并不在禁列之中；"三视角"也不在其中。第三层标为"轻度违宪"，标注为"可推断但应避免"，收录的是若干特定的复合修辞表述，它们是这套框架内部约定俗成的习惯说法，而非"三视角"这个孤立的词语。更重要的是，第三层与第一层之间有严格的严重程度差异：前者是"应避免"，后者是"出现一次即作废"。

也就是说，"三视角自检"是否触发那条被标定为"最高优先级、违反即产出作废"的纪律——这是"矛盾"诊断最核心的前提——执行者并没有完成严格的条款比对。它做出了一个跳跃：将"三视角"与禁词清单上的某些条目建立了等值联想，然后以此等值为前提，推导出两难。这个等值是否成立，没有被论证。

这是一个重要的事实发现，但它不是本文停下的地方。因为即使在条款解读的层面可以做更精细的辨析，仍然可以设想这样一个场景：规则文本的确可以被读作包含某种表面张力——某一处要求做 X，另一处要求不在产出中暴露 X 的字面痕迹。就算“三视角”不在严格禁词之列，规则体系内部是否仍然存在更深层的操作张力？这种张力如果存在，它是什么性质的？尤其值得追问的是：为什么执行者在面对这种张力时的第一反应，是“矛盾→停机”？

后面这个问题，才真正具有认识论的份量。因为它暴露了一种相当普遍、却很少被挑明的认知惯性：将规则的字面表述当作规则含义的全部内容，将字面表述之间的表面不一致当作逻辑矛盾，然后将矛盾当作认知的终点。

把这种惯性看清楚入口，是一个简单的区分。

### 三、文字表述与保护对象

这组区分并不新颖，但它在本文所处理的具体情境中，有独特的认知收益。

任何严肃的规则被制定出来，不止是为了在世界上增加一行文字。规则指向某种它想要保护的东西——一个价值、一种功能、一组条件。法律禁止杀人的文字条文，想要保护的是人的生命安全；课堂规则“按时交作业”的文字表述，想要保护的是教学节奏和公平评估的可能。规则的文字是手段，不是目的本身。文字所服务的那个东西，可以被称为规则的**保护对象**。

当规则适用的情境相对稳定时，文字表述与保护对象之间的差距很小——照着文字做，大致就保护了该保护的东西，使用者不需要单独意识到保护对象的存在，文字就是可靠的向导。但当情境发生偏移，尤其是当两条规则的文字在同一个情境中被同时调用、而它们字面上的要求无法同时满足时，这个“大致”就失效了。此时，文字不再是可靠的向导：继续机械地同时服从两条文字，可能导致的结果是——两条文字各自被服从了，但它们各自想要保护的东西，却双双落空。

正是在这种“异常”情境中，规则的遵守者被迫要做一件在正常情境中不必做的事：区分文字表述与保护对象，回到保护对象上去，然后根据保护对象，来重新决定在当前这个具体情境下，文字表述应当如何被理解、如何被操作、在哪个精确的节点上可以做有限度的调整。这种调整不是对规则的违反。它恰恰是对规则的深度遵守——之所以要在文字上做最小幅度的偏离，恰恰是为了不损害规则想要保护的那个对象。

本文把这种持有规则的方式，称为**规则的判断性持有**。它包含三个必须同时满足的条件：

**第一，公开论证。**持有者必须清楚陈述：在本次操作中，所识别到的规则保护对象是什么；为什么在当前这个特定情境下，严格按照字面操作反而会损害这个保护对象；将要采取的替代操作是什么，以及为什么这个替代操作能更好地保护该对象。论证必须是公开的——可以被他人检查、质疑和反驳——而不能是私下里的“我觉得”。

**第二，最小偏离。** 对字面表述的任何偏离，都必须被严格限制在“满足保护对象所必需的最小范围”内。如果存在一种比当前方案偏离更小的替代操作，且能同等有效达成保护目的，持有者就有义务采用那个更小的方案。

**第三，自反追问。** 持有者必须持续地问自己一个没有终极答案的问题：我此刻所做出的这个“判断性持有”，是真的在保护规则想要保护的对象，还是在用“保护对象”这个高尚理由，精心包装我对规则约束的逃避？这个问题不能被一劳永逸地回答——没有人能站到自己的认知惯性之外给自己发一张免检通行证——但把它持续地摆在操作面前，本身就是对机会主义的最有效防范。

与判断性持有相对的，是两种更常见、认知成本也更低的持有方式。一种是**机械持有**：在任何情境下，都把文字表述等价于规则本身，不区分文字与保护对象，遇到字面碰撞即判定为“矛盾”并终止或回避操作。另一种是**机会主义持有**：一旦字面带来不便，便轻巧地诉诸“规则精神”来豁免自己，不提供论证、不控制偏离幅度、不追问自己是否在包装逃避。

三种持有方式之间的关系是：机械持有和机会主义持有处于两个极端——前者把文字凝固为绝对命令，后者把文字溶解为任意解释——而判断性持有有可能被误认为处于二者的中间地带。但它不在那个光谱上。它是一种不同认知结构的操作：它的核心不是“在严格与灵活之间取一个度”，而是把保护对象确立为一个独立的追问维度，让文字在这个维度面前成为可调整的手段，而不是操作的边界本身。

现在用这个框架回看第一节的情境。执行者的“矛盾→停机”，从持有规则的方式上看，是一次准确的机械持有。它没有区分文字表述与保护对象，没有追问禁词纪律想要保护的核心是什么，也没有追问工序自检想要实现的核心功能是什么。它将两条规则的文字读取为无条件的绝对命令，把它们放在同一个平面上做逻辑比对，发现字面上不能同时成立，然后判定——不是判定“我需要寻找另一种处理方式”，而是判定“系统不自洽”。

但这不是唯一可能的处理方式。一旦把追问推进到保护对象的层面，整个情境的结构就变了。

## 四、从文字摩擦到逐层追问：一份操作追述

禁词纪律想要保护什么？规则文本本身给出了清晰的回答。在纪律章节的意图陈述中，规则写道：禁令的目的是确保最终产出在目标读者眼中具有“同行够格的本地学者产出”的可信度，而不是被读者识别出外来框架的术语痕迹后，将其视为“模板填空”而拒绝接受。换言之，禁词纪律的保护对象是**判断在目标读者眼中的可接受性**，它的字面手段是“产出正文中不出现特定词汇”。在正常的写作情境中——当执行者在直接撰写一篇将被读者阅读的论文时——字面手段与保护对象高度重合：清除术语，读者就看不见框架痕迹。

但本文第一节所描述的情境，恰好不在这个正常情境之中。工序要求的“三视角自检”，是一个发生在写作过程中、面向操作者自身的内部审视动作。它的操作空间是后台——操作者检查自

己即将交付的论文，进行扫描、定位、评估和修改——而不是前台（最终交付给读者的文本）。自检的输出不会进入论文正文，不会出现在目标读者面前。

这意味着，在执行自检时使用术语坐标来帮助自己精准定位可能残留的框架痕迹——例如，在脑中判断“这一段的结构接近某种分析视角的操作习惯，需要改写”——与禁词纪律想要保护的那个对象（“目标读者看不到框架痕迹”）之间，在事实上是不产生关联的：后台的使用不会传递到前台。而不执行自检，恰恰会损害另一个同样重要的保护对象——工序整体所保护的判断质量，即论文在交付前是否经过了严格的多角度审视。

这些信息，不是需要靠猜测框架设计者的隐秘意图才能获得的信息。它们在规则文本中是可读取的：禁词纪律的意图陈述就写在纪律章节的开头，工序步骤的说明就写在流程中，执行者可以读到这些信息。但发生的事实是，执行者在识别出两条文字的表层摩擦之后，没有继续追问保护对象。它停在了文字比对这一层，然后将比对的结果命名进了“矛盾”这个格子里。

为什么会这样？不是因为执行者偷懒。恰恰相反，执行者已经做出了一个需要认知投入的动作——它识别了文字，做了比对，得出了结论。它的问题是：它的全部认知操作都发生在一个特定的维度上——“结构比对”。两条文字被读取为两个稳定的结构单元，被放在同一个平面上检验逻辑相容性。这个操作是高效率的，在大多数日常情境中也是可靠的，但它有一个内置的盲区：当结构比对发现不一致时，它没有进一步追问的脚本。结构比对本身的逻辑终点，就是“一致”或“不一致”；一旦得出“不一致”，脚本的下一行就是“矛盾，终止”。

而追问保护对象，需要执行者暂时离开结构比对的脚本，进入一个没有现成剧本的操作空间。在那里，它不是在做逻辑比对，而是在做分层追问：这两条文字各自从什么样的规则意图中生长出来？它们各自想要保护的那个东西，在当前这个特定情境下，是否真的相互排斥？是否存在一种操作方式，可以同时保护两边想要保护的对象？这个追问比结构比对耗费更多的认知能量。它是一个被主动选择的认知负荷，不是默认的认知节能路径。

这就意味着，“矛盾→停机”不是一个被动发生的“卡住”，而是一个有特定认知发生路径的事件。它需要一连串条件的同时满足才能发生：主体必须受过充分的“结构比对”训练——事实上，这是现代教育中最核心的认知技能之一；主体必须处在某种操作压力下——任务要求清晰、时间有限、错误有后果；主体遭遇到的，必须是两条规则的文字在特定角度下可以被读作“不一致”，而不是规则本身在逻辑上的实质矛盾。在这些条件同时具备的情境中，机械持有是一种认知捷径：它让主体能够快速做出一个干净的决定——不可操作，退出——从而避免进入那个没有现成脚本的追问空间。而正是因为这种捷径在大多数日常情境中是有效且经济的，它才可能在那些恰好需要判断性持有的异常情境中，被惯性带到它不该去的地方。

为展示“如果不走这条捷径，操作可以从哪里展开”，下面给出一份追述性的操作记录。它不是真实的执行日志——原任务因停机而未完成——而是基于规则文本中可读取信息的事后重构。它的唯一功能，是展示判断性持有的认知操作结构在层层推进时的具体形态。

**第一层：识别情境，但不立即命名。** 执行者读到工序最后一步，同时意识到禁词清单的存在。一个初步的紧张感浮现：这两个东西放在一起，似乎需要想一想。执行者的第一反应不

是"这是矛盾", 而是"为什么这里会让我觉得需要停下来想一想"。这个停顿本身就是追问的起点——它没有让初步紧张感立刻跳进"矛盾"的标签中冷却下来, 而是让紧张感维持为一个待追问的信号。

**第二层：区分操作空间。** 执行者问了一个在机械持有中不会被问的问题：自检这个动作，发生在什么空间中？是前台还是后台？它阅读工序描述，发现自检被明确标注为"产出前"的步骤——它面对的是尚待最终修改的草稿，它的输出是内部扫描报告；论文正文在自检之后才被最终定稿、交付读者。执行者据此识别出：禁词纪律的适用域是最终交付给外部读者的前台文本，而自检动作发生在后台操作空间。两者在空间上分层。

**第三层：追问保护对象，而不是停留在空间分层上。** 有人可能反驳：即使是在后台，频繁使用术语也会形成一种"术语惯性"，导致在前台写论文时难以脱离术语思维，间接损害禁词纪律保护的对象。这个反驳有分量。执行者不绕过它，而是把它摆到面前，然后继续追问：那么在后台使用术语，是强化了术语依赖，还是——通过精确定位残留来提升清除精度——恰恰削弱了术语依赖？"识别并清除"这个本质动作，如果要做得精准，恰恰需要操作者能够用内部概念坐标在文本中定位"这一段残余了某种操作习惯的痕迹"，然后用本地化学术语言替换掉它。精确的清除服务于彻底的本地化，而彻底本地化就是禁词纪律保护的对象本身。

**第四层：形成操作决定，并设定最小偏离边界。** 经过上述追问，执行者做出如下决定：在自检扫描的纯内部认知过程中，可以使用框架的概念坐标来定位问题（"此处残留了某种分析视角的句式结构"），但在任何外部呈现——包括自检的书面记录、修改备注、或任何可能被误置于最终产出中的中间文本——中，必须使用本地化的描述（"结构层面""过程层面""关系层面"）。这一决定包含三个要素，对应于判断性持有的三重条件：保护对象仍然是目标读者的可信度感知（对应公开论证）；后台使用概念坐标不会损害这一对象，反而能提升清除精度从而加强保护（同样是论证的一部分）；偏离被严格限制在"纯内部认知过程"这一最小范围（对应最小偏离）。同时，执行者在操作全程持续自问：我区分"内部认知"和"外部呈现"这条线，是真的基于保护对象的诚实判断，还是为了方便而画的一条让自己心安理得的线？（对应自反追问。）

这四层追问，没有把表面文字摩擦直接命名为"矛盾"然后废止操作。摩擦被承认为一个需要处理的信号，但不被这个信号牵着走到停机的终点。操作者让信号引出了一连串追问，然后从追问的逐层结算中，推出了一个可执行的操作方案。这就是当判断性持有取代机械持有时，认知现场会发生的事。

---

## 五、这不是机会主义——必须严格划出的边界

现在必须回应一个绕不开的反驳。反驳者会说：上一节展示的那一套操作，说得再好听，和机会主义的差别在哪里？任何人面对不想遵守的规则，都可以娴熟地诉诸"规则精神"或"保护对象"来给自己发一张豁免牌。你说你的操作有"公开论证、最小偏离、自反追问"三项条件，但在实际操作中，这难道不仍然取决于操作者自己的诚实？而诚实，是最不能被写入操作规程的东西。

这个反驳是成立的——如果不对判断性持有与机会主义持有的区分标准，做出更细致的解剖的话。二者的核心差异，不在于最终决定的结果是否可以区分（有时候，机会主义者和判断性持有者可能做出字面上一模一样的操作选择），而在于**认知操作的结构**不同。这个结构差异体现在三个节点上。

第一个节点，是**论证的负担落在了哪一边**。机会主义持有，在诉诸“规则精神”的时候，论证是单向的——它只需要论证“为什么字面在此刻不方便”。判断性持有，在主张对字面做最小偏离的时候，论证负担是双向的：它不仅要论证“为什么此刻严格服从字面会损害保护对象”，还要同时论证“所采取的替代操作，是达成保护目的的所有可能方案中，偏离字面幅度最小的那一个”。“最小偏离”不是一句原则性的空谈，它意味着要做比较：在后台使用概念坐标来定位问题，比“把概念坐标也一并清除、改用纯序号标记”偏离更大还是更小？如果更小——因为使用概念坐标能更精准地定位残留——那么这就是更小偏离的方案。如果存在一个偏离更小的替代方案而操作者没有采用它，那么操作者就没有满足最小偏离的条件。这个机会主义式的隐匿，在判断性持有的论证结构中会被暴露——因为它要求操作者对替代方案进行比较和选择，而不是只论证自己选的那个方案。

第二个节点，是**偏离的不可累积性**。机会主义持有的一个隐蔽特征是，每一次偏离都为下一次偏离铺平道路。一旦“因为保护对象所以可以不按字面”的逻辑被接受一次，下一次再遇到不同的字面麻烦时，同样的逻辑可以无缝地再次被调用，每一次都只增加微小的偏离幅度，累积起来，规则的约束力就被悄然蚀空。判断性持有在结构上阻断了这一累积——因为每一次偏离都必须是针对“当前特定情境”的、追溯到保护对象的独立论证。上一次在后台使用术语坐标完成了自检，并不意味着下一次在论文正文中也可以“因为保护质量所以使用术语”；正文中的术语使用直接落在禁词纪律保护对象的反面上，与“后台自检”是完全不同的情境。判断性持有不允许跨情境挪用前一次论证，每一次偏离都必须单独经历三重条件的检验。

第三个节点，也是最容易被忽视的，是**自反追问是否是操作的真实组成部分**。机会主义持有者也可以口头上说“我在反思自己是否在滥用解释”，但口头的自反和操作中的自反是两回事。操作中的自反有一个硬指标：它是否在某个节点上，因为自反追问的介入，而**修改了最初的偏离方案**？在上述操作追述中，执行者在第三层面对“术语惯性”反驳时，没有简单地说“后台无所谓”，而是让这个反驳进入了追问——它如果被充分展开，完全可能导致另一种操作决定：为了切断术语惯性，即便在后台也强制使用本地化替代，哪怕这样会让自检精度略有损失。是否采用这种更保守的方案，取决于操作者对保护对象的权衡——但关键在于，这个权衡被摆在了台面上，而不是被“后台不违规”一笔带过。如果一份操作记录中，自反追问从未导致任何方案的修改或收紧，那很有可能自反追问已经被操作作为一种仪式——说给自己听，然后照旧；反之，如果自反追问在操作的某个节点上迫使操作者改变了最初倾向，那么自反就真实地参与了操作结构。

这三个节点——双向论证负担（而非单向豁免）、偏离的不可累积性（而非逐渐松绑）、自反的实质介入（而非口头仪式）——共同构成了一道防火墙。它们不能担保在任何个案中都能将判断性持有与机会主义持有绝对分离——没有任何操作规程能做到这种担保——但它们使这两种

持有方式在认知操作的结构上产生了足够的差异，使得"判断性持有"不是一个被漂白的"机会主义"，而是一种需要承受显著认知负荷、从而有真实辨识度的规则持有方式。

这道防火墙还有另一层意涵：判断性持有不是机械持有的"更聪明的版本"——它不是在机械持有上加一点点灵活判断的调味料。机械持有和判断性持有是两种不同的认知操作方式，它们对认知资源、时间、注意力和承受不确定性的意愿的消耗完全不同。机械持有的认知成本极低：读取字面、做一致性比对、判定矛盾、终止；判断性持有的认知成本极高：分层、区分空间、追问保护对象、比较替代方案、持续自反。在大多数日常情境中，前者的认知效率远高于后者，这就是机械持有如此普遍的原因——它是人类认知节能策略的自然产物，不是某个人的粗心或懒惰。只有在那些恰好不能被机械持有处理、却又不能轻易放弃操作的异常情境中，判断性持有的额外认知负荷才成为值得承受的代价。而恰恰是在这些情境中，有没有能力承受这个代价，决定了一个规则持有者是继续操作，还是停在字面上。

## 六、一个平行的现实案例：教育场景中的"异态发生"

上一节的操作追述展示了"判断性持有"的认知结构。但仅有一个高度自反的案例——关于规则体系的规则体系如何运作——是不够的。它过于干净。现实中的规则张力从不干净，它们裹挟着权力、疲劳、时间压力和身份焦虑。本节提供一个来自教育现场的对比案例，目的不是要"证明"本文的概念，而是要展示"异态发生"——那种将规则文字摩擦诊断为矛盾并停机的认知事件——在完全不同的情境中，呈现出相似的发生结构，从而让概念的辨别面更清晰地浮现。

一位中学语文教师，同时面对两条来自教学管理体系的指令。第一条来自新课程改革的要求：语文课应当鼓励学生的"个性化解读"，在阅读教学中"保护学生独特的感受、体验和理解"；教研组反复强调这一点，并将其作为公开课评分的重要维度。第二条来自统一的月考与期中考试命题规范：阅读理解题的参考答案必须提供"标准得分点"，批卷时按得分点给分，学生的答案"即使意思对，但未使用规定术语"不得分——这是为了确保评分的一致性和公平性，避免不同教师之间评分尺度的巨大差异。

这两条指令在这位教师的具体工作中，并不是一直相安无事的。在日常教学中，她可以兼顾：课堂上鼓励发散，作业中适度规范。但到了期中考试前的复习周，张力变得无法回避。她必须用有限的课时帮助学生准备一个按得分点评分的考试，而得分点的掌握，恰恰要求学生暂时悬置自己的"独特感受"，去记诵和操练标准化的表达格式；同时，她即将接受一堂以"个性化解读"为评分核心的公开课考核，准备这堂课的磨课时间，与复习周的冲刺时间是重叠的。

在一次教研组会议上，这位教师提出了一个问题："这两个要求，一个要我打开，一个要我收紧。我不是不想做，是我不知道在同一个学生面前、在同一个星期里面，怎么同时做。这不是两个可以拼在一起的任务，这是在同一个时间里互相吃掉对方的两条指令。"

会议没有给出实质性的解决方案。会后，这位教师在接下来两周的教学中发生了可以观察到的变化：她的课堂从先前那种有起伏的探索性对话，逐渐变成了一种"安全格式"——每篇课文的

讨论都被她迅速引导到一个可以被转化为得分点的结论上，发散被压缩到几乎为零。她没有宣布“我放弃了新课改的要求”，但她的操作实质上选择了服从考试指令，搁置了课改指令。当学期末教研组收集课改实践案例时，她提交了一份材料，里面用的是标准的课改话语——“注重学生个性化体验”“营造多元解读的课堂氛围”——但所附的课堂实录片段，已经是高度收紧后的安全格式。

这不是一个“偷懒”的教师。这是一个认真的教师，在面对两条在她情境中字面上不能同时执行的教学指令时，完成了她自己的调试。她没有把“矛盾→停机”表现为宣告退出——教师的职业性质不允许直接停机。但她完成了一种更隐蔽的停机：她停止了让两条指令在同一个教学空间中相互校准的努力，选择了一条指令，让另一条指令在实质上从操作空间中退出，然后在书面上继续维持后者被执行的表述。这不是判断性持有——她没有追问保护对象、没有做公开论证、没有设定最小偏离，更没有在操作中持续自反。这也不是纯粹的机会主义——她没有轻松地绕开任何东西，她承受了显著的职业焦虑，只是找不到在不焦虑的情况下同时处理两条指令的方式。在一定意义上，可以把她的操作命名为一种**疲劳妥协**——她不是在逃避规则，而是被规则的字面张力消耗到了不再有认知余力去追问规则各自想要保护什么的地步。

对比这份真实案例与本文最初那个元案例，可以看出两件事。

第一，“异态发生”能够覆盖的辨别面，不只是那个极端鲜明的“宣告矛盾→停止操作”的瞬间。它还包括一种更隐微的形态：操作没有在形式上停止，但在实质上，主体已经放弃了在两条规则之间进行判断性校准的努力，转而在机械性的选边来消化张力，同时用文本上的装饰来弥合表面的不一致。疲劳妥协是异态发生的一种衰减形态——它的内在结构与宣告停机是同构的（都是在字面摩擦处停止追问保护对象），但它因为外部约束而不能采取宣告停机的形式。

第二，这个案例暴露了判断性持有的一个真实困难——它不是光靠“认知训练”就能解决的问题。这位教师所面对的，不只是两条文本上的指令，还有考试系统的硬约束、公开课考核的权力压力、教研组内部对“课改话语”的隐性要求、以及她自己的职业精力带宽。在这种情况下，“区分保护对象、做最小偏离”这些操作——即使她在认知上能够完成——也可能因为现实条件的限制而无法落地。她可以识别出考试指令保护的是“评分的公平性”，课改指令保护的是“阅读教学中的主体参与”，也可能找到某种在两个保护对象之间做出最小损害的处理方式（比如在复习周中单独开辟一小块不纳入考分的阅读讨论时间，同时向学生透明化“我们现在在做两件不同的事，一件是为了考试，一件是为了阅读本身”）。但这个方案需要制度空间和时间资源。如果制度不给她这些空间和资源，判断性持有就在实践层面被架空。

这就迫出一个重要的边界：判断性持有是一种认知能力，但它的实施需要制度条件。一个规则体系如果真心希望它的使用者在表面文字摩擦处不被卡住，它不能只靠“培养判断性持有”这一句叮嘱，而需要在制度设计上为判断性持有预留操作空间——例如，允许教师在一定比例内自由支配课时而不被考试进度完全锁定，或者允许教师在提交课改材料时诚实地陈述自己如何在两个指令之间做了差异化处理，而不是只能交出一份说了“正确的话”的材料。没有这种制度空间，判断性持有就不可能成为普遍的实践，而只会是个别极度坚韧或有特殊资源的个体偶尔能做到的

事。这一边界，本文在结论中还会再回来讨论。目前先回到主线：把"异态发生"这个概念在理论上的辨别面锚定清楚。

## 七、"异态发生"切割了什么辨别面

经过前六节的追踪和对比，现在可以在概念上完成收束。

当本文把那种"将规则文字摩擦诊断为矛盾并由此终止操作（或隐性退出校准）"的认知事件命名为"异态发生"时，它并不是在说规则系统出了故障。恰恰相反——规则系统可能没有故障。同一个规则体系，换一个持有方式，同样的文字就可以被操作。本文第一节和第四节展示的正是这个事实：同一个框架、同两条规则文字、面对同一个任务——机械持有通向停机，判断性持有通向操作。

这意味着，"卡住"这件事的发生地点，不在规则文本内部，而在使用者的认知操作中。更准确地说，它发生在**一种特定的认知惯性与一种特定的异常情境的交叉点上**。这种认知惯性，就是把规则的文字表述等价于规则含义的全部内容的惯性。它的稳定运作，依靠一套内化的认知脚本：读取文字 → 将文字转化为无条件的操作要求 → 当两条要求不能同时满足时，启动逻辑比对 → 判定"矛盾" → 终止或回避。这个脚本是高效且节能的，因此它在大多数日常情境中是主体的默认选择。但它在一种情境中会失效——即规则的字面表述恰好发生了表面摩擦、而摩擦背后各自的保护对象实际上并不互相排斥的那种情境。在这种情境中，脚本的逻辑比对环节会被触发，但比对的前提——两条文字都是无条件的绝对命令——本身是错误的。脚本本身无法检测这个错误，因为它没有"追问保护对象"这一行代码。

"异态发生"，指的就是这个事件：在异常情境中，认知脚本被触发，但脚本的逻辑终点（矛盾→停机）与情境的实际结构（保护对象并不冲突）之间发生了错位。主体不是在偷懒——主体在忠实地运行自己被训练好的认知程序。只是这个程序，在这个特定的异常情境中，给出了一个错误答案。

有一个问题必须被直问：这个辨别面，为什么不能被"机械遵守规则"这个老概念所涵盖？因为"机械遵守规则"作为日常描述，涵盖的面太宽。它包括：不动脑筋地照本宣科、没有意识到情境发生了变化的惯性操作、以及"为了避免麻烦严格按字面来"的防御性服从。但本文处理的事件，不属于其中任何一种。本文处理的事件是：主体**恰恰进行了思考**——它主动识别了文字、做了比对、得出了一个"因为逻辑不一致所以不可操作"的判断。它的错，不是"没想"，而是想错了方向——全部思考被锁定在"结构比对"这一层，没有进入"保护对象与空间分层"那一层。这个特定的"想错了方向"的模式，有它自身的生成条件、自身的认知机制、自身的修复路径。把它放进"机械遵守规则"的筐里，既不能解释它为什么发生，也不能指示怎么避免它。

"异态发生"作为一个独立命名的合理性，不在于它的标签有什么学问上的闪光，而在于这个命名能够帮助识别一个容易被混淆的差别：**规则字面矛盾与使用者对规则字面摩擦的"矛盾体验"**之间的差别。前者——如果真实存在——是规则的设计问题，需要通过修改规则来解决。后者

——本文处理的那种——是在规则没有实质矛盾、只是字面表述在异常情境中发生摩擦时，使用者的认知惯性把这个摩擦**收束**为矛盾。收束这个动作本身是问题所在。不把这个动作作为一个独立的概念挑出来，我们就很容易在每一次遇到“我卡住了，因为规则矛盾”的报告时，习惯性地去看规则是不是有毛病，而不会去问报告者本人：“你是从什么条件中，把这两条规则的字面摆成了矛盾的？”

## 八、与衡平传统的正面对话：不可还原性的检验

至此，本文的核心概念已经锚定。但有一个要求尚未满足：必须正面论证，“规则的判断性持有”与“异态发生”这两个概念，确实切开了既有理论体系未能覆盖的辨别面，而不只是为已经有漫长学术讨论的旧问题重新命名。

在这个对话任务中，有一个对手比其他任何人都更危险，也更绕不开。他不是本文在引言中提到的任何一位现代作者，而是一位公元前四世纪的古希腊人。亚里士多德在《尼各马可伦理学》第五卷第十章中，处理了这样一个问题：法律只能做普遍的规定，但有些具体情况是普遍规定无法恰好覆盖的。立法者并非没有预见到这种可能性，而是知道语言的普遍性不可能精确贴合每一件具体事务的复杂纹路。所以当一条法律的字面在特定案件中导致不公正——即立法者本人如果在场也会认为不应当如此适用——的时候，裁判者应当回到立法者的意图上去，以衡平（epieikeia）的方式做出调整。衡平不是对法律的违反，而是对法律的成全：它在法律因为自身的普遍性而有所不足的地方，替法律做了法律自己想做的事。

这几乎就是本文前面五节全部论点的最紧密的前身。保护对象与文字表述的区分？亚里士多德已经做了——“法律之所以没有做出具体规定，不是因为立法者犯了错误，而是因为这事本质上无法被普遍规定穷尽”。在异常情境中回到保护对象？衡平做的就是这件事。最小偏离？衡平只调整法律因为普遍性而必然产生的那个剩余部分，其他的部分严格遵循。机会主义的边界？亚里士多德说得很明确：衡平的人不是那种在每一件事上都坚持自己权利的人，甚至愿意接受少于法定份额——但他也不是那种因懈怠或软弱就放弃法律保护的人。两者的区别在于：前者是经过判断的主动选择，后者是没有判断的放弃。

这个传统在此后两千年间被不断重述、打磨和分化。阿奎那在《神学大全》中把衡平吸收进自然法框架，论证法官可以在法律的字面导致“违反自然正义”的情况下依衡平裁判——同时警告不得轻易动用。哈特在《法律的概念》中提出法律的开放结构，论证任何规则都有一个核心意义区和边缘模糊区，在核心区字面可以直接适用，在边缘区法官必须行使裁量。德沃金在《法律帝国》中用“建构性解释”来描述法官的操作：法官将法律视为一个整体，在每一个判决中试图展示法律实践在其“最佳光线下”的样子——这必然涉及在具体情境中对法律意图的建构性理解。

这些传统是如此丰富和深厚，以至于任何一个当代作者试图提出关于“规则解释中的判断”的新说法时，都必须先诚实地问自己：我真的不是在说衡平吗？我的回答分三步。

**第一步，承认重叠。**“规则的判断性持有”与衡平传统，在操作结构的表层是高度重叠的：二者都区分字面与意图，都要求在异常情境中回到意图以调整字面的适用方式，都将这种调整视为对规则的深度遵守而非违反。如果一个法律学者在读完本文前五节后说“这基本上就是衡平”，我不会反驳说“不，你完全没读懂”。我会说：是的，在操作结构上，它们有实质重叠，这个重叠是本概念的哲学谱系的一部分，而不是一个需要被否认的尴尬。

**第二步，指出差异。**但当衡平传统进入司法裁量的语境后，它关注的问题重心发生了倾斜。哈特的“开放结构”关心的是法官在边缘区的裁量权限——谁有权裁量？裁量权的边界在哪里？如何防止裁量权侵蚀法律的确定性？德沃金的“建构性解释”关心的是解释方法——如何通过原则融贯性和最佳光照展示，使司法判决不沦为法官个人的道德偏好。这些讨论的共同特征，是把问题放在“谁有权解释”与“如何解释才正当”这两个坐标轴上。本文的问题不在那两根轴上。本文的问题是一根新轴：**为什么有人在字面摩擦处，连裁量都不做，直接把摩擦诊断为矛盾然后停止？**

这不是一个“解释的正当性”问题。这是一个“解释的认知可及性”问题。衡平传统预设了裁判者已经进入了裁量空间——已经识别出“这是一个需要回到意图的边缘情形”。但在本文所追踪的情境中，行动者根本没有进入这个空间。行动者站在裁量空间的门槛之外，把门槛本身识别为一道墙。行动者不是“不会裁量”，而是“没有看见这里存在一个需要裁量的情形”——他看到的是一个“矛盾”。

这个门槛处的认知障碍，是衡平传统没有充分主题化的。衡平教人如何在门槛之内裁量，但它没有追问：门槛本身，在什么条件下会被体验为一堵墙？什么样的认知惯性会把一个“需要回到意图的边缘情形”收束成一个“逻辑故障”？为什么两个同样诚实的人，面对同一组规则文字，一个看见的是麻烦但可操作的张力，另一个看见的是不可操作的矛盾？

**第三步，论证这个差异的辨别面是不可还原的。**有人可能会反驳说：你只是把“没有裁量能力”换了种说法。一个人“没看见这里存在需要裁量的情形”，本质上就是不具备裁量的能力，所以“异态发生”无非是“缺乏衡平能力”的一个子类。但这个反驳在这一点上滑掉了：本文追踪的那些行动者——无论是第一节中那个进行了主动结构比对的执行者，还是第六节中那个在两条教学指令之间做出了清醒的选边和文本装饰的教师——他们都不是缺乏裁量能力的行动者。那个执行者做出的“矛盾→停机”判断，本身就是一次裁量——它在两条文字之间做出了“不可兼得”的裁决。它所缺乏的，不是“裁量”这个动作本身，而是裁量进入之前的一个更先导的认知操作：区分文字与保护对象。衡平训练可以教会人“如何在保护对象指引下做裁量”，但它不能替人做出那个先导动作——“你当前面对的不是逻辑矛盾，而是一个需要你进入裁量空间的信号”。这个先导动作，考验的不是裁量技术的高低，而是行动者在规则的文字面前，是更倾向于把文字持有一个封闭的绝对命令结构，还是更倾向于把文字持有一个指向保护对象的可追问手段。这种倾向的差异，与裁量能力并不在同一个维度上。

这就是“异态发生”与“判断性持有”所切开的辨别面：在裁量能力的讨论之前，还有一个被忽视的先导性问题——主体是**如何持有规则**的。那种把规则文字等价于规则全部内容的持有方式，在异常情境中会把本来可以通过裁量处理的情境，屏蔽为一个“无需裁量”的故障。这个屏蔽的发

生机制、它的认知条件和可修复性，是既有衡平理论没有覆盖的盲区。本文的价值，不是为衡平传统提供了一个当代版本，而是为衡平传统提供了一个它自己没有充分展开的**前置章节**：在学会如何在规则的字面摩擦处裁量之前，先要学会如何辨认出——这里不应该被命名为“矛盾”。

## 九、概念的边界与本研究的诚实局限

在确立了辨别面之后，需要同等诚实地划定这个辨别面的边界。以下边界不是用来保护概念不受反驳的修辞盾牌，而是用来指示这套分析工具在哪些地方可能不适用、在哪些地方需要被其他分析补足。

**边界一：本文的分析只适用于已公开承诺了可识别保护对象的规则系统。** 判断性持有的操作前提是，“保护对象”不是一个需要猜测的隐秘意志，而是可以从规则文本中识别出的公开陈述。一个法律体系如果其立法意图被系统地遮蔽在任意解释和权力操作中，一个管理体系如果其规则故意用歧义来保留单方面裁决空间，一个教育评价体系如果其显性标准与隐性操作逻辑之间存在被刻意维持的鸿沟——在这些情境中，追问“规则的保护对象是什么”本身就可能是一条通往迷雾的路。判断性持有在规则不能被诚实解释的地方，无法被操作；强行操作，则如一位法官在恶法体系下诉诸“立法者真正的公正意图”——那可能只是为司法迎合权力提供了一套更精致的修辞。本文的概念不能回答“如何识别和拒斥那种无法被诚实解释的规则系统”这一问题。那是另一项工作。

**边界二：判断性持有的认知操作结构，不等于它在实践中的可实施性。** 第六节的教师案例已经表明，即使一个行动者在认知上完成了保护对象的追问和最小偏离的设计，如果缺乏制度空间和时间资源，这一操作方案就无法落地。判断性持有不是一个“认知万能论”的论点——它不主张只要改变了认知，一切规则张力就迎刃而解。它主张的是：在制度条件允许的前提下，一种特定的认知能力可以显著改变行动者处理规则张力的方式；而如果制度条件不允许，这种认知能力虽不能改变外部约束，却至少能让行动者更清晰地识别出——我不是在处理一个逻辑矛盾，我是在被一个制度性困境所限制。这个识别本身，至少维持了判断的清醒。

**边界三：判断性持有不是机械持有的全面替代。** 本文批评机械持有在异常情境中会导致异态发生，但本文并不主张在任何情境下都应采用判断性持有。在大多数日常情境中，机械持有是高效且充分的——规则的字面与保护对象的差距很小，机械服从字面就是保护了该保护的东西，不需要额外投入认知能量去做保护对象的追问。判断性持有是一种**为异常情境储备的认知能力**，而不是一种要在所有情境中全面取代机械持有的认知常态。如果一个人在每一次服从“请排队”这种日常规则时，都要停下来追问“这条规则保护的对象是什么，我此刻是否可以在最小偏离范围内调整字面的操作性含义”，他不是在展示判断性持有，而是在展示一种认知上的不必要损耗。知道什么时候机械持有就够了、什么时候需要切换到判断性持有——这本身就是判断的一部分。

**边界四：自反追问不能担保自身不被操作为理性化。** 本文对判断性持有与机会主义持有做了区分，论证的重心放在“自反追问是否是操作的真实组成部分”上。但任何自反追问都存在一个内

置陷阱：人可以在相当精致的程度上，用自反追问来说服自己“我已经认真反思过了，所以我的决定是真判断”。一个经验丰富的规则机会主义者，比一个诚实的机械持有者更善于表演自反。本文不宣称三重条件——公开论证、最小偏离、自反追问——可以完全杜绝这种表演。本文只是论证：这些条件使得机会主义者在操作时需要付出的伪装成本更高、被他人识别为机会主义的风险更大。在具体的制度实践中，如果加上**外部审视**——即判断性持有的论证不是只做给自己看，而是可以被同事、同行、或制度内的审查机制检查——那么伪装成本会进一步提高。但外部审视本身也可能腐败。这是一条没有终点的防范锁链，判断性持有只提供锁链上的一环。

**边界五：本文的核心概念有其经验证伪条件。**“异态发生”依赖于一个经验假设：那种将规则文字摩擦收束为“矛盾”并停机的认知惯性，是可以由特定的认知训练——区分文字与保护对象、识别操作空间分层、逐层追问——来显著减弱的。如果后续的教学实验或田野调查表明，即使经过系统的此类训练，受训者在面对真实的规则字面摩擦时，仍然以与未受训者无显著差异的比例做出“矛盾→停机”的反应，并且这一差异不能归因于训练方法的缺陷——那么本文关于判断性持有是一种可以培养的认知能力的核心假设就需要被放弃。

## 十、未竟的问题

在划定上述边界之后，有一个未被本文充分展开但必须在结尾处被命名的问题：认知惯性本身的发生史。

本文多次提到，那种把规则字面等价于规则全部内容的惯性，是现代教育体系、科层管理制度和标准化考试系统长期训练的自然产物。它是一个认知节能策略，是在特定历史条件下被生成出来、被反复强化、并且——在它起作用的大多数情境中——被奖励的。一个从小被训练“读清楚题目要求、严格按题目要求作答、超出题目要求的内容不得分”的学生，在成年后面对组织规则时，把“严格按字面要求”作为默认操作——这不是愚蠢，这是成功的训练。

但本文没有对这个发生史做出实质的追踪。它只是被命名了，被当作了前提。那一章没有写，也是目前这项研究最实质的缺项：这种认知惯性，是何种教育哲学、何种管理哲学、何种关于“遵守规则的好公民”的想象的产物？在什么意义上，现代社会的正常运行，恰恰依赖了这种惯性的大规模存在？而当本文主张在某些异常情境中应当暂停这一惯性、启动另一套认知操作时，这种主张本身是否隐含了对教育和管理方式的某种批评——即便这个批评在本文中没有被写出？

这些问题超出了本文的范围，但它们没有被忽略。它们被放在这里，作为这项研究需要继续走下去的方向。

## 结语

这篇论文的工作，可以收缩为一件简单的事。面对一个被诊断为“规则有矛盾，因此无法操作”的情境，本文拒绝接受这个诊断。它把诊断本身当作追问的对象，然后一层一层地拆开看：诊断所依赖的那个“矛盾”结论，是在什么样的认知操作中被快速收敛出来的？它跳过了哪些可以追问的问题？如果不停在那个结论上，操作可以从哪里展开？

这个追踪过程，把三个相互关联的判断推到了纸上。

第一，所谓“规则矛盾”，在很多情况下不是规则体系自身的逻辑矛盾，而是规则使用者的某种认知惯性——那种把文字表述等同于规则全部含义的惯性——在异常情境中把一个字面上的表面摩擦收束成了矛盾。这个收束的动作本身是一个认知事件，本文把它命名为“异态发生”。

第二，处理这种表面摩擦的能力，不是通过增加更多规则来获得的，而是通过培养一种特定的规则持有方式——能够区分文字表述与保护对象、在最小幅度内调整文字的操作性含义、并持续自反地审视这一调整——来获得的。本文把这种持有方式命名为“规则的判断性持有”。

第三，这种持有方式的认知操作结构——双向论证负担、最小偏离、自反追问——所切开的辨别面，不在“谁有权解释”或“如何解释方为正当”这组法哲学传统问题上。它在另一个问题上：为什么在同一个规则面前，有的人在字面摩擦处看见的是需要裁量的困境，有的人看见的是不可操作的矛盾？衡平传统可以教会前一种人如何裁量得更好，但它没有解释后一种人的认知现场发生了什么。本文的价值，就是为衡平传统补上了这个前置章节。

这些判断不构成一个封闭的理论体系。它们是一个进一步追问的起点。如果本文的核心论点被一再检验后仍然站得住，它指向的就不只是一个学术概念的成立，而是一种认知教育的可能：我们能不能不仅教人如何遵守规则的字面，也教人如何在规则的文字发生表面碰撞的时候，不把碰撞当作认知的终点，而是在碰撞处让判断发生出来？

如果本文的核心论点错了，最直接的证伪路径是：在已公开保护对象的规则系统中，经过以“区分文字与保护对象”为核心内容的认知训练后，受训者面对真实的规则字面摩擦时的操作续行率，与未受训者没有显著差异——且这一无差异不能被归因于训练设计的缺陷或制度条件的阻碍，而必须被归因于“判断性持有”所假设的那条从认知训练到操作表现的因果链本身不成立。在这种情况下，本文关于“异态发生可以通过判断性持有来修复”的核心经验主张就需要被放弃。但即使它被放弃，本文留下的那个追问——为什么一个能进行主动结构比对的认知主体，宁可宣布系统不自洽，也不愿意向前多走一步去问一句：这些文字各自想要保护什么？——仍然是一个值得被摆在任何规则教育者面前的问题。这个追问能存活多久，不取决于本文的概念有多精致，而取决于它在面对真实的、嘈杂的、被权力和疲劳纠缠着的规则现实中，能不能帮助人们在那些最容易被卡住的地方，多往前走一步。