

当消化失败成为基本功

论可译性幻觉与判断被逐级改写为技术问题的过程

作者 金华 · SDE 学员 · 约 2.1 万字 · 发表于2026年7月29日 · 深化增补版

摘要

为什么高度智识化的组织会集体错过事后看来清晰无比的致命信号？既有组织理论将这些失明归因于某种能力的丧失或压制——目标置换、防御性常规、越轨正常化——它们共享一个未被审查的预设：健康的组织本应内建一种对自身前提进行质询的能力，此能力一旦失效，灾难便接踵而至。本文论证这一预设本身正是问题的根源。通过对挑战者号发射前夜决策过程的微观解剖，本文揭示了一个此前未被准确指认的机制：当一个既有的认知格式被连续的成功加固到无需检验的程度，指向该格式适用性本身的异质质疑，会经一系列日常工序，被一步步重新分类为格式内的技术执行问题，再纳入常规的改进流程，最终在系统的自我完善叙事中湮灭。本文以“操作漂移”命名这一三步渐变位移，并论证它不需要压制反对意见、不需要恶意或愚蠢——它只需要系统正常、理性、甚至优秀地运转。本文将驱动这一漂移的深层动力命名为“可译性幻觉”，并将此概念与尼克拉斯·卢曼的操作封闭理论以及博尔坦斯基与泰弗诺的“测试格式”理论进行逐层对质，证明“可译性幻觉”提供的是一种现有的组织理论无法完整容纳的辨别维度：它解释的不是系统为何只能翻译环境刺激，而是系统如何消化对翻译格式本身的挑战，并在此消化中生产出一种自我增强的叙事闭合。本文进而论证，抵抗操作漂移所需要的警觉，不可能从系统内部精密化中生长出来；它只能来自一种在性质上不同于“格式内分析”的打断行动——置身判断——而这一行动本身，必须在发生学上被历史地理解为其生成与退化皆受特定制度条件限制的脆弱产物。全文收束于核心命题的可证伪条件。

关键词 可译性幻觉；操作漂移；认知格式；组织失明；置身判断；卢曼

一、一种被反复目睹却从未被准确命名的事实

1986年1月27日夜，莫顿·瑟奥科尔公司的工程师罗杰·博伊斯乔利在电话会议上竭力阻止挑战者号航天飞机次日的发射。他依据密封圈在低温下弹性丧失的物理特性以及此前几次发射后烧蚀加剧的原始数据，认定零下的气温将导致密封失效。他后来在总统委员会作证时描述了自己当时的状态：几乎是嘶吼着请求管理层不要发射。

NASA管理层并未直接否定他的结论。马歇尔太空飞行中心的一位主管问道：“你的意思是，你不能用定量数据证明发射是不安全的？”这句问话以技术澄清的面目出现，其实际功能却是将举证责任颠倒过来：质疑者必须在系统承认的唯一认知货币中支付他的反对。博伊斯乔利无法支付。他对材料物理特性的直觉把握、对非标准化数据之“不对劲”的感觉——这些不属于格式合规的证据。

会议结论：发射按计划进行。理由是“没有充分的数据表明发射不安全”。次日，挑战者号升空七十三秒后爆炸，七名宇航员全部遇难。

这件事已经被反复分析过了。黛安·沃恩在其经典研究《挑战者号发射决策》中将其诊断为“越轨正常化”（Vaughan, 1996）：组织在一次次接受微小异常的过程中，逐步将曾被视为不可接受的风险重新定义为可接受的常态。组织沟通研究者将其归结为管理层与工程师之间的权力距离与沟通壁垒。这两种解释各自贡献了深刻的洞见，但合在一起，仍未能触及一个更深的事实：那一夜，博伊斯乔利的声音没有被切断。他的话被所有人听见了。沟通渠道是畅通的。然而听见了，不等于被接收——因为“接收”在工程师那里的定义是“不要飞”，在管理层那里却变成了“请提供标准格式的风险评估数据”。

这不是压制。压制预设反对意见的实质被看见了，然后被有意地排除。但这里发生的是另一件事：反对意见在未被看见其实质之前，其身份就被重新定义了。它从“对既有论证框架在此是否适用的前提性质询”，被无声地重新分类为“框架内一项在证据格式上暂时不达标输入”。博伊斯乔利事后说了一句精准到令人心寒的话：“我们以为自己在谈论生命安全，而他们以为我们在谈论证据格式。”两种性质完全不同的行动——一个指向框架的适用边界，另一个在框架内等待充分证据——被系统的认知语言统一收编为同一个词：“讨论”。

这件事不是挑战者号的孤例。2008年全球金融危机前夕，少数交易员和风险管理者发出了关于次级抵押贷款关联违约风险的警告。他们在说的不是“模型参数需要调整”，而是“模型框架本身在结构性产品关联违约的情境下不成立”。这些警告的命运与博伊斯乔利的嘶吼如出一辙：被系统重新分类为“模型的参数灵敏度分析”待办事项。系统没有把它们扔掉；系统把它们纳入了改进流程——然后那些银行继续运转。他们不是在评估风险之后决定冒险；他们是在自己最精密的风险评估工序中，把“你会算不到”这个提醒处理成了“请你再算一遍”的请求，然后算了，然后继续了。COVID-19暴发初期，一位中国临床医生在同事群中发出关于非典型病毒性肺炎聚集的警告，其格式是社交媒体上的一句口语——“大家别去那里，可能有SARS”。这不是流行病学风险评估的标准格式。在既有监测体系内，这条预警在等待实验室确认的标准流程中被搁置。这里的悲剧性不（仅）在政治权力的介入——更基础的一层是：预警的格式恰好在系统被训练去接收的频道之外。系统在等待一个它认识的声音，而这个声音说的是它自己的话；当警告以外语喊出时，它被算作需要等待验证的“未经确认的信息”，而非此刻就需要做出反应的警告。

这三个案例的共同结构在于：不是系统未能处理风险；而是系统在把风险翻译成内部可处理格式的过程中，耗尽了警告原本携带的穿透力——警告原本在说“你此刻赖以判断安全的那套认知格式可能已经不适用了”，但系统将其听成了“请帮助我更好地使用这套格式”。这个操作不是失误。它是一种被连续训练得过于纯熟的系统性能力。

面对这一再重现的事实，既有组织理论给出的解释共享一个隐蔽的预设。目标置换理论假设组织本该把手段放在目的之下，但它“丧失了”这一能力。防御性常规理论（Argyris & Schön, 1978）假设组织本该具备质疑自身根本假定的学习能力，但它“防御性地”压制了公开质疑。越轨正常化理论假设组织本该维持一个不可逾越的安全底线，但它在一系列微小越轨中“漂移”了这一底线。这些解释各自捕捉了组织认知闭合的一个侧面，但它们共同默认：健康组织的认知结构中，理应内建一种能够对自身前提进行质询的自我免疫功能。当这个功能坏掉、被压制或漂移时，失明就出现了。恢复健康的路，因此就指向修复这个受损的功能——打破防御、优化指标、重建组织学习回路。

本文试图论证：这个预设是错的。并且这个错误不是枝节问题。它将整整一个家族的补救方案都锁在了同一个天花板之下。

系统逻辑的运转，在根本上，从来不做一件事。它从来不对自身的前提进行质询。这不是因为它“丧失”了这个能力，而是因为它运作的方式——把环境信号翻译为内部可操作的编码，再根据编码执行操作——从来就没有、也不需要为“我此刻赖以运转的这个前提本身是否还成立”这个动作，预留任何执行步骤。从识别输入到执行输出，中间没有任何一站叫“停下来，重新审视我们当初为什么把这一类信号判定为可接受异常”。每一步都在把环境中的不确定性转化为内部可处理的确定性——从来如此，没有例外。因此，当一套认知格式被连续的成功加固到无需检验的程度，当外部变化的信号恰好在内部格式无法还原其刺穿力的位置出现时，系统不是在犯错。系统是在最健康地运转——它在忠实地、高效地、无可指摘地执行它被要求执行的每一道工序。

这个判断迫使我们翻转追问的方向。我们不需要继续问“系统为什么会丧失自我质疑的能力”——因为它从来就没有拥有过这种能力。我们需要问的是另一个问题：那些至今仍能有效预警灭绝性风险的组织，其警觉的源头在何处？它不可能来自系统内部精密化本身——每一次精密化都是同一种翻译工序的延伸，只会让信号处理得更光滑、分类得更干净、论证得更自洽。警觉必须来自一个不属于系统内部逻辑、却能在关键时刻介入并叫停翻译动作的行动。它不是“更好的系统思维”，而是一个判断者在感觉到既有认知格式在此处可能已经失效时，越过了格式继续改进的请求，直接说出的那一声“停”。

下文将完成以下论证工作。第二节以挑战者号前夜为文本，追踪一种此前未被精确指认的三步渐变过程——操作漂移——并在此追踪中，让它的机制从具体的决策节点中被逼出，而非被逻辑推演先行给出。第三节将驱动这一漂移的深层动力正式命名为“可译性幻觉”，并对之做出严格的概念界定。第四节将可译性幻觉与卢曼的操作封闭理论以及博尔坦斯基与泰弗诺的“测试格式”理论进行严肃对质，证明本文提供的是一种不可被既有理论还原的新辨别维度。第五节处理抵抗操作漂移的唯一行动类型——置身判断——并追问它自身的发生学条件。第六节回应两个必须正视的理论反驳。第七节简短收束，给出全文核心命题的可证伪条件。

二、一场漂移的解剖

挑战者号发射前夜电话会议的过程，在沃恩的《挑战者号发射决策》和总统委员会报告中已有详细记录（Vaughan, 1996; Rogers et al., 1986）。本节的独特工作不在于增加新史料，而在于追踪这一夜之间，一个指向既有认知格式适用性的异质判断，如何在系统正常的决策工序中被一步步重新分类、消解、乃至湮灭。

发射前夜约八点半，NASA马歇尔太空飞行中心与莫顿·瑟奥科尔公司的管理层和工程师之间召开电话会议，讨论次日低温发射的安全性。博伊斯乔利及其同事依据前几次发射后密封圈烧蚀的原始数据和低温下材料弹性丧失的物理特性，坚决反对发射。此刻必须仔细辨认他在说什么，而不是简单地将他的立场标注为“清醒的工程师提出了警告”。博伊斯乔利当时赖以做出判断的依据，是一组非标准化的材料性能直觉、对烧蚀加剧趋势的不安，以及一个跨越多年的工程经验的凝结：密封圈的弹性在低温下不是线性的。他在说的不是“你的概率风险评估中参数估计有误”；他越过了这一层。他在

说的是：用来论证密封圈可以承受低温的那一整套基于已有发射数据的概率安全论证逻辑，在零下的气温条件下可能根本不适用。

这是一个在性质上不同于常规工程技术讨论的判断。常规的工程争议——参数设定、测试充分性、数据解释——共享同一个前提：概率风险评估框架本身是适用的，只是需要被更好地执行。博伊斯乔利跨过了这个前提。他的判断不是框架内的反驳，而是对框架本身在此特定条件下的适用边界的提醒。

接下来的漂移发生在一小时内。NASA管理层在听取工程师的反对意见后，发出了那句被精心包装成技术澄清的质询：“你的意思是，你不能用定量数据证明发射是不安全的？”在一个被二十四年安全论证格式强化的组织认知装置中，这句话的实质功能是一道格式审查闸门：请使用本系统承认的唯一认知货币来支付你的质疑。如果你无法支付，那么你所说的——无论其内容是什么——在系统的信息处理流程中，不构成“有效输入”。

博伊斯乔利无法支付。他握有的判断依据——非标准化的物理直觉、对一组烧蚀数据之“不对劲”的感觉——在格式审查的第一秒就被判定为“证据格式不达标”。这里发生的不是压制。压制仍然预设反对意见的内容本身被看见了，然后被有意排除。这里发生的事情更隐蔽：反对意见在尚未被看见其实质内容之前，其分类身份就被系统重新定义了。它从“对框架适用性的提醒”，被无声地重新归类为“框架内一项在证据格式上暂时不达标的输入”。质疑者本人也从“对框架的条件提出警示的一方”，被重新指派为“框架内暂时无法完成举证的一方”。举证责任在此刻被颠倒：不再是发射方需要证明框架在低温下依然适用，而是质疑方需要在框架内证明发射不安全。

在会议结论中，博伊斯乔利的嘶吼被静静地归档为一项在技术讨论中被提出、但未能达到可操作证据门槛的担忧，列入“后续监测”项。此前整整二十四次成功发射所加固的安全论证格式，没有被挑战——它只是接受了一项“仍需改进”的内部建议。

这里必须暂停，仔细辨认正在发生的渐变。这并不是一个组织以单一“决策”来压制异质判断的事件。它是一连串分散在多个日常工序中的微观位移被累积起来的产物。格式审查的闸门完成了第一笔位移：将性质与边界的问题，重新分类为证据与合规的问题。文件的归档与纪要完成了第二笔位移：将证据与合规的问题，列入常规的监测与改进流程。这两笔位移一旦完成，整个事件的“问题性质”便被不可逆地锁定了：此后组织内部的所有行动，都将在“这是一个已被识别的、正在被处理的技术改进项”的总括名目下展开。质疑者最初那句指向根本适用性的警告，已经被格式审查与待办归档两道工序，成功地翻译为一句关于如何改进参数校准或加强后续监测的常规工作语言。而组织上下，包括NASA管理层和瑟奥科尔管理层在内，无需任何一个“恶意的压制者”。他们只是各尽其职地完成了自己被要求完成的任务：格式审查者要求合规证据，会议记录者准确归纳待办项，项目管理者推进后续改进。每一个人都做了对的事——而正是这些“对的事”合在一起，完成了对异质判断的消化。

在这里还可以辨识出第三步位移的预览，尽管它在挑战者号案例中因为爆炸而被中断。如果博伊斯乔利的反对被顺利纳入下季度的参数校准任务，并被结案归档——不管参数后来是否真的被调整了——组织的自我叙事将获得一次安静的增强。如果参数被调整了，叙事是：“我们识别了改进点并完善了安全论证。”如果参数未被调整，叙事是：“经审慎评估，现有框架足以覆盖该风险。”在这两种情况下，叙事都绕开了同一个要害：用来评估风险的那套安全论证格式本身，是否还成立。每一次

参数调整所加固的，恰恰被认为可能已经不适用于新条件的那个框架本身。它向整个组织证明：框架是这样好用，以至于它的边界问题从不值得被严肃质疑，只需要被更好地校准。

这三步连在一起，构成一个完整的、闭环的、无需任何环节“出错”的自我维系装置。每一步单独取出都是组织理性的典范——处理反对意见、将争议转化为待办项、把待办项纳入改进流程——合在一起，它们联合完成了一件任何单一环节都不觉得自己在做的系统性工作：阻止组织听见那些无法被既有认知格式所接收的信号。

本文以“操作漂移”命名这一三步位移：异议被接纳（而非压制），但在接纳的同时，其性质被重新分类（从前提性质询转化为格式内争议）；重新分类后的异议被纳入常规优化流程；此流程的完成反过来加固了被质疑的框架本身，形成叙事上的自我增强。这不是一个“坏苹果”或“病态部门”造成的局部故障。这是整个系统在正常运转中的联合产出。它不需要恶意，不需要越轨，不需要任何人有意识地欺骗。它只需要每个节点都忠实地做着设计它时被要求做的事：审查格式、记录待办、推进改进。

这一机制在高可靠性组织（High Reliability Organization）研究中长期未被识别。研究高可靠性组织的学者已经注意到了组织如何在复杂环境中避免灾难，并将这一能力归因于一种“正念”（mindfulness）——对微弱的异常信号的警觉捕捉能力，以及对专业直觉的制度性尊重（Weick & Sutcliffe, 2001）。但“正念”这个概念本身，可以被操作漂移的透镜重新质询：当一套“正念”实践被组织连续巩固、写入流程手册、纳入培训大纲之后，它自身就会成为组织内部新的认知格式。下一次指向这个新格式的异质信号抵达时，它会经历同样的命运吗？本文的判断是：会。因为“正念”一旦被制度化为一套可操作的程序，它就不再是它初始宣称要做的事——对微弱异常保持开放的警觉——而变成了“按要求执行正念检查清单”。这种转变不需要组织犯任何错误：它是制度化本身的必然效果。可译性幻觉并不区分“坏的格式”与“好的格式”。任何格式，包括那些为抵抗闭合而设计的格式，只要被成功经验硬化到足够程度，都能成为下一轮操作漂移的跑道。

2008年金融危机为操作漂移的跨域重现提供了一个发生在完全不同制度场域中的切片。当时全球最精密的银行内部风险模型并非没有“看见”次贷风险。它们用VaR（在险价值）和压力测试反复评估了这一敞口，结论是：尾部损失在置信区间内。少数交易员和风险管理者口头警告——“这些东西的关联违约一旦发生，组合的损失会远超过模型预测”——却无法在正式风险报告的格式中找到立足之地。这些警告不是框架内的参数争议；它们在说“这套模型框架在结构性产品关联违约的情境下不成立”。但风险管理部部门的日常工序只被设计去接收和处理一种信号：可以被表述为模型输入参数的定量数据。口头警告在这个工序的入口处，便被重新分类为“尚未被量化的担忧”和“需要更多实证支持的假设”——然后被列入“待进一步研究”的议程。那些银行并没有在哪一次独立的决策会议上“决定”无视风险。风险是在一系列正常的、分散的、各尽其职的微观行动中被漂移掉的：模型验证团队审阅了参数，合规部门确认了流程，风险评估委员会接收了报告中关于“模型局限性”的标准条款（那正是被系统安全容纳的异见）。各个节点都完成了自己的工作，而所有工作的联合产出是：系统继续运转，直到崩盘。

这个案例还凸显了操作漂移区别于经典科层制理论的独特之处。罗伯特·默顿在一九四〇年代对“科层制人格”的经典分析，已经揭示了规则如何从手段异化为目的本身（Merton, 1940）：科层制中的行动者将遵守规则本身视为最高价值，从而失去对规则所服务之外部目的的判断力。米歇尔·克

罗齐埃在六十年代对“科层制现象”的研究则进一步揭示，规则的僵化并非系统的故障，而是科层组织内部各群体利用规则的间隙进行权力博弈的产物（Crozier, 1964）。这两个分析传统至今仍是理解组织僵化的基准。但操作漂移与二者存在两笔关键差异，这些差异使得它无法被还原为规则僵化的一个子类。

第一笔：操作漂移不需要行动者“病态地遵守规则”。在挑战者号案例中，马洛伊要求博伊斯乔利提供定量证据，这并非一个被规则异化了的头脑在机械地重复流程要求——这是一个在工程理性上完全站得住脚的技术提问。在金融危机案例中，要求风险管理者提供量化的模型输入，恰恰是风险管理被设计来要做的事。操作漂移的可怕之处不在于规则被坏掉了，而在于规则在正常运转中就可以产生闭合。第二笔：操作漂移的终点不是系统僵化，而是系统自我增强的叙事闭环。科层制理论的经典图景是：系统在自身的规则中变得越来越僵化、越来越不能应对环境变化，最终遭遇危机。操作漂移的图景则更为隐蔽：系统发生了一场操作漂移，信号被消解，决策被正常执行，此后如果不爆发灾难，系统会从“没有出问题”的经验中收获一笔额外的自我确认——这笔确认会进一步加固那个从未被质疑过的认知格式。系统不是僵化地走向危机，而是自信地走向危机。它越成功，越自信；越自信，越不需要听见下一次的嘶吼。

操作漂移由此从一个描述性概念指向一个在既有组织理论中尚不存在的辨别维度。它不是系统的故障，不是规则的异化，不是权力的压制。它是在系统正常运转中，经由一连串分离的、各自理性的微观操作，完成的一项将“对格式本身的挑战”转化成“格式内的改进项”的系统性工作。而这些微观操作的每一个执行者，都在忠实地、专业地、甚至可嘉地，做着自己的本职工作。

三、可译性幻觉：一个贯穿概念的正式确立

上一节追踪了操作漂移如何在一个特定案例中发生。本节转身处理它的生成条件：是什么样的深层动力，使得这些各尽其职的微观操作被反复组织起来，联合产出一种无人蓄意的闭合？

我们将这个深层动力命名为可译性幻觉。这个称谓选择本身，需要立即被加以精确限定。“幻觉”二字在这里不是对组织成员的道德指控，也不是一个认识论上的贬义词。它不假定存在一个“正确的”认知位置，而系统邪恶地或愚蠢地偏离了它。本文使用的“幻觉”，是一个结构性的术语：它指称那种被系统自身的连续成功所硬化、从而使系统在特定边界处系统性地看不见自身运作方式的认知习惯。它不是可以被一纸备忘录纠正的“错误信念”，而是一种在组织的日常运转中被再生产的、被成功经验不断滋养的、去个性化的认知状态。它的反面不是“真实”，而是“对自身前提的警觉”。

可译性幻觉的核心命题可以表述为：系统在连续成功中将自身认知格式的有限适用性，无意识地等同于格式的普适性。这一幻觉由三个相互支撑的信条构成，这三个信条不是组织在书面文件中明确宣布的原则，而是沉淀在组织的日常工序、人员培训、绩效评估与话语习惯中的实践信条——它们在每一个具体节点的微观行动中被默会地执行，而不是被言明地信仰。

第一信条：所有重要信号都应该可以被翻译为系统内部的认知格式。注意“应该”一词的规范性分量。它不只是一个中立的观察——“能被翻译的信号更易被处理”——而是一个规范性的预期：如果一个信号无法以系统认可的格式呈现，那不是信号本身有其不可还原的异质性，而是信号的持有者

尚未完成他的认知工作。博伊斯乔利面对的格式审查闸门——“你能用定量数据证明吗？”——其预设正是：如果你不能，那就是你还没有准备好把你的担忧变成真正的信息。

第二信条：如果某个信号暂时无法被翻译，那是因为翻译工具还不够好。这一信条将系统在面对异质信号时的应对方向，稳定地引向同一个出口：精化翻译工具。投资更好的模型，设立更敏锐的预警指标体系，培训更聪明的人来操作格式。这个出口每一次被使用，都是一次对格式本身的加深加固。它向全组织传递了一个无声却有力的信息：框架没有问题，它只需要更好的执行。

第三信条：因此，应对新威胁的最佳策略，就是不断精化翻译工具。这一信条构成了一个自稳的闭环。第一信条说“应该可译”，第二信条说“暂时不可译是因为工具不精”，由此推得第三信条——“所以继续精化工具就好”。在这个闭环中，没有任何一个环节指向了那个被绕开的问题：翻译这个动作本身，在此处可能是无效的。可能存在一些信号，其“信号性”恰在于它指向了既有翻译格式的不适用性，而一旦它被翻译、被重新表述为格式可以处理的形式，它的“异质性”——也即它对格式边界的提醒——就会在翻译过程中被损耗殆尽。

可译性幻觉之所以冠以“幻觉”之名，要害便在于此。它不是在翻译工具的问题上犯了错，而是在一个更根本的问题上建立了一种系统性的视盲：它运行在格式内部，却以为自己在俯瞰格式。每一次格式内部的成功，都强化了这个幻觉——因为从内部看，格式确实越来越精密，处理能力越来越强。但从外部看（如果允许这个透视），系统的每一次精化都在把接收异质信号的音量旋钮往更小处拧——因为它越来越擅长把听起来不对劲的信号处理成自己听得懂的东西。它把听不见做得越来越像听见。

这个幻觉不是一代管理人员或某个组织的专利。它是系统性认知本身的伴生物。只要一个系统在长期运转中形成了一套稳定的认知格式，并被这套格式的连续成功所验证和加固，可译性幻觉就倾向于在那套格式的边界处生长出来。它的生长不需要任何人的蓄意；它只需要人们忠实地、认真地、日复一日地做好自己的本职工作。

这一判断区分于几个表面相似的既有概念。第一个需要区分的，是查尔斯·佩罗在一九八四年提出的“正常事故”理论（Perrow, 1984）。佩罗已经论证了某些技术系统中的灾难不是功能失灵的结果，而是系统结构性特征的正常产物——具体地，交互复杂性与紧密耦合这两个结构特征使得某些事故在概率上不可避免。这一定位与本文有家族相似：二者都将事故归因于系统正常运转而非局部故障。但二者的区分是实质性的。佩罗的论证锚定在技术系统的物理结构上——谁连着谁、哪个阀门会影响哪个反应器——而本文的论证锚定在组织认知的翻译工序上。佩罗要求我们审视管道的连接方式；本文要求我们审视“连接”这个隐喻本身——管道只能连接管道，系统只能连接那些已经进入了它认知格式的东西。可译性幻觉解释了为什么在某些系统被充分警告的情况下，警告仍然无从进入。这一笔是佩罗的结构分析所不能直接提供的。

第二个需要区分的，是卡尔·韦克关于组织意义建构的经典论述。韦克以其曼恩峡谷山火事故分析奠定了这一传统：组织不是在既有环境中做出决策，而是通过回溯性地将模糊的环境线索建构为可理解的叙事，来“演出”其面临的环境本身（Weick, 1993, 1995）。线索被转化为意义，从来离不开组织的解释格式。在这个意义上，韦克的“意义建构”已经包含了“翻译”——环境信号必须经过组织的意义建构过程，才能成为对组织“存在”的信息。但本文要推进的关键一笔在于：韦克的意义建构模型很好地解释了组织如何从模糊中建构清晰，但没有解释组织如何反复错过已经足够清晰的警告。

在挑战者号案例中，博伊斯乔利的嘶吼不是模糊的线索，等待意义建构来赋予它形状；它已经是一个高度清晰的判断，只是这个判断的质地组织的认知格式中没有任何可以被接收的形态。可译性幻觉处理的不是“建构意义”的过程，而是当一个判断已经带有自身完整的意义时，系统如何通过格式审查，把它重新变成“尚待建构”的一团模糊——然后把它纳入意义建构的常规流程，慢慢处理掉。这便是操作漂移与意义建构之间的分野：前者不是在缺乏意义处建构意义，而是在意义已经站立的地方，用一个标准化的翻译工序，把它拆去骨、刮去肉，再重新捏成它承认的形状。

由此可以给出更严格的定义。认知格式，在本文中特指一套被组织正式或非正式地制度化了的规则、证据标准与分类程序，它规定了什么样的陈述在此组织内被承认为“有效输入”，以及该陈述应如何被进一步转化为执行操作。这一定义将认知格式与我们日常所说的“范式”“框架”“话语”区分开来：认知格式不是一般性的思想倾向或文化氛围，而是一套操作性的、在组织的日常工序中可被逐项执行的选择与分类装置。它同时具有认识论的性质（它裁决什么算作“信息”而非噪音）和行政的性质（它将有效输入分配至特定的处理流程）。

可译性幻觉，是指这样一整套被硬化的认知格式在连续成功的滋养下所产生的一种系统性认知状态：在此状态下，系统将自身格式的有限适用性，无意识地等同于格式的普适性；将指向格式适用边界的异质信号，系统性地在翻译工序中重新分类为格式内的待优化项；并将每一次由此产生的“成功处理经历”转化为对格式自身的叙事加固。操作漂移则是这一幻觉在日常运转中的具体执行过程：异议被接纳、重新分类、纳入改进、叙事增强的三步位移。

这一组概念的提出，使得一个被既有组织理论反复触及但从未被单独指认的现象，第一次有了自己的名字和边界。它回答了一个问题——系统如何在未压制异议、未关闭渠道、未发生局部故障的情况下，将异议从“对系统前提的威胁”处理成“对系统自我完善的贡献”——而这个问题，是目标置换、防御性常规、越轨正常化各自单独无法完整回答的。下一节将进行最关键的理论对质：证明可译性幻觉不是一个可以被卢曼或博尔坦斯基与泰弗诺的既有概念轻易替换的重述。

四、与最近邻对质：为什么可译性幻觉不可还原

本文的工作如果要成立，必须通过一场严苛的二阶碰撞。在此节中，本文将“可译性幻觉”与两位最强最近邻——尼克拉斯·卢曼的操作封闭理论，以及博尔坦斯基与泰弗诺的“测试格式”理论——进行逐层对质。目标是证明：既有理论各自握住了“翻译/测试”这一代理变量的一侧，但它们都没有完整地解释系统如何将“对翻译格式本身的质疑”消化为“格式内的优化项”，并在此消化中生产出自我增强的叙事闭合。可译性幻觉所提供的辨别维度，在这些理论的夹缝中生长出来，不能被还原为其中任何一个。

卢曼是本文最危险的理论占位者——危险到如果本文不与卢曼完成严肃对质，整篇论证就可能被一句话打发掉：“这不就是卢曼说的操作封闭吗？”因此必须在一开始就承认卢曼已经精确论证了什么，再精确指出他未曾论证的那一笔。

卢曼的社会系统理论的核心判断是：社会系统——法律、经济、科学、政治——是通过特定二值代码进行沟通的操作性封闭系统。经济系统以“支付/不支付”为代码，法律系统以“合法/非法”为代码。系统不能直接触及环境；每一次环境对系统的影响，都必须先被系统自身的代码重新编码，才能

成为系统内的沟通事件。这个操作，卢曼称之为“结构耦合”——系统与环境之间不是没有关系，而是这种关系只能以系统内部接收和重新建构的方式发生（Luhmann, 1984/1995）。环境可以“刺激”系统，但它永远无法直接“告诉”系统任何东西。系统只能听见自己代码允许它听见的声音。在这个意义上，卢曼已经准确揭示了“可译性幻觉”的第一层面：系统确实只能翻译，它的操作封闭确实是其存在的基本方式。本文在这一点上与卢曼毫无分歧。

但卢曼在此精确地停住了——而本文的论证恰是从他停住的地方开始的。卢曼系统的代码是预定的：法律系统识别合法与非法，经济系统识别支付与不支付。代码执行的是“是/否”的判断，这道闸门一旦通过，系统就知道该如何处理——判例、合同、支付指令。但本文追踪的操作漂移，发生在代码本身不明确的情境中。安全评估——在挑战者号夜间的那个具体关头——不是一套预定的法律代码或经济代码。它是组织在历史经验中逐渐形成的、由特定的证据标准和分类程序构成的一套认知格式。它的“代码”不是“安全/不安全”，而是“是否可以用概率风险量化论证来陈述”。这个格式在它的日常运作中，容纳了技术分歧、参数争议、方法争辩。但它从未被设计去处理一种特定类型的输入：一种指向它自身适用边界的异质判断。当博伊斯乔利说“这套概率论证逻辑在此可能不成立”时，他不是框架内提出一个可以被框架处理的技术问题——他是在框架的代码本身的适用性面前，说了它在自己的词典里找不到的话。

而卢曼的理论从不处理这类信号。卢曼讨论的是环境对系统的“刺激”——一种始终带着混沌和未编码形态的外部信息。他的经典图景是：环境中有大量潜在的刺激，系统只能选择性地将其一部分编码为内部沟通。但博伊斯乔利的判断已经不属于这个图景。他的判断不是混沌的“刺激”——它已经是一个高度清晰的、结构化的判断，只是这个判断的身份是“对编码本身的质询”。在卢曼的模型中，这样一个判断在抵达系统的编码闸门时，模子不对——它无法被“合法/非法”或“支付/不支付”处理，于是被系统作为无关刺激筛出。卢曼可以很冷峻地说：是的，系统就是这样运作的，它无法处理那些不能落入其编码的事件。

但本文要指出的是：卢曼的这一描述恰恰漏掉了最危险的那个机制。博伊斯乔利的判断并没有被筛出。它被非常认真地接收了。它被召开了专门的电话会议，被置于NASA和承包商的联席讨论中，被纳入正式的风险评估流程。它被告知：你的担忧很重要，但请你用我们可以接收的方式来表述它。系统不是在“代码”处把它筛掉了——系统是把它迎入了门内，但在门内，系统用格式审查这第一道工序，将它的身份从“异质判断”转化为“格式内不成熟输入”，然后再把被转化后的它，一步步推向待办归档与叙事增强。操作漂移完成的不是卢曼式的“筛出”，而是更隐蔽的“迎入后重新命名”。系统不是听不见；系统是听见了，但把它听成了别的东西。

由此可以定位卢曼与本文的精确分野。卢曼的操作封闭理论解释了系统为什么只能翻译——因为它的代码就是它的存在方式。但他没有解释，当系统翻译不了某个信号、却又不直接筛出它时，系统会对这个信号做什么。他缺乏一个关于消化的理论。消化不是筛出。消化是：将异质信号纳入系统中，通过一连串分散在日常工序中的微观操作，改变它的身份，使之从“对代码的威胁”变为“代码内的待办”，然后处理掉。可译性幻觉，恰恰是使得消化可能发生的那种系统性认知状态。它解释了系统如何能够在不违反自身代码的前提下，在表面上“认真对待异议”的流程中，将异议的异质性消解殆尽。

第二组对质指向博尔坦斯基与泰弗诺的“辩护”（justification）与“测试格式”（test format）理论。在《论辩护》一书中，博尔坦斯基与泰弗诺系统论证了当争议出现时，行动者并非简单地以权力压制对方，而是将对方纳入一个共同的“证成格式”之中，要求争议双方在一个公认的“等价原则”下展示证据、比较论证力度；由此，异议不是被压制，而是被一个公认的测试程序所消化和裁决（Boltanski & Thévenot, 1991/2006）。这与本文描述的操作漂移确实存在结构上的同构：异议被迎入一个格式，并在格式内被转化。如果本文不能证明操作漂移与“测试”在机制上存在不可消解的差异，那么“可译性幻觉”就只是“证成格式”的一种特殊情况——又一个漂亮的重述。

但二者的差异是实质性的。在博尔坦斯基与泰弗诺的模型中，测试格式预设了争议双方共享一个更高层级的“城邦”（cité）——也即一组关于什么是“共同福祉”、什么构成有效论证的规范性共识。一个市场城邦中的测试是价格竞争，一个公民城邦中的测试是权利平等的论证。争议之所以能被测试所裁决，是因为双方都承认这个城邦的测试是合法的、适用的。换言之，测试所能处理的争议，只能是同一个城邦内部的争议。

操作漂移所要处理的情形，却已经越出了这个前提。博伊斯乔利与NASA管理层在那一刻的冲突，不是一个市场城邦中的竞争性论证，也不是一个工业城邦中的效率论证。它是一个触及了城邦本身适用边界的情形：共同福祉是什么——是“程序合规”还是“生命安全”——在二者之间已经产生了不可靠城邦本身裁决的分歧。格式审查闸门——“请用定量数据证明”——在此处的实质功能，是单方面地将争议界定为必须在效率城邦的技术格式内被解决。这不是城邦内对证据的测试，而是用一个城邦的测试格式，去覆盖一个它已经不足以处理的城邦间冲突。它向异质判断说：“进入我这个测试，否则你的担忧不算信息。”而这个要求本身——要求对生命安全的前提性质询必须以概率风险量化的形态呈现——恰恰是冲突的要害所在。

由此，操作漂移不是城邦内的测试。它是将城邦间冲突，强行降格为城邦内测试，并在此降格中完成对冲突性质的暗中置换。博尔坦斯基与泰弗诺的理论可以解释一种异议如何被城邦内部的测试所裁决，但它没有解释一种异议如何被测试格式的单方面强加所漂移——因为这后一种情形，恰恰越出了他们理论所预设的“共同城邦”的条件。可译性幻觉提供的辨别维度就在这一裂缝中生长出来：它捕捉了系统如何在不被争议双方都承认的共同城邦存在的情况下，仅凭自身的格式惯性，单方面地将城邦间冲突重新分类为城邦内待测项。这个辨别维度，无法被“测试格式”理论所覆盖——因为后者的有效性恰以共同城邦为前提。

对卢曼和博尔坦斯基与泰弗诺的这两组对质，共同指向了同一个理论收益：可译性幻觉不是对“系统只能翻译”或“争议需要测试”的重述，而是对系统如何消化那些指向翻译格式本身的刺穿信号的独立命名。这个消化过程不是系统失灵的证据；它恰恰是系统在自身的成功中所训练出的最精密的功夫。而它所生产出的那个闭环——异议被迎入、重新分类、纳入改进、叙事增强——在卢曼的理论中不可见（因为卢曼只看到了代码筛出），在博尔坦斯基与泰弗诺的理论中不可见（因为他们预设了争议双方共享城邦），在佩罗的正常事故理论中也不可见（因为佩罗关注的是技术耦合而非认知翻译）。可译性幻觉由此占据了一个这些理论各自边界的交汇点，却无法被其中任何一个所替换的位置。

五、置身判断：打断的性质

如果前文的论证成立，抵抗操作漂移所需要的警觉，就不可能从系统内部精密化中生长出来。那么它来自何处？

让我们先处理一个几乎必然会被提出的反驳。博伊斯乔利难道不是系统内部的人吗？他不是一个被雇用的工程师、一个组织成员、一个拿着工资为NASA的承包商工作的人吗？如果系统逻辑只生产操作漂移，他怎么成了例外？这里没有悖论。博伊斯乔利确实是系统内部的人。但他那一夜的判断——这是整个案例中最微妙的一笔——不是来自“更好地执行概率风险评估格式”这一系统内部行为。他那时在做的事，他的二十多年材料工程经验为他积攒的那种判断质地，以及他在电话会议上被逼到墙角时仍然坚持说出的那句话，不是一个被格式训练出来的精化产物。它是另一种认知动作——本文称之为置身判断。

置身判断在性质上不同于格式内的分析。格式内的分析——无论被做得多么精密、反思性多么强——都共享同一个操作前提：分析所依赖的认知格式是合法的、适用的。它的所有改进都在格式内部进行：优化参数、纳入更多变量、校正指标、更新数据。但置身判断恰恰越过了这一前提。它不是一个“更好的格式内分析”。它是一个判断——它判断的是：格式本身在这个特定境况下，可能已经不再适用。

这不是缄默知识（tacit knowledge）。迈克尔·波兰尼的缄默知识概念——我们知道的多于我们能说出的（Polanyi, 1966）——已经在组织理论中占据了一个稳固的位置，用来解释那些无法被标准化流程所涵盖的专家直觉与技能（Nonaka & Takeuchi, 1995）。但缄默知识仍然是框架内的知识：它服务于格式，作为格式无法明言化却不可或缺的补充部分被系统倚赖。一个手工艺人的手感使他能做出机器做不到的活，这不构成对机器逻辑的挑战——它只是补充机器逻辑的局限。系统可以安全地容纳缄默知识，把它圈在一个“专业直觉保护区”里，既否定它，也不让它威胁格式的前提。事实上，许多高可靠性组织已经这样做了——它们设立了“专家直觉热线”，在标准流程之外为模糊担忧提供了一条绿色通道。这个安排看起来很明智。但本文要指出的是：这条绿色通道，本身也是一道格式审查。它的输入格式——“专业直觉”——已经被预先定义了。一个担忧必须能够被表述为“专业直觉”，才能进入这条通道。而博伊斯乔利那一夜的判断，其要害不在于它是直觉而非数据；要害在于它指向了将“数据”与“直觉”对立起来的那套分类格式本身在低温条件下可能已经站不住脚。他的判断不是在为缄默知识争取一席之地。他的判断是把整张桌子——数据这边、直觉那边——都质疑了。

这不是双环学习。克里斯·阿吉里斯与唐纳德·舍恩的双环学习概念（Argyris & Schön, 1978）是组织学习理论中用来描述“对主导变量的质疑”的术语。单环学习是在给定目标与规范下纠正行动；双环学习是对目标与规范本身进行重新审视。从表面看，双环学习与本文的置身判断似乎指向同一件事——都是在问那个更基础的问题。但“双环学习”这个概念，正如它通常被实践的那样，已经被安全地纳入了系统可以操作的格式。它是组织的一次内部倡议：一个学习型组织设置了专门的双环学习工作坊，在季度战略复盘会上留出“质疑根本假设”的议程，委任了一个反思性领导力小组。这些安排使双环学习成为一项系统内部的制度化活动——它被执行、被记录、被评估，然后被归档为一个“已完成”的组织学习环节。而操作漂移从一开始就不排斥质疑；它只是把质疑转化为改进项。当一个系统把双环学习设为一项定期议程时，它可能正在为下一轮操作漂移铺设跑道——下一次指

向“双环学习议程本身是否还适用”的异质判断，会以同样的格式审查被迎入、被重新分类、被纳入“优化学习议程”的待办项。可译性幻觉不区分“单环”与“双环”。它消化一切可以被格式化的东西。双环学习一旦被格式化，就变成了单环的延长。

置身判断的独特质地在于：它不属于任何一种学习循环。它不是“更好的学习”。它是一个学习循环无法容纳的中断。它说“停”——不是在循环内部说“我们该转到双环了”，而是越过格式说“此刻不要在这个格式里讨论，不管单环双环”。它的发生不是来自系统内部精密化；它是被逼出来的。它发生在系统在其认知格式的缝隙处遭遇到一个无法被消解的张力时，被那个张力逼出的一个临界产物。

在挑战者号之夜，这个张力是博伊斯乔利二十多年的材料工程经验——那些长期观察烧蚀数据所累积的物理直觉——与NASA安全论证格式在低温条件下的边界之间的无法消解的对峙。他不是在选择缄默知识而非正式证据；他是在两者都无法被安全地折叠进对方体内时，被逼着做出了一个不依赖任何既有格式授权的判断。这个判断在说出“不能飞”的那一刻，不是在执行格式；它是在执行判断者自身的责任。

这一定位将置身判断从根本上区别于一个常见的误解：它不是将“专家直觉”或“实践智慧”浪漫化为某种神秘的认识论优越者。博伊斯乔利可能错。置身判断不是靠其内容“正确”而成立；它是靠其行动的不可还原性与责任承担而成立。它的认识论地位并不高于格式内的分析——它只是在一个格式不再可能继续适用、而时间又不等人时，越过格式的直接行动。它与格式内分析之间不存在认识论等级制：它们是在不同条件下、承担不同风险、回应不同责任的不同类型的行动。格式内分析在正常条件下是理性的巅峰；置身判断在格式边界处是唯一还能开口说话的方式。

但置身判断自身也有发生学的问题。它不是从天空中降下的永恒救星。它是被发生的——在特定的条件下，经特定的路径，成为特定时刻的一个显露的判断。博伊斯乔利二十多年的材料工程经验（一种特定的专业知识土壤）、他在瑟奥科尔公司中的特定位置（不是NASA雇员，而是承包商工程师，这给了他一个稍远的距离）、以及他参与此前多次发射后烧蚀分析的特定经验积累，在那一夜电话会议的压力和格式审查的逼迫中，联合将他推到了那个说出“不能飞”的临界点上。他的判断不是一个“纯然的真理之音”；它是一个在多方矛盾中被逼出的脆弱产物。如果他的专业经验有所不同，如果公司的进退权衡有所不同，如果会议的气氛稍有改变，这个判断可能就不会以那种形态显露出来。

更重要的是，这个判断在说出之后，已经经历了我们追踪过的消化过程——它已经开始被漂移，只是爆炸打断了漂移。这提醒我们：置身判断不是站在系统外面俯视系统的超越行动。它是一个在系统内部发生、但试图以不可被系统格式消化的方式作用于系统的打断。它的生成需要特定的制度土壤、经验积累和伦理承重，而它的维持则随时面临被系统重新格式化（“罗杰后来加入了安全改进委员会，从此他的担忧就在正常的流程中被处理了”）或退化（“经历了那次爆炸，罗杰再也不在公开会议上反对管理层了”）的命运。一个发生学的结论，不是呼唤“让更多的置身判断发生吧”——那是把判断当成了可以被调用的管理工具，恰恰落入了可译性幻觉。一个发生学的结论是：置身判断只能在特定的条件下被发生出来，并且这些条件本身，也在操作漂移的笼罩之下。

六、两个必须正视的反驳

本文至此已经绕开了两个必须正面回答的反驳。本节不再回避，将它们从阴影中拖出。

第一个反驳来自卢曼传统内部的一种自我防御。它可以说：你把操作漂移描述为系统“消化”了异质判断，但在卢曼的理论中，系统本来就没有“异质”这回事——任何进入系统的沟通都已经是系统自身的产物，谈不上“身份被改变”。你所谓的“漂移”，在系统理论看来，只是系统自身持续运作的正常形态。你用一种规范性的语言（“异质性被剥夺了”）描述了一个系统理论用纯粹功能描述已经说清楚了的事实。

这个反驳在逻辑上是自治的。但它的自治性本身，就是问题的一部分。卢曼的系统理论为了维持其描述的纯粹性，付出了将“信号被系统改变身份”这个判断本身的规范性分量全部悬置的代价。它可以说：博伊斯乔利的嘶吼在被格式审查重新分类后，就不再是同一个事件了；在系统内部，它从一开始就只作为“一个未能提供合规证据的反对意见”而存在。这个描述在卢曼自己的理论内部是准确的——但它不提供任何资源，去回头质询这套系统运作本身的合理性。它不能问：系统把那个判断变成了另一个东西，这是一个需要被严肃对待的事实吗？它能说的是：系统就是这样运作的。这不是反驳了本文的论证；这是以退位的代价避开了本文的追问。本文的论证并不要求卢曼承认“异质性”——本文的论证要求的是：在卢曼的系统理论止步之处，存在着一个它无法命名的现象：系统在正常运转中，将一种特定类型的信号——对自己认知格式适用性的判断——通过日常工序重新构造成了对格式的加固物。卢曼的理论可以描述这一过程的操作，但它不能将这个过程指认为一种需要被抵抗的东西。而本文恰恰要完成这个指认。

第二个反驳与迈克尔·鲍尔关于风险管理和审计社会的分析直接相关。鲍尔已经深刻论证了现代组织中的风险管理往往不是对风险的实质性控制，而是一套管理“声誉风险”的仪式——组织通过执行标准化的风险管理流程，来生产关于自身“在管控风险”的合法性叙事，而非实际降低风险（Power, 2004, 2007）。在这个意义上，鲍尔的分析已经触及了本文第三节所描述的叙事增强。他可以说：你把“可译性幻觉”当作一个新发现，但鲍尔已经展示了风险模型如何沦为组织的自我叙事工具——这不是一模一样的吗？

两个论证确实存在重叠：二者都关注到了系统如何在自身流程中生产自我确认的叙事。但鲍尔的论证对象是被有意识设计的管理工具——风险管理框架、审计标准、合规体系——及其意外后果（被管理层用作声誉管理的手段）。本文的论证对象则更深一层：不是特定的管理工具，而是认知格式本身——决定什么算作“有效信息”的基本规则——在系统运转中所产生的消化效应。鲍尔的风险管理批判仍然预设：如果我们能遏制管理层的“仪式性”使用，让风险管理真正聚焦于实质风险，系统就能更好地预警。他的批判仍然在格式内：它要求的是对管理工具的更诚实的使用。但本文指出的是：即使最诚实的使用，也无法免除可译性幻觉——因为格式审查这个动作，不需要任何人的不诚实就能产生闭合。不是有人把风险模型用歪了；是风险模型作为一套认知格式，在它被设计出来的那一刻，便只能接收能落入它的变量口径的信号。诚实的使用只会让这个局限被更忠实地执行。可译性幻觉是一层比仪式主义更深的岩石。

七、结语：可证伪的条件

本文从头到尾做了一件事：将一个被反复目睹但从未被准确命名的事实，借由“可译性幻觉”这个新概念，从既有组织理论的夹缝中单独指认出来。这件事如果有价值，不是因为本文替既有理论换了一套说法，而是因为它打开了一个此前未被打开的辨别维度。在这个辨别维度建立之后，我们不再只能问：这个组织的自我免疫功能是否失效了？我们可以进一步问：它的认知格式是否已经被自身的成功硬化到了这样一个程度——以至于对格式本身的质疑，无法不被格式的日常工序所消化？

这个追问引向的结论不是管理的改良方案。它是一份清醒。它说：系统在其最优秀、最诚实、最理性的运转中，并不内建听见自身极限的能力。那种能力——如果有的话——必须来自另一种性质的行动。而那种行动的发生条件，本身就是脆弱的。这听起来悲观。但悲观只是它所拒绝的廉价乐观的另一个侧面。它在二者之间选择的是诚实——诚实地面对一个事实：警觉不能从系统中生长出来，但可以在系统的缝隙中被一些特定的人，以他们自己也无法长久保持的方式，艰难地护卫。博伊斯乔利那一夜的嘶吼，不是系统成功的证据；它是系统在一个罕见的瞬间没能成功消化掉的例外。

全文收束于此。一个负责任的论证必须给出其自身的可证伪条件。本文的核心命题——组织在其最健康运转中，通过操作漂移系统性地消化指向其认知格式适用性的异质判断——在三种情况下可被证伪。第一，如果有人展示一个组织的认知格式在长期连续成功中被加固后，仍然能够在不设立制度性打断保护的情况下，仅仅通过内部精密化，有效接收并认真对待指向该格式自身适用性的刺穿信号，本文的论证便需要被大幅修正。第二，如果操作漂移的三步被证明可以分解为其他更基础的已知机制的组合而无任何剩余——即“重新分类”“纳入改进”“叙事增强”被分别完整还原为卢曼的代码筛出、阿吉里斯的防御性常规、和鲍尔的仪式主义——那么可译性幻觉就只是一个组合概念，而非一个有独立辨别力的命名。第三，如果置身判断被证明可以通过系统内部的方法论革新而被常规化、制度化，并且在被制度化之后不发生新一轮的操作漂移，那么本文关于打断行动之不可消化性的判断便告失效。

这些条件，是本文留给其批评者的具体通道。一个真正的批评者不应只说“这不就是卢曼吗”——而应拿出真实的经验材料或理论论证，证明本文所指认的辨别维度确实不存在。我们都在等待那一场对质。而在对质到来之前，本文所辨析的这个新的辨别维度——它居于卢曼的操作封闭与博尔坦斯基与泰弗诺的测试格式之间、佩罗的结构耦合与鲍尔的仪式主义之上、韦克的意义建构与阿吉里斯的防御性常规之外——有理由被暂时地保留在组织理论的对话桌面上，作为一个应当继续被讨论的提案。

深化增补·全球近邻补切（编辑增补）

说明：以下各节为编辑在审读时增补，不改动作者正文与核心判断，只补上本文原稿未正面处理的最近邻理论，并把分离线写成可操作的判别。每一条都给出“若观察到什么，本文即错”的对照预测。

术语的正式处理：本文的「操作漂移」与斯努克的 practical drift 撞名而义相反

本文原稿以卢曼、韦克、阿吉里斯与沃恩为主要对话者，并使用「操作漂移」指称质疑被逐级改写的过程。此处必须做一次术语上的正式切割：斯努克（Snook，2000）在黑鹰直升机误击事件研究中提出的 practical drift，中文文献通行译作「实践漂移」或「操作漂移」，指松耦合系统中的行动者在局部合理性驱动下，使实践悄悄脱离书面规程，且位移在系统层面不被察觉。

二者的位移方向恰好相反。斯努克的漂移是**实践脱离纸面**；本文的位移**全程发生在纸面之内**——每一步改写都严格合规、都被记录为一次技术改进，没有任何行动者偏离规程。真正发生位移的不是行为，而是「**这属于哪一类问题**」这个分类动作。

为避免持续的混淆，本文此后以**分类位移**指称自身所追踪的机制，「操作漂移」一词保留给斯努克意义上的用法。这一改称不影响原稿的任何论证，只使它与既有文献可以正确地相互定位。

由此得到一条诊断上的对照：斯努克式漂移的系统里，调查者会发现大量未被记录的变通与心照不宣的偏离；分类位移的系统里，记录反而异常干净——所有质疑都在档案中以「已受理、已转为技术改进项、已闭环」的形态存在。

与特纳「灾难孕育期」及德克尔「漂入失败」的分离线

特纳（Turner，1978）的灾难孕育理论指出，重大事故之前存在一段长期的孕育期，其间警示信息分散在组织各处而未被汇聚成一个可被识别的整体图景。德克尔（Dekker，2011）的漂入失败则强调，在竞争与资源压力下，一连串局部理性的小决策使系统缓慢滑向边界，且无人拥有全局视角。

两者的共同前提是**信息问题**：要么信息没有汇聚，要么没有人看得见全局。本文的判断与此不同：在挑战者号发射前夜的关键会议上，信息**已经完整到场**，质疑被明确提出、被认真对待、被以标准流程处理——问题不在信息未到，而在它**被正确地处理进了错误的类别**。把「这个格式是否还适用」听成「这个参数是否达标」，随后一切后续动作都在参数这一层上无可指摘地展开。

这条分离线给出可判决的检验：若在这类案例中，能证明关键质疑在会议记录与往来文件中**始终保持为格式适用性质疑的类别**、只是被职权或进度压力压下，那么本文的分类改写机制不成立，权力解释与信息解释已经足够。本文的成立条件是：档案中可以看到质疑**逐步换类**的痕迹——从「我们不知道低温下密封是否还能成立」到「请补充低温密封的测试数据」，再到「已列入下一轮改进计划」。

证伪条件的收紧

综合上节，本文的核心命题可被以下任一观察否定：其一，在多起同类事故的原始档案中，异质质疑均未出现类别迁移，而是以原样被否决；其二，在实验条件下，若给决策会议引入一名不掌握权力、但被明确授权只做一件事——每当一项质疑被转写为技术项时出声指认——而这一介入并不改变最终决策的分布，则说明类别迁移并非致命环节；其三，若在质疑被转写之后，组织仍能在后续任一环节自行回到格式适用性这一层，则本文所称的不可逆性不成立。

本次增补所核验的文献

- Dekker, S. (2011). *Drift into Failure: From Hunting Broken Components to Understanding Complex Systems*. Ashgate.
- Snook, S. A. (2000). *Friendly Fire: The Accidental Shootdown of U.S. Black Hawks over Northern Iraq*. Princeton University Press.
- Turner, B. A. (1978). *Man-Made Disasters*. Wykeham Publications.
- Vaughan, D. (1996). *The Challenger Launch Decision: Risky Technology, Culture, and Deviance at NASA*. University of Chicago Press.