

借来的理论：前沿技术的概念如何解释一家投行的日常

Borrowed theories: how frontier-technology concepts explain an investment bank's daily work

庚辛研究院 / Glacier Institute^{*†}

出自 维特根斯坦研究所 · The Wittgenstein Institute

庚辛研究院 (Glacier Institute) | 通讯论文 Working Paper GI-WP-2026-P11 | 2026-09-22 | 中文正文，附英文标题与摘要

建议引用格式 Cite as:

Glacier Institute (2026). Borrowed theories: how frontier-technology concepts explain an investment bank's daily work. Glacier Institute Working Paper GI-WP-2026-P11. <https://glacier.mba/research/GI-WP-2026-P11.html>

庚辛研究院 (2026). 《借来的理论：前沿技术的概念如何解释一家投行的日常》. 庚辛研究院通讯论文 GI-WP-2026-P11. <https://glacier.mba/research/GI-WP-2026-P11.html>

要点 Key Takeaways

- 1 客户的理论解释得了我们，多半不是巧合，是我们只听得懂自己已经会的那部分。
- 2 说「懂一个行业」，可检验的含义是：借得到它的概念，并且借对了。
- 3 借来的概念，四成是六十年前的老东西，只是被重新用了一遍。
- 4 一个类比读起来越顺，越像是被证明了，其实只是越像。

摘要

一家长期服务前沿科技公司的中介机构，会反复发现一件事：客户用来解释自己技术的概念，恰好也解释了这家机构自己的组织与作业方式。本文把这件事当成一个可研究的现象，而不是一句修辞。

本文的底本是一份公开文本：这家机构在官网上列出了十二对「他们的理论」/「我们的现象」，每一对左边是一个来自前沿技术领域的概念，右边是自己的一项日常做法。十二对全部逐字抄录，标出在公开纯文本里的行号。

机制部分用五组文献：类比推理与概念迁移（结构映射理论、多重约束理论）、吸收能力（一个机构只学得会它已经有相关先验知识的东西）、组织惯例与能力的连续积累、专业服务机构里的知识编码与默会知识，以及一组边界文献——相似性被当作证据的代表性启发、隐喻在组织研究

里的作用与危险、探索与利用的误读。

本文的学术贡献在逐对分析的第三段：官网写了理论是什么、迁移到什么上，没有写这个迁移在什么条件下成立、什么条件下不成立。十二对逐一补上。其中三对在原理论的前提下并不成立，本文逐条说明为什么，并给出仍然可用的那一部分。

在此基础上给出一件工具：一张九行的类比自查表，把一个外来概念用到自己组织之前要过的判据写成「判据 / 通过的样子 / 不通过的样子」三格。再用这张表对十二对逐对打分，给出一张自己算出来的表。本文的主张写成一个可以被推翻的命题，并说明什么证据能推翻它。文末给出五类边界，其中两类是必须写的：类比讲得通不等于做得对，以及事后解释的诱惑。

本文不出现任何客户公司的名字，不出现可反推到具体公司的描述，不出现这家机构自身的任何经营数字。

关键词：类比推理；结构映射；概念迁移；吸收能力；组织惯例；专业服务机构；一级市场

JEL 分类 JEL Classification: D83 (搜寻、学习与信息), O33 (技术变迁：选择、后果与扩散), L84 (专业与商业服务), G24 (投资银行与创业投资)

* 利益披露 Disclosure: 庚辛研究院是庚辛资本的研究机构，本文由庚辛研究院署名。庚辛资本在其业务中担任部分科技企业的财务顾问，并以自有资金参与部分企业的股权投资；这类关系可能与本文讨论的行业范围重合。本文不涉及任何具体在投项目，不使用任何未公开信息；文中引用的企业、行业与数据信息一律取自公开来源并逐条注明出处。作者未因本文获得任何第三方报酬。

† 免责声明 Disclaimer: 本文为方法讨论性质的通讯论文，仅代表作者在写作时点的分析。本文不构成投资建议，不构成任何证券、基金份额或其他权益的要约或要约邀请，亦不构成对任何公司估值、融资结果或投资回报的承诺或预测。本文未经同行评议，内容可能在后续版本中修订。

Abstract

A financial intermediary that spends years working with frontier-technology companies keeps noticing the same thing: the concepts its clients use to explain their own technology also explain the intermediary's own organization and way of working. This paper treats that observation as a phenomenon worth studying rather than a figure of speech.

The source text is public. The firm lists twelve pairs on its website under the heading "their theory, our phenomena": on the left a concept from a frontier-technology field, on the right one of its own daily practices. All twelve are reproduced verbatim, with line numbers in the public plain text.

The mechanism section draws on five bodies of work: analogical reasoning and conceptual transfer (structure-mapping theory and multiconstraint theory), absorptive capacity (an organization can only take in what it already has related prior knowledge for), organizational routines and the cumulative growth of capability, knowledge conversion and tacit knowledge in professional service firms, and a set of boundary references — the representativeness heuristic that treats similarity as evidence, the role and danger of metaphor in organization studies, and misreadings of exploration and exploitation.

The contribution lies in the third paragraph of each pairwise analysis. The website states what the theory says and what it is mapped onto; it does not state the conditions under which the mapping holds and the conditions under which it fails. This paper supplies those for all twelve. Three of the twelve do not hold under the premises of the source theory; the paper says why, and what part remains usable.

From this we build one tool: a nine-row analogy audit, each row written as criterion, what passing looks like, what failing looks like. The twelve pairs are then scored against it, producing a table computed here rather than quoted. The central claim is stated as a falsifiable proposition together with the evidence that would refute it. Five boundary cases close the paper, including two that must be stated: an analogy that reads well is not the same as a decision that is right, and the pull of after-the-fact explanation.

No client company is named, no description that could identify one is used, and none of the firm's own financial figures appear.

Keywords: analogical reasoning; structure mapping; conceptual transfer; absorptive capacity; organizational routines; professional service firms; primary market

1. 问题：客户的理论，为什么解释得了我们自己

一家精品投行八年只做硬科技。八年里，它读的材料、开的会、见的人，绝大部分内容是别人的技术。技术的名词它未必全懂，但那些用来解释技术的概念，它听了成千上万遍：训练与推理是两回事，误差会在段与段之间累积、决定上限的不是数据量是数据质量、智能需要一个世界模型做先验。

然后发生了一件事。这些概念开始解释它自己。

这件事被写在了官网上。庚辛研究院在官网第 18 章列了十二对卡片，左边是「他们的理论」，右边是「我们的现象」。导语只有两句：

「庚辛长期与最前沿的公司一起工作，八年下来发现：他们用来理解世界的理论，恰好也解释了我们自己。理论来自这些公司与行业，现象是庚辛的日常。」

——庚辛官网 /institute.html [1]；同文亦见首页 /index.html [2]

章末一句：

「向自己陪伴的公司学习，是这份工作的特权。」

——庚辛官网 /institute.html [1]

本文回答两个被真实问过的问题。

第一个来自客户与投资人，问法通常是：一家中介机构说自己「懂」某个技术行业，这句话有没有可检验的含义？懂到什么程度算懂？

第二个来自内部，问法通常是：把客户的理论借过来解释自己，这是洞察，还是好听？两者怎么分辨？

本文认为这两个问题是同一个问题。一个机构能借到什么概念，取决于它已经会什么（第 3.2 节）；而一个借来的概念是洞察还是修辞，取决于它映射的是结构还是表面相似，以及它能不能指导一个具体动作（第 3.1 节与第 5 节）。第 9 节直接给出回答。

本文的路线：第 2 节把底本逐字摆出来，十二对一对不漏；第 3 节说明概念迁移在什么机制下发生；第 4 节逐对分析，每对三段，第三段是本文补的；第 5 节给出工具并用它给十二对打分；第 6 节把主张写成可以被推翻的命题；第 7、8 节划边界。

说「懂一个行业」，可检验的含义是：借得到它的概念，并且借对了。

表 1 / Exhibit 1 官网第 18 章「他们的理论」/「我们的现象」十二对逐字抄录。

序号	他们的理论	理论句（官网原文）	我们的现象	现象句（官网原文）	行号
1	训推分离（来自具身智能的工程实践）	「训练时靠多个辅助任务调优参数；到了推理，把所有的枝剪掉，只留一个输出。」	这解释了「找钥匙」。	「准备一场融资，我们穷尽技术、财务、结构的每一条辅助线；见投资人的那三十秒，剪掉全部枝蔓，只递一把钥匙。」	L3868–3873
2	专家模型融合（来自基础模型的前沿研究）	「不同的专家模型融合成一个更强的模型；每一代不必推倒重训，智能连续增长。」	这解释了「理解的规模化」。	「每个项目组都是一个专家模型，每周复盘把它们融合进同一个组织。八年，庚辛没有推倒重训过——能力是连续长出来的。」	L3874–3879
3	端到端（智能驾驶与深度学习的传统）	「从输入到输出，一个模型端到端负责；分成多段的流水线，误差会在段与段之间累积。」	这解释了「全托管」。	「从恢复事实到交割收口，同一个系统端到端负责，六个维度不换手——误差没有地方累积。」	L3880–3885
4	数据质量与教师模型（智能驾驶行业的公开工程实践）	「决定模型上限的不是数据量，是数据质量；要有教师模型持续评估：缺哪一类泛化数据，就补哪一类场景。」	这解释了我们的复盘制度。	「复盘会就是那个教师模型——评估上一周的判断质量、找出缺口，指挥下一周去补缺的那类场景。」	L3886–3889、3919
5	世界模型是先验（世界模型研究的共识）	「任何模型最终都需要一个世界模型做先验和约束，泛化能力才不靠运气。」	这解释了北坡、价格带与 4D。	「它们是庚辛的资本世界模型：每一笔交易进入执行之前，先被这套先验约束过一遍。」	L3920–3925
6	泛化与效率成反比（来自具身智能的场景分析）	「结构化场景要效率，非结构化场景要泛化；真正把模型训出来的，是非结构化场景的数据。」	这解释了我们复杂局面的偏好。	「标准轮次是结构化场景，方法早已成熟；复杂局面才是高质量数据。我们据此选择走进非结构化的局面：每解开一个，组织就多一分泛化。」	L3926–3929、3959
7	范式的接力（基础模型行业的公开叙事）	「一个范式到了效率下降的节点，增长不会停——接力棒交给下一个范式。」	这解释了「第八年」。	「八年，庚辛两次交棒：从单点服务到操作系统，从中国到全球——增长的曲线没有断过。」	L3960–3963、3984
8	世界先于画面（来自空间智能与世界模型研究的公开路线之争）	「屏幕上的一切呈现都只是底层世界的投影；智能应当建模投影之前的世界本身，而不是投影之后的像素。」	这解释了「恢复事实」。	「BP、Memo、路演，都是投影；投影之前，先有那个真实的世界。所以每个项目的第一步，是像企业考古学家一样重建业务、技术、客户、交易历史与关键风险——先建世界，再渲染画面。材料只是这个世界的一次采样。」	L3985–3990
9	解耦生长，统合交付（来自模块化学习与复杂系统工程公开思想）	「复杂能力不该由一个整体硬扛：按性质拆成层，让每一层独立生长，最后面向同一个目标重新耦合成一次交付。」	这解释了「二十七个子系统」。	「4D、价格带、四层旋涡、GQW——每个子系统各居其层、各自成章；到了一单交易，再被重新编排成一次交付。拆开是为了各自变强，合拢是为了共同成事——名单会过期，系统不会。」	L3991–3994、4063
10	少原则，多演绎（来自强化学习与自我博弈的公开传统）	「聪明的行为不是规定出来的，是演绎出来的：只写下第一性原理级的少数原则，最优策略交给大规模的真实实践去推演。」	这解释了「只做减法的清单」。	「我们写下的是边界，不是打法：踩过的坑写成禁令，复盘会全员重读；规则写得越细，越有人卡规则的缝，所以模糊处交给被所有人信任的人。方法可以变，边界只加不减——具体的打法，在一场场真实交易里长出来。」	L4064–4067、4131
11	智能出于多样（来自《心智社会》以来的认知科学经典）	「强大的智能来自大量异质心智的协作互补，而非某个单一完美的单体——多者成社会，社会即心智。」	这解释了「球面上的组织」。	「庚辛不是一张层级表，是同一个球面上一组尽量拉开的方向：方向铺得越开，覆盖越大、冗余越小；哪个角度空了，就有人转过去。组织的力量不靠某个孤胆英雄，靠异质方向的互补——挤在一起，就都成了平均值。」	L4132–4135、4160
12	观测易得，干预稀缺（来自因果推断的教科书共识）	「因果之梯上，《看见》低于《去做》：真正稀缺的不是海量观测，而是『在什么状态、做了什么』。后每样干预数据。」	这解释了「把钱投进去」。	「看一百个案子是观测；把自己的钱投进深度服务过的公司，是干预。只有带反馈的干预，才教得会状态如何转移。」	L4161–4164、4193

2. 底本怎么取的
十二对「理论」与「现象」的逐字对照。底本是公开网页，任何人可以复现。取法：[curl](https://glacier.mba/institute.html) 拉取 <https://glacier.mba/institute.html> 与 <https://glacier.mba/index.html>，依次剥掉 HTML 注释、<script>、<style>，再把块级标签换成换行、去掉

其余标签、反转义实体、去空行，得到纯文本。抓取日期 2026-09-22。下文行号指 [institute.html](https://glacier.mba/institute.html) 纯文本的行号。纯文本（抓取 2026-09-22），行号为该纯文本行号。理论侧来源为官网第 18 章在 [institute.html](https://glacier.mba/institute.html) 的纯文本第 3839 行到第 4194 行。首页 [index.html](https://glacier.mba/index.html) 也有这一章，但只放了十二对里的两对（专家模型融合、世界模型是先验），完整十二对只在 [institute.html](https://glacier.mba/institute.html) [1][2]。

表 2 / Exhibit 2 十二对的理论侧来源，逐字取自官网「来源」行 [1]，本文补上可解析的 DOI 或 arXiv 地址（2026-09-22 逐条请求核可达，结果见写作说明）。

序号	官网「来源」行原文	可访问地址
T1	Jaderberg et al. (2017), Reinforcement Learning with Unsupervised Auxiliary Tasks, ICLR 2017, Figure 1	arXiv:1611.05397
T2	Jacobs, Jordan, Nowlan & Hinton (1991), Adaptive Mixtures of Local Experts, Neural Computation 3(1), Figure 1	doi:10.1162/neco.1991.3.1.79
T3	Bojarski et al. (2016), End to End Learning for Self-Driving Cars, arXiv:1604.07316, Figure 4	arXiv:1604.07316
T4	Hinton, Vinyals & Dean (2015), Distilling the Knowledge in a Neural Network, arXiv:1503.02531	arXiv:1503.02531
T5	Ha & Schmidhuber (2018), World Models, arXiv:1803.10122, Figure 8 · CC BY 4.0	arXiv:1803.10122
T6	Wolpert & Macready (1997), No Free Lunch Theorems for Optimization, IEEE Trans. Evolutionary Computation	doi:10.1109/4235.585893
T7	Kuhn (1962), The Structure of Scientific Revolutions, University of Chicago Press	press.uchicago.edu 书页
T8	Dawid & LeCun (2023), Introduction to Latent Variable Energy-Based Models: A Path Towards Autonomous Machine Intelligence, arXiv:2306.02572, Figure 9 · CC BY 4.0	arXiv:2306.02572
T9	Simon (1962), The Architecture of Complexity, Proceedings of the American Philosophical Society 106(6)	jstor.org/stable/985254
T10	Silver et al. (2017), Mastering the Game of Go without Human Knowledge, Nature 550	doi:10.1038/nature24270
T11	Minsky (1986), The Society of Mind, Simon & Schuster	openlibrary.org OL2714978M
T12	Pearl & Mackenzie (2018), The Book of Why, Basic Books	basicbooks.com 书页

来源与说明：标号 T1-T12 只在本文内使用。

表 3 / Exhibit 3 十二条理论侧来源的载体与年代分布。

维度	分档	条目	条数	占比（分母 12）
载体	同行评议期刊论文	T2、T6、T9、T10	4	33%
载体	会议论文与预印本	T1、T3、T4、T5、T8	5	42%
载体	专著	T7、T11、T12	3	25%
年代	1962–1997（经典期）	T2、T6、T7、T9、T11	5	42%
年代	2015–2023（当代期）	T1、T3、T4、T5、T8、T10、T12	7	58%
领域	机器学习／人工智能	T1、T2、T3、T4、T5、T6、T8、T10	8	67%
领域	认知科学、科学哲学、复杂系统、因果推断	T7、T11、T9、T12	4	33%

来源与说明：本文按表 2 逐条归类计数，分母 12。归类规则写在写作说明 8.1。

一条抄录说明：纯文本里理论名与来源标签之间的「——」和多余空格，是剥标签时两个相邻元素接缝产生的，不是原文标点。本文抄录时把它还原成「理论名（来源标签）」的形式，理论句与现象句一字未改。

2.2 十二对，逐字

2.3 理论侧来源：官网自己标了

十二对每一对下面都有一行「来源」，指向一篇具体的论文或专著。这一点值得单独说一句：借概念的人写下了自己从哪儿借的。这让第 4 节的逐对分析成为可能——可以回到原文，看这个概念在它自己的领域里到底说了什么。



图 1 / Exhibit 4 「他们的理论 → 我们的现象」这条迁移关系的结构。

来源与说明：图为本文归纳；映射关系结构而非表面相似出自本文第 3.1 节，上游那道吸收能力的闸出自第 3.2 节，下游九条判据与两个出口出自第 5 节表 4。图中的条件与结果是占位，不指任何具体一对。

2.4 这些来源本身透露了什么

把表 2 按载体与年份归类，能看出一件不在官网文字里的事。下表是本文从表 2 的十二条来源数出来的。

归类规则：载体按官网「来源」行标注的载体判（T1 官网标 ICLR 2017，计会议论文）；年代按来源行标注的年份判；领域按该文发表时所在的学科社群判。每一维度三档或两档合计均为 12，无重复计数。

表 3 里有两个数字值得停一下。

第一，67% 的概念来自机器学习与人工智能。这不是一个中立的分布。如果一家机构借概念是随机的、跟着热度走的，来源会散在更多领域。集中在一个领域，说明借的不是热词，是它每天真的在读的东西。这一点在第 3.2 节会被当成吸收能力的证据，也在第 6 节被写成可检验的推论。

第二，42% 的来源是 1997 年以前的。最老的一条是 1962 年的两篇（T7、T9）。也就是说，这些「前沿技术的概念」里有接近一半是六十年前的老东西，只是被当代的工程实践重新用了一遍。借概念的人借到的，往往是老概念的新用法。

借来的概念，四成是六十年前的老东西，只是被重新用了一遍。

3.

机制：概念为什么会从客户跑到自己身上

第 2 节摆的是现象。这一节问的是：这种迁移是怎么发生的，为什么会发生在这家机构而不是别处，以及它什么时候可靠、什么时候不可靠。

3.1 类比推理：映射的是结构，不是像不像

认知科学里对类比有一套成熟的理论。Gentner 的结构映射理论 [3] 给出了核心判据：一个好的类比，映射的是关系结构，不是对象属性。把太阳系映射到原子，有价值的不是「都是圆的」，是「中心的东西吸引外围的东西，外围的东西绕着转」这一组关系。她把这条原则叫做系统性原则（systematicity）：高阶关系（关系之间的关系，尤其是因果关系）优先被映射，孤立的表面属性不被映射。

Gentner 与 Markman [4] 把这条原则推广到相似性判断本身：人判断两样东西像不像，做的其实就是一次结构对齐。Holyoak 与 Thagard [5] 从另一条路给出一个多重约束的框架：一个类比映射同时受三种约束拉扯——结构一致性、语义相似性、以及使用者当下的目标。第三条约束是关键：目标会把映射拉偏。一个人想证明某件事，类比就会往那个方向长。后来他们把这套框架写进综述 [7]，也把这条警告一起写了进去。

Gick 与 Holyoak [6] 用实验证明了迁移有多难：给被试读一个结构相同的故事，多数人在没有提示时想不起来用

它解决面前的问题；给了提示，解决率大幅上升。这个实验的含义是反直觉的——类比不是自动发生的，需要有人提示、需要检索线索。一家机构每天泡在一个领域的概念里，等于每天被提示，这正是迁移会发生在它身上而不是别处的原因。

对本文的用处是一把尺子：判断一个借来的概念是洞察还是修辞，先看它映射的是关系结构还是表面相似。第5节的自查表第一行就是这条。

3.2 吸收能力：你只学得会你已经会的东西

为什么是这家机构借到了这些概念，而不是任何一家读同样新闻的机构？

Cohen 与 Levinthal [8] 的答案是吸收能力：一个组织识别、消化并应用外部新知识的能力，取决于它已有的相关先验知识。他们强调这种能力是累积的、路径依赖的——今天能吸收什么，取决于昨天积累了什么；而且在某个领域停止投入，会让未来在该领域的吸收能力下降。Zahra 与 George [9] 后来把它拆成两段：获取与消化（潜在吸收能力），转化与利用（实现的吸收能力）。两段之间有落差——很多组织获取了知识却没有转化成行为。

这两条直接对上了表3的两个数字。67% 的概念集中在机器学习领域，是先验知识在起作用：八年只做硬科技，读的材料集中在那儿，所以能识别的概念也集中在那儿。而 Zahra 与 George 的两段划分给出了本文最重要的检验点：把概念写上墙是「获取与消化」，改变了某个具体动作才是「转化与利用」。第5节自查表的第五行、第6节的主命题，都建在这条上。

吸收能力还有一个不太被引用的推论，本文要把它写出来：吸收能力也是一个过滤器。它决定了什么进得来，也决定了什么进不来。一个机构借不到的概念，不一定是那些概念没用，可能只是它没有对应的先验。这是第7节反例三的来源。

3.3 惯例与能力：为什么「没有推倒重训」是一句需要小心的话

表1第2对说：「八年，庚辛没有推倒重训过——能力是连续长出来的。」

演化经济学对这句话有现成的理论，也有现成的警告。Nelson 与 Winter [10] 把组织的能力安放在惯例（routines）上：惯例是组织的记忆，是能力的存放形式，它的变化是渐进的、路径依赖的、带惯性的。这套理论一方面支持「能力连续长出来」这个描述，另一方面提醒：连续性与惰性是同一件事的两面。惯例不会自己退出，旧惯例会在环境变了以后继续运行。

Feldman 与 Pentland [11] 补上了另一半：惯例不是死的。他们把惯例拆成 ostensive（抽象的、大家口中的那个流程）与 performative（每一次真的做出来的那个动作）两个面向，指出两者之间的张力正是惯例内生变化的来源——每一次执行都可能偏离口头版本，偏离累积起来就是变化。

对本文的用处有两条。第一，「没有推倒重训」这句话在理论上成立，但它同时意味着承担路径依赖的成本，这一点官网没写，本文在第4.2节补。第二，Feldman 与 Pentland 的两个面向给了一个可操作的量：口头版本与实际动作之间的差距。一个借来的概念如果只改了口头版本（大家开始这么说了），没有改实际动作（每次真的做出来的事没变），那它停在了 Zahra 与 George 的第一段。

3.4 知识怎么在一家专业服务机构里传下去

Nonaka [12] 的组织知识创造理论把知识分为默会与显性两类，把组织里的知识转化画成四步循环：社会化（默会到默会）、外化（默会到显性）、组合（显性到显性）、内化（显性到默会）。Polanyi [13] 更早给出了默会知识的根本命题：我们知道的比我们说得出来的多。

这两条解释了第18章这件事本身是什么类型的动作。一家投行的判断力主要是默会的——怎么看一家公司、什么时候该慢下来、哪句话不能说。把这些默会的东西对上一个来自客户领域的显性概念，是 Nonaka 说的外化：给说不出的东西找到一个能说出来的名字。

外化有明确的收益：命名之后可以传授、可以讨论、可以批评。八年的判断力如果只留在几个人身上，新人要重新踩一遍坑；给它一个名字，新人至少能从名字开始问。

外化也有明确的代价，本文在第7节反例四写：名字一旦好听，就不再被质疑。外化把默会知识变成显性知识，同时也把它变成了一句口号。

3.5 边界文献：类比什么时候会骗人

前面四小节讲的是迁移为什么发生、为什么可靠。这一小节讲它什么时候不可靠。

Tversky 与 Kahneman [14] 的代表性启发是最根本的一条：人判断 A 是不是属于 B，靠的是 A 像不像 B，而不是概率。相似性被当成证据用。类比的危险正在这里——一个类比读起来越顺，越像是被证明了，其实只是越像。

Gilovich [15] 做过一个更贴近本文的实验：让被试对一个假想的国际危机做判断，材料里插入与二战或与越战相呼应的无关表面线索（比如会议地点是不是「慕尼黑大厅」、难民是不是坐「火车」），被试的政策倾向随之

系统性偏移。这个实验说明：表面相似会在人不知情的情况下改变决策。把这条搬到本文的处境：一个来自客户领域的概念，哪怕只在表面上像，也会影响这家机构怎么描述自己、进而怎么行动。

March [16] 的探索与利用是另一类错误的来源。他的模型说明，个体与组织代码之间的相互学习速度会影响长期绩效：社会化太快，多样性被吞掉，组织收敛到一个平庸的共识。这条常被误读成「要多探索」这样一句口号。真实的结论是有条件的——最优的学习速度取决于环境变化的速度与组织的规模。本文在第 4.11 节用它给「智能出于多样」那一对划边界。

最后两条是关于隐喻本身的。Morgan [17] 指出组织理论各个流派其实建在不同的根隐喻上（机器、有机体、文化、政治系统），每个隐喻照亮一部分、遮住另一部分。Tsoukas [18] 进一步区分了隐喻的两种用法：作为字面描述的替代品（装饰性的，可以被去掉），与作为生成性的认知工具（改变了你能看见什么，去不掉）。他给出的判据对本文特别有用：一个隐喻如果去掉之后什么都没少，它就是装饰。

五组文献合起来给出的判据，正是第 5 节那张表的九行。

一个类比读起来越顺，越像是被证明了，其实只是越像。

4. 逐对分析：十二对，每对三段

体例：每一对三段。第一段说这个理论在它自己的领域里说的是什么，给来源。第二段说迁移到了什么上，引官网原句。第三段说这个迁移在什么条件下成立、什么条件下不成立——这一段官网没有，是本文补的。

4.1 训推分离 → 找钥匙

理论说什么。T1 的做法是在强化学习的主任务之外加一组无监督辅助任务，让它们共享同一套表征去做梯度更新；辅助任务的作用是在奖励稀疏的时候提供额外的学习信号，训练完成后推理只走主策略那一条。要点有两个：辅助任务和主任务共享同一套参数，以及推理时剪掉的只是计算路径，不是它们训练时留下的影响。

迁移到什么上。官网写的是「找钥匙」：

「准备一场融资，我们穷尽技术、财务、结构的每一条辅助线；见投资人的那三十秒，剪掉全部枝蔓，只递一把钥匙。」[1]

什么条件下成立。三个前提。第一，辅助线必须真的改变主线结论——技术、财务、结构的工作要影响那把钥匙长什么样，否则它们不是辅助任务，是并行产出一堆没人读的材料。第二，接收端的带宽确实稀缺；三十

秒是硬约束时剪枝才有意义。第三，剪枝要可恢复：被追问时能把枝接回来。

什么条件下不成立。原理论里推理时剪枝不丢信息，是因为信息已经编码进权重。一份材料没有权重。剪掉的枝蔓不会自动留在钥匙里，除非有人记得。所以这个迁移有一个隐含的载体：把辅助线做完的那个人必须在场。如果做辅助线的人和递钥匙的人不是同一批人，训推分离就退化成信息丢失。还有一个反向的失效：在尽调阶段，对方要的恰恰是枝蔓，这时候继续剪枝就不是提炼，是隐瞒。

可检验的形式。被剪掉的辅助线，能否在被追问后的一个工作日内复原成可交付的形式。能，说明它真的进了权重；不能，说明它从来没有影响主线。

4.2 专家模型融合 → 理解的规模化

理论说什么。T2 是混合专家模型的原始论文：若干个专家网络加一个门控网络，门控决定每个输入交给谁；专家之所以会分化，是因为门控制造了竞争，让每个专家去负责自己更擅长的那一类输入。这里要说一句准确的话：T2 讲的是分工加路由，不是「融合成一个更强的模型」。把多个模型合成一个的那一支（模型合并、集成蒸馏）是另一支文献，官网理论句写的是后者，标注的来源是前者。

迁移到什么上。

「每个项目组都是一个专家模型，每周复盘把它们融合进同一个组织。八年，庚辛没有推倒重训过——能力是连续长出来的。」[1]

什么条件下成立。混合专家要工作，需要两样东西：一个门控，以及一个共同的损失函数。组织里的门控是存在的——决定哪类问题交给哪个组的那个人或那套规则。共同的损失函数是问题所在：不同项目组服务不同的客户、面对不同的局面，各自的成败不共享同一个目标函数。缺了共同损失，「融合」就没有梯度可用，只剩下信息通报。所以这个迁移成立的条件是：复盘必须改变其他组的下一次动作，而不是让其他组知道发生了什么。

什么条件下不成立。第 3.3 节的惯例理论给出另一侧的限定：「没有推倒重训」在描述上成立，但它同时是路径依赖的另一种说法。神经网络不推倒重训的代价是灾难性遗忘与旧表征的束缚；组织不推倒重训的代价是旧惯例在环境变了以后继续运行 [10]。这一句在官网上是正面表述，在文献里是一句中性的描述加一条成本提示。本文的补充是：连续积累的能力需要一个配套装置，用来定期让某条旧惯例退休，否则连续性会慢慢变成惰性。

可检验的形式。取任意一次复盘里得到的教训，看它是否在另一个组接下来的动作里出现；以及过去一年里有没有任何一条明确废止的旧做法。前者检验融合，后者检验惰性。

4.3 端到端 → 全托管

理论说什么。T3 训练一个网络直接从摄像头像素映射到方向盘转角，不再人工划分车道检测、路径规划等中间阶段。它的论证是：分段流水线的每一段各自优化自己的目标，段与段之间的目标并不一致，误差会在接缝处累积。

迁移到什么上。

「从恢复事实到交割收口，同一个系统端到端负责，六个维度不换手——误差没有地方累积。」[1]

什么条件下成立。端到端要赢过分段，需要两个条件：一是有足够的端到端训练信号；二是误差确实主要产生在接缝处，而不是产生在单个环节内部。第二个条件在跨机构协作里通常满足——换手处的信息损失是真实的、可观察的。

什么条件下不成立。第一个条件在这里是弱的。端到端的学习信号在自动驾驶里是海量驾驶数据，在一笔交易里是交割结果，样本稀疏且反馈很慢。更重要的是，端到端有一个公认的代价：不可解释与难以定位。分段流水线的好处恰恰是出问题能归因到某一段。一个组织如果「六个维度不换手」，就少了换手时天然产生的那次交叉检查。所以这个迁移成立的前提是配一个外部评估者——正好是下一对里的教师模型。第3对与第4对必须成对使用，单独用第3对是危险的。

可检验的形式。出问题时，能否在不追问个人的前提下定位到是哪一個维度先出的问题。不能，说明端到端的代价已经发生。

4.4 数据质量与教师模型 → 复盘制度

理论说什么。T4 是知识蒸馏：先训练一个强的教师模型，让学生模型去拟合教师输出的软概率分布，而不是硬标签。它的核心洞察是「暗知识」——教师在错误类别之间给出的相对概率，携带了硬标签里没有的信息。学生学的不只是「答案是什么」，还有「错在哪里、次优是什么」。

迁移到什么上。

「复盘会就是那个教师模型——评估上一周的判断质量、找出缺口，指挥下一周去补缺的那类场景。」[1]

什么条件下成立。蒸馏要工作，教师的输出必须是软的。翻译到复盘：复盘的产出必须包含「当时第二好的选项是什么」「这个判断错在哪一档」，而不只是成或败。这是本文认为整章里最能直接指导动作的一条：在复盘纪要里加一栏「次优选项」，就是把硬标签换成软标签。

什么条件下不成立。两条。第一，蒸馏预设教师比学生强。组织里的教师与学生常常是同一批人——自己评估自己上周的判断。自蒸馏在机器学习里有效，但它放大已有的偏差而不是纠正它，因为教师的错误会被完整地传给学生。第二，T4 里的教师是先独立训练好的。复盘会如果只用本周自己产生的材料，它就不是一个独立的教师，是同一个模型的回声。

可检验的形式。随机抽十份复盘纪要，数其中有多少份写了「当时还有哪些选项、为什么没选」。这个比例就是软标签的比例。

4.5 世界模型是先验 → 北坡、价格带与 4D

理论说什么。T5 把智能体拆成三块：视觉编码器 V 把观测压成潜变量，记忆模块 M 学会预测下一步，控制器 C 在这两者之上做决策。论文里最有名的一段是：控制器可以完全在学到的世界模型里训练（在「梦」里），然后迁移回真实环境。论文自己也报告了这条路的陷阱：智能体会去利用世界模型的缺陷，找到在梦里得分很高、在现实里不成立的策略。

迁移到什么上。

「它们是庚辛的资本世界模型：每一笔交易进入执行之前，先被这套先验约束过一遍。」[1]

什么条件下成立。世界模型要当先验用，必须持续被现实纠错——M 的预测误差要回流去更新 M。一套框架如果每年都在被某次真实结果改写一两条，它是先验；如果十年没改过，它是教条。

什么条件下不成立。原论文的陷阱在组织里有完全对应的形式：用内部框架自我论证。一个项目如果只在自家的框架里算得通，而没有一次外部现实的检验，那就是在梦里得高分。这不是一个比喻式的担心，是原论文实测过的失败模式。

可检验的形式。过去十二个月里，这套先验有没有任何一条因为一次真实结果被改写。数目为零，先验已经变成教条。

表 4 / Exhibit 5 类比自查表。

序号	判据	通过的样子	不通过的样子	依据
1	映射的是结构还是表面相似	能写出两边各自的因果关系，并逐条对上	只能说「都是…」 「很像…」，说不出关系怎么对	[3][4]
2	原理论的适用前提在我们这儿成立吗	能列出原理论的前提，逐条判成立或不成立	从没找过原理论的前提，只看结论	[3][5]
3	可证伪吗	能说出「看到什么就算这个类比不成立」	怎么都说得通	[14]
4	有没有反例	至少能举一个这个类比明显失效的场合	想不出反例，也没找过	[15]
5	迁移后能不能指导具体动作	能改掉或加上一个可观察的动作	只改了大家的说法，动作没变	[9]
6	是不是只在事后解释	有一条做法是先有理论、后做出来的	全部做法都在借用之前就存在	[14][15]
7	会不会掩盖真实成本	能说出这个做法的代价是什么，谁承担	这个类比只讲好处	[10][16]
8	换个理论能不能同样解释	找过替代解释，并说明为什么这个更好	没找过替代解释	[17][18]
9	去掉这个概念，还剩什么	去掉后确实少了一样看得见的东西	去掉后什么都不少，只是不好听了	[18]

来源与说明：一个外来概念用到自己组织之前要过的九条判据。判据来源标在末列：[3]-[7] 类比与迁移，[8][9] 吸收能力，[10][11] 惯例，[12][13] 知识转化，[14]-[18] 边界文献。本表为本文归纳。

4.6 泛化与效率成反比 → 对复杂局面的偏好

理论说什么。这里要先做一个区分。官网标的来源 T6 是没有免费午餐定理，它的命题是：在所有可能的目标函数上平均，任何两个优化算法的期望表现相同。它说的不是「泛化与效率成反比」，它说的是「在所有问题上平均，没有算法更强」。官网图注画的锥面与等长向量忠实于 T6；理论名那一句则是具身智能工程实践里的经验总结，是另一个命题。这是十二对里第二处理论名与标注来源不完全对齐的地方。

迁移到什么上。

「标准轮次是结构化场景，方法早已成熟；复杂局面才是高质量数据。我们据此选择走进非结构化的局面：每解开一个，组织就多一分泛化。」[1]

什么条件下成立。没有免费午餐定理的前提是问题在所有可能目标函数上均匀分布。现实里不均匀，所以专用化是有价值的——这恰恰支持了「选一类局面并在它上面变强」这个选择。迁移成立的真正条件在别处：非结构化局面之间要共享结构。解开一个复杂局面之后，下一个复杂局面要能用上这次的解法，「多一分泛化」才成立。

什么条件下不成立。如果每个复杂局面都是一次性的、彼此不共享结构，那么解开它只产生消耗，不产生泛化。这在实务里是个真问题：复杂常来自特异的历史遗留

，而特异的东西按定义不复用。这一对因此需要一个量：复用率。

可检验的形式。取过去若干个被判为复杂的局面，看后一个局面里有没有用到前一个的解法。有，泛化成立；没有，那是消耗。

4.7 范式的接力 → 第八年

理论说什么。T7 的核心命题恰恰不是平滑接力。Kuhn 说的是：常规科学在一个范式里解谜，异常累积到一定程度产生危机，然后发生革命；新旧范式之间不可通约——没有共同的尺子可以度量两者的成就，旧范式的一些成果在新范式里甚至不再被当成问题。

迁移到什么上。

「八年，庚辛两次交棒：从单点服务到操作系统，从中国到全球——增长的曲线没有断过。」[1]

什么条件下成立。「一个范式到了效率下降的节点，增长不会停——接力棒交给下一个范式」这句，在描述多条 S 曲线相继接续这个形状上是成立的，这也是这个意象在产业叙事里流行的原因。

什么条件下不成立。按 Kuhn 的原意，范式更替伴随的是断裂而不是连续：旧的衡量标准被放弃，旧的成就变得不可比。官网说「曲线没有断过」——这更接近 Kuhn 说的常规科学内部的累积，而不是革命。所以这一对里，借来的是「接力」这个意象，没有借到原理论

表 5 / Exhibit 6 用表 4 的九条判据给表 1 的十二对逐对打分。

序号	理论	1 结构	2 前提	3 可证伪	4 反例	5 指导动作	6 非事后	7 不掩盖成本	8 替代解释	9 去不掉	得分
1	训推分离	✓	△	✓	✓	✓	×	△	✓	✓	7.0
2	专家模型融合	△	×	✓	✓	✓	×	×	△	✓	5.0
3	端到端	✓	△	✓	✓	✓	×	×	✓	✓	6.5
4	数据质量与教师模型	✓	△	✓	✓	✓	×	△	✓	✓	7.0
5	世界模型是先验	✓	✓	✓	✓	✓	×	△	✓	✓	7.5
6	泛化与效率成反比	△	×	✓	✓	△	×	△	△	✓	5.0
7	范式的接力	×	×	△	✓	△	×	×	×	×	2.0
8	世界先于画面	✓	✓	✓	✓	✓	×	△	✓	✓	7.5
9	解耦生长， 统合交付	✓	△	✓	✓	✓	×	△	✓	✓	7.0
10	少原则，多 演绎	△	×	✓	✓	✓	×	△	×	×	4.0
11	智能出于多样	△	△	✓	✓	△	×	×	△	✓	5.0
12	观测易得， 干预稀缺	✓	×	✓	✓	✓	×	✓	✓	✓	7.0

来源与说明：✓ = 通过，△ = 有条件通过（条件写在第 4 节对应小节），× = 不通过。打分由本文做出，依据是第 4 节的逐对分析；官网没有对应评价。未列为通过数（✓ 计 1，△ 计 0.5）。

的核心命题。这不必然是错误，但要说清楚：它是一个装饰性隐喻，不是一个生成性隐喻 [18]。

它仍然能贡献什么。Kuhn 理论里有一个可以直接落地的部件：异常。范式更替的先行指标不是增长放缓，是用现有方法解释不了的失败在累积。翻译成组织里的动作：维护一张「解释不了的失败」清单。这张清单比任何增长指标都更早告诉你该换方法了。这条是本文从原理论里取出来、官网没有的部分。

4.8 世界先于画面 → 恢复事实

理论说什么。T8 是隐变量能量模型这条路线的介绍，背后是一场公开的路线之争：智能应当建模产生观测的那个潜在世界，还是直接建模观测本身（像素）。主张前者的理由是，在潜空间里预测可以忽略不可预测的细节，而在像素空间里预测会被迫去建模噪声。

迁移到什么上。

「BP、Memo、路演，都是投影；投影之前，先有那个真实的世界。所以每个项目的第一步，是像企业考古学家一样重建业务、技术、客户、交易历史与关键风险

——先建世界，再渲染画面。材料只是这个世界的一次采样。」 [1]

什么条件下成立。前提是那个「底层世界」确实存在，而且可被重建到足以约束材料。财务事实、技术事实、交易历史大体可以重建，因为它们留下了痕迹。这一对是十二对里结构映射最干净的一对：潜变量与观测的关系，对上事实与材料的关系，两边的因果方向一致。

什么条件下不成立。两处。第一，意图、关系、未来计划不是可重建的痕迹，它们只存在于人身上，重建会变成猜测。第二，也是更要紧的一处：把「投影」一律当噪声会丢掉真实的信息。市场对一家公司的叙事本身就是事实的一部分，它影响定价，也影响别人的行动。潜变量模型里像素是果，在一级市场里，画面有时候是因。可检验的形式。在一次事实重建里，被判为「材料噪声」而丢弃的东西，事后有多少被证明是影响了结果的。



图 2 / Exhibit 7 九条判据各自在十二对上的通过分布。

来源与说明：图为本文归纳，计数逐列取自本节表 5，每行合计十二。本图只看每条判据的分布，不排名、不合计逐对得分，也不标出任何一对的名次。

4.9 解耦生长，统合交付 → 二十七个子系统

理论说什么。T9 提出近可分解系统：复杂系统往往由层级构成，层内的相互作用远强于层间；短期内每个子系统的动力学近似独立，长期看只有子系统的聚合行为才影响其他子系统。Simon 用一栋三房八格建筑的热扩散矩阵做例子，格与格之间的传导系数差着一到两个数量级——这正是「近可分解」的定量含义。

迁移到什么上。

「4D、价格带、四层旋涡、GQW——每个子系统各居其层、各自成章；到了一单交易，再被重新编排成一次交付。拆开是为了各自变强，合拢是为了共同成事——名单会过期，系统不会。」[1]

什么条件下成立。Simon 的条件是可以量化的：层内耦合必须远强于层间耦合。组织里的对应物是接口稀疏且稳定。如果两个子系统之间每个月都要重新约定怎么衔接，近可分解性不成立，拆分带来的是协调成本而不是并行生长。

什么条件下不成立。当一单交易需要二十七个子系统同时高带宽协商时，这个结构反而更慢。近可分解系统的好处在于它把改动局部化；一旦改动无法局部化，层级结构就只剩下开销。

「名单会过期，系统不会」这句要限定。系统同样会过期，只是慢一档。惯例理论 [10][11] 说的正是这件事：能

力存放在惯例里，惯例有惰性。更准确的表述是：名单以月计过期，系统以年计过期。本文不改官网原句，只补上这条限定。

可检验的形式。数任意两个子系统之间的接口在过去一年里被修改了几次。次数低，近可分解成立；次数高，拆分正在产生成本。

4.10 少原则，多演绎 → 只做减法的清单

理论说什么。T10 是自我对弈强化学习：不用人类棋谱，只给规则，靠自我对弈与搜索生成训练信号，网络与搜索互相提升。它的成功有三个前提，缺一不可：规则固定且完全已知；胜负在每盘结束时立即明确；可以无限次自我对弈。

迁移到什么上。

「我们写下的是边界，不是打法：踩过的坑写成禁令，复盘会全员重读；规则写得越细，越有人卡规则的缝，所以模糊处交给被所有人信任的人。方法可以变，边界只加不减——具体的打法，在一场场真实交易里长出来。」[1]

什么条件下不成立——这一对要先说不成立的部分。三个前提在一级市场里全部不成立：规则不固定（法域、市场、参与者都在变）；反馈慢且含混（一轮几个月，真正的结果要等到退出）；不能自我对弈（不能凭空生成一笔交易来练手）。这一对是十二对里读起来最顺、

前提落空最彻底的一对，本文在第7节反例一把它当成主要例子。

那还剩下什么。T10 里有一条结论与那三个前提关系较弱：注入人类先验知识（棋谱、手工特征）在长程上会限制上限。它对应的组织命题是：把规则写得过细会限制适应。而官网原句「规则写得越细，越有人卡规则的缝」说的其实是另一件事——规则细化引发的是规则套利，这一点用 Goodhart 式的指标失效或制度经济学解释得更直接，不需要自我对弈。

这一对的价值在哪。它的价值不在论证，在表述。「写边界不写打法」是一条能立刻执行的规则，它本身是对的，只是它的正当性不来自 T10。这正是第5节自查表第八行（换个理论能不能同样解释）要抓的情形：做法对，理由借错了。做法可以留，理由该换。

4.11 智能出于多样 → 球面上的组织

理论说什么。T11 的命题是：心智由大量本身不智能的 agent 组成，智能出现在这些 agent 的组织方式里；同一个节点对上 is agent，对下是 agency。这本书强调的是层级与分工的结构，多样性是手段。

迁移到什么上。

「庚辛不是一张层级表，是同一个球面上一组尽量拉开的方向：方向铺得越开，覆盖越大、冗余越小；哪个角度空了，就有人转过去。组织的力量不靠某个孤胆英雄，靠异质方向的互补——挤在一起，就都成了平均值。」^[1]

什么条件下成立。一个不能跳过的差异：T11 的 agent 是简单且不智能的，它们不需要互相理解，力量来自结构。组织里的成员不是简单 agent，他们之间要沟通，而沟通成本随多样性上升。所以多样性的收益在两个条件下成立：问题空间足够宽（需要多个角度才能覆盖），且任务可分解（不同角度的产出可以合起来用）。

什么条件下不成立。March [16] 给出转折点在哪儿：个体与组织共识之间的相互学习速度决定长期表现。社会化太快，多样性被吞掉，组织收敛到平庸的共识——这正是官网说的「挤在一起，就都成了平均值」。但反过来也成立：环境稳定时，慢社会化是浪费。所以「方向尽量拉开」不是一个无条件的好，它的最优程度取决于环境变化的速度。官网这句里缺的正是这个条件。

这一对缺一个量。「尽量拉开」拉开多少算够，原句没给。本文提一个可测的代理指标：同一个问题上，两个人独立给出不同判断的比例。这个比例太低，说明方向挤在一起了；太高，说明缺共同基础。它是可以按季度数出来的。

4.12 观测易得，干预稀缺 → 把钱投进去

理论说什么。T12 的因果之梯把推断分三级：关联（看见）、干预（去做）、反事实（设想）。核心命题是高级的问题不能只用低一级的数据回答。要注意的是，Pearl 说的干预有严格含义： $do(x)$ 指的是外生地把变量设成某个值，切断它原本的所有输入。

迁移到什么上。

「看一百个案子是观测；把自己的钱投进深度服务过的公司，是干预。只有带反馈的干预，才教得会状态如何转移。」^[1]

什么条件下不成立——这一对在原理论的框架里恰好不成立。投与不投是基于对结果的预期而做的选择，这是内生的，不是 $do(x)$ 。它是教科书里典型的选择偏差：只投看好的，然后观察结果，得到的是被选择过的样本，而不是干预数据。在 Pearl 的定义下，自选择的行动不是干预。这一对是十二对里唯一一处理论的核心定义与迁移直接冲突的。

那这句话在说什么。它说的其实是另一件真的、也重要的事：承担后果会改变你收到的反馈。不投只能看到自己看对没看对；投了会收到持续的、带成本的、无法回避的反馈。这在文献里更接近有切身利害的学习，而不是 do-calculus。这个修正不削弱那句话，它让那句话变得可以检验。

本文的修正，以及它指向的动作。因果之梯真正稀缺的那一级是第三级——反事实。「如果当时不投会怎样」

「如果当时按住不发会怎样」这类问题，需要的数据是被否掉的那些项目的后续。这几乎没有机构系统地记录，因为它不带来看得见的成果、也不好看。本文认为这才是最值得建的一份台账：记录被否掉的项目，而不只是记录做成的项目。这也是本文对整章唯一一条「官网的做法应当加一件」的建议。

记忆点 · KEY POINT

真正稀缺的不是干预，是反事实——被否掉的那些项目，谁在记。

5. 工具：类比自查表

5.1 九行判据

把第3节的五组文献压成一张表，给要借概念的人第二天就用。用法：一行一行过，任何一行落在「不通过的样子」那一列，这个借用就不能当论证用，只能当表述用。

第9行是 Tsoukas [18] 那条判据的直接落地：生成性的隐喻去不掉，装饰性的隐喻去掉之后什么都不少。它放在最后，因为它最狠。

5.2 用这张表给十二对打分

三个读法。

第一，第6列整列是×。十二对的现象句全部以「这解释了…」开头，被解释的做法在借用之前就已经存在。这不等于这些迁移没有价值——外化本身有价值（第3.4节）——但它意味着整章目前是解释性的，不是生成性的。要变成生成性的，需要出现至少一条相反方向的例子：先有理论，后做出来。第6节把这一条写成可检验的推论。

第二，得分最低的是第7对（范式的接力，2.0）与第10对（少原则，多演绎，4.0）。两者的共同点是：句子最顺，前提最不成立。这是第7节反例一的数据形式。

第三，得分最高的三对是第5、第8（各7.5）与第1、4、9、12（各7.0）。它们的共同点是两边的因果方向一致，而且都能落到一个可以数出来的动作上：先验被改写过几次、被丢弃的材料有多少后来证明有用、接口一年改几次、被否掉的项目有没有人记。能落到一个可以数的东西上，是这张表里最硬的信号。

记忆点 · KEY POINT

借来的概念要能改掉一个具体动作，否则它只是个名字。

6. 一个可以被推翻的命题

主命题。如果一个外来概念对一家机构的借用是结构映射（映射的是关系与因果结构，而不是表面相似），那么它应当能产生至少一条该机构在借用之前没有做、借用之后开始做、且可被第三方观察到的动作变更。只解释既有做法、不改变任何动作的借用，是修辞不是机制。

什么证据能推翻它。在一组机构里，把它们借来的概念按判据1（结构还是表面）分为两组，然后统计「借用后十二个月内出现新的可观察动作变更」的比例。如果两组比例没有差异，主命题不成立——那说明结构映射与表面相似在后果上没有分别，判据1就不是一条判据。

推论一（分布检验，本文已做了一次）。如果借概念靠的是吸收能力[8]，那么一家机构借来的概念应当集中在它长期服务的领域，而不是随机分布在当时的热门概念上。表3显示十二条来源里67%来自机器学习与人工智能，与这家机构八年只做硬科技的先验一致；另有42%是1997年以前的经典，说明借的不是热度。推论一在这

一个样本上没有被推翻。如果换一家机构，发现它借的概念散布在与业务无关的多个热门领域，推论一被推翻，吸收能力就不是这里的机制，追热点才是。

推论二（方向检验，尚未检验）。如果这些借用里至少有一部分是生成性的，那么应当存在至少一条做法，是先有概念、后做出来的。检验方法是时间顺序：概念写进内部文本的日期，早于该做法第一次执行的日期。表5第6列显示，仅凭公开文本无法判断先后，因为公开文本只给了现在的状态。这一条需要内部时间戳才能检验，本文做不了，明确留空。

推论三（配对检验）。第4.3节主张第3对（端到端）必须与第4对（教师模型）成对使用，因为端到端的代价是难以归因，而教师提供外部评估。可检验的形式：在只用了端到端、没有配外部评估的组织里，问题定位所需的时间应当更长。如果观察到没有外部评估的组织定位一样快，这条配对主张不成立。

记忆点 · KEY POINT

说不出「看到什么就算它不成立」，那它不是一个判断，是一句话。

7. 边界与反例

一份只有正面例证的分析不值得信赖。以下五类，是这套做法已知会失灵的地方。

反例一：类比讲得通，不等于做得对。这是最重要的一条。表5里得分最低的两对（范式的接力、少原则多演绎）恰恰是读起来最顺的两对。顺，是因为它们的意象强、句子短、画面清楚；而前提在一级市场里并不成立（第4.7、4.10节逐条列过）。Tversky 与 Kahneman [14] 解释了为什么会这样：人用相似性代替概率，读起来像，就被当成了被证明。Gilovich [15] 的实验进一步显示，表面线索会在人不知情的时候改变政策判断。所以判据不是句子顺不顺，是前提成不成立。一个类比可以留在墙上当表述，但不能进论证链——这是两种用途，必须分开。

反例二：事后解释的诱惑。表5第6列整列是×。每一句都是「这解释了…」，被解释的做法在借用之前就存在。事后解释有两个特征，都很难自觉：它不会失败（总能找到一个概念去解释已经发生的事），它让人觉得自己理解了（相似性被当作理解）。分辨方法只有一个，就是时间顺序：先有理论后有做法，才是机制；先有做法后有理论，是命名。命名有价值（第3.4节说过：可以传授、可以批评），但它不是解释。这一节之所以

必须写进论文，是因为第 18 章整章的句式天然把命名说成了解释，而写的人不会自己发现这件事。

反例三：只看看得见被写上墙的那些对。分母不可知。一个概念如果试过、发现对不上、于是被丢掉了，它不会出现在官网上。所以表 5 的得分分布是一个幸存者样本，不能读成「这家机构的类比质量」。真正的质量指标是丢弃率——试过多少对、留下多少对。这需要一份被丢弃的类比清单，而没有人会天然去建这份清单。这一条与第 4.12 节的结论是同一条：记下被否掉的，分母才存在。

反例四：借来的词会变成免检通行证。外化 [12] 把说不清的判断变成一个能说出口的名字，收益是可传授，代价是名字一旦好听就不再被质疑。一个做法挂上一个前沿概念之后，讨论它的人会先去论证这个概念，而不是先去检查这个做法。Tsoukas [18] 与 Morgan [17] 讲的正是这件事：根隐喻照亮一部分、遮住另一部分，而被遮住的部分恰恰是不再被问的那部分。表 4 第 9 行（去掉之后还剩什么）是针对这一条的解药。

反例五：反向污染——挑那些能被自己的理论解释的客户。吸收能力 [8][9] 是个过滤器：决定什么进得来，也决定什么进不来。一个机构长期用某一组概念理解世界之后，会更容易理解那些用同一组概念说话的公司，也更难理解不用这套语言的公司。理解得快，容易被当成判断得准。这一条没有现成的解药，只有一个提醒：理解顺畅本身不是一个判断依据。

先有理论后有做法，才是机制；先有做法后有理论，是命名。

8. 这个分析在什么范围内成立

一、对象范围。本文分析的是一份公开文本里的十二对概念迁移，出自一家专业服务机构。结论适用于「长期深度服务某个技术领域的中介机构如何借用该领域的概念」这一类情形。它不外推到技术公司之间的概念迁移，也不外推到不以知识工作为主的组织。

二、证据的性质。表 1、表 2 是逐字抄录，行号可复现。表 3 是从表 2 数出来的，规则写在写作说明。表 5 是本文的判断，不是测量。九条判据里有几条（尤其是第 5、第 6 行）需要内部信息才能准确判定，本文只用公开文本判，因此这些格子是保守估计——倾向于判 × 或 △。换一个评分者，个别格子会不同；本文的三个读法依赖的是整体形状（第 6 列整列、最低两对、最高几对），不依赖单个格子。

三、单样本。全文只有一个机构、一份文本。主命题（第 6 节）需要跨机构样本才能检验，本文做不了。推论一在这一个样本上没有被推翻，这不等于被证实。

四、没有做因果推断。本文不主张借用这些概念导致了这家机构的任何结果。第 4 节讲的是映射在什么条件下成立，不是借用带来了什么绩效。任何把表 5 的得分读成绩效指标的做法都是误读。

五、原理论的表述。第 4 节对十二篇来源的概括，是为了判断迁移是否成立而做的简化，不能替代原文。三处指出理论名与标注来源不完全对齐的地方（第 4.2、4.6、4.7 节），指的是概念与来源之间的对应关系，不是说原论文有问题。

六、一个诚实的空缺。本文没有做的是：拿到内部时间戳，检验有没有哪一条做法是先有概念后做出来的（第 6 节推论二）。这需要内部记录的时间顺序，公开文本给不了。在这项工作完成之前，本文对「解释性还是生成性」的判断是一个可被推翻的工作假设。

七、本文不做什么。不评价这家机构的任何业务结果，不构成投资建议，不给出任何可回溯到具体交易或具体公司的信息。

表 5 是判断不是测量；能推翻它的是内部时间戳，不是更好的措辞。

9. 结论

一家长期服务前沿科技公司的机构，会反复发现客户用来解释技术的概念也解释了自己。本文把十二对这样的概念逐字摆出来，逐对问了一个官网没问的问题：这个迁移在什么条件下成立。

回到第 1 节的两个问题。

「懂一个行业」有没有可检验的含义。有。可检验的含义是：借得到这个行业的概念，并且借对了。借得到，是吸收能力的结果——表 3 显示十二条来源里三分之二集中在一个领域，且四成是六十年前的经典，这说明读的是原文不是热搜。借对了，看的是表 4 那九行，尤其是第 1 行（映射结构还是表面）与第 5 行（能不能改掉一个具体动作）。一家机构如果借来的概念散落在当季所有热词上，并且没有一条改变过任何动作，那它不是懂，是跟读。

借客户的理论解释自己，是洞察还是修辞。两者都有，可以分开。本文给出的分法是表 4 的第 9 行：去掉这个概念，还剩什么。剩下一个可以数出来的动作（先验被改写过几次、复盘里写没写次优选项、接口一年改几次、

被否掉的项目有没有人记），是洞察；什么都不少，只是不好听了，是修辞。修辞不是罪，但它不能论证链。

本文对整章提了一条建议，只有一条：记录被否掉的项目。第 4.12 节的分析说明，因果之梯上真正稀缺的不是干预，是反事实；而反事实需要的数据，是没做成的那些。第 7 节反例三说明，类比的质量也要靠被丢弃的那些才能衡量。两条指向同一件事，本文认为这是整章最值得加的一件。

三句话可以概括全文：

1. 客户的理论解释得了我们，多半不是巧合，是我们只听得懂自己已经会的那部分。
2. 类比读起来越顺，越像是被证明了，其实只是越像——判据是前提，不是句子。
3. 一个借来的概念要能改掉一个具体动作，否则它只是个名字。

去掉这个概念，还剩什么——剩一个能数的动作，是洞察；什么都不少，是修辞。

参考文献 / References

体例：作者（年份）. 题名. 期刊/出版方, 卷(期), 页码. DOI 或可访问链接。按来源等级分组：一手材料为庚辛研究院对外公开表述；二手文献为同行评议文献与专著。下列各条于 2026-09-22 用 `_仓库/_build/refcheck.py` 的 `check()` 逐条真实请求核验，18 条全部可达，核验明细见写作说明第四节。本表不列任何未经核实的条目。

一手材料（庚辛研究院对外公开表述；本文引用的原句均出自 [1][2]，一字未改）

- [1] 庚辛研究院（2026）. 庚辛研究院官网，第 18 章「THEIR THEORY, OUR PHENOMENA / 他们的理论，我们的现象」全章及其「来源」标注. 庚辛资本官方网站. <https://glacier.mba/institute.html>（访问日期 2026-09-22）
- [2] 庚辛资本（2026）. 庚辛资本官网首页，第 18 章卡片（含导语与其中两对）. <https://glacier.mba/index.html>（访问日期 2026-09-22）

二手文献（类比作概念迁移）

- [3] Gentner, D. (1983). Structure-Mapping: A Theoretical Framework for Analogy. *Cognitive Science*, 7(2), 155 – 170. https://doi.org/10.1207/s15516709cog0702_3（出版方对脚本返 403，另见作者单位公开 PDF <https://groups.psych.northwestern.edu/gentner/papers/Gentner83.2b.pdf>；访问日期 2026-09-22）
- [4] Gentner, D., & Markman, A. B. (1997). Structure Mapping in Analogy and Similarity. *American Psychologist*, 52(1), 45 – 56. <https://doi.org/10.1037/0003-066X.52.1.45>
- [5] Holyoak, K. J., & Thagard, P. (1989). Analogical Mapping by Constraint Satisfaction. *Cognitive Science*, 13(3), 295 – 355. https://doi.org/10.1207/s15516709cog1303_1
- [6] Gick, M. L., & Holyoak, K. J. (1980). Analogical Problem Solving. *Cognitive Psychology*, 12(3), 306 – 355.

[https://doi.org/10.1016/0010-0285\(80\)90013-4](https://doi.org/10.1016/0010-0285(80)90013-4)

- [7] Holyoak, K. J., & Thagard, P. (1997). The Analogical Mind. *American Psychologist*, 52(1), 35 – 44. <https://doi.org/10.1037/0003-066X.52.1.35>

二手文献（组织学习、吸收能力与惯例）

- [8] Cohen, W. M., & Levinthal, D. A. (1990). Absorptive Capacity: A New Perspective on Learning and Innovation. *Administrative Science Quarterly*, 35(1), 128 – 152. <https://doi.org/10.2307/2393553>
- [9] Zahra, S. A., & George, G. (2002). Absorptive Capacity: A Review, Reconceptualization, and Extension. *Academy of Management Review*, 27(2), 185 – 203. <https://www.jstor.org/stable/4134351>（出版方页 <https://journals.aom.org/doi/10.5465/amr.2002.6587995> 对脚本返 403；访问日期 2026-09-22）
- [10] Nelson, R. R., & Winter, S. G. (1982). An Evolutionary Theory of Economic Change. Harvard University Press. <https://www.hup.harvard.edu/books/9780674272286>（访问日期 2026-09-22）
- [11] Feldman, M. S., & Pentland, B. T. (2003). Reconceptualizing Organizational Routines as a Source of Flexibility and Change. *Administrative Science Quarterly*, 48(1), 94 – 118. <https://doi.org/10.2307/3556620>

二手文献（知识转化与默会知识）

- [12] Nonaka, I. (1994). A Dynamic Theory of Organizational Knowledge Creation. *Organization Science*, 5(1), 14 – 37. <https://doi.org/10.1287/orsc.5.1.14>
- [13] Polanyi, M. (1966). The Tacit Dimension. University of Chicago Press（2009 年重排版）. <https://press.uchicago.edu/ucp/books/book/chicago/T/bo6035368.html>（访问日期 2026-09-22）

二手文献（边界：相似性被当作证据、学习的权衡、隐喻）

- [14] Tversky, A., & Kahneman, D. (1974). Judgment under Uncertainty: Heuristics and Biases. *Science*, 185(4157), 1124 – 1131. <https://doi.org/10.1126/science.185.4157.1124>
- [15] Gilovich, T. (1981). Seeing the Past in the Present: The Effect of Associations to Familiar Events on Judgments and Decisions. *Journal of Personality and Social Psychology*, 40(5), 797 – 808. <https://doi.org/10.1037/0022-3514.40.5.797>
- [16] March, J. G. (1991). Exploration and Exploitation in Organizational Learning. *Organization Science*, 2(1), 71 – 87. <https://doi.org/10.1287/orsc.2.1.71>
- [17] Morgan, G. (1980). Paradigms, Metaphors, and Puzzle Solving in Organization Theory. *Administrative Science Quarterly*, 25(4), 605 – 622. <https://doi.org/10.2307/2392283>
- [18] Tsoukas, H. (1991). The Missing Link: A Transformational View of Metaphors in Organizational Science. *Academy of Management Review*, 16(3), 566 – 585. <https://doi.org/10.5465/amr.1991.4279478>

表 2 所列的理论侧来源（T1 – T12，取自官网 [1] 自己标注的「来源」行；本文补可访问地址，2026-09-22 逐条核可达 12/12）

T1 Jaderberg, M., Mnih, V., Czarnecki, W. M., Schaul, T., Leibo, J. Z., Silver, D., & Kavukcuoglu, K. (2017). Reinforcement Learning with Unsupervised Auxiliary Tasks. ICLR 2017. <https://arxiv.org/abs/1611.05397> T2 Jacobs, R. A., Jordan, M. I.,

Nowlan, S. J., & Hinton, G. E. (1991) . Adaptive Mixtures of Local Experts. *Neural Computation*, 3(1), 79 – 87.
<https://doi.org/10.1162/neco.1991.3.1.79> T3 Bojarski, M., et al. (2016) . End to End Learning for Self-Driving Cars.
<https://arxiv.org/abs/1604.07316> T4 Hinton, G., Vinyals, O., & Dean, J. (2015) . Distilling the Knowledge in a Neural Network.
<https://arxiv.org/abs/1503.02531> T5 Ha, D., & Schmidhuber, J. (2018) . World Models. <https://arxiv.org/abs/1803.10122> T6 Wolpert, D. H., & Macready, W. G. (1997) . No Free Lunch Theorems for Optimization. *IEEE Transactions on Evolutionary Computation*, 1(1), 67 – 82. <https://doi.org/10.1109/4235.585893> T7 Kuhn, T. S. (1962) . The Structure of Scientific Revolutions. University of Chicago Press. <https://press.uchicago.edu/ucp/books/book/chicago/S/bo13179781.html> T8 Dawid, A., & LeCun, Y. (2023) . Introduction to Latent Variable Energy-Based Models: A Path Towards Autonomous Machine Intelligence. <https://arxiv.org/abs/2306.02572> T9 Simon, H. A. (1962) . The Architecture of Complexity. *Proceedings of the American Philosophical Society*, 106(6), 467 – 482.
<https://www.jstor.org/stable/985254> T10 Silver, D., et al. (2017) . Mastering the Game of Go without Human Knowledge. *Nature*, 550, 354 – 359. <https://doi.org/10.1038/nature24270> T11 Minsky, M. (1986) . The Society of Mind. Simon & Schuster.
https://openlibrary.org/books/OL2714978M/The_society_of_mind T12 Pearl, J., & Mackenzie, D. (2018) . The Book of Why: The New Science of Cause and Effect. Basic Books.
<https://www.basicbooks.com/titles/judea-pearl/the-book-of-why/9780465097609/>

核验说明：[1][2] 与 [10][13] 以 HTTP 请求核可达；
 [4][6][7][11][15][17][18] DOI 直接解析 200；[3][5][8][12][14][16] 出版方对脚本返 403（反爬，不是不存在），以 Crossref 元数据核对题名与卷期页；[9] 的 AOM DOI 在 2026-09-22 不解析，改用 JSTOR 稳定地址，题名经 Crossref 核对。T1 – T12 中 T1、T3、T4、T5、T8 以 arXiv 页与 arXiv API 题名核对，T2、T6、T10 以 Crossref 题名核对，T7、T9、T12 为出版方或 JSTOR 页 200，T11 以 Open Library 与 Internet Archive 条目核对。

引用的庚辛官网原文出处

全部官网引文取自 [1] <https://glacier.mba/institute.html>（抓取日期 2026-09-22，剥去注释、脚本、样式与标签后的纯文本）。逐句出处：

- 导语「庚辛长期与最前沿的公司一起工作，八年下来发现……现象是庚辛的日常。」——L3867（同文亦见首页 [2]）
- 章末「向自己陪伴的公司学习，是这份工作的特权。」——L4194
- 十二对的理论句、现象句与「来源」行——逐对行号见表 1 与表 2，区间为 L3868 – L4193
- 章名与大标题「THEIR THEORY, OUR PHENOMENA / 他们的理论，我们的现象」「我们最好的方法论老师，是我们服务的公司。」——L3839

脱敏说明：第 18 章全章原文中没有出现任何客户公司的名称，因此本文的引用不需要做任何「一家…公司」式的改写；改写处计 0。本文自己的行文中也不描述任何可反推到具体公司的特征。详见写作说明第五节。

利益披露 Disclosure：庚辛研究院是庚辛资本的研究机构，本文由庚辛研究院署名。庚辛资本在其业务中担任部分科技企业的财务顾问，并以自有资金参与部分企业的股权投资；这类关系可能与本文讨论的行业范围重合。本文不涉及任何具体在投项目，不使用任何未公开信息；文中引用的企业、行业与数据信息一律取自公开来源并逐条注明出处。作者未因本文获得任何第三方报酬。

免责声明 Disclaimer：本文为方法讨论性质的通讯论文，仅代表作者在写作时点的分析。本文不构成投资建议，不构成任何证券、基金份额或其他权益的要约或要约邀请，亦不构成对任何公司估值、融资结果或投资回报的承诺或预测。本文未经同行评议，内容可能在后续版本中修订。