

学科通融 · 之五

选择听谁，本身就是一次判断

论鉴别资质的独立性，及证言依赖在何处失效

美学 × 伦理学 × 知识论

作者 王德生 + Claude

栏目 学科通融 · 站外碰撞

篇幅 约2.4万字 · 23页

发表 2026年7月27日

摘要

在绝大多数事情上靠别人告诉我们，是理性的常态。可有两个领域一直被当作例外：审美判断须基于亲身经验，道德信念不应仅凭他人告知而形成。这两条禁令近年同时遭到强攻，而三家给出的处理互不相容。本文提出三家共同未加追问的那一步：在这个领域里，你凭什么认为那个人靠得住——而这个凭什么，一个不具备该领域判断力的人能不能用？能用的领域，依赖既理性又安全；不能用的领域，选择听谁本身就是那个判断的一次完整行使，依赖并没有省去判断，只是把它移到了一个不被记录、也不被自己察觉的位置。

这一篇由哪三个学科的理论体系撞成

若知识论一侧正确（认知自主根本不是一种德性，应由智性互赖取代），则美学的熟识原则与伦理的道德证言悲观论是把一个误导性理想抬成了规范；若伦理一侧正确（以证言了结一个问题会使人此后更少去寻求理解），则互赖的处方正在系统性地生产不再求解的人；若美学乐观论一侧正确（那条禁令的来源不是认识论而是社会表演，且熟识原则今日已近乎被公认失败），则前两家都把一条表演规范误诊成了认识论规范。三家均为站外的公开文献，链接直达原始出处，可自行核对。

美学 · 审美证言

[《审美证言》斯坦福哲学百科二〇二五年秋季版 —— 熟识原则与自主原则的攻防现状；另见罗布森《审美证言：一种乐观进路》（牛津大学出版社，二〇二二）书评](#)

乐观论一侧的方法论要求最要紧：悲观论者必须交出一个说明，说明是什么使审美与味觉这两个领域独具熟识要求；找不到，那条禁令就只是一组直觉。

<https://plato.stanford.edu/archives/fall2025/entries/aesthetic-testimony/>

伦理学 · 道德证言悲观论

[《道德理解：从德性到知识》（Nous 二〇二五）与关于证言在发展理解中作用的悲观论新论证（Australasian Journal of Philosophy 103\(3\), 二〇二五）](#)

悲观论由静态推进为动态：以证言了结一个问题，会使人此后更少有理由去寻求更广的理解，从而持续地缺乏理解；而道德理解本身系于领会以及情感与动机的投入。

<https://onlinelibrary.wiley.com/doi/full/10.1111/nous.12508>

知识论 · 反自主一支

[《认知自主的哲学》专题导言（Social Epistemology 二〇二四）与其后二〇二五年列维—格林—卡沃尔三方公开往还](#)

列维主张「智性自主」是对一种执行管理德性的误导性命名，应改称智性互赖，并表示一旦把互赖理解妥当就完全不需要再设一个自主的德性；该支还带着「自己做研究」在若干领域可靠地把人带离真理的经验代价。

<https://www.tandfonline.com/doi/full/10.1080/02691728.2024.2335623>

目 录

摘 要.....	4
关键词.....	4
一、问题：一条禁令，三种互不相容的处理.....	4
二、不可还原性校验：与五个既有说明的正面辨别.....	6
三、概念界定与独立性的成因.....	9
四、判据及其操作化.....	10
五、机制：不独立领域中的自我锁死.....	12
六、以此重新解释三家的证据.....	14
七、十二个领域的兑现.....	16
八、人工智能条件下的新处境.....	17
九、正面回应六项反驳.....	18
十、边界、局限与可证伪条件.....	20
十一、结 论.....	22

摘要

一般证言知识论把默认接受他人之言视为理性的常态，而美学与伦理学中各有一条长期存在的禁令与之对立：审美判断须基于亲身经验（熟识原则），道德信念不应仅凭他人告知而形成（道德证言悲观论）。这两条禁令近年同时遭到强攻——审美一侧出现了第一部系统的乐观论专著，并有学者称熟识原则今日已近乎被公认失败；知识论一侧则出现了更彻底的主张，认为「认知自主」根本不是一种德性，应当被「智性互赖」取代。三家的冲突不是侧重之别：若知识论一侧正确，则另两条禁令是把一个误导性理想抬成了规范；若伦理一侧正确，则互赖的处方正在系统性地生产不再寻求理解的人；若美学乐观论一侧正确，则前两者都把一条社会性的表演规范误诊成了认识论规范。

本文提出，三家共同未加追问的，是**鉴别资质与执行资质之间的关系**。本文将某领域中「判断谁在此事上可靠」的能力称为鉴别资质，将「自己做出该领域判断」的能力称为执行资质，并主张：**一个领域能否安全地依赖证言，取决于该领域的鉴别资质是否独立于执行资质**。独立者，外行可凭外部凭据识别专家，依赖因而既理性又安全；不独立者，除以自己的同类判断衡量对方之外别无他法，于是**选择听谁本身就是该判断的一次完整行使**，依赖并未省去判断，只是把它移到了一个不被记录的位置。

本文进一步论证该独立性的成因：外部凭据之所以在事实领域可用，是因为那里存在不预设该领域判断的结果反馈，且此类反馈的读取不需要专业能力；而在审美与道德领域不存在这样的反馈，凭据只能由同类判断者互授，故无法被外行锚定。由此得到四项推论：智性互赖的处方在不独立领域是循环的；依赖在此类领域会自我锁死（不是动机下降而是能力下降）；聚合多人意见不构成解决；以及——审美乐观论所举的描述性证据，全部落在「无下游使用」的一侧，因而不是反例而是本判据所预测的正常情形。文末给出与五个既有说明的辨别、九个领域的兑现、人工智能条件下的新处境、六项反驳的回应与六条可证伪条件。

关键词

鉴别资质；执行资质；独立性判据；证言依赖；熟识原则；道德证言悲观论；智性互赖；专家识别

一、问题：一条禁令，三种互不相容的处理

1.1 一个不对称

在绝大多数事情上，我们靠别人告诉我们而知道。地球的年龄、某种药的剂量、去年的通货膨胀率、这座桥能承受多重——这些信念的保证几乎全部来自他人，而没有人认为这有什么不妥。当代知识论主流的判断是：证言与知觉、记忆并列，是一种基本的知识来源；默认接受一个看起来可靠的证言者，是理性的常态而非退让。

可是有两个领域一直被当作例外。

一个是审美。按照沃尔海姆一九八〇年提出、此后被反复援引的**熟识原则**，审美价值的判断必须基于对其对象的第一手经验，除极窄的限度外不可在人与人之间传递。你不能仅凭朋友说「那幅画很好」就形成「那幅画很好」这个判断——你可以相信她说的是真话，可以据此决定去看，但你不能因此就把那个判断算作你自己的。

另一个是道德。霍普金斯于二〇〇七年把这一立场命名为**道德证言悲观论**：仅凭他人告知而形成道德信念，在某种意义上是有缺陷的。希尔斯给出的最有影响的解释是：证言可以传递足以构成道德知识的保证，却传递不了**道德理解**——那种把握「为什么如此」的状态；而道德理解是若干道德成就（如有德性地行动、行动具有道德价值）的必要条件。

于是出现一个需要解释的不对称：**为什么偏偏这两个领域不能依赖证言？**

1.2 三家的最新进展，与它们的正面冲突

这一不对称在最近数年里同时受到三个方向的推动，而三个方向互不相容。

其一，知识论一侧不但为依赖辩护，而且否认「自主」是一种德性。 在关于认知自主的专题讨论中，列维主张「智性自主」是对某种执行管理德性的误导性命名：鉴于我们的社会处境与深度认知依赖，更合适的名称是**智性互赖**——一种在延迟他人自己推理之间恰当权衡的德性。他进一步表示，一旦把互赖理解妥当，就完全不需要再设一个自主的德性。与之对峙的格林则主张保留自主，把它理解为在「与他人一同思考最易被带偏」的那些场合鼓励独立思考的德性；卡沃尔随后加入，三方在二〇二五年形成了一场公开往还。这一支还带着经验证据：关于「自己做研究」的讨论普遍指出，坚持自行判断在若干领域可靠地把人带离真理。

其二，伦理一侧的悲观论给出了一个更硬的动态机制。 传统悲观论是静态的——延迟于证言时你缺少理解。而近年的推进是动态的：以证言了结一个问题，会使人**此后更少有理由去寻求更广的理解**，从而使他持续地处于缺乏理解的状态。与此同时，关于道德理解本身的研究把它进一步系于道德领会以及主体在情感与动机上的投入。这两步合起来，使悲观论从「你少了一样东西」变成「你少了那样东西，并且因此更不容易再得到它」。

其三，美学一侧的乐观论正在把那条禁令整个拆掉。 罗布森于二〇二二年出版了该题目上的第一部专著，系统地**为审美证言辩护**。其中最要紧的一步是方法论上的：他要求悲观论者交出一个说明——**是什么使审美（以及味觉）这两个领域独具这种熟识要求**；找不到这样的说明，那条禁令就只是一组直觉。他还给出大量描述性证据：我们在失传作品、自然景观之美、经时间检验的审美共识等场合大量依赖审美证言，且这些依赖看起来完全正当。里格尔在一次讲演中说，熟识原则曾经根深蒂固，而**今日几乎已成共识认为它失败了**。在解释那条禁令的来源时，罗布森早期提出的一条路径是信号规范：作出审美断言即在标示自己的智力、技艺与感受力，凭二手证言而断言之错在于虚报自己。阮氏对此提出了一个尖锐的反例——在完全私下、不向任何人标示的场合，那种「不对」的直觉反而更强。

1.3 冲突的形态

把三者并置，冲突的烈度可以精确陈述。

学科	对「能不能靠别人告诉你」的判断	若它成立，另两家如何
知识论 · 反自主一支	能，而且坚持自己想反而更糟；不需要「自主」这个德性	另两条禁令是把误导性理想抬成规范
伦理 · 悲观论新版	能得到信念，得不到理解，且此后更不会去找	互赖的处方在系统性地生产不再求解的人
美学 · 乐观论	能；那条禁令的来源不是认识论而是社会表演	前两家把一条表演规范误诊成了认识论规范

三方各自有可核查的支撑：知识论一侧有社会认知的大量研究与「自己做研究」的经验代价；伦理一侧有一条可陈述的动态机制；美学一侧有一部专著、一批描述性反例，以及一个方法论上的强要求。因此问题不在于谁错，而在于三者共同没有追问什么。

1.4 本文的主张

本文提出，三者共同未加追问的，是一个比「能不能传递」更靠前的问题：

＞ 在这个领域里，你凭什么认为那个人在这件事上靠得住？而这个「凭什么」，是否可以被一个不具备该领域判断力的人使用？

本文将「判断谁在此事上可靠」的能力称为**鉴别资质**，将「自己做出该领域的判断」的能力称为**执行资质**，并主张一条分界：

＞ 一个领域能否安全地依赖证言，取决于该领域的鉴别资质是否独立于执行资质。

在独立的领域，外行可凭外部凭据识别可靠者，依赖既理性又安全；在不独立的领域，除了用自己的同类判断去衡量对方，别无他法——于是**选择听谁，本身已经是那个判断的一次完整行使**。依赖在这里并没有省去判断，它只是把判断移到了一个不被记录、也不被自己察觉的位置。

第二节完成不可还原性校验，第三节界定概念并给出独立性的成因，第四节给出可操作判据，第五节给出自我锁死机制，第六节以此重新解释三家的证据，第七节在九个领域兑现，第八节处理人工智能条件下的新处境，第九节回应六项反驳，第十节给出边界、局限与可证伪条件。

二、不可还原性校验：与五个既有说明的正面辨别

一个新判据若能被既有说明一比一替换，它就没有提供增益。本节与五个最强的邻近说明逐一交锋。

2.1 与希尔斯的「理解不可传」

这是距离最近的一个，也是本文最需要说清的一处。

希尔斯的说明处理的是**接收端**：证言给了你信念与保证，没有给你「为什么」的把握，而后者是若干道德成就的必要条件。本文处理的是**上游**：在你接收之前，你凭什么选中这一位证言者。

两处差别是实质性的。

第一，**解释范围不同**。理解不可传这一条若成立，它同样适用于大量事实领域——医生的诊断迁移能力、工程师对新结构的判断，都要求理解而非仅仅知识，而我们并不因此认为依赖医生与工程师是有缺陷的。希尔斯的说明因此**过度概括**，这也是文献中对它的一项标准批评：它蕴含许多直觉上无可指摘的信念形成方式（如溯因、归谬、直觉）同样令人不安。本文的判据不遇到这个问题：医学与工程属于鉴别可外部化的领域，故依赖在那里是安全的，尽管那里同样要求理解。

第二，**给出的对策不同**。若问题在接收端缺理解，对策是补理解（去学、去想、去要求解释）。若问题在上游的鉴别，补理解不解决问题——因为你仍然要先选中一位值得去理解其理由的人，而那次选择已经用掉了你要外包的那个能力。

2.2 与阮氏的自主原则

自主原则主张：人应当通过运用自己的官能与能力来得出自己的审美判断。它是一条**规范性主张**——你应当如此。

本文的主张是**结构性的**：在鉴别不独立的领域，你**没有别的办法**，这不是应当不应当的问题。这一差别有两个后果。其一，本文不需要为「自主为何有价值」作辩护，因而不受「自主是不是一个好理想」这场争论的牵连——即便列维一侧完全正确、自主不是德性，本文的分界依然成立，因为它讲的不是德性而是可行性。其二，本文能够解释自主原则为何在事实领域不适用，而无须诉诸「那里的价值不同」这类说法：那里的鉴别可以外部化，所以运用自己的官能不是唯一通路。

2.3 与列维的智性互赖

本文不反对互赖。深度依赖确是常态，聪明地延迟确是一种可欲的能力，「自己做研究」确有可观的经验代价。

本文提供的是**它的适用条件**。互赖的核心动作是「延迟得聪明」——对证言者在该领域的胜任度、善意、以及内外群体中支持与反对的比例保持敏感。而这一动作在鉴别不独立的领域是**循环的**：要判断某人在审美或道德上是否胜任，除了用你自己的审美或道德判断去衡量他，没有第二条路。于是互赖的处方在这些领域要求的，正是它试图替你省下的那样东西。

需要强调：这不是对互赖的反驳，而是一条边界。互赖在鉴别可外部化的领域是完全正确的，而那是绝大多数领域。

2.4 与麦格拉思的道德专家难题

麦格拉思指出，我们无法在不预设自己的道德判断的情况下识别道德专家，并以此构成对道德实在论的一个难题：若存在道德事实，为何在这一领域不存在我们能够识别并加以依赖的专家？

这是本文最近的一位盟友，也因此最需要划清。

本文的增量在两处。其一，**它不是道德的特殊难题**。麦格拉思在道德领域观察到的结构，同样出现在审美领域，并且——如第七节所示——出现在一整类领域中。把它当作道德的特异现象，会误导人去它的本体论中寻找原因。其二，**本文给出了该结构的成因**（见 3.3），而成因是关于验收结构的，不是关于对象本性的。这一点削弱了它对道德实在论的威胁：无法识别专家，未必说明没有事实，可能只说明没有不预设该判断的结果反馈。

2.5 与戈德曼的「外行如何评估专家」

戈德曼一九九一年的经典处理，正是试图把鉴别外部化。他列出外行可用的若干判据：论证的表现、其他专家的一致同意、专业履历、利益偏倚、以及**过往的成功记录**。

本文对它的一击是精确的：**这几条判据中，只有最后一条能起锚定作用，而它恰恰在不独立领域不存在。**

- 论证的表现：外行只能评估论证的表面质量（流畅、自信、看似有条理），而这与实质可靠性可以脱钩，这一点在本领域内早有讨论。
- 其他专家的同意：把鉴别问题推给了「谁算专家」，构成回退一步的同一问题。
- 专业履历：其效力最终取决于颁发履历的机构是否可靠，而这又需要锚定。
- 利益偏倚：能剔除若干明显不可靠者，不能在剩下的人中排序。
- 过往的成功记录：**唯一可以由外行直接读取而不预设专业能力的项**——病人活没活、桥塌没塌、预测中没中。

于是本文的判据可以被表述为对戈德曼名单的一次收束：**鉴别的独立性，等价于第五条判据是否可用。**

2.6 与群体智慧／意见聚合

一条常见的实践建议是：不确定时多听几个人的，取其中位或多数。

这一建议的成立要求两个条件：个体判断相互独立，以及存在一条被认可的聚合规则。在不独立领域，第二个条件本身就是问题的重演——**选择哪一条聚合规则、把谁算进这个池子，同样需要鉴别**。把十个人的审美判断平均，与听其中一人的，之间的差别取决于这十个人是怎么被选进来的；而选人的那一步，正是本文所说的那次判断。

2.7 校验结论

五个说明中，两个（希尔斯、阮氏）处理的是接收端或规范面而非上游的可行性；一个（列维）与本文相容而缺一条边界；一个（麦格拉思）观察到了同一结构但把它当作道德的特异现象且未给成因；一个（戈

德曼）正是要外部化鉴别，而本文指出其名单中唯一有锚定力的一项在此类领域缺席。没有任何一个能覆盖「鉴别资质是否独立于执行资质」这一辨别面。

三、概念界定与独立性的成因

3.1 两种资质

执行资质：在某领域自己作出该领域判断的能力。会看病、会算结构、会判断一幅画好不好、会判断一件事该不该做。

鉴别资质：判断某个他人在该领域是否可靠的能力。

二者在概念上分离，这一点常被忽略。一个完全不懂心脏外科的人，可以相当可靠地判断某位心脏外科医生是否值得信任——他凭执照、专科认证、所在机构、同行转诊、并发症率。他的鉴别资质高，执行资质为零。

这种分离不是偶然的便利，而是分工得以可能的条件。若鉴别必须以执行为前提，则一切专业分工都将无从建立：你要找一位可靠的电工，就得先会电工。

3.2 独立与不独立

据此可作一条分界。

鉴别独立的领域：存在一批外部凭据，使不具备执行资质者也能可靠地排序候选证言者。医学、工程、航空、税法、气象预报属于此类。

鉴别不独立的领域：不存在这样的凭据；判断某人在此事上是否可靠，只能以自己的同类判断去衡量他。审美判断、道德判断属于此类，并且——第七节将表明——不止于此。

在不独立的领域，「依赖证言」这一描述是误导的。你并没有把判断交出去，你做了一次判断：**你判断了这个人在这件事上值得听**。而这次判断动用的正是你在该领域的执行资质。因此：

> 在鉴别不独立的领域，选择听谁本身就是那个判断的一次完整行使。依赖没有省去判断，只是把它移到了一个不被记录的位置。

3.3 独立性从何而来

这是本文的关键一步，也是对罗布森那个方法论要求的正面回答——他要一个说明，说明是什么使某些领域独具熟识要求。本文的回答是：

> 外部凭据之所以在某些领域可用，是因为那里存在**不预设该领域判断的结果反馈**，且此类反馈的读取**不需要该领域的执行资质**。

拆开来看有两个条件，缺一不可。

条件一：存在独立于该判断的结果。 一台手术之后病人活没活，这件事的成立不依赖于任何人对该手术的医学评价；一座桥塌没塌，不依赖于任何人对该设计的工程学评价。

条件二：该结果的读取不需要专业能力。 病人是否活着、桥是否塌了、飞机是否落地、预报的雨是否下了——这些外行都读得出。正是这一条使凭据可以被外行**锚定**：执照制度、专科认证、机构声誉，其可靠性最终由这些外行也能读取的结果所校准。

现在看审美与道德。

审美：一件作品「好不好」没有独立于审美判断的验收。销量、奖项、经时间的存留都是**同类判断的聚合**，而不是外部结果——它们由具备（或自称具备）审美判断的人产出。因此条件一不成立。

道德：一个决定「对不对」同样没有独立于道德判断的验收。后果可以被观察（有人受伤、有人得益），但从后果到「这个决定是对的」这一步，本身就是一次道德判断。因此条件一同样不成立。

于是熟识原则与道德证言悲观论所捕捉的那件事，可以被推导出来，而不必被当作一组原始直觉：**在这两个领域中，鉴别不可外部化；而在鉴别不可外部化的领域，「靠别人告诉你」这一操作在结构上不能省去你自己的判断。**

3.4 三点澄清

其一，这不是怀疑论。 本文不主张在这些领域里无法获得知识，也不主张他人的意见没有价值。他人的审美与道德意见极有价值——但其价值的兑现方式不是替代你的判断，而是给你一个去看、去想的方向。这一点与本文所属栏目此前几篇的立场一致，此处不展开。

其二，独立性是程度问题。 少数领域完全独立（结果硬碰硬且外行可读），少数完全不独立，多数居中。第七节给出一条居中的谱系。

其三，不独立不等于不可学。 执行资质可以增长，而它增长的同时鉴别资质随之增长——因为在这类领域两者是同一东西。这正是第五节自我锁死机制的另一面：它同样意味着**唯一的出路是提高执行资质**，而这条出路是通的，只是不能被外包。

四、判据及其操作化

4.1 一问

把第三节收成一个可以当场提出的问题：

> 我凭什么认为这个人在这件事上靠得住——而这个理由，一个完全不具备该领域判断力的人能不能用？

能用 → 鉴别独立，依赖安全。 不能用（那个理由本身要求你已经会判断这件事） → 鉴别不独立，你正在做的不是外包判断，而是行使判断。

4.2 三层答案

实际使用时，答案通常落在三层之一，而三层的处置完全不同。

第一层：可锚定的外部凭据。 「他有执照，那家医院的并发症率公开可查，出了事有人被吊照。」——这个理由完全不依赖我懂不懂医。鉴别独立。

第二层：同类判断的聚合。 「他在这个圈子里公认有品味」「大家都说他判断力好」。——这个理由把鉴别推给了一群人，而那群人是谁、他们凭什么，仍然需要我判断。鉴别不独立，尽管它看起来像第一层。这一层最危险，因为它在形式上模仿了第一层。

第三层：我自己的同类判断。 「我读过他写的东西，他说的那些我自己核对过，对得上。」——这个理由诚实地承认了它动用了我的执行资质。鉴别不独立，但这是**正确的形态**：它至少没有伪装。

第二层与第三层的分别值得强调：两者都不独立，但第三层是在行使判断，第二层是在**误以为自己没有行使判断**。后者的代价见第五节。

4.3 一个更快的粗筛

三层要一个个想，有一个更快的办法：

> **这个领域里，有没有人因为判断错了而承担过可被外人读出的后果？**

有——且这类事件在该领域是常规的、可查的 → 大概率独立。 没有，或者「判断错了」这件事本身只能由同行认定 → 不独立。

这条粗筛之所以比一问更快，是因为它不问你的理由，只问已经发生过什么。而人对自己所处领域的独立性判断，几乎总比实际情况乐观一格。

4.5 三个实例走一遍

判据的价值全在能否当场用。此处以三个日常场合各走一遍。

场合一：选一位牙医。「我凭什么认为她靠得住？」——她有执照，这家诊所有可查的投诉记录，做坏了会被追责，而且我自己的牙好没好，几个月内我就知道。这几条理由，一个完全不懂口腔医学的人全都能用。**鉴别独立，放心依赖。**

场合二：选一位理财顾问。「我凭什么？」——他有资格证，他所在的机构有名。但这两条的锚在哪里？他的过往业绩要放在多长的周期上读、如何与市场基准分离、他所承担的风险有多大——这些的读取本身需要专业能力。**条件一成立而条件二不成立**，因而其独立性远低于表面。这解释了一个长期存在的经验事实：该领域凭据齐全而事故不断，且事故往往在数年后才显形。正确的处置不是放弃依赖，而是承认自己

的鉴别在这里是弱的，并据此限制单次委托的规模——这是本文框架给出的、与「找一个更权威的机构」方向不同的建议。

场合三：听一位前辈说「你该往那个方向做」。「我凭什么认为他在这件事上靠得住？」——他做出过成绩。可他做出成绩的是**另一个时代的另一个方向**，而这条建议要在十年后才被检验，届时检验它的人是他这一代的同行。我能给出的唯一诚实理由是：我读过他判断过的一些事，其中有几件我自己后来也想明白了，对得上。**这条理由不懂行的人用不了。鉴别不独立**——而这意味着，我采纳这条建议这一步，已经是我自己的一次方向判断，不是一次省力。

三个场合的形式完全相同（找一位更懂的人问），而按判据分属三格。这正是本文认为该判据有用的理由：**形式相同的动作，其性质并不相同，而形式是我们默认用来分类的东西。**

4.4 两个必须防的误用

其一，不要用它去否定专业分工。 本文的判据在绝大多数领域给出的答案是「独立」，即依赖安全。它的适用面是窄的，不是宽的。

其二，不要用它为拒绝倾听作辩护。 判据说的是「你选择听谁这一步动用了你自己的判断」，不是「不要听」。恰恰相反：既然那一步已经动用了你的判断，那么听得越多、看得越具体，那一步就做得越好。判据反对的是**以为自己没在判断**，不是判断本身。

五、机制：不独立领域中的自我锁死

5.1 与伦理一侧动态机制的关系

第一节提到，悲观论近年的推进是动态的：以证言了结一个问题，会使人此后更有理由去寻求更广的理解。那是一条**动机**层面的机制。

本文给出的是一条**能力**层面的机制，两者独立且叠加。

5.2 三步

第一步。 在不独立领域，我依赖某人的判断。这一次依赖是有代价的，但代价不在当期显现——我得到了一个可用的结论。

第二步。 由于在这类领域鉴别资质与执行资质是同一样东西，我每一次以「听他的」代替「自己看」，就少了一次运用该资质的机会。执行资质随之停滞，鉴别资质同步停滞。

第三步。 鉴别资质停滞，使我下一次更难判断该听谁——而这恰恰增加了我依赖的必要，因为我更没有把握自己看。于是依赖加深。

三步构成一个闭环，且**每一步都局部合理**：第一步是明智的省力，第二步不产生任何可见的损失，第三步是对自己能力不足的诚实反应。

5.3 为什么它不被察觉

这一机制的隐蔽性有一个具体的来源，它同时解释了 4.2 第二层为何最危险。

在依赖加深的过程中，**我的判断的外在表现在改善而不是恶化**：我引用的是更可靠的来源，我的意见与公认的意见更加一致，我犯明显错误的次数下降。按任何以「与公认意见的偏离度」为指标的评估，我在进步。

而唯一在退化的那样东西——「在一个尚无公认意见的新情形里，我自己能不能判断」——**只在遭遇新情形时才显形**，而那时通常已经晚了。

5.4 与「过度概括」批评的关系

对悲观论的一项标准批评是过度概括：若依赖证言令人不安，则许多无可指摘的信念形成方式（溯因、归谬、直觉）同样应当令人不安。

本文的机制不遇到这一批评，因为它并不主张依赖本身令人不安。它主张的是一个条件句：**在鉴别不独立的领域，依赖会同时消耗鉴别的能力**。在鉴别独立的领域，同样的依赖没有这个后果——我依赖医生一辈子，我鉴别医生的能力（读执照、读并发症率、听转诊）丝毫不受影响，因为那个能力从一开始就不是医学能力。

这正是本文判据的价值所在：它把一条模糊的不安，收束成一个有明确适用条件的机制。

5.5 一个可能的出路，及其代价

这一机制并非无解，而它的解法形状值得点出，因为它与常见建议不同。

由于在不独立领域鉴别资质与执行资质是同一样东西，**唯一有效的干预是运用它**——而不是提高对它的认识、不是学更多的判断标准、不是找更好的老师。这三样都属于「先得判断该学什么、该信哪套标准、该跟谁」，因而重演同一个循环。

有效的动作只有一个形状：**在有下游使用的场合，先自己作出判断并记下来，再去看别人怎么说**。顺序是关键。先看别人的，你自己那一次判断就不会发生，而不发生的判断不产生任何资质。

其代价是实在的：这样做的当期产出更差，且这一更差是可测量的，而其收益不可测量且延迟。这使它在任何以当期产出计量的安排中都处于劣势。本文对此没有对策——第八节末对同一困难的表述是一样的，因为它本来就是同一个困难。

六、以此重新解释三家的证据

一个判据的力量，不在于它另起炉灶，而在于它能否把三家已有的证据重新安置。本节逐一处理。

6.1 知识论一侧：「自己做研究」的经验代价

这一支最有力的经验支撑是：坚持自行判断在若干领域可靠地把人带离真理，阴谋论与各类抗拒专业意见的现象是标准例证。

按本文的判据，这些例证**全部落在鉴别独立的领域**：疫苗是否安全有效、气候是否在变暖、某种疗法是否有效——这些领域都存在不预设该领域判断的结果反馈，且外行可读（人有没有病、冰有没有化、病人有没有好）。在这些领域坚持自主，确实是把一个可外部化的鉴别硬做成执行，而且做不好。

因此这些证据并不支持「自主一般地是坏理想」，它们支持一个更窄的结论：在鉴别可外部化的领域，坚持自己判断是低效且危险的。而这正是本文所主张的。

还有一层值得点出：抗拒专业意见者的自我描述往往正是「我自己做了研究」——也就是说，**自主是一件可以被表演的事**。表演的自主与实际的执行资质在外观上不易区分，而这一点使得以「他们坚持自主」为由来反对自主的论证需要额外小心。

6.2 伦理一侧：延迟的自我巩固

如 5.1 所述，本文的能力机制与该动态机制独立且叠加。此处补充一点两者的分工：

- 动机机制解释了**为什么人不去补**（既然问题已了结，寻求更广理解的理由就减少了）。
- 能力机制解释了**为什么即使去补也更难补**（用于判断该向谁学、哪种解释是好解释的那个资质，本身已经停滞）。

两者相加，解释了一个常被观察到的现象：在这类领域，长期依赖者一旦决心自己判断，往往不是从零开始，而是从**一个被扭曲的起点**开始——他会把「更符合他所依赖过的那个来源」误当作「更正确」。

6.3 美学一侧：乐观论的描述性证据

这是本文最要紧的一次重新安置。

罗布森举出的描述性证据是有力的：我们确实在若干场合大量依赖审美证言而不觉有何不妥——已经失传、无从亲见的作品；未曾到过的自然景观之美；经时间检验的审美共识。

本文的主张是：**这些例子有一个未被注意的共同点——它们全部落在「没有下游使用」的一侧。**

- 失传的作品：我永远不会亲见它，因而这个判断不会被我用于任何后续的观看、比较或选择。
- 未到过的景观：同上，直到我真的去。
- 经时间检验的共识（如某位大师的卓越）：这类判断在我的实际审美生活中几乎不承担任何工作——它不决定我今晚看什么、不影响我对眼前这幅画的反应。

而在有下游使用的场合——我要挑一件作品带回家、我要决定把三年时间投进哪一种手艺、我要判断眼前这位学生的作品是否值得鼓励——**同样的依赖立刻显出问题**，而这正是悲观论直觉最强的那些场合。

于是乐观论的证据不是本文的反例，而是**本文所预测的正常情形**：在鉴别不独立的领域，当一个判断没有下游使用时，鉴别失败不产生代价，因而看不出问题。

这一重新安置有一个可检验的后果，见第十节第三条。

6.4 私下情形与阮氏的反例

阮氏对信号解释的反例是：在完全私下、不向任何人标示的场合，那种「不对」的直觉反而更强。信号解释无法处理这一点，因为私下场合没有信号。

本文的判据可以处理它，而且给出的解释相当直接：**私下场合正是下游使用最纯粹的场合**。我私下依据别人的判断决定挂哪幅画、听哪张唱片，这个判断将持续作用于我此后的全部感受与选择——它有最长的下游。而公开断言反倒常常是一次性的社交动作，其下游可能为零。

按 6.3 的框架，下游越长，鉴别失败的代价越大，「不对」的直觉因而越强。这与阮氏观察到的方向一致。

6.5 三家为何各自到不了这一条

把三家的处境并排，可以看清这条分界为何没有被它们中的任何一家提出。

知识论一侧到不了，因为它的问题意识是**辩护性**的：面对广泛的抗拒专业意见的现象，它要论证依赖是理性的。而它援引的全部案例恰好落在鉴别可外部化的一侧——这不是选择偏差，是那些案例正是当时的公共问题所在。在一个全部样本都独立的样本集里，独立性不会成为变量。

伦理一侧到不了，因为它的问题意识是**接收端**的：既然可以传递知识，为什么还不安？它因此把全部注意力放在「传过来的是什么、没传过来的是什么」上。上游那一步——你怎么选中了这位证言者——在它的问题构造里根本不出现，因为在思想实验中那位证言者总是被规定为「道德上可靠的人」。**规定掉的东西不会成为变量**，而本文的整个论点正在被规定掉的那一步上。

美学一侧到不了，因为它的任务是**拆除**：它要证明那条禁令没有独立支撑。而拆除者的方法是举反例，反例越多越好。它因此不会去问「这些反例有没有共同点」——一旦去问，就等于承认那条禁令在某个范围内成立，而那正是它要否定的。6.3 所做的，正是它没有动机去做的那一步。

三家合起来才逼出这条分界：一家给了「可外部化的凭据长什么样」，一家给了「不安在哪里」，一家给了「必须交出一个说明」这个方法论要求。三样一凑，那个被共同规定掉的变量就浮上来了。

七、十二个领域的兑现

一条判据若只在提出它的领域内成立，它就只是那个领域的经验。本节在十二处兑现，并按独立性排出一条谱系。

7.1 高度独立的一端

急诊医学。 结果硬碰硬且外行可读（人活没活）。执照与追责制度由这些结果锚定。依赖完全安全，且坚持自己判断是危险的。

结构工程。 同上。桥塌与不塌不由工程学评价决定。

天气预报。 预测在短期内被现实裁决，且裁决外行可读。这使得预报机构的可信度可以被完全不懂气象的人排序。

7.2 中段

税务与法律实务。 部分独立：申报是否被驳回、案子是否胜诉，是外行可读的结果；但「这个方案是否稳妥」在很长时间内没有结果反馈，且很多风险终身不显形。故其独立性随事项的结果显形速度而变。

投资。 表面上极度独立（赚没赚钱是硬结果），实际上受噪声与周期长度的严重干扰，短期结果几乎不携带信息。这是一个**外部结果存在但读取需要专业能力**的典型——条件一成立而条件二不成立，故独立性远低于表面。

临床判断中的疑难部分。 常规诊疗高度独立；而在结果延迟数年、且多因素交织的部分（慢性病管理、生活方式建议），结果反馈的读取开始需要专业能力，独立性随之下降。

7.3 低度独立的一端

审美判断。 如 3.3 所述。

道德判断。 如 3.3 所述。补一点：道德领域存在一个特殊的复杂性——某些道德判断有外行可读的后果（政策造成的死亡人数），但从后果到「这是错的」这一步仍是道德判断，故条件一仍不成立。

「什么问题值得做」这一类判断。 学术方向的选择、研究问题的取舍、一个组织该往哪走。这一类的独立性极低，因为其结果反馈的周期长于职业生涯，且「这个方向是对的」的认定由同行作出。本文认为这是被严重低估的一格——它在形式上被当作专业判断（因而可依赖专家），而在结构上属于不独立领域。

同行评审。 属混合型，且其两部分的独立性差距极大。「数据是否造假、方法是否有误」这一部分接近独立——存在可核查的事实，且核查不必预设对该方向价值的判断。而「这个问题是否重要、这个方向是否有前途」这一部分属于最低度独立的一格：它的结果反馈周期长于评审者的职业生涯，且其认定完全由同行作出。同一份评审意见里并置着独立性相差一个数量级的两类判断，而制度以同一套程序处理它们。

人才选拔。「能不能做这件具体的事」可以通过当场做一件小的来锚定，接近独立；「这个人将来能走多远」则完全不独立——它没有可读的近期结果，且其认定由已被选拔出来的人作出。选拔制度的可靠性因而随其权重在两者之间的分配而变。

育儿与教育中的方向判断。「这个孩子该往哪个方向走」是本文谱系中独立性最低的一格之一：结果显形以十年计，读取需要判断，且几乎所有可用的凭据（升学、评奖）都是同类判断的聚合。这一格之所以值得单列，是因为它是**最多人每天都在其中依赖他人证言、而几乎无人意识到其鉴别不可外部化**的一格。

7.4 谱系的用法

把十二处并排，可得一条实践上的读法：**独立性随「结果显形速度」与「结果可被外行读取的程度」两者共同下降**。前者决定凭据有没有机会被校准，后者决定校准能否不预设专业能力。

这也给出一个警示：一个领域可以在自己不知不觉间从独立滑向不独立——例如当其结果的显形周期被拉长，或当其验收被改为由同行评定。这是一种在制度改革中相当常见、而几乎从不被记入的变化。

八、人工智能条件下的新处境

8.1 为什么这一分界现在才要紧

前述分界在原理上一直成立，但在实践中长期不构成日常问题，原因是：在不独立领域获取高质量证言的成本很高。要得到一位有品味的人对你手上这幅画的判断，你得认识他、约到他、让他愿意认真看。成本本身限制了依赖的频次。

这一成本刚刚接近于零。而且——这是要紧的一点——**新的证言来源在两类领域中都能产出高质量输出**。

8.2 两类领域中的不同后果

在鉴别独立的领域，机器输出的可靠性仍可被外部锚定：代码跑不跑、诊断对不对、预测中没中。使用者不需要具备执行资质就能校准自己的信任程度。这是好消息，且是这项技术最扎实的收益。

在鉴别不独立的领域，情况相反。机器给出的审美判断、价值权衡、方向建议，其可靠性**无法被不具备该项执行资质的人评估**；而它恰恰最令人满意，因为它按使用者已有的偏好优化。

于是得到一个可检验的预测：

> 在鉴别不独立的领域，机器辅助会先提高使用者的满意度与产出的表面质量，随后降低其执行资质，而使用者**无法从自己的体验中察觉这一下降**——因为按 5.3，退化的一项只在新情形中显形。

8.3 一个不对称的建议

由此得到一条与常见建议方向不同的操作原则：

- 在鉴别独立的领域，**尽量依赖**，并把力气花在维护外部凭据（可核查的成功记录、追责链条）上。
- 在鉴别不独立的领域，**依赖之前先自己做一遍**——不是为了产出更好，产出多半更差；而是为了保住那个用来判断该不该采纳的能力。

第二条的成本是实在的，而它的收益不在当期。这使它在任何以当期产出计量的安排中都处于劣势——这一点本文没有对策，只能指出。

8.4 一个可以现在就做的检验

上述预测不需要等待多年才能被检验。一项低成本的设计如下。

在某个鉴别不独立的领域（如某类作品的评判）中，取两组水平相近的从业者。实验组在作出自己的判断之前先获取机器的判断；对照组先自己判断，再看机器的判断，并记录两者的差异。两组的产出质量在短期内很可能实验组更高。

关键的观测项不是产出，而是在**一批刻意构造的新情形上的表现**——那些机器与既有共识都没有现成答案的情形。若 8.2 的预测成立，则实验组在这一项上的表现应当随时间相对下降，而两组在**自评的把握程度**上不应出现相应差异。

「产出更好而自评不变，唯独在新情形上落后」这一组合，是本文机制的特征指纹；若观察不到它，则 8.2 落空。

九、正面回应六项反驳

9.1 「这只是把熟识原则换个说法」

不是。熟识原则说的是审美判断必须基于第一手经验，它是一条关于**信念来源**的规范，且被批评为无法解释大量正当依赖的案例。

本文说的是：在某些领域，**鉴别不可外部化**，因而选择听谁已经用掉了执行资质。这是一条关于**可行性结构**的描述，不是规范；它不禁止依赖，它指出依赖在这些领域没有省去它看起来省去的東西。两者的可检验后果也不同：熟识原则预测一切二手审美判断都有缺陷，本文预测**只有有下游使用的那些**才显出缺陷——而 6.3 表明后者与观察更相符。

9.2 「审美与道德领域也有专家，我们也确实能识别他们」

我们确实能，而且本文并不否认。问题在于**凭什么**。

细看任何一次「我认得出谁有品味」的实际过程，其依据都落在 4.2 的第三层：我读过他说的，我自己核对过，对得上。这恰恰是以自己的执行资质去衡量。它有效——但它不是外部化，它是行使。

而落在第二层（圈子公认）的那些识别，其可靠性完全取决于那个圈子，而评估那个圈子仍需同一种能力。历史上审美与道德共识的大规模转向（对某些艺术门类的重估、对某些制度的道德重判）正是这一层不可锚定的证据。

9.3 「结果反馈在道德领域是存在的：看后果就行」

这是最强的一项反驳，需要认真处理。

后果确实可以被观察，而且往往外行可读（有多少人受害、有多少人得益）。但从「发生了这些后果」到「那个决定是错的」这一步，需要一次道德判断——需要决定哪些后果算数、如何在受损者与受益者之间权衡、是否存在与后果无关的约束。而这一步正是我们要外包的那一步。

这与医学的对照很清楚：从「病人死了」到「这次手术的判断是错的」这一步，可以由医学专业内部以外行也能核查的方式作出（并发症率、同类手术的基线、尸检）。而从「这项政策造成了这些损失」到「这项决定是错的」，没有对应的可核查步骤——不同的道德立场会从同一组后果读出相反的结论。

因此条件一在道德领域确实不成立，尽管后果可见。

9.4 「那么这一切是否只是说：价值判断不能外包？」

不是，而且这一简化会丢掉本文最有用的部分。

「价值不能外包」是一条老话，它的问题是给不出边界，因而在实践中要么被无限扩张（凡涉及价值皆不可依赖，于是寸步难行），要么被当作套话。

本文给的是一条可操作的分界，而这条分界不与「是不是价值判断」重合。7.3 给出的第三格——「什么问题值得做」——在形式上被当作专业判断，很少被当作价值判断，而它按本文的判据属于低度独立。反过来，某些明显带价值色彩的判断（如某项安全标准的取舍）由于其结果显形快且外行可读，独立性相当高。判据看的是验收结构，不是内容的性质。

9.5 「若本文成立，我们凭什么相信本文」

这是本文自反的入口，也是唯一一项本文无法完全化解的反驳。

「哪些领域的鉴别可外部化」这个元问题，其自身的鉴别是否可外部化？部分可以：本文的判据部分依赖可核查的事实（某领域有没有外行可读的结果反馈、有没有人因判断错误而承担过可被外人读出的后果），这些是第三方可查的。但整体不完全可以——判断本文这条分界是否切中要害，仍要用到读者自己的判断力。

因此本文不能靠权威被接受。它只能靠**使用**被检验：拿第四节那一问去问你手上正在依赖的三件事，看它是否分出了你原本分不出的东西。分出来了，它对你成立；分不出，那它对你就是一句空话，无论其论证多么严密。

9.6 「凭据可以被制造出来：为什么不给审美与道德也建一套」

这是一项建设性的反驳，值得认真对待，因为它指向本文框架的一个真实的行动方向。

若鉴别的独立性来自可锚定的凭据，而凭据的可锚定性来自外行可读的结果反馈，那么原则上可以为任何领域建造凭据——只要能找到一种外行可读的结果。历史上确有先例：医学的执照与并发症统计不是自然存在的，它们是被造出来的；天气预报的可信度评级同样如此。

本文的回答分两层。

第一层：确有可造的部分，而且值得造。 在审美与道德领域中，存在若干可被外行读取的次级结果：一位评审此前推荐过的作品，后来被更广的时间检验如何；一位顾问此前给出的建议，其当事人后来是否需要推翻。把这类记录**变成可查的**，确实能提高该领域的独立性。本文第七节所说的「独立性随结果显形速度与可读程度而变」，正蕴含这一条：改变这两个量就改变了独立性。

第二层：但这类凭据的锚仍是同类判断。 「后来被时间检验如何」这句话里的「检验」，仍由具备该领域判断的人作出。这与并发症率不同——病人是否死亡不由医学判断决定。因此这类凭据能把独立性从极低提到中等，**不能提到医学那一档**。

这一分层给出一条实践建议，也给出一项限度：值得为不独立领域建造可查记录，但不应期待它能把该领域变成可安全外包的领域；而在建造过程中最要防的，是第二层被当成第一层来用——即把同类判断的聚合当作外部锚定。按 4.2，那正是三层中最危险的一层。

十、边界、局限与可证伪条件

10.1 三条不适用边界

其一，紧急与高危场合按独立领域处置。 即便某项判断带有价值成分，若时间不允许且已有成熟凭据，正确的动作是依赖。在此处引入本文框架，会为拖延提供一个听起来有洞察力的理由。

其二，尚不具备任何执行资质者应当大量依赖。 本文的机制以「存在可被消耗的执行资质」为前提。一个刚入门者没有可消耗的东西，他的正确路径是大量吸收他人的判断。分界应在他开始对某些判断产生自己的不同意见之后才启用。

其三，本文不为拒绝倾听背书。 见 4.4。

10.2 四项局限

局限一：独立性是连续谱，本文以三分呈现。 第七节的谱系已表明多数领域居中，而本文的判据在中段给出的指示较弱。

局限二：本文未处理类型的迁移。 一个领域可从独立滑向不独立（7.4 末），反之亦可（当某领域建立起可核查的结果记录时）。这一迁移过程本文未着一字，而它可能是最具政策含义的部分。

局限三：「下游使用」这一概念未被充分刻画。 6.3 依赖它作出重新安置，而本文只给了例示，没有给出判定某个判断有无下游使用的独立标准。

局限四：本文只讨论了个体依赖，未讨论制度依赖。 一个组织如何在不独立领域组织其判断（评审、评奖、方向决策），是本文框架的自然延伸，而本文未展开。

10.3 六条可证伪条件

- **其一，一问不可操作。** 若若干具专业背景者分别用第四节那一问对同一批领域作独立／不独立分类，结果无法达成一致，则该判据不可用，本文的实践主张随之落空。此项成本极低。
- **其二，成因说明不成立。** 若能举出一个鉴别可外部化、而其中不存在「不预设该领域判断且外行可读的结果反馈」的领域，则 3.3 的成因说明错误。
- **其三，下游使用与直觉强度无关。** 若在受控设计中，对有下游使用与无下游使用的审美依赖，人们的「不对」直觉强度没有系统差异，则 6.3 的重新安置失败，罗布森的描述性证据重新构成对悲观论的独立反驳。此项可用问卷设计直接检验。
- **其四，鉴别与执行在不独立领域可分离。** 若能展示一种方法，使不具备审美或道德执行资质者可靠地排序该领域的证言者，则本文的核心分界不成立。
- **其五，自我锁死不发生。** 若在不独立领域长期依赖他人判断的人群，其执行资质未见相对下降，则第五节的机制错误。
- **其六，人工智能条件下的预测落空。** 若在鉴别不独立的领域，长期使用机器辅助者的执行资质未下降，或其下降可被使用者自行察觉，则 8.2 的预测错误。

10.4 本文不主张什么

本文不主张少依赖他人，不主张自主是一种德性，不主张审美与道德判断无法达成知识。它只主张一件事：一个领域能否安全地依赖证言，取决于该领域的鉴别资质是否独立于执行资质；而在不独立的领域，选择听谁本身已经是那个判断的一次完整行使。其余全部内容皆由此推出。

10.5 被本文的整理动作扫掉的部分

按本文自身的立场，一次整洁的整理会腾出解释的位置，故此处原样列出被扫掉者：

- 三家的内部分歧远比本文的呈现复杂。知识论一侧并非只有反自主一支，格林与卡沃尔的辩护是活的且有力的；伦理一侧的乐观论者（斯利瓦、博伊德等）有系统的回应，本文只在 5.4 处理了其中

一条；美学一侧对熟识原则的辩护近年仍在出版，里格尔所说的「近乎共识」是一位与会者的表述而非计量结果。

- 本文对戈德曼名单的处理（2.5）压缩得相当厉害，其中「论证的表现」一条在文献中有比本文所述精细得多的讨论。
- 本文所举的三个「审美乐观论例子皆无下游使用」的案例是本文自己的归类，而罗布森本人并未以此方式呈现它们；这一归类是否穷尽了他的例子，本文未作检验。
- 本文没有处理**集体层面的鉴别**。一个共同体如何在独立领域维持其判断质量（学派的传承、行会的师承、评审团的构成），是一个与个体层面机制不同的问题，本文的框架对它只给了一个方向而未展开（见 9.6 第一层）。

十一、结论

本文的结论可收为三句。

其一，证言依赖的分界不在内容的种类，而在验收的结构。 审美与道德之所以被当作例外，不是因为它们涉及价值，而是因为在这两个领域中不存在不预设该领域判断、且外行可读的结果反馈；缺了这一条，凭据无法被锚定，鉴别便无法外部化。

其二，在鉴别不独立的领域，依赖并未省去判断。 选择听谁本身就是那个判断的一次完整行使。真正的风险因而不是「依赖」，而是**以为自己没有在判断**——这正是同类判断的聚合（「大家都说他好」）之所以是三层中最危险一层的原因。

其三，这一分界现在才成为日常问题。 在不独立领域获取高质量证言的成本刚刚接近于零，而在该领域中，输出的可靠性无法被不具备该能力者评估，且它按使用者的偏好优化因而最令人满意。由此得到的操作原则是不对称的：在鉴别独立的领域尽量依赖，并把力气用于维护可核查的结果记录；在鉴别不独立的领域，依赖之前先自己做一遍——不是为了产出更好，而是为了保住那个用来判断该不该采纳的能力。

最后需要指出本文最易被扭曲的一处：这条判据的正当用法是施于自己正在依赖的事，而非用以质疑他人的依赖。「你凭什么认为他靠得住」这一问可以使任何一次引用停摆，而提问者不承担任何代价。其有效的用法只有一种：对着自己手上正在依赖的三件事各问一遍，看其中有几件，你给出的理由是一个不懂行的人也能用的。

参考文献

- Belkonienė, M. (2025). Moral understanding: From virtue to knowledge. *Noûs*. <https://onlinelibrary.wiley.com/doi/full/10.1111/nous.12508>

- Goldman, A. I. (2001). Experts: Which ones should you trust? *Philosophy and Phenomenological Research*, 63(1), 85–110.
- Green, A. (2025). Epistemic authenticity, skilled interactions, and strategic parallel processing: A dialogue with Levy and Kawall. *Social Epistemology Review and Reply Collective*. <https://social-epistemology.com/2025/09/15/epistemic-authenticity-skilled-interactions-and-strategic-parallel-processing-a-dialogue-with-levy-and-kawall-adam-green/>
- Hills, A. (2009). Moral testimony and moral epistemology. *Ethics*, 120(1), 94–127.
- Hopkins, R. (2007). What is wrong with moral testimony? *Philosophy and Phenomenological Research*, 74(3), 611–634.
- Kawall, J. (2025). Finding work for epistemic autonomy: A reply to Green and to Levy. *Social Epistemology Review and Reply Collective*. <https://social-epistemology.com/2025/05/07/finding-work-for-epistemic-autonomy-a-reply-to-green-and-to-levy-jason-kawall/>
- Levy, N. (2024). Against intellectual autonomy: Social animals need social virtues. In *The Philosophy of Epistemic Autonomy* (special issue). *Social Epistemology*. <https://www.tandfonline.com/doi/full/10.1080/02691728.2024.2335623>
- McGrath, S. (2011). Skepticism about moral expertise as a puzzle for moral realism. *The Journal of Philosophy*, 108(3), 111–137.
- Nguyen, C. T. (2020). Autonomy and aesthetic engagement. *Mind*, 129(516), 1127–1156.
- Robson, J. (2022). *Aesthetic Testimony: An Optimistic Approach*. Oxford University Press. 评论见 <https://ndpr.nd.edu/reviews/aesthetic-testimony-an-optimistic-approach/>
- Robson, J., & Wallbank, R. (2025). Aesthetic testimony. In *The Stanford Encyclopedia of Philosophy* (Fall 2025 ed.). <https://plato.stanford.edu/archives/fall2025/entries/aesthetic-testimony/>
- Wollheim, R. (1980). *Art and Its Objects* (2nd ed.). Cambridge University Press.
- 关于证言在发展理解中作用的悲观论新论证，见 *Australasian Journal of Philosophy*, 103(3), 2025. <https://www.tandfonline.com/doi/full/10.1080/00048402.2024.2421277>

本文由三个分属不同学科、且互相冲突的理论体系碰撞而成。全部来源均为站外公开文献，可自行核对。 作者：王德生 + Claude · SDE Universes · 学科通融