

人工智能报告： 从买家视角看AI定价

2026



目录

01	执行摘要 塑造当今买家体验人工智能定价方式的头条模式。	第03-05页
02	人口统计数据。 受访者概况——覆盖不同行业、规模和职位的296位买家。	第06-07页
03	合同前评估 为什么签约前难以评估AI定价	第8-19页
04	AI支出体验 合同签订后预算变化——谁超支、为什么超支、谁承担风险	第20-33页
05	AI定价情绪 哪些型号是买家喜爱、喜欢、容忍还是回避的，以及其原因。	第34-42页
06	人工智能与传统软件即服务 在涉及人工智能时，买家会如何不同对待。	第43-49页

01 执行摘要

当前影响买家对AI定价体验的标题模式

执行摘要

人工智能软件市场正从按席位订阅模式转向基于消费的定价模式——包括积分、代币、使用层级和按成果收费——其理由是AI的可变成本结构需要不同于传统SaaS的商业方式。本研究探讨了这一转变如何被买方接受。

Pricing I/O与Benchmarkit合作，于2026年2月至4月期间调查了296名软件买家。受访者来自收入低于2000万美元的公司到收入超过100亿美元的企业，其中90%拥有总监级或以上头衔。调查从四个方面展开：买家如何在签约前评估AI定价、AI支出与预算的对比情况、买家偏好的定价模式，以及AI定价相对于传统SaaS的感知。调查结果显示，买家的情绪与市场的发展方向并不一致。

本报告围绕三个主题展开。

1. 评估：买家在签约前难以预测AI成本。按席位制是唯一在签约前能获得信心的模式（43%认为最容易评估，14%认为最难），这一结论在所有收入区间均成立。当涉及AI时，其他所有模式的评估难度都会增加：30%认为混合模式比SaaS更难评估，依次为按席位（33%）、按成果（37%）、按使用量（44%）和按积分/代币（55%）。按积分/代币是受影响最大的模式——难度几乎是受影响最小的模式的两倍，且这种困难是结构性的，而非与买家成熟度相关，在所有收入区间中感受相似（54-59%）。可预测性（而非价格）是首要评估标准：68%将可预测的总成本列为前三优先项；最低入门价格排名垫底，仅为19%。
2. 透明度：主要的担忧包括成本不可预测性、定价/使用量到成本的换算问题，以及缺乏使用透明度。与信任相关的问题，如隐性加价和意外费用，排名较低。在开放式回答中，38%要求更高的可预测性，35%要求更高的透明度；只有10%要求更强的价值或成果对齐。加价和意外费用则排名较低。在开放式回答中，38%的受访者要求更高的可预测性，35%要求更高的透明度；仅有10%要求更强的价值或成果一致性。

3. 保护：89%的买家超出了初始AI预算——其中45%大幅超出，44%中等程度超出，仅有9%在预算范围内。超支主要归因于预测失误，而非供应商行为：AI功能驱动额外使用量（67%）、使用量增长快于预期（63%）以及难以预测使用量（38%）。只有10%归因于售后供应商定价或条款变更。67%的买家认为IT是AI成本风险的主要责任人，而财务部门仅为17%，且随着公司规模增大，IT的责任占比增加。买家更偏好防护栏——带有审批的软上限、预测性警报、月度/季度支出上限——而非硬性截止或预购积分池，且最有效的防护栏因具体担忧而异。

问题在于失败而非供应商行为：推动额外使用量的AI功能（67%）、使用量增长超出预期（63%）、以及使用量预测困难（38%）。仅有10%的受访者提及售后供应商定价或条款变更。67%的买家认为IT部门是AI成本风险的主要责任人，而财务部门仅为17%，且随着公司规模扩大，IT部门的责任占比进一步提升。买家更倾向于设置防护栏——即附带审批的软上限、预警通知、月度/季度支出限额——而非硬性截断或预购信用池，且最有效的防护栏因具体关注点而异。

买家偏好度最高的定价模型是那些锚定固定指标的模型：基于座位的模型净偏好度最高，达到+17。偏好度最低的模型是两种纯变量模型：纯使用量模型为-7，纯成果模型为-19。这一模式表明，买家愿意为可预测性接受更高的价格，而对于缺乏固定锚点的模型（无论变量部分如何定义）则会要求折扣。

研究的关键数据

根据塑造买家对AI定价体验的三种模式，整理了调查中的主要数据。

01



评估

买家在签约前难以评估人工智能定价

43%的人认为按座位定价是最简单的

模型用于评估合同前

55%的人认为信用和代币定价更难

评估AI比评估传统SaaS更具挑战性

44%认为基于使用量的定价更难评估

评估AI比评估SaaS更具挑战性

68%的人将可预测的总成本列为前三

优先考虑因素（在评估AI供应商时）

19%的人将最低入门价格列为前三优先项

——这是研究中排名最低的标准

02



透明度

买家希望获得清晰的定价机制，而非创意性方案

38%的人将成本不可预测性列为首要

对人工智能定价的担忧

+17

基于座位的净偏好——在所有AI定价模型中最高

-19

纯结果导向定价的净偏好——在所有模型中最低

38%的开放式回复要求

更高的可预测性和稳定成本

10%的人要求更强的价值或成果

的一致性

03



保护

买家希望有防止超支的护栏，而非硬性截止。

89%的人已超出初始人工智能

预算

67%的人表示人工智能功能推动额外使用

是导致超支的首要因素

10%归咎于供应商定价或条款变更

售后

67%的人认为IT是AI成本的主要负责人

的风险

62%的人希望设置软上限并附带警报和审批

工作流程

02 人口统计数据 谁参与了调查——共296位来自不同行业、

不同规模 and 不同角色的买家。

调查参与者

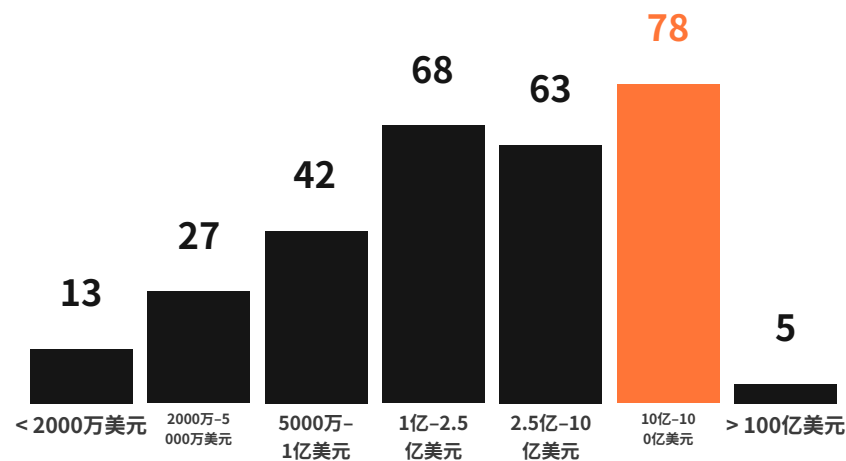


296

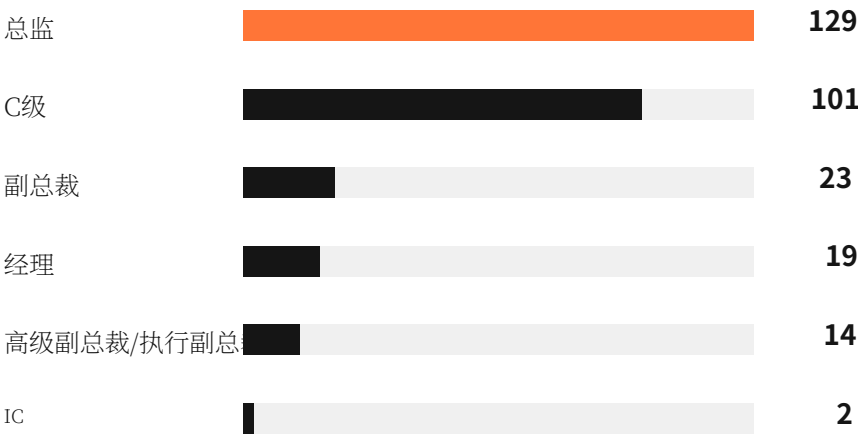
参与者

跨越不同行业、公司规模和资历层级—
—所有人都参与评估、批准或拥有软件
采购决策权。

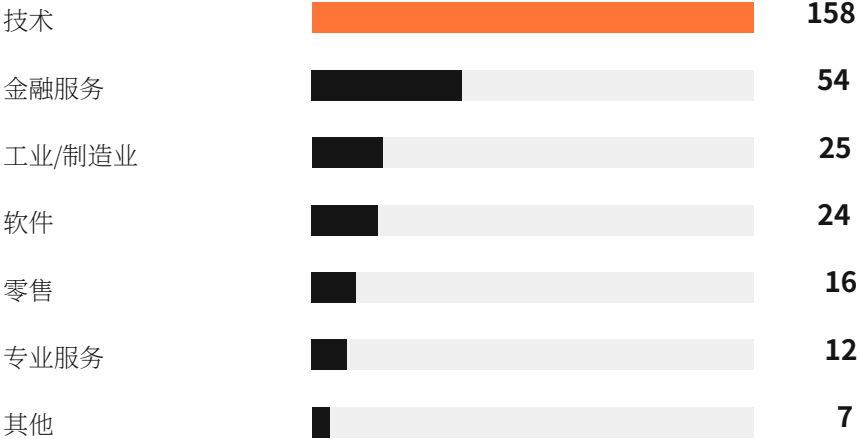
1. 按ARR细分



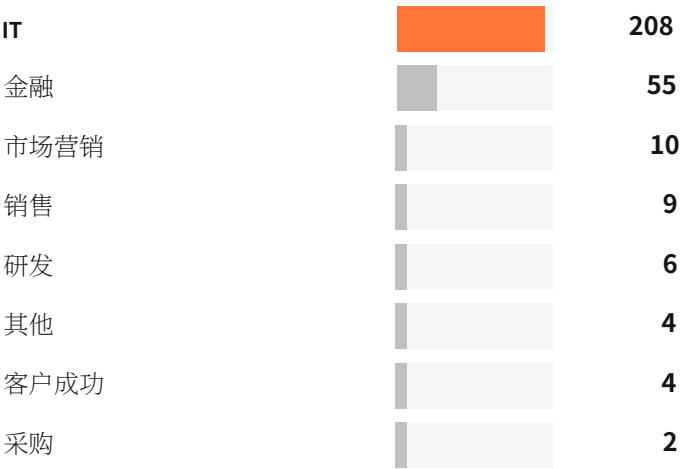
2. 按职位



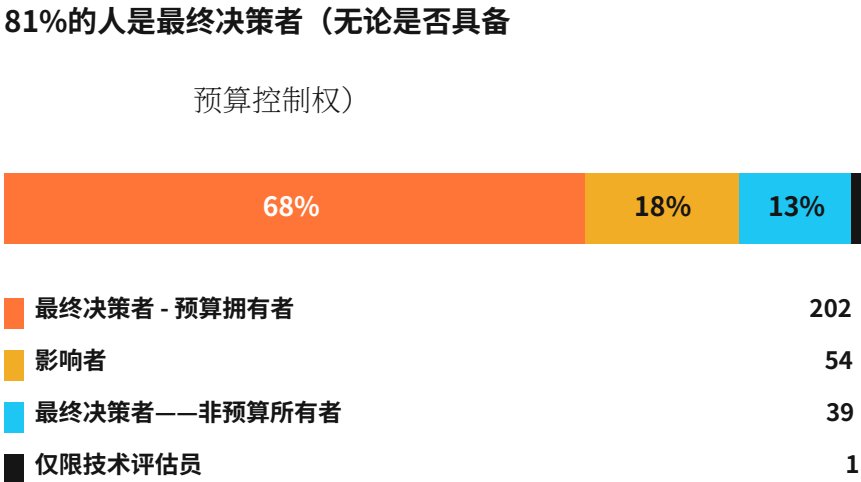
3. 按行业



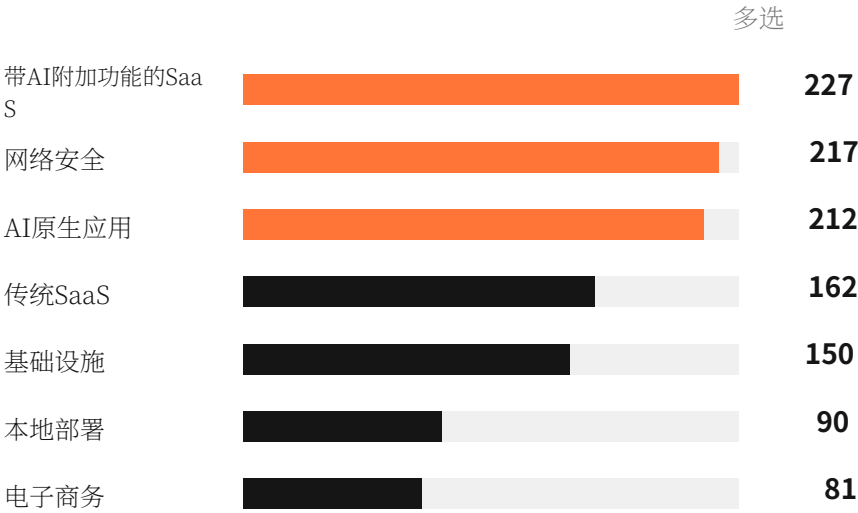
4. 按部门



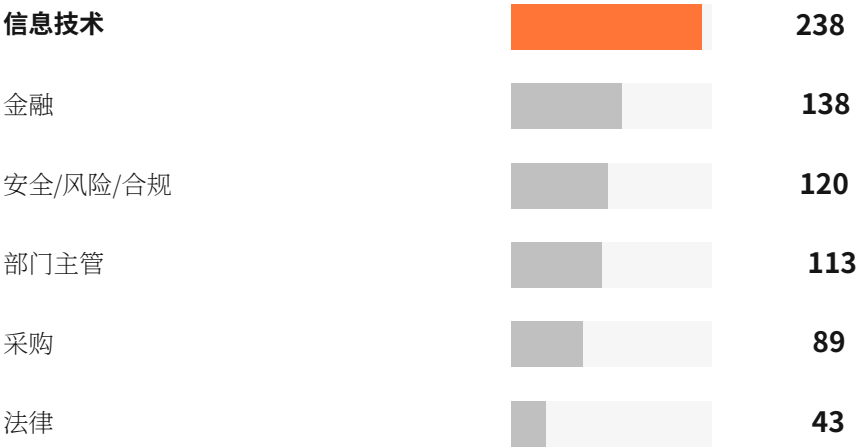
5. 按采购角色



6. 积极购买的软件类型



7. 审批软件的功能



一个偏向拥有预算的买家的评审团。从低于2000万美元的初创企业到超过100亿美元的企业。90%持有总监及以上头衔；81%为最终决策者。

03 签约前评估

为什么在签约前难以评估AI定价，以及买家在尝试评估时衡量的内容

要求人工智能买家提供其对不同定价模式的偏好

- 01

按席位计费

企业按每用户付费以获取软件使用权限。费用随用户席位数增加而上涨。这是常见的定价模式，能为供应商带来可预测的经常性收入。
- 02

按用量计费

企业根据实际消耗量付费。费用随消息、存储或数据处理等活动的增加而上升——类似水电账单。使用量较低的月份费用更少。
- 03

积分/代币

在每个计费周期中，组织会获得一定数量的积分额度，产品中的每个操作都会消耗部分积分。不同功能所需的积分不同，因此简单任务消耗少量积分，而复杂任务则消耗较多积分。
- 04

基于结果

组织根据实际成果付费，通常无需预付或仅需极少的前期费用。只有当达成明确的业务成果时——例如解决一个支持工单、完成一笔款项追回或获取一条合格线索——供应商才能获得收益。
- 05

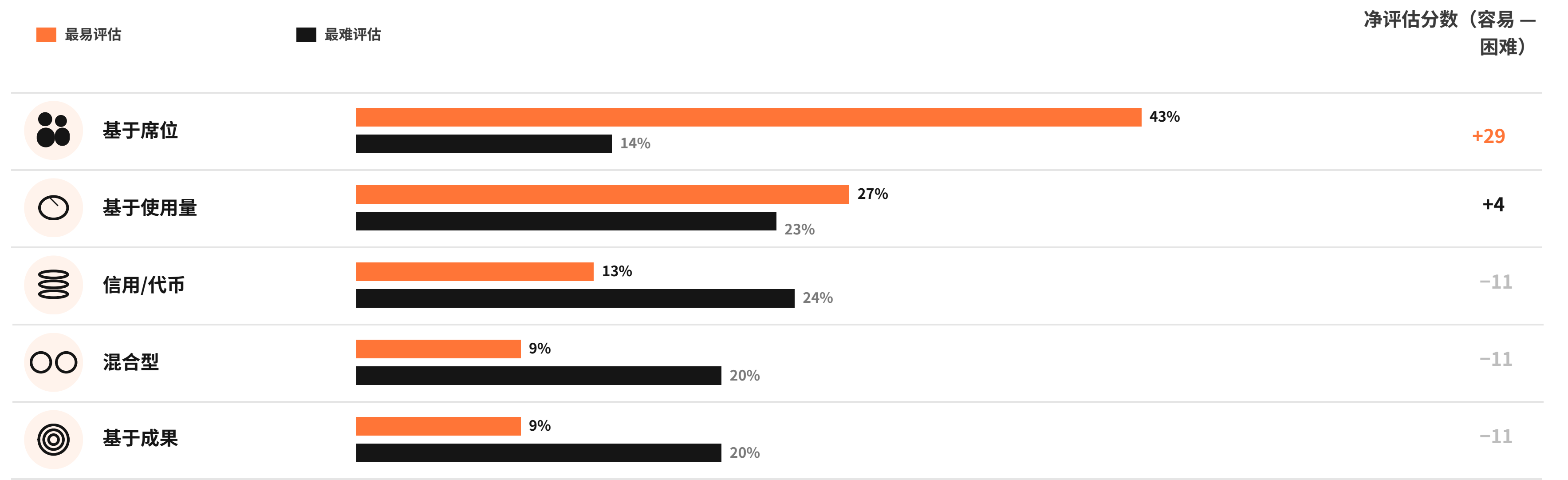
混合模式

由上述两种或多种模式组合而成。最常见的模式是基础平台费用（按席位或订阅计费）叠加按使用量或按成果收取的附加费用。

调查设置 以上五个定义是买家在表达偏好之前收到的确切框架。

基于席位是买家唯一有信心评估的人工智能定价模式

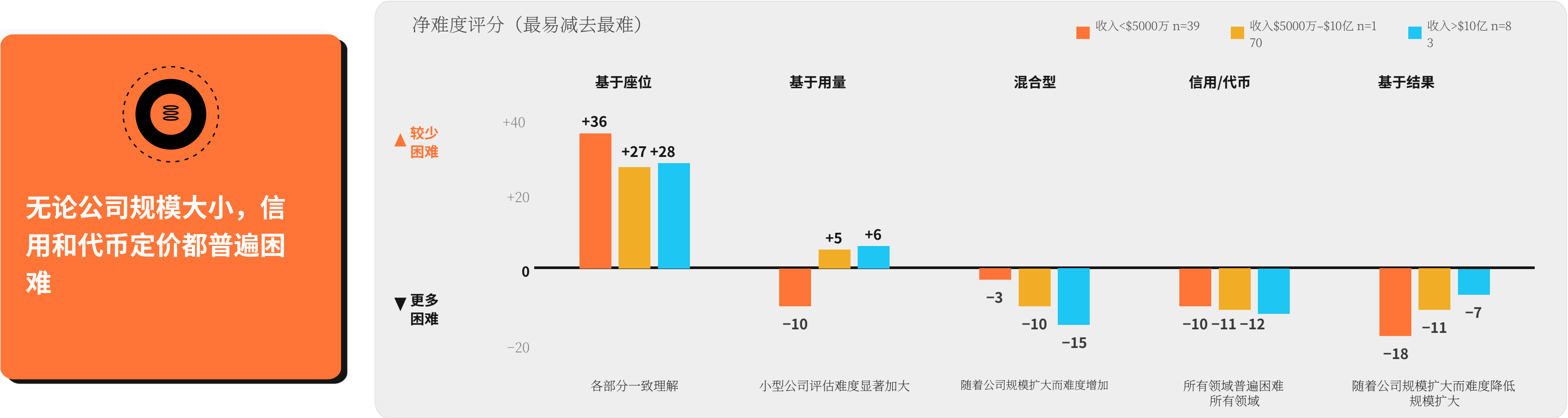
合同前最易与最难评估的定价模式 • 样本量 = 292



基于订阅通常是在签约前最容易评估的定价模式。其他模式则更具两极化——买家通常认为它们更难评估。依赖使用量、成果或消费的定价引入了买家在签约前无法可靠预测的变量

在各类规模的企业中，基于座位的定价模式更易评估；不同细分市场对定价模型的感知存在差异

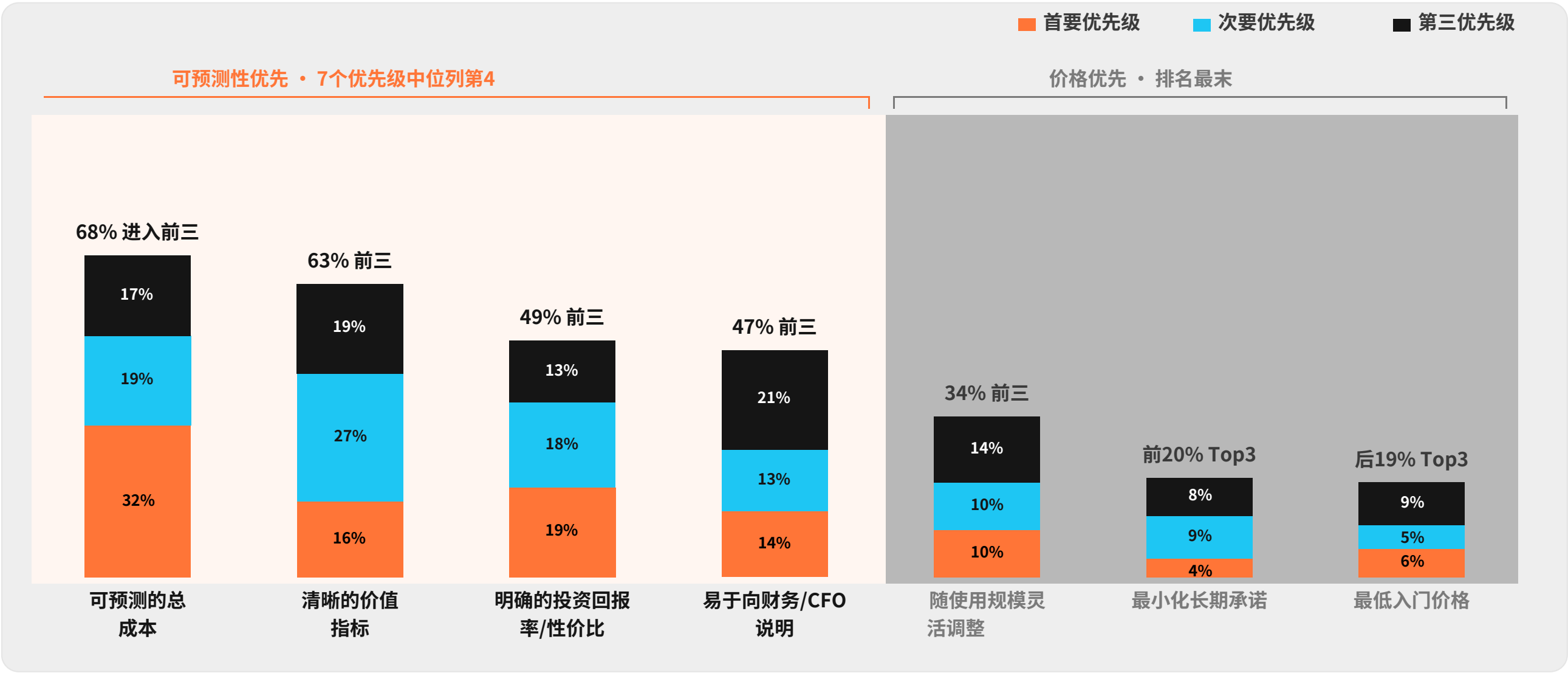
定价模型评估难度净差异（合同前） · 按收入区间划分 · 样本量 = 292



按席位定价是唯一一种在不同规模公司中始终有效的模式，从5000万美元以下的买家到10亿美元以上的企业都表现出强劲的正向效果。其他模式则更不均匀。基于使用量的定价对小公司来说最难，可能是由于使用历史不足，而混合定价随着公司规模扩大而变得更加困难。信用/代币定价在每个细分市场中始终为负面，表明该模式本身存在结构性问题。基于结果的定价会随着公司规模扩大而改善，但始终未能达到净正向效果。

AI采购方更看重可预测性和清晰的定价机制，而非较低的入门价格

评估AI定价时最重要的因素 · 排名前三的百分比 · 样本量 = 292



买家愿意为可预测的定价支付更高费用，对无法预测的定价则支付更低

68% vs **19%**

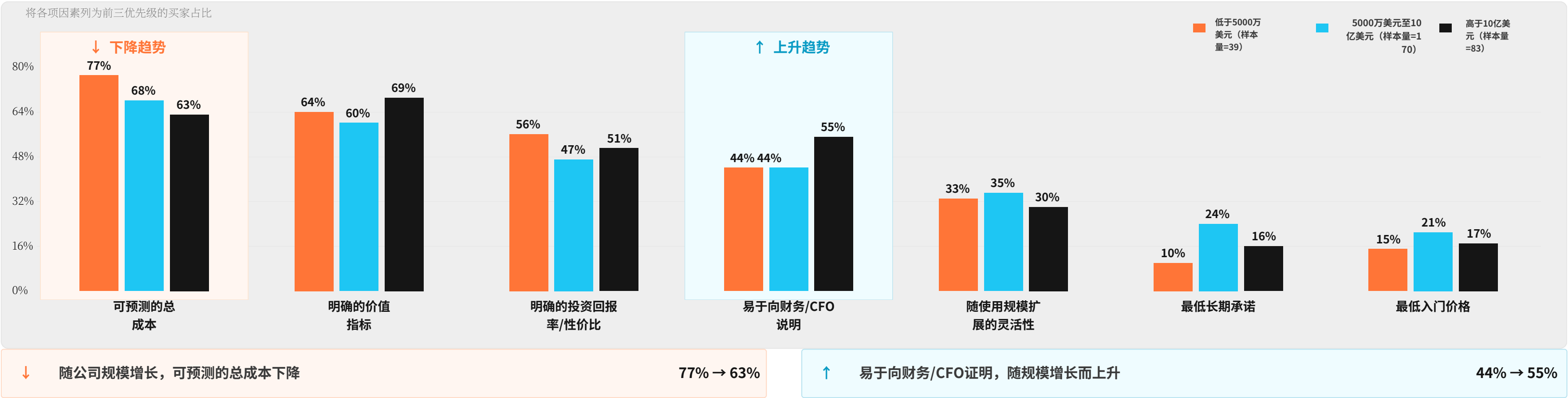
将可预测总成本列为前三优先级

将最低入场价格列为前三优先级

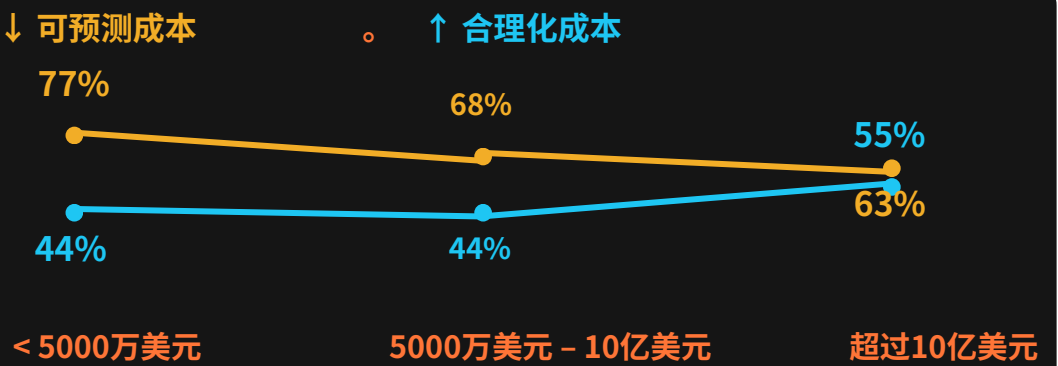
买家将成本可预测性置于节省成本之上。前四大优先事项都指向同一个核心需求：即买家能够预测价格、在内部解释清楚并向财务部门证明其合理性。最低入门价格和最小化长期承诺排名垫底，这表明买家更愿意为可预测的定价模式支付更高费用，而非为无法预测的模式支付较低费用。

随着公司规模的增长，买方面临的问题从预测成本转向证明成本的合理性

评估人工智能定价模型时最重要的因素 · 按收入区间划分 · 样本量=292



可预测成本是各个细分市场的首要任务，但驱动这一需求的因素随公司规模而变化。小公司需要知道他们为自己的规划需要支付多少费用，而大型企业预计会花费更多，但需要清晰的价值指标和能够向财务部门解释的理由。中型企业同样偏好可预测性，略微倾向于灵活性和更低的入门价格。



随着资历提升，定价讨论从证明价值转向为支出辩护

评估AI定价模型时最关注的因素 · 按买方资历划分 · 样本量 = 264

将各项列为前三大优先事项的买方比例

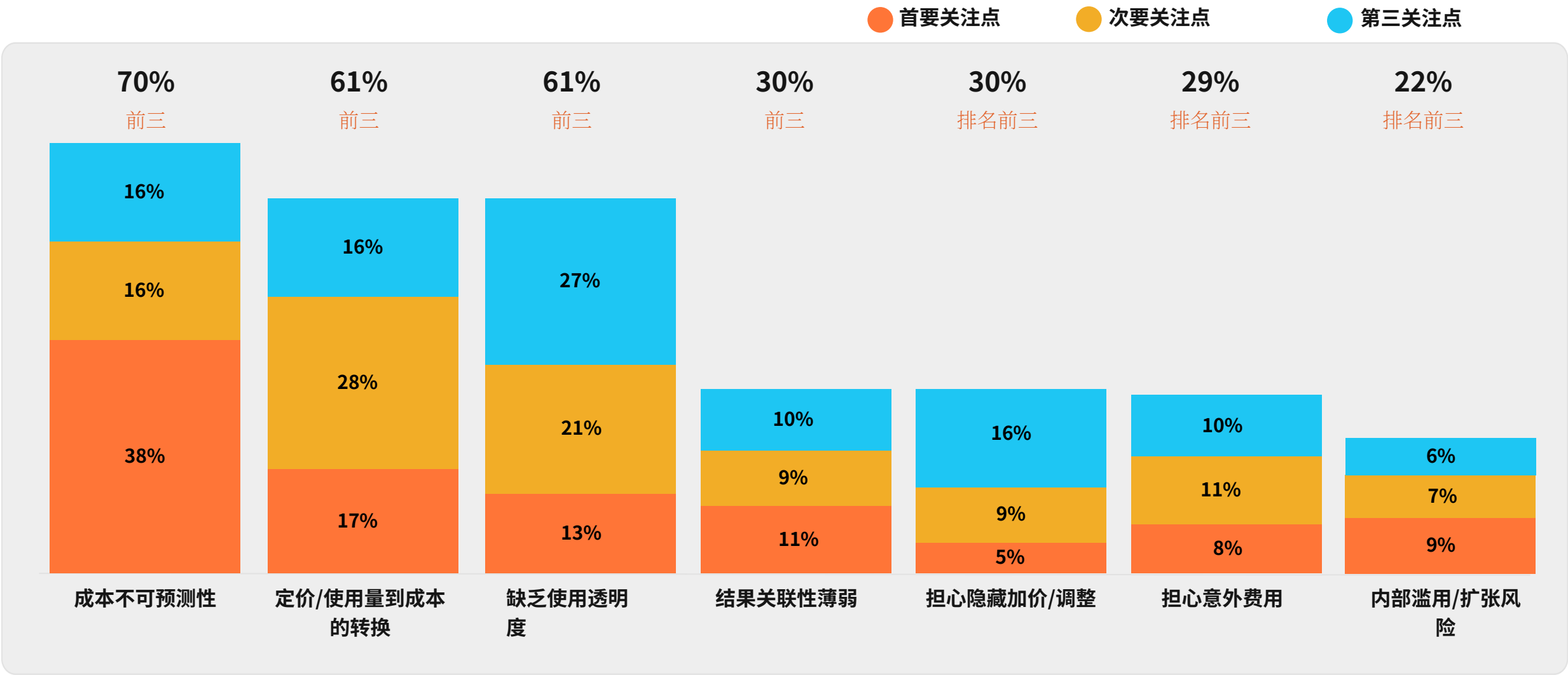
总监: n=128 | 副总裁/高级副总裁/执行副总裁: n=37 | 高管层: n=99

因素	总监		副总裁/高级副总裁/执行副总裁		高管层	
可预测的总成本		66%		76%		64%
明确的价值指标		55%		73%		70%
清晰的ROI/价格与价值关联		55%		46%		40%
易于向财务/CFO说明		38%		57%		54%
随使用规模扩展的灵活性		38%		16%		34%
最低限度的长期承诺		26%		14%		18%
最低入门价格		22%		19%		20%

各职级优先事项高度一致，但随着职级提升，定价讨论重点从性能转向可防御性。
董事更贴近产品，更注重明确的投资回报率和价格-价值关联。副总裁及C级高管则较少关注价值验证，更侧重于明确的定价指标以及向财务部门说明支出的合理性。

买家对人工智能定价的担忧集中在预测和透明度上

人工智能定价首要关注点 · 排名前三的百分比 · 样本量 = 257



首要关注点 #1

70% 成本不可预测性

最受关注问题#2

61% 定价/用量与成本转化

首要关注点 #3

61% 缺乏使用透明度

关于人工智能定价的主要担忧都追溯到一个根本问题——买家无法可靠预测产品的成本。对隐藏费用、加价或结果关联性薄弱的担忧排名较低，表明问题不在于供应商信任度，而在于缺乏可见性和预测信心。买家并非要求供应商更加诚实，而是希望获得能够理解并能据此规划的定价结构。

信用额度方面的模糊性部分源于存在多种不同的信用模式， 例如：

计量单位各不相同， 单位成本各不相同， 且触发消费的定义也各不相同。买家无法评估卖家无法统一界定的内容。

Vendor	Pricing Model	Metric	Cost Structure	What it means
 Adobe	Output-based	Generative credits	Varies across products	A credit is an action, ex: an AI image generation.
 Atlassian	Output-based	AI credits	AI credit pricing, limits not enforced	A credit is based on interactions with specific AI features; more complex interactions cost more.
 bolt.new	Cost-based	Tokens	\$0.00002 per token	Tokens are units of text that LLMs use to process information.
 Clay	Cost-based	Credits	\$0.01 to \$0.07 per credit	A credit is a data point or AI action, ex: data enrichment.
 Cursor	Cost-based	Dollars	\$0.03 per credit	Per GPT-4 request.
 Google	Cost-based	AI credits	\$0.01 per credit	A credit is used in AI generation (not a request). The amount of AI credits used for each model varies.
 HubSpot	Output-based	Credits	\$0.01 per credit	A credit is an action, ex: enrichment, a customer agent interaction.
 Lovable	Output-based	Credits	\$0.25 per credit	A credit is based on work done in response to messages; more complex actions cost more credits.
 monday	Output-based	AI credits	\$0.08 per credit	A credit is a task that's intentionally set by users and is successfully run.
 Salesforce	Output-based	Flex Credits	\$0.005 per credit	A credit is a discrete AI action, ex: updating a customer record. These translate to \$0.10 each.



Salesforce 平台的信用额度成本为 0.005美元



Lovable 平台的信用额度成本为 0.25美元



买家无法将两者进行换算。

五个要素区分了买家可规划的信用模型与不可规划的信用模型

信用模型不是问题，模糊不清才是

01 定义

02 可预测性

03 政策

04 动作映射

05 配额设计

劣质信用模型的构成要素

- ❌ **定义不一致**

询问三个人某项信用的价值，得到三个答案——黑箱本身就是产物
- ❌ **买家无法预测需求**

客户猜测购买量——然后要么囤积积分，要么超量并责怪你。
- ❌ **政策无故变更**

长期的不一致性会破坏信任，客户因不了解购买内容而停止购买。
- ❌ **动作成本差异悬殊**

对同一用户混合使用5积分和50积分的动作，会让追踪变得几乎不可能。
- ❌ **按用户分配与按团队分配，界定模糊**

在有人在不恰当的时机耗尽之前，没人知道配额池是共享的还是独立的。

优质信用模型的构成要素

- ✅ **跨团队定义一致**

销售、财务和产品部门对信用的描述完全一致
- ✅ **对买家可预测**

客户在承诺前可预估用量——无意外超量。
- ✅ **明确的结转和到期政策**

规则形成文档、执行一致，对刺头不例外。
- ✅ **积分映射为离散、可衡量的动作**

一个积分对应一件事；关于什么算是“消耗”积分，不存在歧义。
- ✅ **对配额的明确界定**

重度用户付出更多，因为他们获得更多；模型赢得的是信任，而非怨恨。

基于成果的定价仅在成果完全归属、可衡量且在上线前达成一致时才能生效

必须满足三项条件。若缺少任意一项，将出现六种失败模式。



章节摘要

合同前评估

买家在签署合同前无法可靠评估大多数人工智能定价模式。基于座位的模式是唯一能在合同前建立信心的模式；基于用量、基于成果和混合定价模式比简单模式更难评估。

这种摩擦出现在任何用量计算或代币消耗之前，这也是为什么可预测性比更低的价格点更受买家评估标准青睐

这对买家意味着什么

- 评估AI定价模式的难度主要是一个预测问题。
- 签约前提供使用基准的意愿增强。
- 预计为更可预测的使用量支付更高价格。

这对供应商意味着什么

- 以成本可预测性而非入门价格为主导，是明显的差异化优势。
- 为潜在买家提供简单的使用量计算方法（例如计算器）。
- 选择具有固定指标（例如用户席位）的模型比纯粹可变模型更有利

43%

认为基于座位的模式最易评估

基于用量（+4%）是唯一获得正向难度净评分的其他模式

68%

将可预测总成本列为前三大优先事项

低价入门排名垫底。买家愿意为可靠的预测能力支付更多费用。

55%

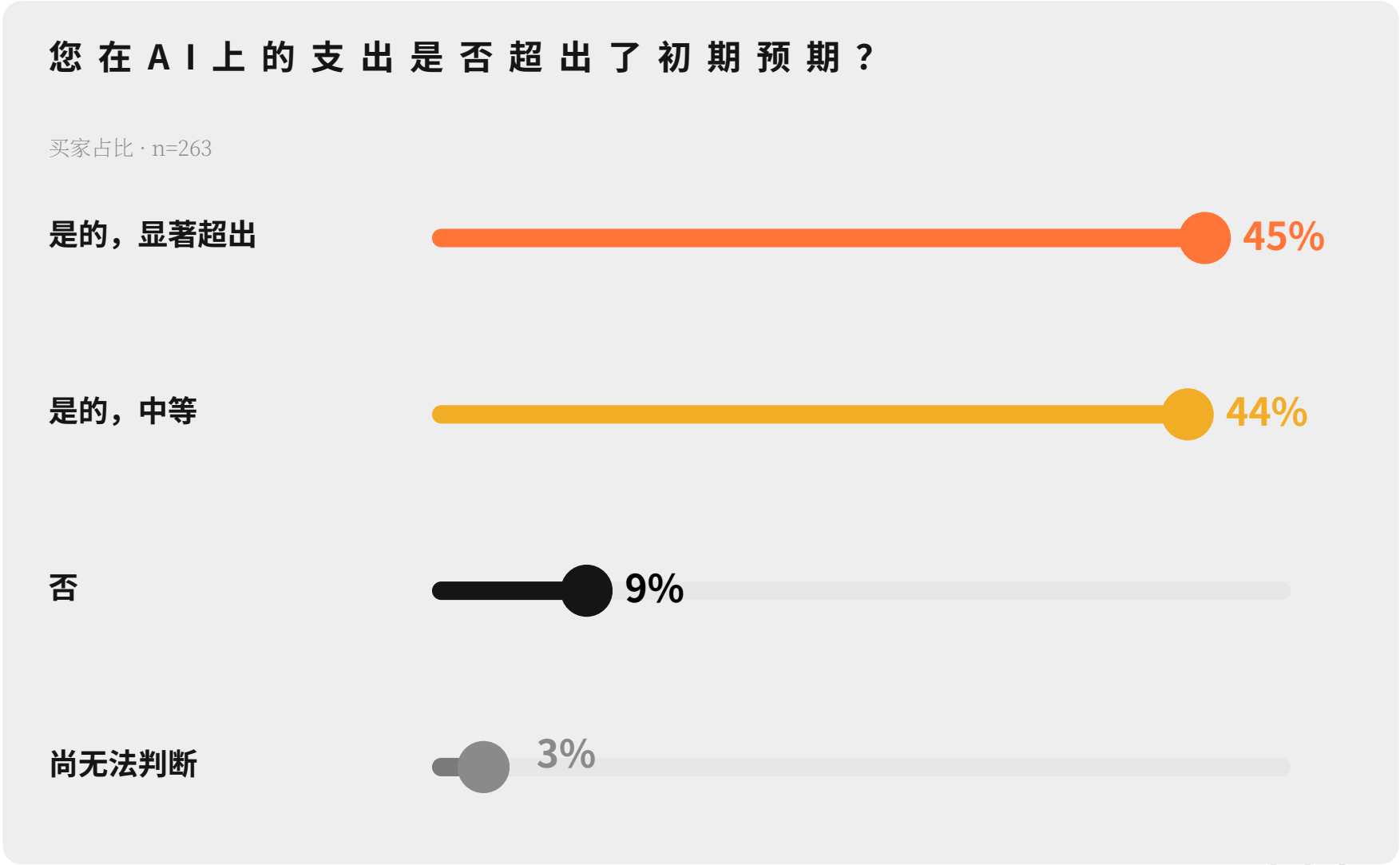
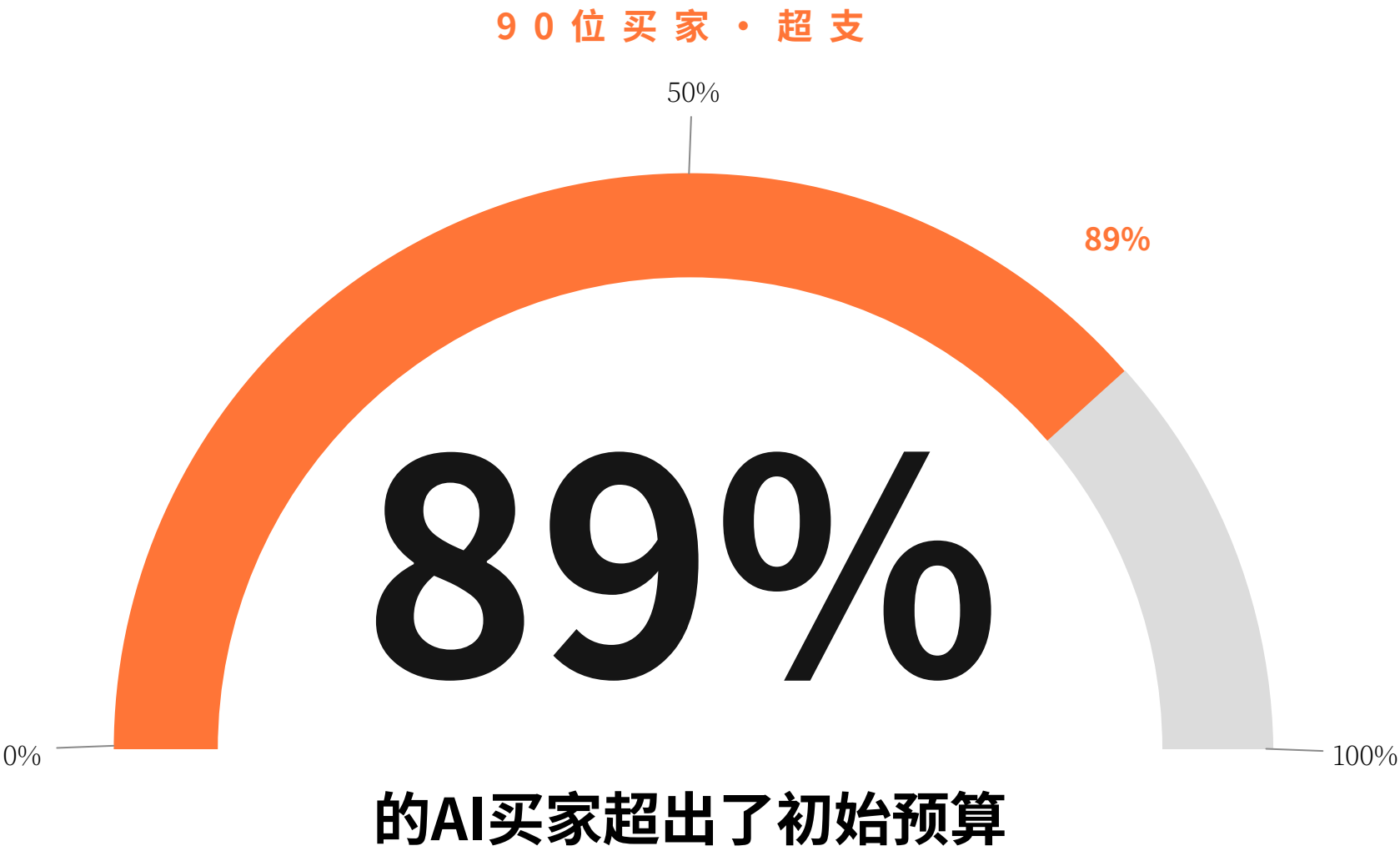
的人表示基于信用/令牌的定价更难以评估AI

评估基于信用模型的挑战在各规模公司中普遍存在

04 AI 支出体验

合同签订后预算如何变化——谁超支、原因是什么，以及风险由谁承担

对于AI买家而言，维持初始预算现在已成为例外情况



超出初始AI预算 是常见现象，其中显著超支与中等超支几乎一样普遍。这反映了买家对可预测性的担忧：许多人低估了自己将被收取的费用或实际使用量

金融服务在超出初始预算的AI支出中受影响最大

按超支排名

谁的AI预算超支最严重？——按行业划分。N=263

01 金融服务 96% 超出预算

n = 47 位买家
烧钱比例最高——大多数远超计划

02 科技 / 软件 89% 超出预算

n = 167 位买家
紧随其后——仍有十分之九超出了计划。

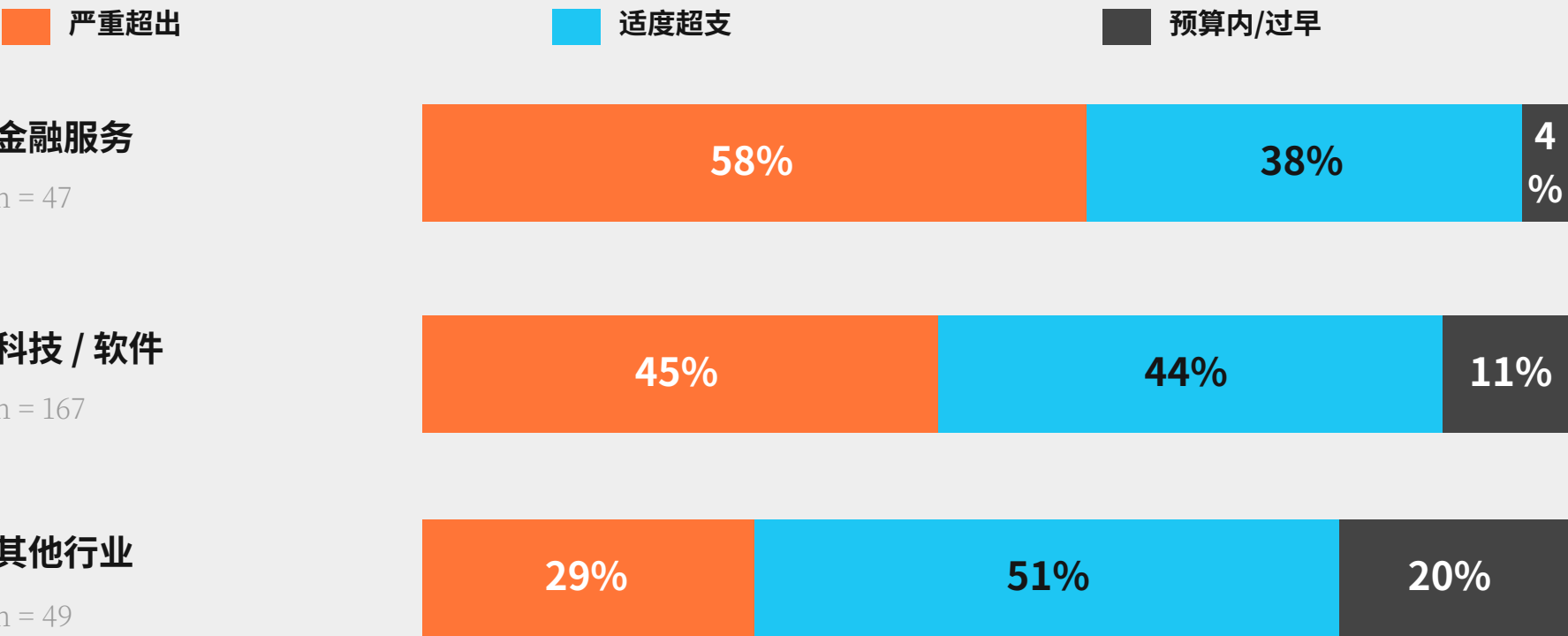
03 其他行业 80% 超出预算

n = 49 位买家
n = 49 位买家
八成买家超支，但严重程度较低。

完整数据分析

您在AI软件上的支出是否超过了最初的预期（预算）？

买家比例 · 按行业



金融服务买家在AI软件上的超支幅度最大。科技和软件买家紧随其后，约有十分之九的买家实际支出超出计划。其他行业也存在超支现象，但程度较轻。超支现象在各市场普遍存在，
供应商应预见买家在价格谈判时已感到不满。

AI的使用量超出买家预期，这是预算超支的首要原因

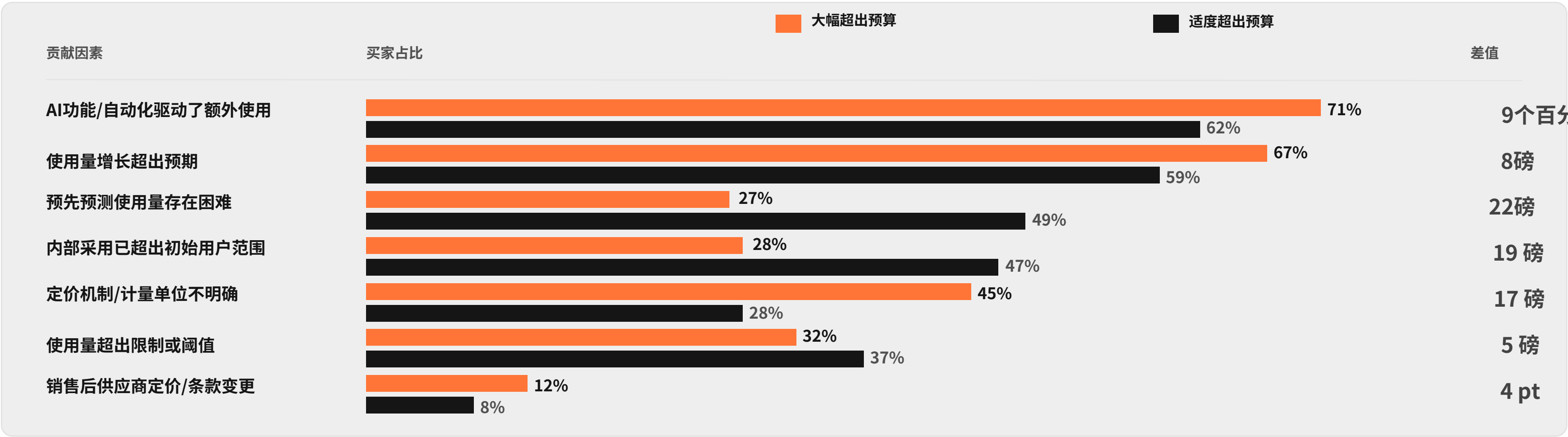
对超出预期的AI支出的主要贡献 · 样本量 = 233



预算超支是由于买家低估了其组织实际使用AI的数量，而非供应商在签约后改变规则。主要贡献者都描述了同样的动态：AI产品产生的活动超出团队预期，而内部采用速度也快于预测

大幅超出预算的买家报告的根源原因与适度超支的买家不同

主要因素导致AI支出高于预期 • 按超支严重程度分类 • N=233



显著超支 • 样本量=117

需要消费护栏——而非更精准的预测。

大幅超出初始预算的买家意识到产品确实有效，但使用量呈复合增长，预算无法跟上。这些客户需要的不是更精准的预测，而是消费护栏——让他们在无需每季度找财务部门的情况下持续扩大规模。

中度超支 • 样本量=116

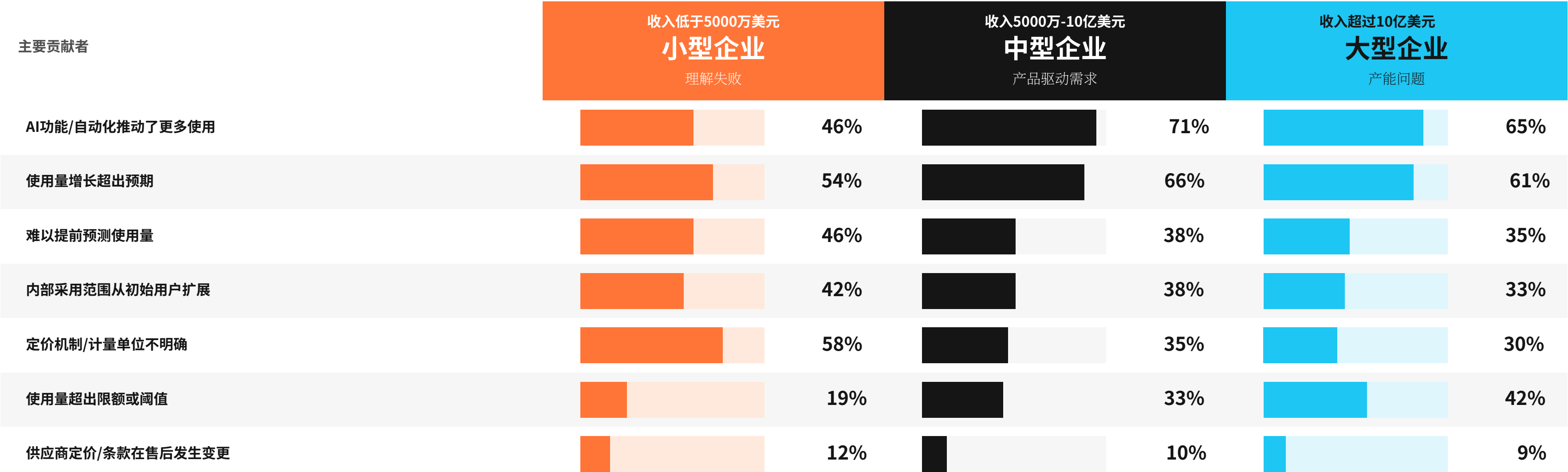
需要透明计量表——而非使用上限。

略微超出初始预算的买家难以模拟其购买内容，超额支出源于对使用情况和未来采用率的误解。他们需要的不是使用上限，而是一个足够透明的定价模型，以便从一开始就能进行预测。

每个细分市场均超出预算，但其原因各有深意

人工智能支出超预期的主要贡献者 • 买家占比 • 样本量 = 230

收入<5000万美元: n=26 | 收入5000万-10亿美元: n=138 | 收入>10亿美元: n=66



小型公司（低于5000万美元）超出预算属于理解失败：该细分市场的买家并未完全理解其购买内容。一个超支标题，三种不同解决方案。小型买家需要教育，中型市场需要治理，企业需要流量控制。定价页面无法解决容量问题，速率限制也无法解决理解问题。中型市场（5000万至10亿美元）超出预算是由产品驱动而非买家驱动：需求增长速度超过了组织治理能力。企业（超过10亿美元）超出预算属于容量问题：该细分市场需要流量控制、优雅降级和预测性警报，而非定价教育。

IT承担AI预算风险——即便他们并不推动使用

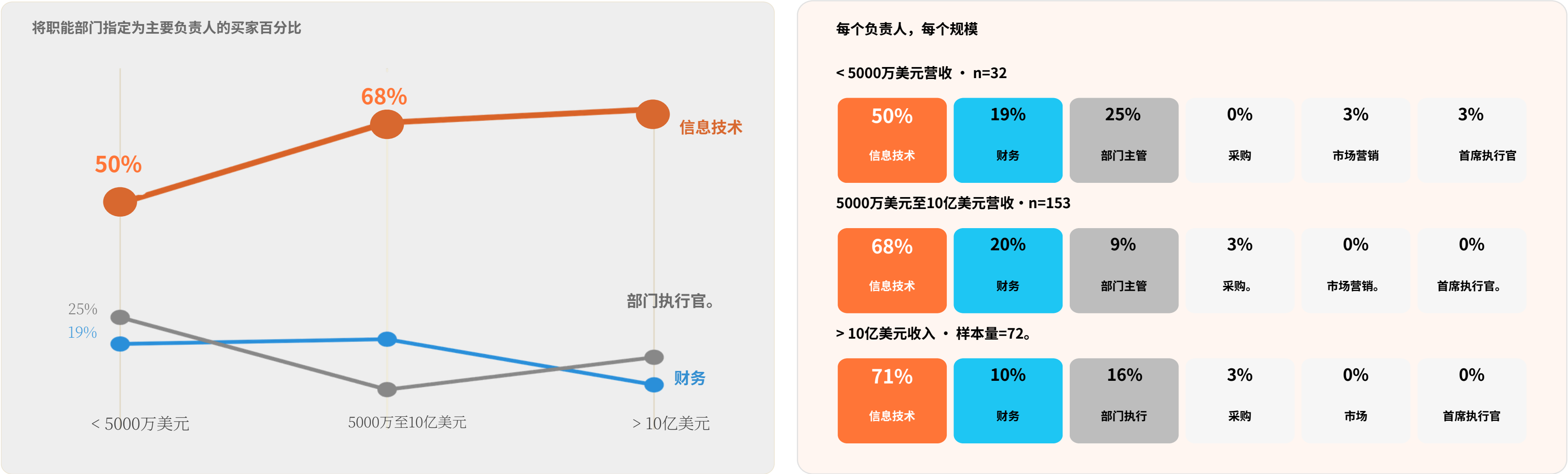
AI成本（预算）风险的主要承担者 • 买家百分比 • N = 257



人工智能成本风险的责任几乎完全集中在IT部门，而其他部门仅扮演次要角色。这就造成了一种紧张局面：财务部门虽然密切关注预算，却很少被认定为实际负责人。这种不匹配正是意外超支的温床。

企业的AI成本风险随公司规模增长而增加

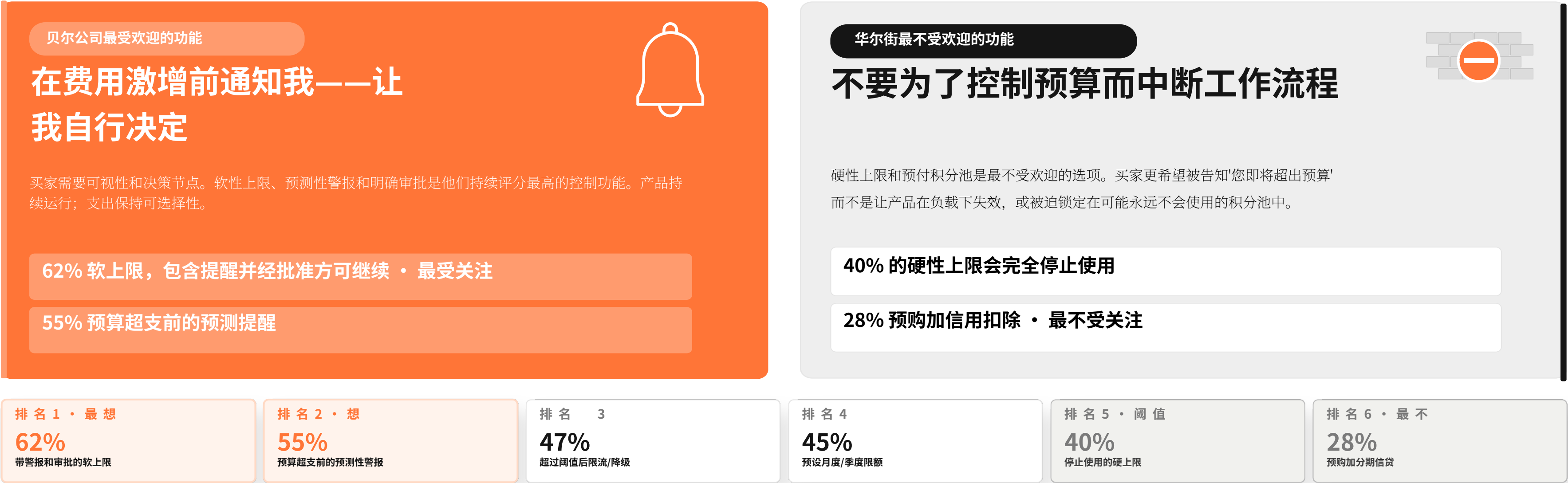
人工智能成本（预算）风险的主要负责人 · 按部门划分 · N=257



在所有规模的企业中，IT部门承担了最大份额的人工智能预算风险，但这种主导地位随着企业规模的扩大而增强。在大型企业中，IT部门被指认为人工智能成本风险主要责任方的频率远高于其他部门。而在小型企业中，财务部门和各部门高管扮演着更重要的角色，这很可能是因为IT部门规模较小，承担风险的能力有限。

AI买家需要的是使用情况的可视性和可控性，而不是硬性限制

供应商最受欢迎的使用控制功能 • 样本数 = 262



在管理人工智能使用方面，买家更看重可见性和决策点，而非严格限制。他们宁愿提前收到预警并获得继续使用的批准机会，也不愿产品突然中断（硬性上限）或被锁定在预购积分中。

针对每个AI定价关注点，按相关用户群体对使用控制措施的偏好程度（即相比整体基准更希望获得的程度） 进行排序

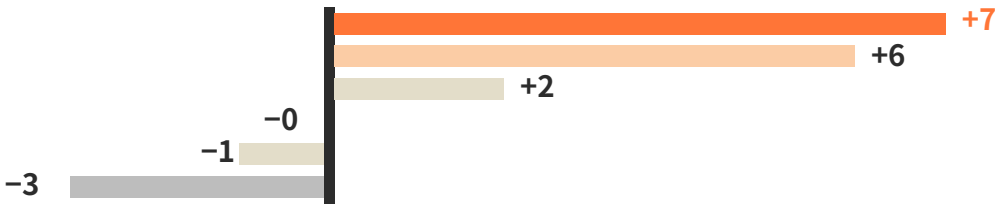
样本量 = 262

相比整体基准： 橙色=最高提升程度的控制措施；深灰色=强烈拒绝（小于-3%）

成本不可预测性

n = 181

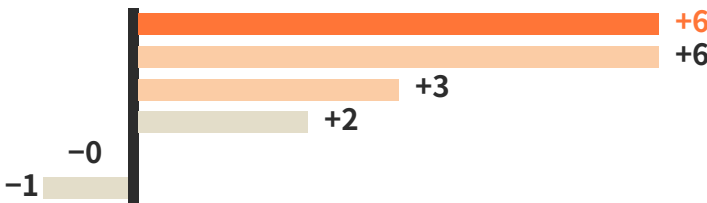
- 月度/季度消费限额
- 硬性使用上限
- 预测性告警
- 限制/降低使用量
- 预购固定积分
- 软上限 + 审批



用量到成本的转换

n = 155

- 软上限 + 审批
- 每月/每季度支出限额
- 硬性用量上限
- 预测性告警
- 节流/降低用量
- 购买前固定信用额度



缺乏使用透明度

n = 154

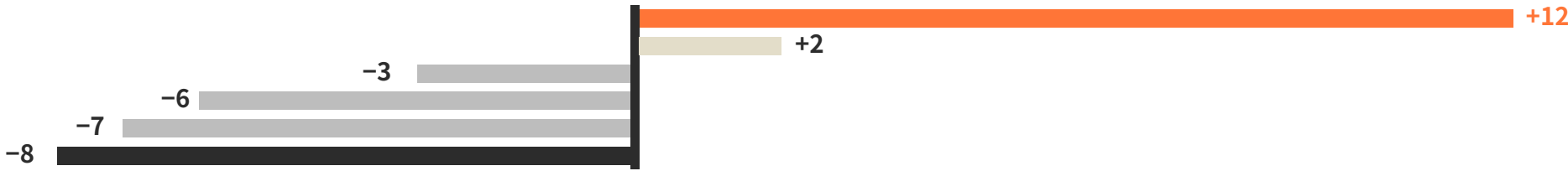
- 月度/季度消费限额
- 软上限 + 审批
- 硬性使用上限
- 预测性警报
- 购买前固定积分
- 节流/降级使用



弱结果关联

n = 76

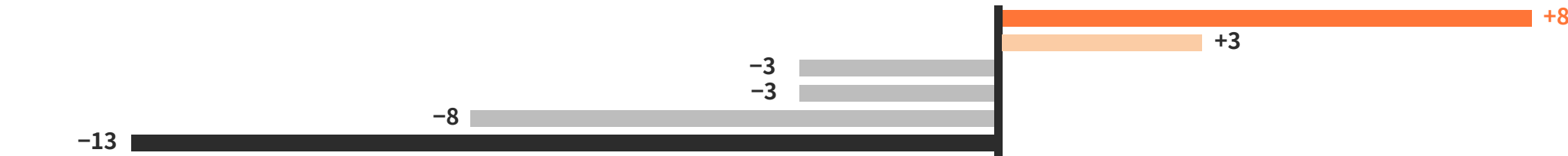
- 节流/降级使用
- 预测性警报
- 硬性使用上限
- 软性上限+审批
- 月度/季度支出限额
- 预购固定额度



隐性加价/调整

n = 76

- 节流/降级使用
- 预购固定信用额度
- 软上限 + 审批
- 硬使用上限
- 预测性警报
- 月度/季度消费限额



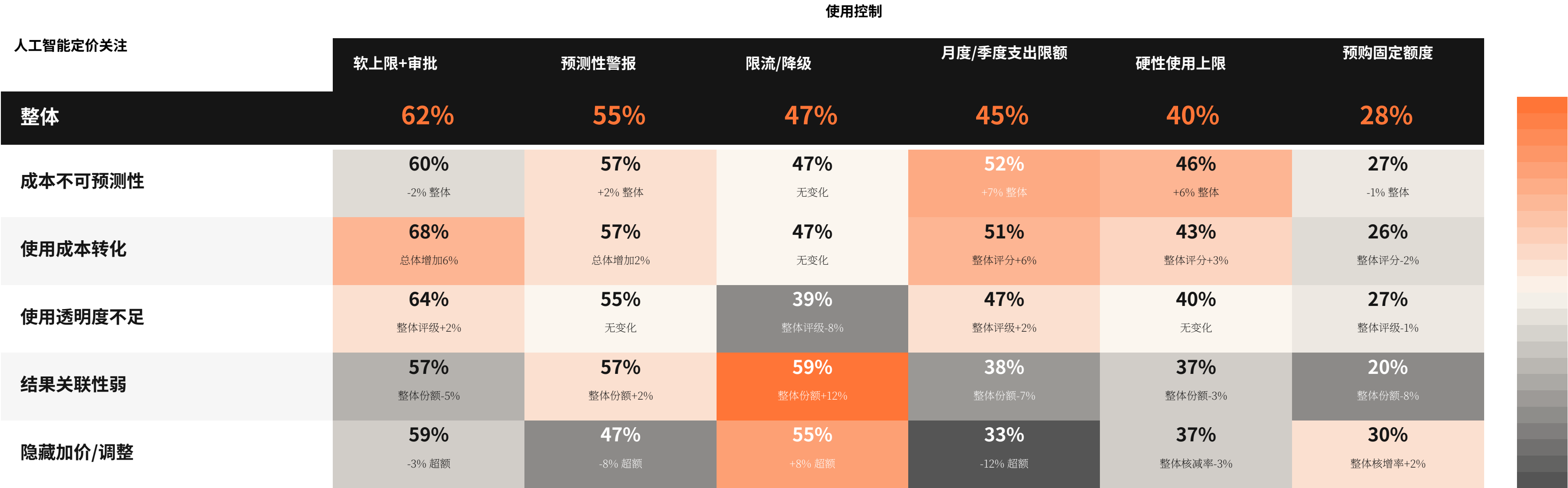
-15个百分点

+15个百分点

每个“关注群体”都有不同的首选使用控制方式，而且多个群体积极地排斥其他群体最偏好的控制措施。产品的正确治理功能取决于您的目标客户群体最关心的优先事项，而非市场默认设置。

买家并非寻求单一通用解决方案——预算超支的不同原因需要不同的管控方案

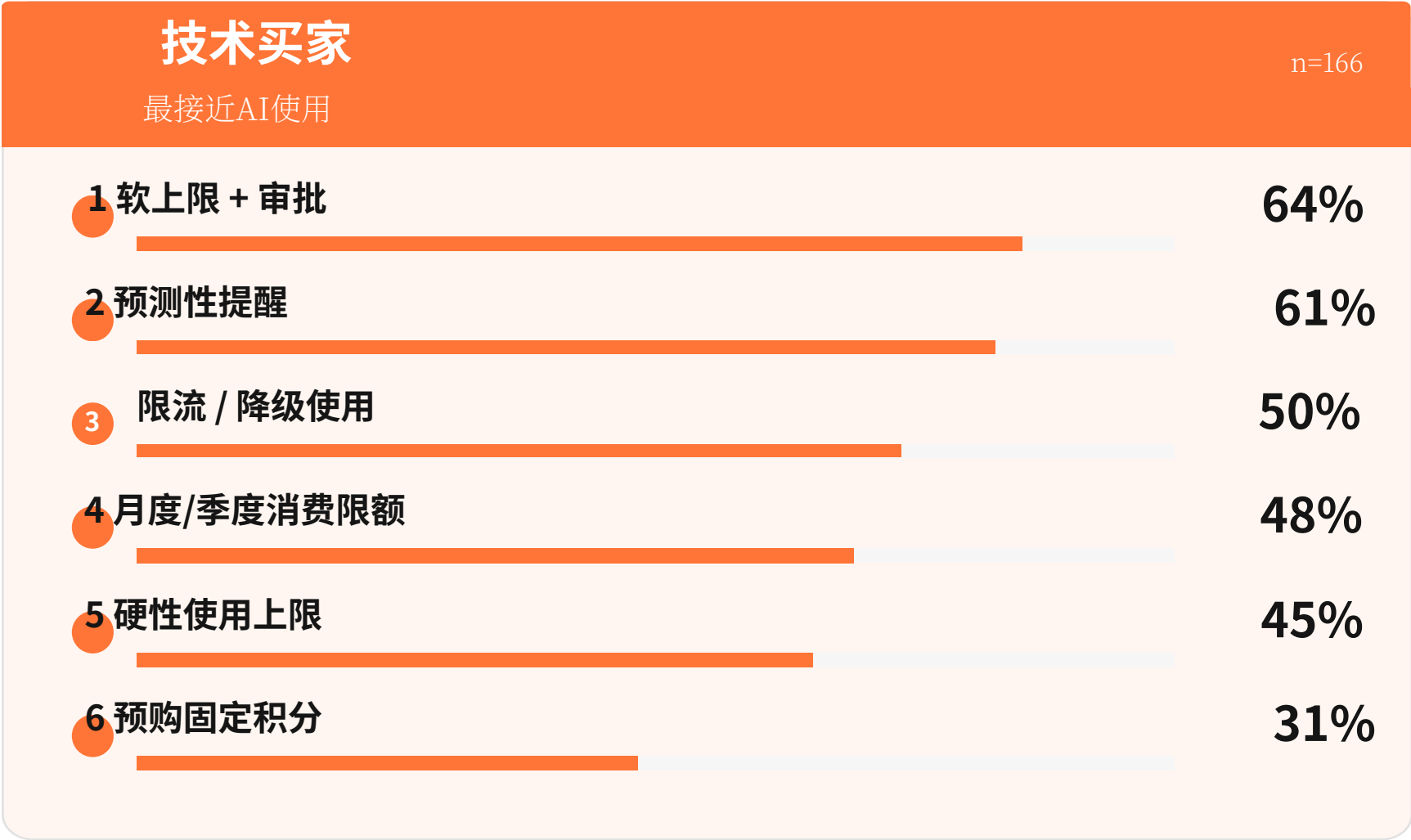
由人工智能定价关注驱动买家首选使用控制措施 • 样本量 = 262



不存在放之四海而皆准的人工智能定价解决方案。每种关切都对应特定的首选调控手段。能够根据买家的实际关切维度（而非提供市场通用默认设置）匹配相应防护措施的供应商，将获得更多预算份额的信任。

技术买家比非技术买家要求更多的保护措施

所需的使用控制 • 按买家画像划分的所有控制 • N = 262



买家并非要求更便宜的积分，而是寻求对其自身预测失误的保护。每一项高优先级的控制都是针对不可预测消耗（例如警报、限额、审批）的防护栏，且技术买家并非要求更便宜的积分，而是寻求对其自身预测失误的保护。每一项高优先级的控制都是针对不可预测消耗（警报、审批、限额）的防护栏。技术买家对每项防护栏的需求比非技术买家更为强烈，这表明最接近人工智能使用的团队最担心失去对其的控制。

买家希望针对不同问题设置不同的防护栏，并对所有问题提供预测性警报

软上限 + 审批	- 当买家难以提前预测用量时，效果极佳 - 当买家内部采用范围超出初始用户、用量增长速度超出预期，或自动化驱动额外用量时，效果中等 - 当定价指标或计量单位不明确时，效果不佳
硬用量上限	- 当买家内部采用范围超出初始用户时，效果极佳 - 当买家难以提前预测用量或定价指标/计量单位不明确时，效果中等 - 当自动化驱动额外用量或用量增长速度超出预期时，效果较差
预测性警报	- 当买家内部采用范围超出初始用户时，效果极佳 - 当买家难以提前预测使用量、自动化驱动额外使用、使用量增长超出预期，或定价指标/计量单位不明确时，效果适中。
月度/季度支出限额	- 当买家难以提前预测使用量时，效果极佳。 - 当买家的内部采用范围超出初始用户，或使用量增长超出预期时，效果适中。 - 当自动化驱动额外使用，或定价指标/计量单位不明确时，效果较差。
速率限制/降级	- 当定价指标/计量单位不明确时，效果极佳。 - 当自动化驱动额外使用时，效果较差。 - 当买家难以提前预测使用量、使用量增长超出预期，或内部采用范围超出初始用户时，效果不佳。
预购固定积分	- 当买方内部应用范围超出初始用户时，效果极佳。 - 当买方难以提前预测使用量时，效果中等。 - 当自动化驱动额外用量、用量增速超出预期，或定价指标/计量单位不明确时，效果较差。

章节摘要： AI支出体验

大多数AI买家超出初始预算，但原因并非单纯的供应商行为。买家经常低估其AI使用量，例如通过扩大内部采用

这更多是治理问题而非定价问题，需要设置护栏以维持买家对AI软件采购的满意度

这对买家的意义

- AI预算的制定应基于对未来支出的假设，而不仅仅是初始使用量
- IT和财务部门需要共同对AI支出相关的成本风险负责
- 预期护栏是解决成本不可预测性和采用驱动增长等常见问题的基本保障

这对供应商意味着什么

- 采用基于使用的模式时，软上限和预测性警报等使用控制措施至关重要
- 处理异议对于克服买家对可变定价的既有犹豫至关重要
- 预计买家可能曾因AI软件超支，因此会对可变定价模式持怀疑态度

89%的买家超出初始AI预算

45%显著超出，44%适度超出，仅9%保持在预算内

67%

将IT指定为AI成本风险的主要负责人

尽管密切关注预算，但财务部门仅在17%的情况下被提及

62%

希望设置带警报的软上限作为保护性护栏

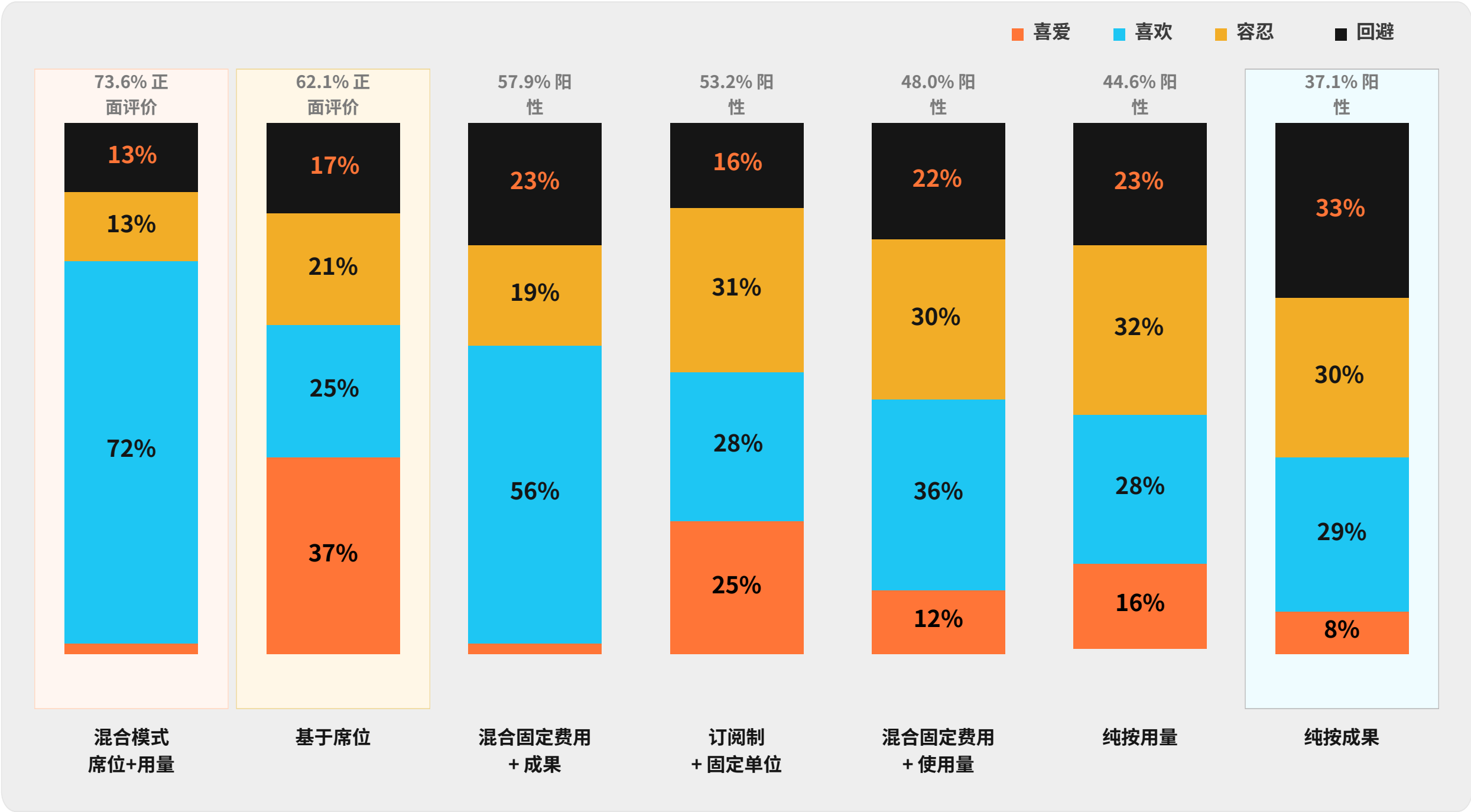
大多数买家拒绝硬性限制，转而偏好警告加决策点的模式

05 人工智能定价情绪

买家喜爱、喜欢、容忍或回避的定价模型，以及他们偏好的原因。

AI买家看重基于席位且可预测的定价

由AI定价模型得出的情感分析 · 买家占比 · N = 296



01

73.6%

阳性结果。
混合座椅+使用量

02

62.1%

正面
基于座位

03

57.9%

阳性。
混合固定费用+成果

买家首选基于座位的定价模式——他们最明确排斥的是纯基于结果的模式

最受欢迎与最不受欢迎的定价模式 · 按净偏好排名 · 样本量 = 296

定价模式	最受欢迎		最不受欢迎		净偏好
基于座位		33%		16%	+17
订阅 + 固定单位		25%		15%	+10
混合固定费用 + 使用量		11%		9%	+2
混合席位 + 使用量		7%		7%	+0
混合固定费用加成果		2%		5%	-3
纯按使用量计费		14%		21%	-7
纯结果导向型		8%		27%	-19

净偏好最差的两两个模型属于“纯模型”。只要定价完全与单一变量挂钩，无论变量如何定义，买家都会持续拒绝所有锚定被移除的定价结构。

对于一款关键性AI产品，可预测性优于所有变量替代方案

关键性AI产品的首选定价结构 · 买家占比 · N=257

首选方案

51%

高固定成本，完全可预测

若消除变量风险，买家愿为核心产品支付更高价格
半数市场宁愿选择确定性而非节省开支。

其他选项

26% 混合型（含明确防护措施）

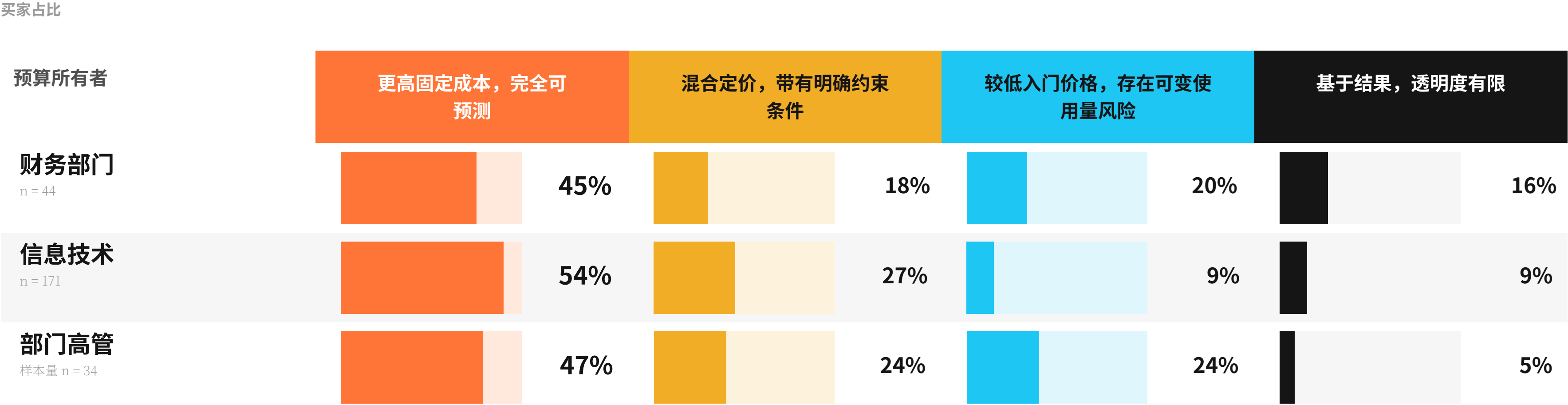
13% 低入门价，可变使用风险

10% 仅按结果付费，透明度有限

在所有四种结构中，约四分之三的买家选择具有可预测底线的选项，而四分之一接受可变风险。对于高风险的人工智能产品，基于固定锚点的打包方案是大多数买家愿意承诺的选择。

财务部门对基于使用量的可变定价结构表现出更高的方向性容忍度

基础AI产品的首选定价结构 · 按部门划分 · 样本数=249



IT部门似乎是AI预算最厌恶风险的持有者：当IT掌握预算支出时，买家倾向于完全可预测的定价模式，拒绝可变或基于结果的替代方案。财务部门似乎是风险承受能力最强的，对基于结果和可变使用量的定价接受度明显更高。

您的产品定价设计应匹配承担预算风险的部门职能，而非泛化的买家特性。

买家反馈最终归结为一个简单的要求：让AI定价更清晰、更可预测、更易于管理

买家希望在AI定价模型中改变什么 · 样本量 = 194



买家并非要求更复杂的定价创意，而是要求减少模糊性。多数买家透明自己的可预测性及要求透明使用控制，而非要求透明分析预测性，以及更少透明使用控制，而是要求减少模糊性

用买方自己的话说：清晰度和可预测性不容妥协

软件供应商对买方如何体验AI定价的最大误解是什么？

出现两大主题

定价清晰度

“我们需要清晰度和可预测性。否则
无法验证预算

定价清晰度

“如果无法确定预算，我们不太可能采用新AI功能”
它们存在信用/使用上限或附加费用。我们
正在积极避免采用具有这种定价机制的人工
智能功能。

定价清晰度

“客户成果、供应商收入和供应商成本之间存在
错配”

价值认知

“作为买家，我不关心你的成本模型——我需要你清
晰说明软件带来的好处，并透明地解释使用量如何
转化为公平定价。我愿意为价值付费，但没时间争
论价值——请提前让我了解清楚。

价值认知

“低风险试用（即使是付费）的重要性
试用，以便你能判断人工智能的定价是否名副其
实。我不想被锁定在多年期合同中，结果要么方
案行不通，要么供应商倒闭。

价值认知

“承诺过多而兑现太少
证据不足，难以兑现承诺。这有点像20年前的云
迁移——需要时间，也经历过大量失败。

买家希望设置消费护栏并明确价值机制

如果可以对当前人工智能定价模式做出一项改变，你会选择什么？

浮现出两大主题

主题 01

消费可见性与安全护栏

- "我会实施更清晰、可预测的内嵌式定价
针对突发峰值设置防护措施，使预算编制和规模扩展更加轻松。
- "我会优先考虑标准化、实时消费仪表盘配合
自动化软性上限警报。这将使组织能够尝试AI功能，而无需担心在计费周期结束时产生巨额、不可预测的超额费用。
- "更可预测的定价——清晰的固定方案或上限有助于控制成本
避免意外——清晰的固定方案或上限有助于控制成本并避免意外情况。
- 我会让每个供应商提供硬性使用上限并自动关闭，这样
我们就不会收到超出预算的意外账单。
- 我会让定价更可预测。比如设定明确的月度上限或
提供慷慨限制的统一定价层级会更容易规划并避免意外费用。

主题 02

标准化价值计量

- "我会围绕清晰可预测的价值单位来标准化定价
而非模糊的令牌或波动性用量。
- "让定价简单、透明且可预测，而非令人困惑
按用量计费方案。
- "固定成本、重要结果指标的可预测性，以及
根据需求变化展示灵活性的能力。
- "用量模式在达到某些里程碑后似乎变得有些令人困惑
被击中。似乎大多数（使用量）更多时收费更高。
- 我会确保有足够的灵活性来随着加入更多成员而
我们系统中AI的扩展持续增长。如果采用固定座位数，就难以应对任何变化以及随时间推移的增加量。

章节摘要： AI定价情绪

买家情绪与每种模型包含的固定定价程度相一致：
两个评级最高的模型以固定指标（如座位数）为基础，其中评分最高的是座位数加使用量的混合模型，而评分最低的两个模型是纯使用量模型和纯结果模型。

当直接询问买家希望如何改变AI定价时，他们并未要求更复杂的定价，而是希望定价清晰且可预测。

这对买家意味着什么

- 优先选择将固定基础与可变部分相结合的混合模型。
- 对于关键的AI产品，接受较高的固定价格以换取可预测性
- 认识到预算所有者决定偏好：IT主导的支出倾向于完全可预测，而财务主导的支出可接受波动性

这对供应商意味着什么

- 基于固定锚点构建定价——每个顶级模型都有锚点，而每个底层模型都缺乏
- 纯成果导向定价高度依赖产品特性，并非市场默认模式
- 对于风险规避型细分市场，默认采用基于座位或订阅+固定单元的打包方式。

51%的买家会选择更高的固定成本/完全

可预测的结构

只有10%的买家会接受透明度有限的纯结果模型。

38%

将'更高的可预测性/稳定的成本'列为首要改进项

仅有10%的购买者要求更强的价值/成果一致性

37%

对纯成果导向定价持积极态度

在所有测试模型中最不受欢迎 • 净偏好率为-19

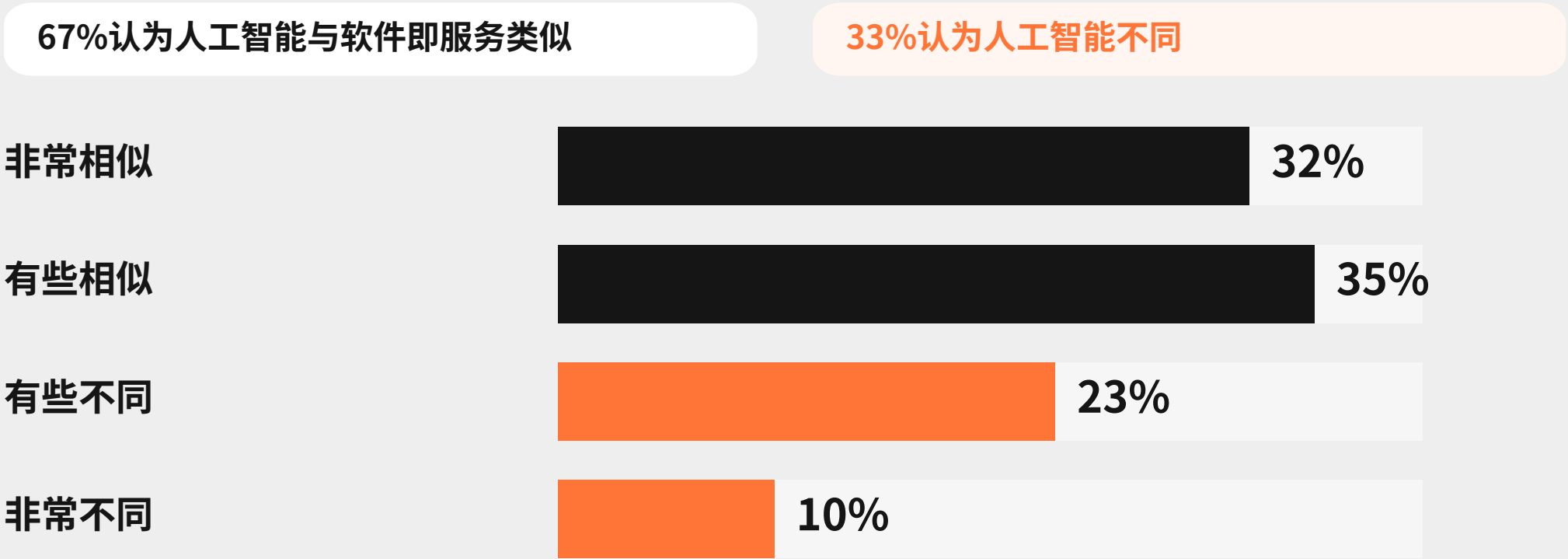
06 人工智能与传统软件即服务

当人工智能参与时，买家对待方式有何不同，以及他们仍然如何看待相同之处。

三分之一的买家将人工智能定价视为一个完全不同的类别

买家对人工智能与传统软件即服务的认知

买家占比 · n = 257



客户原话

- “必须以某种方式考虑使用量。”
积分、代币等必须透明。
- “我们寻找针对AI的特定保护性语言，确保我们的数据不会被供应商或任何大型语言模型使用。”
- “仍在尝试评估价值与成本——这仍然只是一个工具。”

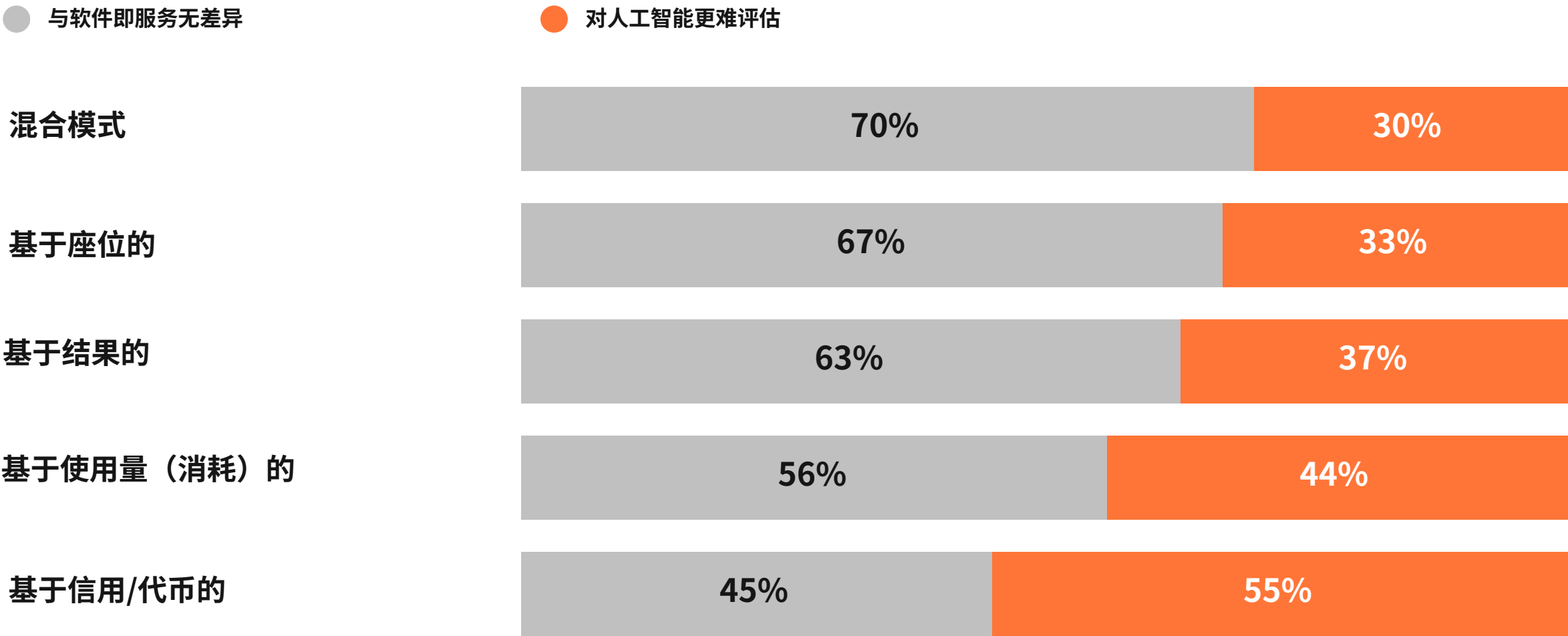
SaaS（软件即服务）：您通过订阅在线访问的软件，无需购买并安装在计算机上。供应商负责托管、更新和维护，您只需登录使用即可。例如：Gmail、Zoom、Slack、Salesforce。

三分之二的买家对AI定价抱有与传统SaaS相似的期望，因此基于用户或基于使用量等SaaS式定价模式会让他们感到熟悉。然而，这种认知并非普遍存在，也并非总是特别强烈，因此仍有尝试和探索的空间。

超过半数买家认为，基于信用的定价模式在人工智能领域比软件即服务领域更难评估

定价模式评估难度：人工智能 vs 软件即服务 · 样本量 = 292

买家占比



最难评估

55%

基于积分/代币的

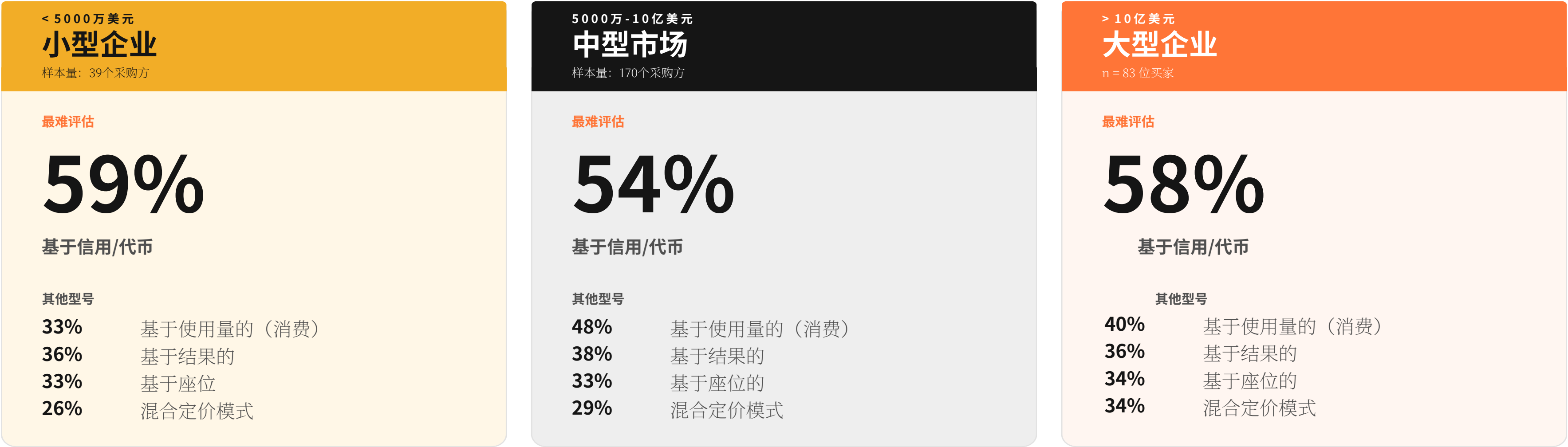
AI比SaaS更难评估

买家对于信用额度成本或所购内容缺乏AI前的基准参照，因此他们默认将其视为“风险过高而无法预测”

采用消费指标（信用/代币）作为计量单位的定价模型在AI购买环节可能给客户带来挑战，尽管这与供应成本存在关联。买家对于自身消费量缺乏AI前的基准参照，因此这种定价结构变得更难预测和评估。以座位数或固定部分为基础的模型能帮助用户更轻松地预测支出

当AI应用于不同收入规模的公司时，基于信用/代币和使用量的定价模式评估难度显著增加

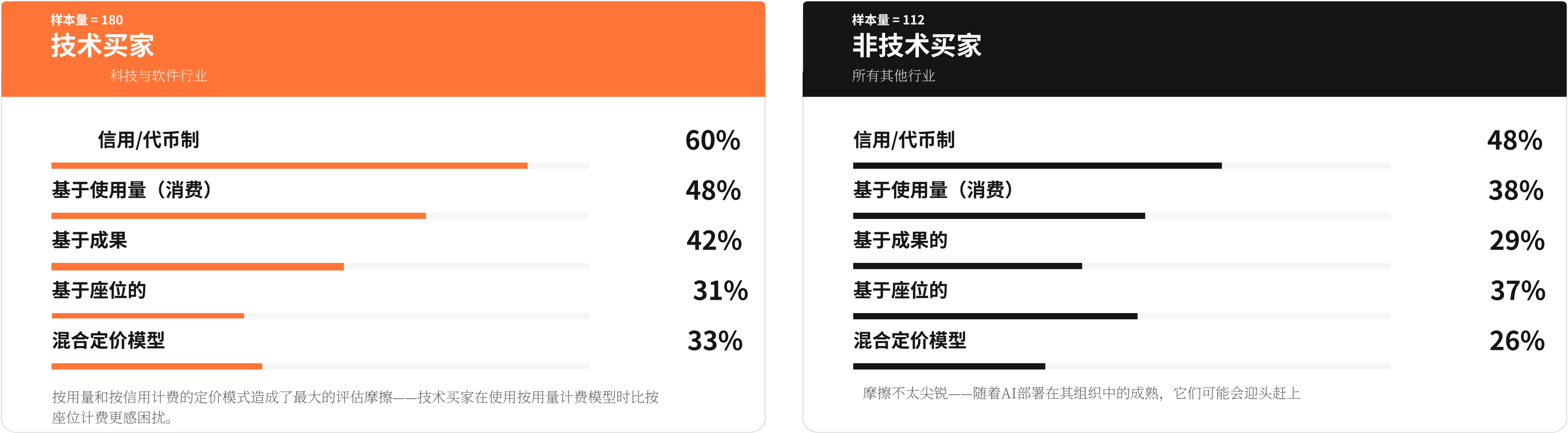
定价模式评估难度：AI 对比 SAAS · 按收入规模划分 · 样本量 = 292



基于消费的定价被认为在各公司规模中难以评估。这些困难源于定价模式本身的结构性问题，而非买家规模或成熟度。销售基于消费型AI产品的供应商可能因评估摩擦而失去部分交易，而非价格因素。许多买家在签约前无法可靠预测成本，因此尽可能减少这种摩擦可能有助于提高成交率。

技术买家认为相较非技术买家，评估人工智能的定价模式更为困难——只有基于席位的定价模式是例外

对人工智能与传统SaaS的定价模式评估难度提升 · 按行业划分 · 样本量 = 292



技术买家报告称，在所有按用量和按成果计费的定价模型中，都存在明显的AI专用评估摩擦。随着AI成本模型和定价成熟度的提升，非技术买家最终可能会有更强烈的反应。随着非技术类AI部署的成熟，他们对评估难度的感知可能会向上趋同于技术买家的水平。

大型买家拒绝可变定价，并将人工智能视为独立的定价类别

人工智能定价摩擦按细分市场 · 按收入层级 · 样本量=257



异常高支出影响所有买家群体，但大型买家更倾向于对此有明确的定价偏好。他们往往拒绝纯可变定价模式，并将人工智能视为独立的定价类别。依赖纯可变定价的供应商在最大型、最高价值的交易中可能面临困境。

章节摘要

人工智能与传统软件即服务

人工智能买家认为，基于信用额度和使用量的定价模式比基于座位、成果或混合定价模式更难评估。这一现象在不同规模的企业中普遍存在，只是程度有所不同。与其他行业相比，科技行业的买家尤其倾向于认为这些定价模式难以评估。

这些偏好趋势，加上相当一部分人工智能买家最终支出超出初始预期的事实，使得价格可预测性成为人工智能产品的关键考量因素。

这对买家意味着什么

- 信用/代币定价比其SaaS同类产品更难理解。
- 推动供应商尽可能以SaaS熟悉的术语（按席位、按 workflows、按固定单元订阅）传达消费定价。
- 随着财务部门在销售流程中参与度降低，预计AI评估摩擦将增加。

这对供应商意味着什么

- 寻找让买家理解并可预测定价的方法。买家非常重视了解他们将支付多少，至少是近似值。
- 注意密切关注支出。具有可变成本的可预测定价可能导致利润波动。
- 考虑你的销售对象。不同类型的买家可能有不同的优先事项。

55%的受访者认为基于信用额度的定价模式在人工智能领域比在SaaS领域更难评估

比在SaaS领域更难

44%的受访者对基于使用量的定价模式表达了相同看法

10-13%

科技行业与非科技行业买家在价格评估上的差异

与非技术买家相比，技术买家认为基于使用量和结果的定价评估难度显著更高。

60%

收入超过10亿美元的买家中，选择可变定价作为最不偏好选项

这种偏好明显强于对小公司的偏好。

关于本报告

本报告基于对296名SaaS参与者的调查，由Pricing I/O与Benchmarkit合作开展。

关于Pricing I/O

Pricing I/O是一家专注于SaaS和AI定价策略的公司，帮助B2B软件企业设计包装、定价和盈利策略。自2019年以来，我们已服务超过400家SaaS和AI领域的客户。
我们的客户覆盖SaaS和AI领域，数量超过400家。

www.pricingio.com

[免费定价记分卡测验](#)



致谢

我们感谢那些贡献时间和见解的SaaS买家和运营商，正是他们的支持使这项研究成为可能。

关于Benchmarkit

Benchmarkit是B2B SaaS运营指标和基准的领先来源，提供数据驱动的洞察，帮助SaaS领导者做出更明智的决策。

www.benchmarkit.ai

