

2026–2027 学年短学期
课程综合实践 II：数据要素交易基础

Lec 10: 数据定价的其他视角

吴一航 yhwu_is@zju.edu.cn

浙江大学计算机科学与技术学院



CONTENT

目录

1. 隐私与外部性视角

2. 信息提取机制

3. 诚实数据获取机制

考虑个体 i 拥有一条隐私数据（例如是否患有某种疾病等），他在纠结是否以一定价格出售自己的这条隐私数据。假设数据买方是一个可信的数据中心，承诺对 i 使用满足 ε_i -差分隐私的机制 M 保护其隐私。

- i 关心出售的价格是否能补偿隐私泄漏带来的自身效用降低，因此隐私泄漏带来的效用降低自然也就成为了隐私成本。

考虑个体 i 拥有一条隐私数据（例如是否患有某种疾病等），他在纠结是否以一定价格出售自己的这条隐私数据。假设数据买方是一个可信的数据中心，承诺对 i 使用满足 ε_i -差分隐私的机制 M 保护其隐私。

- i 关心出售的价格是否能补偿隐私泄漏带来的自身效用降低，因此隐私泄漏带来的效用降低自然也就成为了隐私成本。

为了衡量隐私泄漏带来的效用降低，首先需要定义 i 的效用函数 u_i ；

- 自然地，可以将 u_i 定义在所有未来可能发生的事件 $a \in A$ 上，其中 A 表示未来可能发生的事件集合；
- 机制 M 的输出会影响未来事件发生的概率，令 $f : \text{Range}(M) \rightarrow \Delta(A)$ 表示从机制输出到事件发生概率的映射。因此 $f(M(D))$ 可以表示当机制作用于数据库 D 时，每个事件 $a \in A$ 发生的概率。

举一个例子理解上述定义：

- 设 D_1 是包含 i 是否患有某种疾病的隐私数据的数据库， D_2 是不包含该条数据的数据库；
- A 表示未来 i 可能接受的医疗保险的价格；
- 机制 M 保证数据库上的查询满足 ϵ -差分隐私，保险公司在看到机制 M 返回的结果后会决定 i 的医疗保险价格；
- 因此 $f(M(D_1))$ 的含义就是当保险公司在看到数据库 D_1 的查询结果后为 i 提供不同保费的概率分布。

举一个例子理解上述定义：

- 设 D_1 是包含 i 是否患有某种疾病的隐私数据的数据库， D_2 是不包含该条数据的数据库；
- A 表示未来 i 可能接受的医疗保险的价格；
- 机制 M 保证数据库上的查询满足 ε -差分隐私，保险公司在看到机制 M 返回的结果后会决定 i 的医疗保险价格；
- 因此 $f(M(D_1))$ 的含义就是当保险公司在看到数据库 D_1 的查询结果后为 i 提供不同保费的概率分布。

在这一场景下， i 关心的问题是出售自己的患病数据能否补偿他在未来医疗保险中可能遭遇的价格上升带来的效用下降。形式化地，而当 i 不出售自己的数据时（即数据库为 D_2 ），其（期望）效用为

$$v_i = \mathbb{E}_{a \sim f(M(D_2))} [u_i(a)] = \sum_{a \in A} u_i(a) \cdot \mathbb{P}(f(M(D_2)) = a),$$

当 i 出售自己的数据时（即数据库为 D_1 ），其（期望）效用为

$$v_i' = \mathbb{E}_{a \sim f(M(D_1))} [u_i(a)] = \sum_{a \in A} u_i(a) \cdot \mathbb{P}(f(M(D_1)) = a).$$

假设 M 是一个满足 ε_i -差分隐私的机制，根据差分隐私的后处理性质， $f \circ M$ 也是满足 ε_i -差分隐私的，故有

$$\begin{aligned} v_i' &= \sum_{a \in A} u_i(a) \cdot \mathbb{P}(f(M(D_1)) = a) \\ &\leq \sum_{a \in A} u_i(a) \cdot \exp(\varepsilon_i) \mathbb{P}(f(M(D_2)) = a) \\ &= \exp(\varepsilon_i) \cdot v_i \end{aligned}$$

同理也可以证明 $v_i \leq \exp(\varepsilon_i) \cdot v_i'$ 。

假设 M 是一个满足 ε_i -差分隐私的机制，根据差分隐私的后处理性质， $f \circ M$ 也是满足 ε_i -差分隐私的，故有

$$\begin{aligned} v_i' &= \sum_{a \in A} u_i(a) \cdot \mathbb{P}(f(M(D_1)) = a) \\ &\leq \sum_{a \in A} u_i(a) \cdot \exp(\varepsilon_i) \mathbb{P}(f(M(D_2)) = a) \\ &= \exp(\varepsilon_i) \cdot v_i \end{aligned}$$

同理也可以证明 $v_i \leq \exp(\varepsilon_i) \cdot v_i'$ 。

- 由此可见，只要数据中心承诺使用满足 ε_i -差分隐私的机制 M ，那么数据卖家 i 出售自己数据后的效用 v_i 受到的影响 $(v_i - v_i')$ 不会超过 $(1 - \exp(-\varepsilon_i))v_i$ ；
- 因此可以将隐私成本函数写为 $c(v_i, \varepsilon_i) = (1 - \exp(-\varepsilon_i))v_i$ ；
- 当 ε_i 较小时，则可以更进一步地近似为 $c(v_i, \varepsilon_i) = \varepsilon_i v_i$ 。这是一个线性函数，表明用户的隐私成本随着隐私泄漏的增加而线性增加。

与数据价值分类相同，隐私成本也可分成**内在成本**和**工具成本**两部分：

- 此前定义隐私成本的角度是工具成本的角度，即泄露隐私后可能会因为价格歧视或其他原因带来多少损失；
- 然而在很多场景下，个人不愿意透露一些信息并不出于披露信息后会导致自己遭受经济损失，而是出于内在的不情愿的心理（内在成本）。

与数据价值分类相同，隐私成本也可分成**内在成本**和**工具成本**两部分：

- 此前定义隐私成本的角度是工具成本的角度，即泄露隐私后可能会因为价格歧视或其他原因带来多少损失；
- 然而在很多场景下，个人不愿意透露一些信息并不出于披露信息后会导致自己遭受经济损失，而是出于内在的不情愿的心理（内在成本）。

最后，读者可以回想一下自己的真实经历：我们真的有这么重视隐私成本吗？答案对于很多人而言是否定的：**即使每个人都会表面上认为个人数据的隐私保护非常重要，都表现出对自己的隐私数据被滥用的担忧，但事实上很多人都在无时无刻地泄露自己的隐私信息，这一对隐私的态度和行为不匹配的现象被称为隐私悖论”。**

一个很生动的例子是，考虑如下两个问题：

1. 给你一个披萨，你是否愿意提供自己的邮箱地址；
2. 某一平台已经拥有你的邮箱地址，平台要求你支付一个披萨的价钱来保证你的邮箱不会被进一步公开，你是否愿意支付。

非常现实地，很多人会在第一个问题上选择愿意，而在第二个问题上选择不愿意。

假设你是一个数据分析师，希望收集数据统计某一人群中的吸烟率、肥胖率或患某种疾病的比例，你希望以尽可能低的成本获得比较高的准确率。

数据分析师有可能的两个目标：

1. 有预算约束，希望最大化估计的准确度；
2. 有估计准确度的目标，希望在实现目标的前提下最小化支付。

如何为这两个目标设计相应的机制？

假设你是一个数据分析师，希望收集数据统计某一人群中的吸烟率、肥胖率或患某种疾病的比例，你希望以尽可能低的成本获得比较高的准确率。

数据分析师有可能的两个目标：

1. 有预算约束，希望最大化估计的准确度；
2. 有估计准确度的目标，希望在实现目标的前提下最小化支付。

如何为这两个目标设计相应的机制？

困难：每个人的隐私成本和隐私数据是相关的：

- 例如吸烟 / 过于肥胖的人会认为自己的数据隐私成本很高，如果只用固定价格获取数据，会导致统计偏差；
- 如果隐私成本的先验信息完全未知，目标很难达到；如果隐私成本先验分布已知，可以设计相应的机制（见教材相关章节）。

- 微观视角：数字平台机制设计
 - 好处：平台使用隐私数据可以**优化推荐算法**，提高用户福利
 - 坏处：平台使用隐私数据可以**实施价格歧视**（现实中不多见？）
 - 权衡：如何设计隐私数据使用政策权衡上述利弊

- 微观视角：数字平台机制设计
 - 好处：平台使用隐私数据可以**优化推荐算法**，提高用户福利
 - 坏处：平台使用隐私数据可以**实施价格歧视**（现实中不多见？）
 - 权衡：如何设计隐私数据使用政策权衡上述利弊
- 宏观视角：非竞争性与数据产权
 - 数据具有非竞争性，可以同时被多个企业使用；
 - 自动驾驶车厂的例子：自动驾驶车厂需要驾驶员驾驶习惯等数据增强算法，增强算法后可以提高产出，产出提高后可以收集到更多数据，从而数据产生了**乘数效应**；
 - 产权：如果数据产权归厂商所有，厂商会担心其他公司拥有后竞争加剧，因此数据共享少；如果产权归个人所有，个人可以权衡自己的隐私和卖数据的收益，产生更多的数据共享；
 - 数据和 idea 的比较：都属于信息产品，都有公共物品属性（非竞争性），**但数据相比于 idea 更容易传播（idea 需要教育），但数据的传播也更容易受到管控（加密等）。**

外部性问题通常分为两个角度讨论：

- 买家视角
 - 不同买家之间具有竞争关系，例如是两个自动驾驶车厂，购买数据增强自动驾驶算法，此时一方购买数据增强算法后可能会影响另一方的市场占有率；
 - 不同买家是股票投资者，购买某些股票的调研报告后会导致供求关系改变，股票价格改变从而影响其他人的收益；

外部性问题通常分为两个角度讨论：

- 买家视角

- 不同买家之间具有竞争关系，例如是两个自动驾驶车厂，购买数据增强自动驾驶算法，此时一方购买数据增强算法后可能会影响另一方的市场占有率；
- 不同买家是股票投资者，购买某些股票的调研报告后会导致供求关系改变，股票价格改变从而影响其他人的收益；

- 卖家视角

- 假设两个消费者隐私数据强相关（他们年龄、性别等背景类似，习惯也类似），那么其中一个人出售自己的隐私数据后，另一个人的隐私数据就不值钱了；
- 这会导致隐私数据均衡价格下降，以及隐私数据使用的泛滥。

感兴趣的同学可以参考教材相关内容。

CONTENT

目录

1. 隐私与外部性视角

2. 信息提取机制

3. 诚实数据获取机制

接下来我们的讨论目标是如何让数据提供者诚实地提供自己的数据。我们首先研究更一般的情况，即**如何得到参与人对一个问题的诚实答复**。

- 自然想到此前关于拍卖诚实性的讨论。但拍卖（或一般的机制设计）**更关注参与人的汇报与自身效用相关的场景**，例如拍卖中参与人的汇报会决定其是否赢下拍卖以及参与人的支付（**内生激励**）；
- 本节更关注能否**设计一个奖励机制**，使得能根据参与人的回答决定给每个参与人的奖励从而让参与人诚实回答问题，也就是说回答问题给出不同答案对自己的内部效用没有影响，但不同的答案会导致获得的外部奖励不同（**外生激励**）。

为方便后续的讨论，首先需要介绍一些信息论的基础知识。首先介绍最基本的“熵”，也称为“香农熵”。

定义

设离散随机变量 X 的取值集合为 $\mathcal{X}$ ，则其熵（**entropy**）定义为

$$H(X) = - \sum_{x \in \mathcal{X}} \mathbb{P}(X = x) \log \mathbb{P}(X = x).$$

熵衡量的是随机变量 X 的不确定性：当 X 几乎确定时，熵较小；当 X 更“分散”时，熵更大。后续讨论只涉及离散情形，因此在此不展开连续型随机变量的微分熵。

若有两个随机变量 X, Y ，则自然可以定义联合熵与条件熵：

定义

联合熵 (joint entropy) 定义为

$$H(X, Y) = - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) \log \mathbb{P}(X = x, Y = y).$$

条件熵 (conditional entropy) 定义为

$$\begin{aligned} H(Y \mid X) &= \sum_{x \in \mathcal{X}} \mathbb{P}(X = x) H(Y \mid X = x) \\ &= - \sum_{x \in \mathcal{X}} \mathbb{P}(X = x) \sum_{y \in \mathcal{Y}} \mathbb{P}(Y = y \mid X = x) \log \mathbb{P}(Y = y \mid X = x) \\ &= - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) \log \mathbb{P}(Y = y \mid X = x). \end{aligned}$$

联合熵和条件熵之间满足基本的链式法则：

命题

$$H(X, Y) = H(X) + H(Y \mid X).$$

证明：

$$\begin{aligned} & H(X, Y) \\ &= - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) \log \mathbb{P}(X = x, Y = y) \\ &= - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) (\log \mathbb{P}(X = x) + \log \mathbb{P}(Y = y \mid X = x)) \\ &= - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) \log \mathbb{P}(X = x) + H(Y \mid X) \\ &= H(X) + H(Y \mid X). \end{aligned}$$

□

互信息 (mutual information) 是信息论中另一个非常重要的概念。

定义

两个随机变量 X, Y 的互信息定义为

$$I(X; Y) = H(X) - H(X | Y) = H(Y) - H(Y | X).$$

等价地,

$$I(X; Y) = \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} \mathbb{P}(X = x, Y = y) \log \left(\frac{\mathbb{P}(X = x, Y = y)}{\mathbb{P}(X = x) \mathbb{P}(Y = y)} \right).$$

- 互信息衡量的是：知道 Y 之后，关于 X 的不确定性减少了多少。若 X 与 Y 独立，则互信息等于 0；
- $I(X; Y) = I(Y; X)$ ，因此互信息是对称的；
- $I(X; X) = H(X)$ ，因此熵也被称为自信息。

互信息还有链式法则和非负性，下面逐一展开讨论。首先是链式法则，为了给出链式法则的表述，需要先给出条件互信息的定义：

定义

给定随机变量 Z 时， X 与 Y 的**条件互信息**定义为

$$I(X; Y \mid Z) = H(X \mid Z) - H(X \mid Y, Z).$$

互信息链式法则

$$I(X, Y; Z) = I(X; Z) + I(Y; Z \mid X).$$

证明：由定义与熵的链式法则，

$$\begin{aligned} I(X, Y; Z) &= H(X, Y) - H(X, Y \mid Z) \\ &= H(X) + H(Y \mid X) - H(X \mid Z) - H(Y \mid X, Z) \\ &= I(X; Z) + I(Y; Z \mid X). \end{aligned}$$

□

下面讨论互信息的非负性，为证明这一性质，首先引入**相对熵 (relative entropy)**（也称为**KL 散度 (Kullback-Leibler divergence)**）的定义：

定义

设 p, q 是同一支撑集 $\mathcal{X}$ 上的两个概率分布，则它们的 KL 散度定义为

$$D_{\text{KL}}(p \parallel q) = \sum_{x \in \mathcal{X}} p(x) \log \left(\frac{p(x)}{q(x)} \right).$$

KL 散度的非负性

$D_{\text{KL}}(p \parallel q) \geq 0$ ，当且仅当 $p = q$ 时取等号。

证明：由于对数函数是凹函数，由 Jensen 不等式，

$$-\sum_{x \in \mathcal{X}} p(x) \log \left(\frac{q(x)}{p(x)} \right) \geq -\log \left(\sum_{x \in \mathcal{X}} p(x) \frac{q(x)}{p(x)} \right) = -\log 1 = 0.$$

等号成立当且仅当 $\frac{q(x)}{p(x)}$ 为常数。由于 p, q 是概率分布，故只能 $p = q$ 。 $\square$

根据上述命题，KL 散度自然可以用来**衡量两个概率分布之间的差异**。但 KL 散度并不是一个距离函数，因为它**不满足对称性和三角不等式**，但在机器学习等领域，KL 散度仍然被广泛用于衡量两个概率分布之间的差异。

证明：由于对数函数是凹函数，由 Jensen 不等式，

$$-\sum_{x \in \mathcal{X}} p(x) \log \left(\frac{q(x)}{p(x)} \right) \geq -\log \left(\sum_{x \in \mathcal{X}} p(x) \frac{q(x)}{p(x)} \right) = -\log 1 = 0.$$

等号成立当且仅当 $\frac{q(x)}{p(x)}$ 为常数。由于 p, q 是概率分布，故只能 $p = q$ 。 $\square$

根据上述命题，KL 散度自然可以用来**衡量两个概率分布之间的差异**。但 KL 散度并不是一个距离函数，因为它**不满足对称性和三角不等式**，但在机器学习等领域，KL 散度仍然被广泛用于衡量两个概率分布之间的差异。

不难看出互信息可以写成 KL 散度的形式：

$$I(X; Y) = D_{\text{KL}}(p(X, Y) \parallel p(X)p(Y)).$$

因此立即得到 $I(X; Y) \geq 0$ 。

在信息论介绍最后，给出一个非常重要的结论：

数据处理不等式 (data processing inequality)

若随机变量满足马尔可夫链 $X \rightarrow Y \rightarrow Z$ ，即在给定 Y 的条件下， X 与 Z 条件独立，则

$$I(X; Y) \geq I(X; Z).$$

直观上， Z 是由 Y 再加工得到的信息，因此它不会比 Y 含有更多关于 X 的信息。

证明：由马尔可夫链条件，有 $I(X; Z | Y) = 0$ 。由互信息链式法则，

$$I(X; Y, Z) = I(X; Y) + I(X; Z | Y) = I(X; Z) + I(X; Y | Z).$$

因为条件互信息非负， $I(X; Y) = I(X; Z) + I(X; Y | Z) \geq I(X; Z)$ 。 $\square$

在信息论介绍最后，给出一个非常重要的结论：

数据处理不等式 (data processing inequality)

若随机变量满足马尔可夫链 $X \rightarrow Y \rightarrow Z$ ，即在给定 Y 的条件下， X 与 Z 条件独立，则

$$I(X; Y) \geq I(X; Z).$$

直观上， Z 是由 Y 再加工得到的信息，因此它不会比 Y 含有更多关于 X 的信息。

证明：由马尔可夫链条件，有 $I(X; Z | Y) = 0$ 。由互信息链式法则，

$$I(X; Y, Z) = I(X; Y) + I(X; Z | Y) = I(X; Z) + I(X; Y | Z).$$

因为条件互信息非负， $I(X; Y) = I(X; Z) + I(X; Y | Z) \geq I(X; Z)$ 。 $\square$

这一不等式与后续讨论直接相关：不诚实报告本质上就是先把自己的真实信息做一次加工再汇报，而数据处理不等式说明：**这种加工不会增加你和别人报告之间的相关信息。**

考虑如下场景：向一个专家询问对随机事件 X 的预测，例如“明天是否下雨”，“某只股票是否上涨”等。我们希望设计一个奖励机制，使得专家诚实地报告自己的真实概率判断。

在这一情况下，真实结果会在事后出现，因此可以通过“预测”和“结果”比较来决定专家的奖励（即评分规则）：

1. 专家报告一个概率分布 $p \in \Delta(X)$ ；
2. 未来真实结果 $x \in X$ 出现；
3. 平台按函数 $S(p, x)$ 给奖励。

若专家的真实信念是 $q \in \Delta(X)$ ，则其期望奖励为

$$S(p; q) = \mathbb{E}_{x \sim q}[S(p, x)] = \sum_{x \in X} q(x) S(p, x).$$

定义

若对任意真实信念 q 和任意报告 p 都有

$$S(q; q) \geq S(p; q),$$

则称 S 为**适当评分规则 (proper scoring rule)**。若当 $p \neq q$ 时严格不等，则称 S 为**严格适当评分规则 (strictly proper scoring rule)**。

因此，适当评分规则的本质是：**诚实报告自己的真实信念可以获得最大的期望奖励**。注意常数奖励（即无论专家报告什么都给同样的奖励）是适当评分规则，但不是严格适当评分规则。

定义

若对任意真实信念 q 和任意报告 p 都有

$$S(q; q) \geq S(p; q),$$

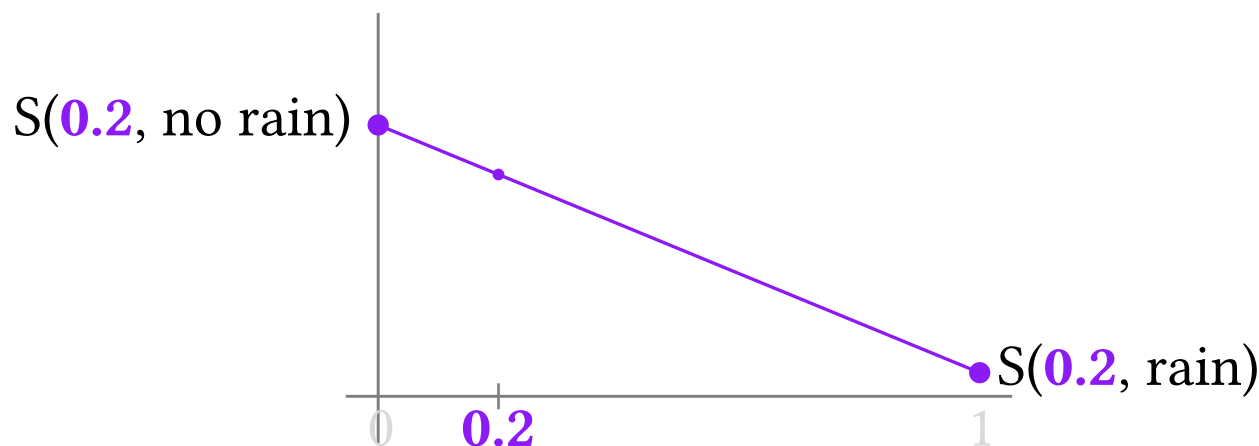
则称 S 为**适当评分规则 (proper scoring rule)**。若当 $p \neq q$ 时严格不等，则称 S 为**严格适当评分规则 (strictly proper scoring rule)**。

因此，适当评分规则的本质是：**诚实报告自己的真实信念可以获得最大的期望奖励**。注意常数奖励（即无论专家报告什么都给同样的奖励）是适当评分规则，但不是严格适当评分规则。

为了构造适当评分规则，可以先做一个启发式分析。考虑“是否下雨”的例子，要求专家报告 $p \in [0, 1]$ 表示明天下雨的概率。此时真实信念为 $q \in [0, 1]$ 。考虑专家报告 $p = 0.2$ 的情况。根据期望得分的定义，

$$\begin{aligned} S(0.2; q) &= q \cdot S(0.2, \text{rain}) + (1 - q) \cdot S(0.2, \text{no rain}) \\ &= q \cdot (S(0.2, \text{rain}) - S(0.2, \text{no rain})) + S(0.2, \text{no rain}). \end{aligned}$$

因此，固定报告 $p = 0.2$ 后， $S(0.2; q)$ 是关于 q 的一条线段：

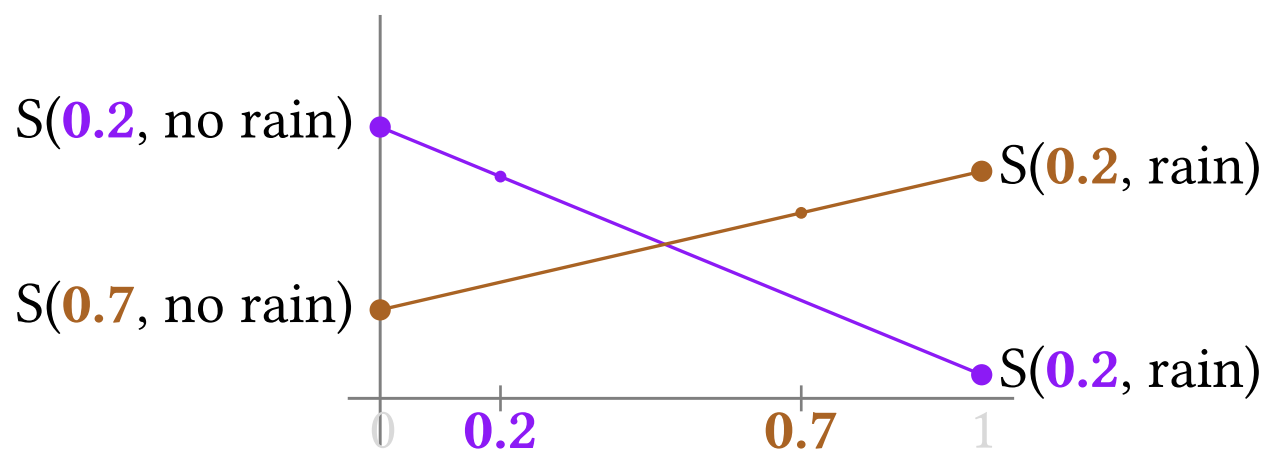


再看报告 $p = 0.7$ 的情况。我们希望评分规则是适当的，因此至少要满足

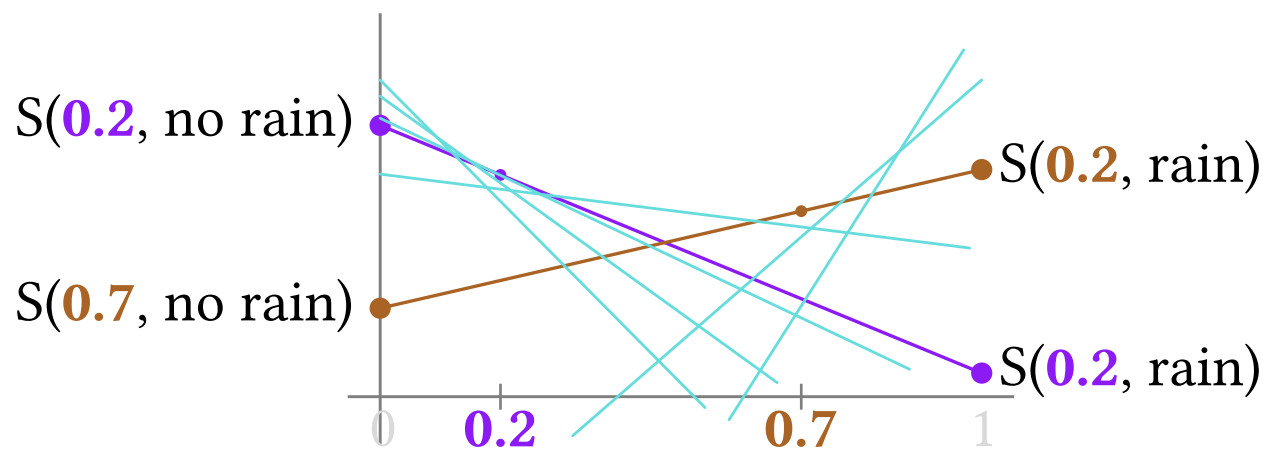
$$S(0.7; 0.7) \geq S(0.2; 0.7),$$

$$S(0.2; 0.2) \geq S(0.7; 0.2).$$

这意味着当真实信念是 0.7 时，报告 0.7 的期望得分不能低于报告 0.2；当真实信念是 0.2 时，报告 0.2 的期望得分不能低于报告 0.7。因此， $S(0.7; q)$ 可以合理地画成另一条线段：



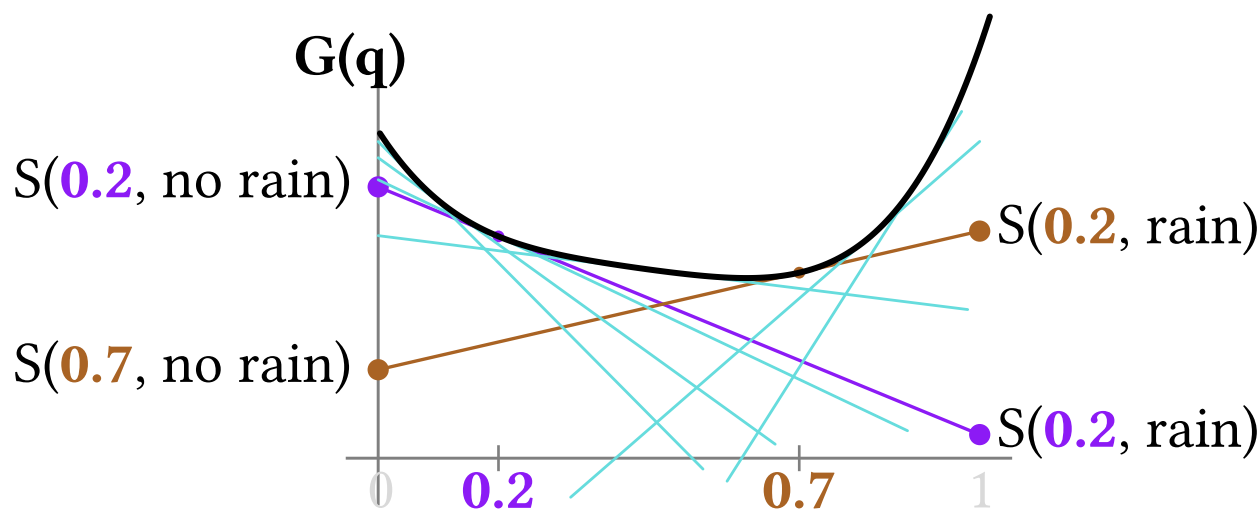
类似地，不断改变报告值 p ，就会得到许多条关于 q 的线段。它们共同描述了“不同报告在不同真实信念下的期望得分”：



现在看这些线段的上包络。实际上这条上包络线在任意 q 点的取值等于使得 $S(p; q)$ 最大的 p 对应的 $S(p; q)$ 的值，结合 S 是适当评分规则可知，这条上包络线就是

$$G(q) := S(q; q) = \max_{p \in [0,1]} S(p; q).$$

由于 G 是多个线性函数取最大得到的，而线性函数是凸函数，因此 G 也是凸函数。



基于上述启发式推导的直观，每条线段 $S(p; q)$ 在 $q = p$ 处与 G 的函数图像相切，因此可以将 $S(p; q)$ 写成

$$S(p; q) = S(p; p) + (S(p; q) - S(p; p)) = G(p) + \nabla G(p) \cdot (q - p).$$

其中 $\nabla G(p)$ 是 G 在 p 处的梯度。由此可以给出下面关于（严格）适当评分规则等价条件的命题：

命题

评分规则 S 是（严格）适当的当且仅当存在一个（严格）凸函数 $G : \Delta(X) \rightarrow \mathbb{R}$ ，使得

$$S(p, x) = G(p) + \nabla G(p) \cdot (\delta_x - p),$$

其中 δ_x 是在 x 处取值为 1，其他取值为 0 的向量。

证明：“仅当”的部分延续上述启发式推导的过程的直观即可，下面证明“当”的部分。假设存在一个（严格）凸函数 G 使得上述等式成立，则

$$\begin{aligned} S(p; q) &= G(p) + \nabla G(p) \cdot (q - p), \\ S(q; q) &= G(q). \end{aligned}$$

由凸函数的一阶条件，

$$G(q) \geq G(p) + \nabla G(p) \cdot (q - p) = S(p; q),$$

因此 $S(q; q) \geq S(p; q)$ 。若 G 严格凸，则当 $p \neq q$ 时不等式严格成立。 $\square$

最常用的（严格）适当评分规则是对数评分规则。

命题

对数评分规则 $S(p, x) = \log p(x)$ 是严格适当评分规则。

证明： 设专家真实信念为 q ，则

$$\begin{aligned} S(q; q) - S(p; q) &= \sum_{x \in X} q(x) \log q(x) - \sum_{x \in X} q(x) \log p(x) \\ &= \sum_{x \in X} q(x) \log \left(\frac{q(x)}{p(x)} \right) = D_{\text{KL}}(q \parallel p) \geq 0, \end{aligned}$$

并且当且仅当 $p = q$ 时取等号。 □

当然也可以找到对数评分规则对应的凸函数 G 。由定义，不难验证 $G = \sum_{x \in X} p(x) \log p(x)$ 是对数评分规则对应的凸函数。

在评分规则的讨论中，我们要求专家给出对一个预测问题的诚实回答，这一预测问题在未来终会给出一个答案，因此可以通过比较答案和专家预测给出奖励。然而，有的时候问题并非预测问题，也没有标准答案，例如要求一群人参与问卷调查：

例

设两位同学对浙大食堂的真实评价分别为 $A, B \in \{\text{好}, \text{差}\}$ ，联合分布为

$$\mathbb{P}(A = \text{好}, B = \text{好}) = 0.5, \quad \mathbb{P}(A = \text{差}, B = \text{差}) = 0.02,$$

$$\mathbb{P}(A = \text{好}, B = \text{差}) = 0.24, \quad \mathbb{P}(A = \text{差}, B = \text{好}) = 0.24.$$

浙大新宇应当如何设计奖励，使两位同学都愿意诚实回答？

在上述例子中，我们并非希望每个人给出对一个问题的预测，而是希望每个人给出对一个问题的诚实回答。由于没有标准答案，因此无法通过比较答案和参与人的回答来提供奖励。这一问题称为**无验证的信息提取 (information elicitation without verification)** 问题。

在给出这一问题的解之前，首先从上述例子中抽象出更一般的问题：

- 两个参与人参与一项调查，两人可能的汇报记为 $A = \{A_1, \dots, A_m\}$, $B = \{B_1, \dots, B_n\}$;
- 每个参与人内心有一个诚实的答案，如 $A = A_i, B = B_j$ ，但**奖励机制设计者不知道具体哪个答案是诚实的，只知道诚实答案的先验分布，且先验分布是共同知识**；
- 记两个参与人在调查中分别报告 $\bar{A}, \bar{B}$ 时，两个参与人能获得的奖励为 $r_1(\bar{A}, \bar{B}), r_2(\bar{A}, \bar{B})$ ，并且**奖励函数是公开的**；
- 针对奖励函数，**两个理性的参与人可以设计报告策略** $\sigma_1(A_i)$ 和 $\sigma_2(B_j)$ （分别表示第一个参与人真实类型为 A_i 时的报告和第二个参与人真实类型为 B_j 时的报告）。诚实报告的策略为 $\sigma_1(A_i) = A_i$ 和 $\sigma_2(B_j) = B_j$ 。

以上定义了两个参与人在奖励机制下的**不完全信息博弈**，下面就可以讨论如何**设计奖励机制使得诚实报告是一个贝叶斯均衡**（注意，接下来的讨论均只涉及纯策略均衡）。

一个简单的想法是，当两个人报告相同的答案时，给两个人奖励 1，否则不给予奖励，即

$$r_1(\bar{A}, \bar{B}) = r_2(\bar{A}, \bar{B}) = \begin{cases} 1 & \text{if } \bar{A} = \bar{B} \\ 0 & \text{otherwise} \end{cases}.$$

但这并不诚实。假设同学 1 诚实报告，而同学 2 的真实类型是“差”，则

$$\mathbb{P}(A = \text{好} \mid B = \text{差}) = \frac{0.24}{0.24 + 0.02} = \frac{12}{13}.$$

同学 2 若诚实报告“差”，其期望收益为 $0 \cdot \frac{12}{13} + 1 \cdot \frac{1}{13} = \frac{1}{13}$ ，若谎报“好”，其期望收益为 $1 \cdot \frac{12}{13} + 0 \cdot \frac{1}{13} = \frac{12}{13}$ 。所以一致性奖励会诱导少数迎合多数，而不是诚实报告。

回忆适当评分规则保证了参与者诚实报告时能获得最大的期望奖励，但适当评分规则需要应用于可以验证真实结果的情况，在无验证信息提取的问题中并没有真实结果。

回忆适当评分规则保证了参与者诚实报告时能获得最大的期望奖励，但适当评分规则需要应用于可以验证真实结果的情况，在无验证信息提取的问题中并没有真实结果。

然而，一个非常巧妙的想法是，对于参与者 1，可以将其任务视为预测参与者 2 的汇报，反之对参与者 2 也可以将其任务视为预测参与者 1 的汇报，这样就产生了一个有真实结果的预测问题。具体奖励机制如下：

1. 参与者 1 报告 $\bar{A}$ ，参与者 2 报告 $\bar{B}$ ；
2. 计算在 $\bar{A}$ 的条件下参与者 2 报告 $\bar{B}$ 的概率分布 $p_{\bar{A}}(B)$ ，以及在 $\bar{B}$ 的条件下参与者 1 报告 $\bar{A}$ 的概率分布 $p_{\bar{B}}(A)$ ；
3. 令 $r_1(\bar{A}, \bar{B}) = S(p_{\bar{A}}(B), \bar{B})$ ， $r_2(\bar{A}, \bar{B}) = S(p_{\bar{B}}(A), \bar{A})$ ，其中 S 是严格适当评分规则。

回忆适当评分规则保证了参与者诚实报告时能获得最大的期望奖励，但适当评分规则需要应用于可以验证真实结果的情况，在无验证信息提取的问题中并没有真实结果。

然而，一个非常巧妙的想法是，对于参与者 1，可以将其任务视为预测参与者 2 的汇报，反之对参与者 2 也可以将其任务视为预测参与者 1 的汇报，这样就产生了一个有真实结果的预测问题。具体奖励机制如下：

1. 参与者 1 报告 $\bar{A}$ ，参与者 2 报告 $\bar{B}$ ；
2. 计算在 $\bar{A}$ 的条件下参与者 2 报告 $\bar{B}$ 的概率分布 $p_{\bar{A}}(B)$ ，以及在 $\bar{B}$ 的条件下参与者 1 报告 $\bar{A}$ 的概率分布 $p_{\bar{B}}(A)$ ；
3. 令 $r_1(\bar{A}, \bar{B}) = S(p_{\bar{A}}(B), \bar{B})$ ， $r_2(\bar{A}, \bar{B}) = S(p_{\bar{B}}(A), \bar{A})$ ，其中 S 是严格适当评分规则。

对于参与者 1，假设参与者 2 选择诚实汇报，根据严格适当评分规则的性质，只有当 $p_{\bar{A}}(B)$ 的确能表示参与者 1 对参与者 2 汇报的真实预测时，参与者 1 的奖励才能最大化，因此诚实报告自然成为了参与者 1 应对参与者 2 诚实报告的最优反应，参与者 2 同理。因此上述奖励机制保证了参与者诚实报告是一个贝叶斯均衡。

总结上述奖励机制的核心思想，就是利用一方的汇报来评估另一方的汇报，因此这一方法也被非常形象地称为**同伴预测 (peer prediction)**。这个机制非常简洁漂亮，但也有三个局限：

1. 它需要先验分布且是共同知识；
2. 诚实报告可能不是唯一均衡，例如在某些情况下可能存在“所有人都始终报好”的合谋均衡；
3. 上面只讨论了两人的机制；
 - 当然推广到多人（大众点评）虽然不难，只需随机匹配“同伴”即可，但前面两个问题仍然存在。

实际上，先验分布在同伴预测机制中的作用主要是用于计算参与人之间的后验概率分布，进而用于适当评分规则。因此一种自然的解决方案是，**让参与人直接报告自己的后验概率分布**。考虑 n 个人参与一项调查的情况，所有参与人可能汇报的集合一致，记为 S 。每个人都有一个真实信息 $S_i \in S$ ，考虑如下**贝叶斯吐真剂 (Bayesian truth serum, BTS)** 机制：

1. 每个参与人 i 报告 $\bar{S}_i$ 和对其他参与人报告估计的概率分布 $\bar{p}^i \in \Delta(S)$;
2. 计算所有汇报的概率分布的几何平均 $\bar{p}_s = \left(\prod_{i=1}^n \bar{p}_s^i \right)^{\frac{1}{n}}$ ，其中下标 s 表示 $s \in S$ 的概率；
3. 根据参与人汇报的信息 $\bar{S}_i$ 计算出参与人信息的经验分布 $\bar{\lambda} \in \Delta(S)$
4. 为每个参与人 i 提供奖励（其中 G 是严格适当评分规则）

$$\mathbb{E}_{S \sim \bar{\lambda}}[G(\bar{p}^i, S)] + \log \left(\frac{\bar{\lambda}_{\bar{S}_i}}{\bar{p}_{\bar{S}_i}} \right).$$

理解上述机制的关键就在于理解这一奖励表达式。

1. 第一项与同伴预测机制的思想类似，当参与人的数量 $n \rightarrow \infty$ 时，分布 λ 会趋向于真实的信息分布（如果其他参与人都诚实汇报），因此第一项在参与人 i 诚实报告对其他人的报告的预测时实现最大化；

1. 第一项与同伴预测机制的思想类似，当参与人的数量 $n \rightarrow \infty$ 时，分布 λ 会趋向于真实的信息分布（如果其他参与人都诚实汇报），因此第一项在参与人 i 诚实报告对其他人的报告的预测时实现最大化；
2. 对第二项，不难发现当 $\bar{\lambda}_{\bar{S}_i}$ 相比于 $\bar{p}_{\bar{S}_i}$ 越大时，参与人 i 的奖励越大。更具体地说，当信息 $\bar{S}_i$ 真实出现的概率高于大家认为这一信息会出现的概率时（或者说 $\bar{S}_i$ **出乎意料的普遍 (surprisingly common)** 时），对 $\bar{S}_i$ 的奖励会更高；

1. 第一项与同伴预测机制的思想类似，当参与人的数量 $n \rightarrow \infty$ 时，分布 λ 会趋向于真实的信息分布（如果其他参与人都诚实汇报），因此第一项在参与人 i 诚实报告对其他人的报告的预测时实现最大化；
2. 对第二项，不难发现当 $\bar{\lambda}_{\bar{S}_i}$ 相比于 $\bar{p}_{\bar{S}_i}$ 越大时，参与人 i 的奖励越大。更具体地说，当信息 $\bar{S}_i$ 真实出现的概率高于大家认为这一信息会出现的概率时（或者说 $\bar{S}_i$ **出乎意料的普遍 (surprisingly common)** 时），对 $\bar{S}_i$ 的奖励会更高；
 - 这背后的直观在于**激励拥有小众答案的参与人诚实报告**，例如问卷调查一个人是否拥有过 20 段感情经历，显然大部分人通常没有 20 段感情经历，并且认为别人有 20 段感情经历的概率很低，从而此时真的拥有 20 段感情经历的人诚实报告会获得比较高的奖励，因此会激励这部分小众人群诚实报告。

1. 第一项与同伴预测机制的思想类似，当参与人的数量 $n \rightarrow \infty$ 时，分布 λ 会趋向于真实的信息分布（如果其他参与人都诚实汇报），因此第一项在参与人 i 诚实报告对其他人的报告的预测时实现最大化；
2. 对第二项，不难发现当 $\bar{\lambda}_{\bar{S}_i}$ 相比于 $\bar{p}_{\bar{S}_i}$ 越大时，参与人 i 的奖励越大。更具体地说，当信息 $\bar{S}_i$ 真实出现的概率高于大家认为这一信息会出现的概率时（或者说 $\bar{S}_i$ **出乎意料的普遍 (surprisingly common)** 时），对 $\bar{S}_i$ 的奖励会更高；
 - 这背后的直观在于**激励拥有小众答案的参与人诚实报告**，例如问卷调查一个人是否拥有过 20 段感情经历，显然大部分人通常没有 20 段感情经历，并且认为别人有 20 段感情经历的概率很低，从而此时真的拥有 20 段感情经历的人诚实报告会获得比较高的奖励，因此会激励这部分小众人群诚实报告。

BTS 在参与人数量 $n \rightarrow \infty$ 时能保证诚实报告是贝叶斯均衡（均衡不一定唯一），证明技术性很强，因此略去。

回答问卷通常不需要付出努力，但很多现实任务不仅要诚实，还要努力，例如机器学习数据打标签和 MOOC 互评作业。在这些场景下，一个平台有一系列任务，可以通过分发给一群人来完成，这一过程被称为这类场景称为**众包 (crowdsourcing)**。此时平台希望同时实现 **(1) 参与人诚实报告自己的评估；(2) 参与人愿意投入足够努力，提升评估质量。**

- 注：众包的场景下的任务通常有标准答案，例如标签是否打对，作业好坏是否评判准确。但由于任务过多，平台获得任务标准答案的成本过高，因此众包问题也属于无验证的信息提取问题。

回答问卷通常不需要付出努力，但很多现实任务不仅要诚实，还要努力，例如机器学习数据打标签和 MOOC 互评作业。在这些场景下，一个平台有一系列任务，可以通过分发给一群人来完成，这一过程被称为这类场景称为**众包 (crowdsourcing)**。此时平台希望同时实现 **(1) 参与人诚实报告自己的评估；(2) 参与人愿意投入足够努力，提升评估质量。**

- 注：众包的场景下的任务通常有标准答案，例如标签是否打对，作业好坏是否评判准确。但由于任务过多，平台获得任务标准答案的成本过高，因此众包问题也属于无验证的信息提取问题。
- 基本思想：如果只奖励同一任务上的一致性，所有人都始终报告同一结果也能拿高分，因此最终的机制是：**奖励不同参与人在同一任务上的一致性，同时惩罚不同参与人在不同任务上的一致性。**
- 由于问题形式化涉及符号较多，此处不展开，仅讨论机制设计大致思想，感兴趣的同学可以参考教材相关内容。

前面已经看到几种不同的同伴预测机制：经典同伴预测、BTS、众包。他们的思想很直观，但构造相当不平凡。自然想法是寻找一个统一框架，把这些机制都解释成同一种思想，这就需要求助于数据处理不等式。

基于数据处理不等式的思想，可以**随机为每个参与人抽取一个参考评分者，然后将两者之间的互信息作为参与人的奖励**。直观上此时参与人不会选择不诚实汇报，因为不诚实汇报相当于对自己报告的随机变量做了“数据处理”，会导致与参考评分者之间的互信息降低，从而导致奖励减少。

形式化问题：假设有 n 个参与人参与一项问卷调查，所有人的可能答案都在同一集合 Σ 中。参与人 i 的诚实答案是其私有信息，构成随机变量 Ψ_i 。为了描述一般框架，首先需要假设随机变量 $\Psi_1, \dots, \Psi_n$ 的联合分布 Q 是共同知识，在之后运用这一框架到具体场景时可以放宽这一假设。

不难理解两个同支撑的随机变量之间的转化可以用一个转移矩阵 M 来表示。基于转移矩阵的定义，下面定义符合奖励机制设计目标的信息单调互信息 (**information-monotone mutual information**)：

定义

称一个互信息度量 I 是信息单调互信息，如果它满足：

- 对称性： $I(X; Y) = I(Y; X)$ ；
- 非负性： $I(X; Y) \geq 0$ ，且独立时为 0；
- 数据处理不等式：对任意转移矩阵 M ，都有 $I(M(X); Y) \leq I(X; Y)$ 。

基于上述讨论，我们可以进一步定义**互信息范式 (Mutual Information Paradigm, MIP)**：

定义

假设每个参与人 i 的真实信息的分布为 Ψ_i ，在调查中报告的分布为 $\hat{\Psi}_i$ 。使用均匀分布抽取一个参考参与人 $j \neq i$ ，其报告为 $\hat{\Psi}_j$ ，则参与人 i 的奖励为

$$I(\hat{\Psi}_i; \hat{\Psi}_j).$$

其中 I 是信息单调互信息。上述奖励机制被称为互信息范式。

利用信息单调互信息的定义，可以证明 MIP 保证诚实报告是占优策略均衡：若参与人 i 采用策略 M 处理自己的真实信号，则其期望奖励为

$$\sum_{j \neq i} \frac{1}{n-1} I(M(\Psi_i); \hat{\Psi}_j) \leq \sum_{j \neq i} \frac{1}{n-1} I(\Psi_i; \hat{\Psi}_j).$$

事实上，普通互信息以及更一般的 f -互信息都可以满足信息单调互信息的定义；这部分证明较技术，这里不展开。

尽管 MIP 拥有我们希望的性质，但需要注意的是，MIP 在严格意义上说并不是一个合理的奖励机制：

- MIP 要求参与人汇报自己信息的概率分布，但真正的机制只能允许参与人汇报自己的信息的实现值；
- 然而事实上，只需要设计机制使得参与人汇报的奖励是 MIP 给出的奖励的无偏估计量即可实现与 MIP 相同的性质；
- 下一节我们将针对诚实数据获取机制给出一个可实施的（只基于报告信息的实现值）奖励机制。

CONTENT

目录

1. 隐私与外部性视角

2. 信息提取机制

3. 诚实数据获取机制

考虑如下场景：一个统计学家希望获取一些数据集用于统计任务（训练贝叶斯机器学习模型），统计学家希望这些数据集有较高的质量。

- 一种很简单的测试数据集质量的方法是设计一个测试数据集，通过获取的数据训练的模型在测试数据集上的表现来判定数据集的优劣；
 - 这一思想将在下一节课进一步展开阐述；
 - 然而这一思想可能会导致一个微妙的问题：如果数据提供者知道了测试集，那么数据提供者可能会通过操纵自己的数据集使其过拟合测试集，这会导致统计学家得到的数据质量下降；
- 因此本小节的目标是基于同伴预测思想，设计机制使得数据提供者诚实地给出自己的数据集。

- 共有 n 个数据提供者。数据提供者 i 拥有数据集 $D_i = \{d_i^1, d_i^2, \dots, d_i^{N_i}\}$;
- 统计学家希望利用这些数据训练一个机器学习模型，例如线性回归模型： $y = \theta^\top x + \varepsilon$;
- 假设可能的参数空间 Θ 有限，并且在给定 θ 的条件下，不同数据提供者的数据集彼此独立：

$$p(D_1, D_2, \dots, D_n \mid \theta) = p(D_1 \mid \theta)p(D_2 \mid \theta) \cdots p(D_n \mid \theta),$$

- 例如在线性回归中，若不同卖家的特征与噪声都相互独立，那么在给定参数 θ 的条件下，各卖家的数据集就是独立的；
- 统计学家的目标是，在预算约束 B 下，为每个数据提供者设计支付规则 $r_i(\tilde{D}_1, \dots, \tilde{D}_n)$ ，从而激励每个人诚实报告 $\tilde{D}_i = D_i$;
- 此外，假设先验分布 $p(\theta)$ 是共同知识，统计学家可以根据任意数据集 D_i 计算后验分布 $p(\theta \mid D_i)$ ，并且统计学家能识别身份，因此一人不能反复冒名多次领取支付。

定义

支付规则是**诚实的**，如果对于任意分布 $p(\theta, D_1, \dots, D_n)$ 和每个数据提供者 i 的数据集 D_i 的任意实现，当其他人如实报告时，对任意 i, D_i, D_i' ，都有如实报告 D_i 能带来最高的期望支付，即

$$\mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)}[r_i(D_i, D_{-i})] \geq \mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)}[r_i(D_i', D_{-i})].$$

- 注意上述定义**不要求数据提供者实际知道条件分布**，因为它保证无论在什么分布下诚实报告都构成均衡；
- 上述定义的一个缺陷是**诚实报告不一定构成唯一的均衡**，因此希望提出一个比诚实性更强的性质要求，下面的**敏感度（sensitivity）**定义就是一个更强的要求。

定义

支付规则是敏感的，如果对于任意分布 $p(\theta, D_1, \dots, D_n)$ 和每个数据提供者 i 的数据集 D_i 的任意实现，当所有数据提供者 $j \neq i$ 报告 $\tilde{D}_j$ 且其给出的后验分布是准确的，即 $p(\theta \mid \tilde{D}_j) = p(\theta \mid D_j)$ 时，有

1. 诚实报告 D_i 能带来最高的期望支付：对任意的 D_i' 有

$$\mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)} [r_i(D_i, \tilde{D}_{-i})] \geq \mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)} [r_i(D_i', \tilde{D}_{-i})];$$

2. 报告后验分布不准确的数据集 D_i' （即 $p(\theta \mid D_i') \neq p(\theta \mid D_i)$ ）严格劣于报告一个后验分布准确的数据集 $\tilde{D}_i$ ：

$$\mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)} [r_i(\tilde{D}_i, \tilde{D}_{-i}(D))] > \mathbb{E}_{D_{-i} \sim p(D_{-i} \mid D_i)} [r_i(D_i', \tilde{D}_{-i})].$$

敏感度定义保证在均衡中，数据提供者报告的数据集必须给出正确的后验分布 $p(\theta \mid \tilde{D}_j) = p(\theta \mid D_j)$ ，否则期望支付会严格降低。

自然的，更理想的性质是数据提供者 i 报告任何 $D_i' \neq D_i$ ，其期望支付都严格降低。而满足敏感度要求的机制可以被视为朝着这一理想目标迈出的重要一步：

- 一方面，此时唯一可能保持支付最大化的数据报告策略是报告一个具有正确后验分布 $p(\theta \mid \tilde{D}_j) = p(\theta \mid D_j)$ 的数据集 $\tilde{D}_i$ ，而数据提供者找到不等于 D_i 的满足要求的 $\tilde{D}_i$ 可能具有挑战性；
- 另一方面，下面的引理表明，当统计学家获取的数据能保证得到正确后验分布时，最终能获得对 θ 的正确估计。

命题

当 $D_1, \dots, D_n$ 在给定 θ 的条件下独立时，对于任意 $(D_1, \dots, D_n)$ 和 $(\tilde{D}_1, \dots, \tilde{D}_n)$ ，如果对任意的 i 都有 $p(\theta \mid \tilde{D}_i) = p(\theta \mid D_i)$ ，则 $p(\theta \mid \tilde{D}_1, \dots, \tilde{D}_n) = p(\theta \mid D_1, \dots, D_n)$ 。

证明就是不断应用条件概率公式，感兴趣的同学可以参考教材。

除了诚实性与敏感性，支付规则还应满足三条基本要求：

1. 对称性：不因卖家身份不同而区别对待，即对于任意 n 元置换 $\pi(\cdot)$ ，对任意的数据提供者 i 都有

$$r_i(D_1, \dots, D_n) = r_{\pi(i)}(D_{\pi(1)}, \dots, D_{\pi(n)});$$

2. 个人理性：支付不能为负，即对任意的数据提供者 i 和任意的数据集 $D_1, \dots, D_n$ 实现都有 $r_i(D_1, \dots, D_n) \geq 0$;
3. 预算约束：统计学家的总支付不超过 B ，即对任意的数据集 $D_1, \dots, D_n$ 实现都有 $\sum_{i=1}^n r_i(D_1, \dots, D_n) \leq B$ 。

首先考虑一种简单的情况：若 $p(D_i | \theta)$ 是共同知识，则可以计算

$$p(D_{-i} | D_i) = \sum_{\theta \in \Theta} p(D_{-i} | \theta) p(\theta | D_i),$$

从而可以直接套用上一节的经典同伴预测机制。如果这一条件不成立，则需要通过互信息范式来设计支付规则。其关键在于设计奖励机制使其期望等于信息单调互信息。为此，定义两个数据集之间的**点互信息 (point-wise mutual information, PMI)**：

定义

对数据集 D_i 和 D_{-i} ，定义

$$\text{PMI}(D_i, D_{-i}) = \sum_{\theta \in \Theta} \frac{p(\theta | D_i) p(\theta | D_{-i})}{p(\theta)}.$$

基于此定义 **log-PMI** 支付规则： $r_i = \log(\text{PMI}(\tilde{D}_i, \tilde{D}_{-i}))$ 。

下面的命题表明，当使用 log-PMI 规则时，期望支付等于 $\tilde{D}_i$ 与 $\tilde{D}_{-i}$ 之间的互信息，从而可知使用 log-PMI 支付规则满足诚实性要求。

命题

在 log-PMI 规则下，数据提供者 i 的期望支付等于 $\tilde{D}_i$ 与 $\tilde{D}_{-i}$ 之间的互信息： $\mathbb{E}[\log(\text{PMI}(\tilde{D}_i, \tilde{D}_{-i}))] = I(\tilde{D}_i; \tilde{D}_{-i})$ 。

证明：利用条件概率公式以及给定 θ 时的数据条件独立性，

$$\begin{aligned} \sum_{\theta \in \Theta} \frac{p(\theta | \tilde{D}_i) p(\theta | \tilde{D}_{-i})}{p(\theta)} &= \sum_{\theta \in \Theta} \frac{p(\tilde{D}_i | \theta) p(\tilde{D}_{-i} | \theta) p(\theta)}{p(\tilde{D}_i) p(\tilde{D}_{-i})} \\ &= \sum_{\theta \in \Theta} \frac{p(\tilde{D}_i, \tilde{D}_{-i} | \theta) p(\theta)}{p(\tilde{D}_i) p(\tilde{D}_{-i})} = \frac{p(\tilde{D}_i, \tilde{D}_{-i})}{p(\tilde{D}_i) p(\tilde{D}_{-i})}. \end{aligned}$$

由 log-PMI 定义，期望支付等于 $\sum_{\tilde{D}_i, \tilde{D}_{-i}} p(\tilde{D}_i, \tilde{D}_{-i}) \log\left(\frac{p(\tilde{D}_i, \tilde{D}_{-i})}{p(\tilde{D}_i) p(\tilde{D}_{-i})}\right)$ ，于是期望支付就是互信息的定义式。 $\square$

原始 log-PMI 还有两个实际问题：(1) 若 $\text{PMI} = 0$ ，则对数无定义；(2) 原始支付既不保证非负，也不自动满足预算约束。

处理方法很直接：(1) **只在 $\text{PMI} > 0$ 的情况下计算对数**；(2) 设所有有效报告上的 $\log(\text{PMI})$ 介于 L 和 R 之间，把它**线性缩放**到 $[0, \frac{B}{n}]$ 。于是归一化后的支付规则可写为

$$r_i(\tilde{D}_1, \dots, \tilde{D}_n) = \begin{cases} \frac{B}{n} \cdot \frac{\log(\text{PMI}(\tilde{D}_i, \tilde{D}_{-i})) - L}{R - L} & \text{if } \text{PMI}(\tilde{D}_i, \tilde{D}_{-i}) > 0. \\ 0 & \text{otherwise} \end{cases}$$

命题

上述归一化后的支付规则满足诚实性、对称性、非负性和预算约束。

证明并不困难：(1) 归一化只是对原始 log-PMI 做正线性变换，不改变最优报告；(2) 每个人的支付都落在 $[0, \frac{B}{n}]$ ，因此总支付不超过 B ；(3) 所有卖家使用同一规则，因此是对称的。

事实上敏感性也满足，证明技术性较强此处略去，可以参考教材了解。

- 希望这门课能带给同学们一些启发，能让各位有“学到了一些有价值的内容”的体验
 - 十年之后你还能记得什么？看待问题的视角，工具理性.....
- 在 LLM 时代，在“疯狂内卷”的时代，仍然保持对知识的敬畏
- 个人研究兴趣：**Econ for LLM**（LLM 的社会经济影响，LLM 偏好对齐，公平性等），**LLM for Econ**（Agent-based Modeling），**复杂系统**（动态视角的经济系统、多智能体系统以及 Deep Learning Theory）
 - 对上述方向，或更广泛的计算经济学或理论计算机方向感兴趣的，欢迎联系我
 - 可以钉钉联系，或者加微信：yhwu_is，或者直接线下聊都可以