

ALLOCATING COMMON-VALUE GOODS

Hiroto Sato
Nagoya University

Ryo Shirakawa
MIT

INTRODUCTION & MODEL

INTRODUCTION

- ▶ Information is **dispersed** across individuals, generating **informational externalities**.
 - ▶ Usually, no single agent possesses all the information about the value of an option.
 - ▶ Hence, information held by others affects an agent's beliefs and decisions.
- ▶ How can society benefit from **leveraging** such informational externalities ?
 - ▶ To study this, we construct a minimal model of **allocating common-value goods**.
 - ▶ What is the optimal design, and what are the implications of its design?

MODEL (1/3): COMMON VALUE GOODS

- ▶ Homogeneous **goods** and ex ante homogeneous **agents** $i \in \{1, 2, \dots, n\}$.
- ▶ Goods have **uncertainty** about **common value** $\omega \in \{-1, +1\}$, equally likely.
 - ▶ (e.g., vaccines, investment opportunities / quality, effectiveness, degree of side effect, etc.)
- ▶ Agent obtains payoff ω from the good, and 0 from not receiving it (normalization).
 - ▶ Agents want to receive a good if their beliefs about $\omega = +1$ exceed $1/2$.
 - ▶ Remark 3: Can be generalized to payoff $\omega + \varepsilon_i$, where ε_i is an idiosyncratic private value.
- ▶ No resource constraint (discussion).

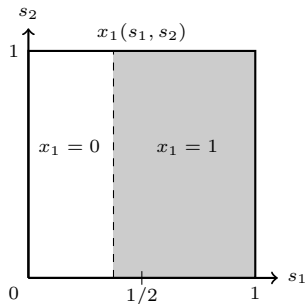
MODEL (2/3): PRIVATE INFORMATION

- ▶ Each agent i has **private signal (type)** $s_i \in S$ regarding common value ω .
 - ▶ Formally, it is drawn from a state-dependent distribution $\mathbb{F}_\omega \in \Delta(S)$.
 - ▶ Signals $s_1, \dots, s_n$ are independent conditional on ω (but, ex ante correlated).
 - ▶ $\text{LR}(E) = \frac{\mathbb{P}[E|\omega=+1]}{\mathbb{P}[E|\omega=-1]}$: **likelihood ratio** for event E ; $\text{LR}(E) \geq 1 \Leftrightarrow \mathbb{P}[\omega = +1 \mid E] \geq 1/2$
posterior belief exceeds 1/2
 - ▶ Normalize $\underset{\text{private signal}}{s_i} = \underset{\text{private belief}}{\mathbb{P}[\omega = +1 \mid s_i]} \in [0, 1]$: (e.g., $s_i = 0.6 \Leftrightarrow \mathbb{P}[\omega = +1 \mid s_i] = 0.6$)
- ▶ Assumption: CDF $\mathbb{F} = \frac{\mathbb{F}_{-1} + \mathbb{F}_{+1}}{2}$ admits a smooth and full support density f on $[0, 1]$.

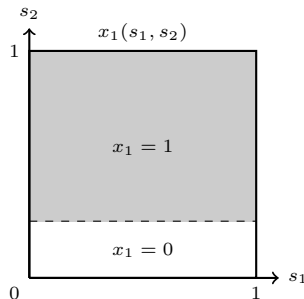
MODEL (3/3): MECHANISM

- ▶ A direct mechanism is an **allocation rule** $x_i(s_1, \dots, s_n) \in [0, 1]$ such that
$$(\mathbf{P}) : \mathbb{E}[\omega \cdot x_i(s_i, s_{-i}) \mid s_i] \geq 0 \text{ and } (\mathbf{IC}) : \mathbb{E}[\omega \cdot x_i(s_i, s_{-i}) \mid s_i] \geq \mathbb{E}[\omega \cdot x_i(\hat{s}_i, s_{-i}) \mid s_i].$$
- ▶ x_i may depend on others' signals, and thus shapes the informational externality.
 - ▶ Screening and aggregating private signals, and leveraging them for allocation.
- ▶ Designing allocation rules may effectively reduce **discard rates** and **inefficiency**.
 - ▶ We focus on a direct mechanism maximizing **allocation probability** $\mathbb{E}_\omega \mathbb{E}_s [\sum_i x_i(s)]$.
 - ▶ Remark 1: Can be generalized to a convex combination with **utilitarian sum**.

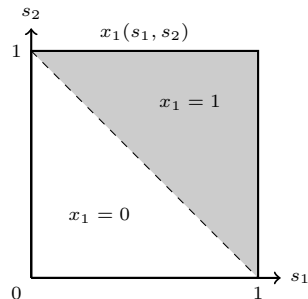
ILLUSTRATIVE EXAMPLES



- ▶ $x_1(s_1, s_2) = x_1(s_1)$.
- ▶ Violating (P) and (IC).



- ▶ $x_1(s_1, s_2) = x_1(s_2)$.
- ▶ Violating (P).



- ▶ $x_1(s) = \mathbb{I}\{\text{LR}(s) \geq 1\}$.
- ▶ Feasible and **Efficient**

LITERATURE

- ▶ Mechanism design with **interdependent values**:
 - ▶ Crémer and McLean (1985, 1988), McAfee and Reny (1992), Lopomo *et al.* (2022).
 - ▶ No transfer: Kattwinkel (2020), Niemeyer and Preusser (2022), Kattwinkel and Winter (2024), etc.
- ▶ Majorization and duality in mechanism design:
 - ▶ Indirect utility and duality: Daskalakis *et al.* (2017), Kleiner (2022)
 - ▶ Extreme points: Kleiner *et al.* (2021), Yang and Zentefis (2024), Augias and Uhe (2025), etc.
- ▶ Mechanism design with **money burning**:
 - ▶ McAfee and McMillan (1992), Hartline and Roughgarden (2008), Condorelli (2012), etc.
 - ▶ Noda and Okada (2024), Tokarski (2024, 2025), Yang (2025), Yang *et al.* (2025), Dworczak (2026).

RESULTS

OPTIMAL MECHANISM

Theorem 1

The optimal mechanism is a monotone threshold rule: there is a partition $\mathcal{S}_i$ of the signal space $[0, 1]$ into intervals such that, for each interval $S_i \in \mathcal{S}_i$ we have $x_i(s_i, s_{-i}) = \kappa \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau\}$ for all $s_i \in S_i$, for some $\kappa \in [0, 1]$, and $\tau \geq 0$.

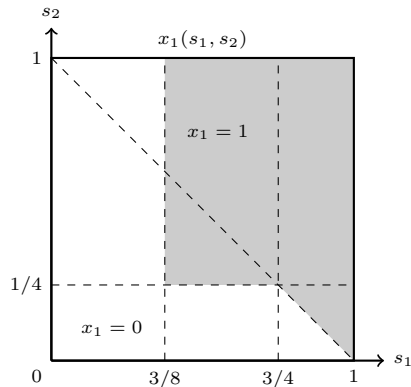
- ▶ The optimal mechanism partitions the signals/types into disjoint intervals.
- ▶ Then, a common allocation rule is applied to all types within the same interval.
- ▶ A common allocation rule consists of two parameters (κ, τ) :
 - ▶ $\kappa \in [0, 1]$ controls the **"volume"** of allocation.
 - ▶ $\tau \geq 0$ determines state-dependent allocation probabilities (**"degree of learning"**).

EXAMPLE: UNIFORM DISTRIBUTION WITH TWO AGENTS

- ▶ No allocation for $s_1 \leq 3/8$.
- ▶ Constant allocation for $s_1 \in [3/8, 3/4]$.
- ▶ Efficient allocation for $s_1 \geq 3/4$.
- ▶ Partition

$$\mathcal{S}_1 = \underbrace{\{[0, 3/8],\}}_{x_1(s)=0}, \underbrace{\{[3/8, 3/4]\}}_{x_1(s)=\mathbb{I}\{\text{LR}(s_2) \geq 1/3\}}, \underbrace{\{s_1 : s_1 \geq 3/4\}}_{x_1(s)=\mathbb{I}\{\text{LR}(s) \geq 1\}}$$

- ▶ Exclusion at the bottom types.



EXCLUSION AT THE BOTTOM

Proposition 1

The optimal mechanism excludes bottom types:

there is an interval $[0, \varepsilon_i]$ such that $x_i(s_i, s_{-i}) = 0$ for all $s_i \in [0, \varepsilon_i]$ and $s_{-i} \in [0, 1]^{n-1}$.

- ▶ Exclusion occurs even though the designer prefers to allocate. (nontrivial !)
 - ▶ Intuition: managing **incentives for downward deviations**.
 - ▶ (P): Allocating to low types is allowed only when others have very positive signal.
 - ▶ Doing so incentivizes a middle type to misreport her belief downward.
- ▶ Exclusion does not arise when there is no information asymmetry about type s_i , and inefficiency does not arise when there is no interdependence or $n = 1$.

PROOF SKETCH OF THEOREM 1

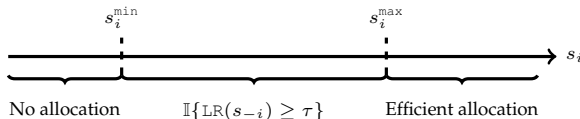
- ▶ Focus on **indirect utility function** $U_i(s_i) = \max_{t_i} \mathbb{E}[\omega \cdot x_i(t_i, s_{-i}) \mid s_i]$ (Rochet, 1987).
- ▶ Envelope formula transforms the objective into a linear functional of U_i .
- ▶ Step 1: Derive necessary conditions for the optimal indirect utility.
 - ▶ U_i is **convex, increasing**, and **bounded by two convex functions** ($\underline{U}_i \leq U_i \leq \bar{U}_i$).
 - ▶ Not a full characterization of feasible indirect utilities, but sufficient under optimality.
- ▶ Step 2: Construct monotone threshold mechanisms that implement them.

OPTIMAL MECHANISM: LOG-CONCAVE CASE

Theorem 2

If f is log-concave, a monotone threshold mechanism with two thresholds is optimal:

$$x_i(s_i, s_{-i}) = \begin{cases} 0 & \text{if } s_i \leq s_i^{\min}(\tau) \\ \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau\} & \text{if } s_i^{\min}(\tau) \leq s_i \leq s_i^{\max}(\tau) \\ \mathbb{I}\{\text{LR}(s_i, s_{-i}) \geq 1\} & \text{if } s_i^{\max}(\tau) \leq s_i. \end{cases}$$

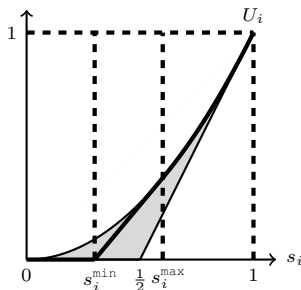


PROOF SKETCH OF THEOREM 2

- ▶ The problem is reduced to maximize a linear functional $\int U_i d\mu_i$ subject to

$$\mathcal{U}_i^* = \{U_i \mid U_i \text{ is convex, increasing, and } \max\{0, 2s_i - 1\} \leq U_i \leq \bar{U}_i\},$$

- ▶ We provide **weak duality**: an upper bound for the objective value.
- ▶ Following form of U_i achieves the upper bound and our mechanisms induce this.



OPTIMAL MECHANISM: LARGE MARKET CASE

- ▶ Consider a market with **many participants** (and many goods).
- ▶ For each market size n , let V_n be the maximum probability of allocation to the agent.

Theorem 3

Fix any agent i . There exists a sequence of mechanisms for agent i of a form

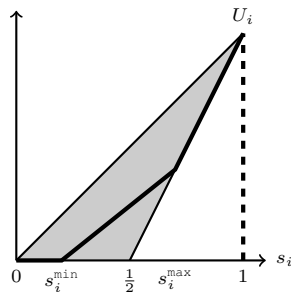
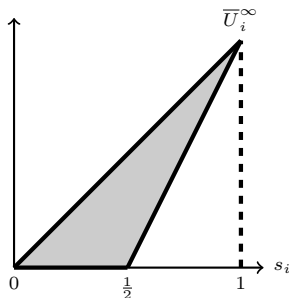
$$x_i(s_i, s_{-i}; n) = \begin{cases} 0 & \text{if } s_i \leq s_i^{\min}(n) \\ \kappa(n) \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau(n)\} & \text{if } s_i^{\min}(n) \leq s_i \leq s_i^{\max}(n), \\ 1 & \text{if } s_i^{\max}(n) \leq s_i, \end{cases}$$

for each market size $n \in \mathbb{N}$, such that the difference $|V_n - \mathbb{E}[x_i(s; n)]| \rightarrow 0$ as $n \rightarrow \infty$.

- ▶ A monotone threshold mechanism with two thresholds is **asymptotically optimal**.

PROOF SKETCH OF THEOREM 3

- ▶ LLN implies that, s_{-i} perfectly reveals the true state in the limit.
- ▶ Thus, in the limit, the first-best payoff becomes $\bar{U}_i^\infty(s_i) = s_i \cdot 1 - (1 - s_i) \cdot 0 = s_i$.
- ▶ The set of U_i over which we optimize and an extremal function are follows:

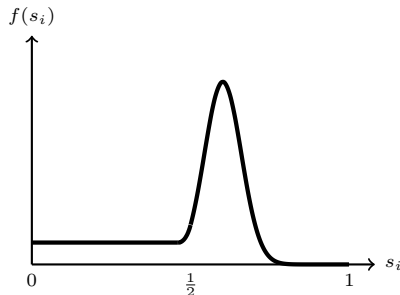


OPTIMAL MECHANISM: LAISSEZ-FAIRE RULE

Proposition 3

If $\log f$ is increasing on $[0, 1/2]$, concave over $[1/2, 1]$, and $\int_{1/2}^1 (1 - t_i) f(t_i) dt_i \geq (1/4) f(1/2)$, then $x_i(s) = \mathbb{I}\{s_i \geq 1/2\}$ is optimal.

- ▶ **Independent rule** is optimal if s_i just above $1/2$ are sufficiently likely.
- ▶ A sufficient condition for the optimality of **empty network** in network design.



PAYMENT DESIGN

PAYMENT DESIGN

- ▶ Pricing is often **infeasible** in many allocation settings (e.g., vaccine allocation).
 - ▶ Thus, we assume no monetary transfers so far.
 - ▶ **Payment-like instruments**, such as **costly effort and waiting time**, may be feasible.
- ▶ Do non-negative payments help improve the allocation rate (and efficiency)?
 - ▶ A mechanism is a profile $(x_i, t_i) : [0, 1]^n \rightarrow [0, 1] \times [0, T]$ such that
$$(\mathbf{P}^*) : \mathbb{E}[\omega \cdot x_i(s_i, s_{-i}) - t_i(s_i, s_{-i}) \mid s_i] \geq 0,$$
$$(\mathbf{IC}^*) : \mathbb{E}[\omega \cdot x_i(s_i, s_{-i}) - t_i(s_i, s_{-i}) \mid s_i] \geq \mathbb{E}[\omega \cdot x_i(\hat{s}_i, s_{-i}) - t_i(\hat{s}_i, s_{-i}) \mid s_i].$$
 - ▶ Payoff: $\mathbb{E}[\omega \cdot x_i(s) - t_i(s) \mid s_i] \Rightarrow$ allocation becomes **less attractive to agents**.
 - ▶ If there is only one agent, positive payments are suboptimal.

OPTIMAL MECHANISM: NON-NEGATIVE PAYMENT

Theorem 4

Suppose that the designer can also design non-negative payments.

The optimal mechanism is a monotone threshold rule: there exists a partition $\mathcal{S}_i$ of the signal space $[0, 1]$ into intervals such that, for each interval $S_i \in \mathcal{S}_i$, we have

$$x_i(s_i, s_{-i}) = \kappa_x \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau_x\}, \quad t_i(s_i, s_{-i}) = \kappa_t \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau_t\},$$

for all $s_i \in S_i$ and $s_{-i} \in [\underline{s}, \bar{s}]^{n-1}$, for some $(\kappa_x, \tau_x) \in \mathbb{R}^2$ and $(\kappa_t, \tau_t) \in \mathbb{R}^2$.

- Technical challenge: **revenue equivalence fails** (x_i does not uniquely pin down t_i).
- (Remark: objective can be extended to $\mathbb{E}[\alpha \cdot x_i(s) + (1 - \alpha) \cdot U_i(s_i)]$ again.)

OPTIMAL MECHANISM: LARGE MARKET WITH MONEY BURNING

Theorem 5

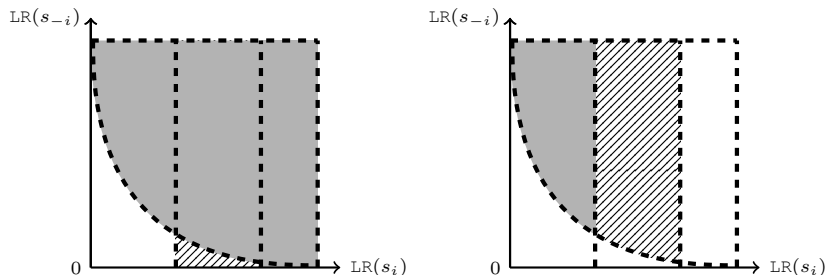
Suppose that the designer can also design non-negative payments and $T \geq 1$.

Then, there exists a sequence of mechanisms for agent i of a form

$$x_i(s_i, s_{-i}; n) = \begin{cases} \mathbb{I}\{\text{LR}(s_i, s_{-i}) \geq \eta(n)\} & \text{if } s_i \leq s_i^{\min}(n) \\ \kappa_x + (1 - \kappa_x) \cdot \mathbb{I}\{\text{LR}(s_i, s_{-i}) \geq \tau(n)\} & \text{if } s_i^{\min}(n) \leq s_i \leq s_i^{\max}(n), \\ 1 & \text{if } s_i^{\max}(n) \leq s_i, \end{cases}$$

$$t_i(s_i, s_{-i}; n) = \begin{cases} \delta(n) \cdot \mathbb{I}\{\text{LR}(s_i, s_{-i}) \geq \eta(n)\} & \text{if } s_i \leq s_i^{\min}(n) \\ \kappa_t(n) \cdot \mathbb{I}\{\text{LR}(s_i, s_{-i}) \geq \tau(n)\} & \text{if } s_i^{\min}(n) \leq s_i \leq s_i^{\max}(n), \\ 0 & \text{if } s_i^{\max}(n) \leq s_i. \end{cases}$$

for each $n \in \mathbb{N}$, such that $|V_n^* - \mathbb{E}[x_i(s; n)]| \xrightarrow{n \rightarrow \infty} 0$, $\delta(n) \xrightarrow{n \rightarrow \infty} 1$, and $\kappa_x, \kappa_t(n) \in [0, 1]$.



□ $x_i(s) = 0$ ▨ $x_i(s) \in (0, 1)$ ■ $x_i(s) = 1$ □ $t_i(s) = 0$ ▨ $t_i(s) \in (0, 1)$ ■ $t_i(s) = 1$

- ▶ The optimal mechanism **no longer excludes bottom types** unlike in Proposition 1.
- ▶ Intuition: positive payments (**money burning**) tighten (P^*) but changes (IC^*).
 - ▶ **Different beliefs** leads to **different expected payments** from the same payment rule.
 - ▶ A payment rule is **decreasing** in type, which makes **downward deviations less attractive**.

CONCLUSION

CONCLUSION

- ▶ We consider the allocation of **common-value goods** with uncertainty.
 - ▶ The allocation rule shapes **informational externalities**, which in turn affects incentives.
 - ▶ The optimal mechanism is a threshold rule based on “**volume**” and “**degree of learning**.”
- ▶ The optimal mechanism **excludes** low-belief agents to deter downward deviations.
 - ▶ Improving **allocation rates** and **efficiency** may require additional policies.
- ▶ Introducing **costly effort** or **waiting time** improves the allocation rate and efficiency.
 - ▶ Non-negative payments (decreasing in types) deter downward deviations.

THE END.

APPENDIX

DETAILED SKETCH

- By normalization, $s_i = \mathbb{P}[\omega = +1 \mid s_i]$, we have

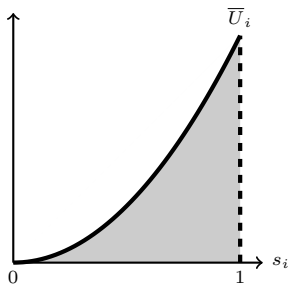
$$U_i(t_i; s_i) = s_i \cdot \underbrace{\sum_{\omega \in \{-1, +1\}} \mathbb{E}[x_i(t_i, s_{-i}) \mid \omega] - \mathbb{E}[x_i(t_i, s_{-i}) \mid \omega = -1]}_{\text{slope} \in [0, 2]}.$$

- This is **linear in the private signal** s_i .
- Thus, any feasible U_i is **convex, increasing, and 2-Lipschitz**.

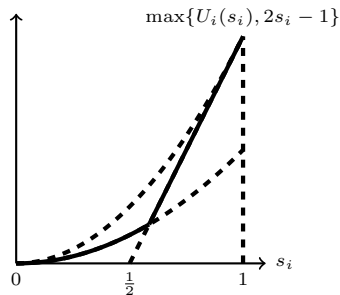
► Additionally, any feasible indirect utility U_i must be **"not too large"**.

► The **first-best payoff** $\bar{U}_i$ when allocation is efficient: $x_i(s) = \mathbb{I}\{\text{LR}(s) \geq 1\}$.

► Any feasible U_i should lie in the shaded region:

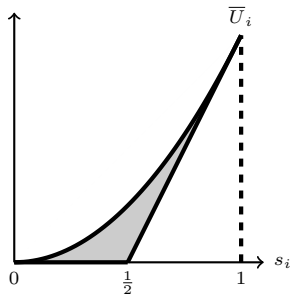


- ▶ There is, however, a function in the triangle that cannot be implemented.
 - ▶ Lemma 6: Any **optimal** indirect utility must be above $\underline{U}_i(s_i) = \max\{0, 2s_i - 1\}$.
 - ▶ Proof: Replace $x_i(s_i, s_{-i}) \in [0, 1]$ with 1 if $U_i(s_i) \leq 2s_i - 1 = \mathbb{E}[\omega \mid s_i]$.



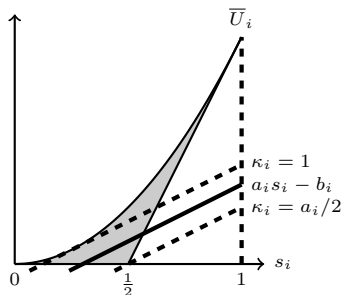
- ▶ (This argument works for a convex combination with utilitarian sum.)

- Thus, optimal indirect utility U_i is **sandwiched by functions**, $\bar{U}_i$ and $\underline{U}_i$.



- In addition, optimal U_i is also convex, increasing, and 2–Lipschitz.

- Finally, consider any linear function $a_i s_i - b_i$. Is it implementable?
 - A **threshold mechanism** $\kappa \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau\}$ induces a linear utility U_i .
 - For each $\kappa \geq a_i/2$, we can find $\tau(\kappa)$ such that U_i has slope a_i .
 - If $a_i s_i - b_i$ intersects with the shaded region, some $\kappa \in [a_i/2, 1]$ induces intercepts b_i .



Summary:

- ▶ Mechanisms $\kappa \cdot \mathbb{I}\{\text{LR}(s_{-i}) \geq \tau\}$ implement any linear segment in the shaded region.
- ▶ We already know that any optimal U_i is convex and in the shaded regions.
- ▶ Therefore, it is an upper envelope of such lines.
- ▶ A **monotone threshold mechanism** can implement it.

QED.

EXTENSIONS OF THEOREM 1

- ▶ Objective can be anything as long as it is increasing in agents' payoffs over $s_i \geq 1/2$.
 - ▶ It can be a **convex combination with utilitarian sum** (allocation rates and efficiency).
- ▶ Distribution $\mathbb{F}_i$ can be heterogeneous across agents $i \in \{1, \dots, n\}$.
- ▶ Can restrict to **queue-based allocations** $x_i(s) = x_i(s_1, \dots, s_i)$.
 - ▶ It does not use information gathered from those who will arrive after i .
 - ▶ **Corollary:** The optimal queue-based allocation mechanism is monotone partitional.
- ▶ Can be generalized to payoff $\omega + \varepsilon_i$, where ε_i is an idiosyncratic **private value**.
 - ▶ Let $t_i = \frac{(1+\varepsilon_i)s_i}{(1+\varepsilon_i)s_i + (1-\varepsilon_i)(1-s_i)} \in [0, 1]$. Then, the interim payoff is proportion to $\propto U_i(t_i(s_i, \varepsilon_i)) = t_i \mathbb{E}[x_i(s_i, \varepsilon_i, s_{-i}) \mid \omega = +1] - (1 - t_i) \mathbb{E}[x_i(s_i, \varepsilon_i, s_{-i}) \mid \omega = -1]$.
 - ▶ The analysis goes through by focusing on the effective type t_i and others' signals s_{-i} .

PROOF SKETCH OF THEOREM 4

- ▶ Any positive payment **only reduces** expected payoffs.
- ▶ As in Theorem 1, any optimal indirect utility U_i lies between $\bar{U}_i$ and $\underline{U}_i$.
- ▶ For fixed t_i , replace x_i with a monotone threshold rule preserving $\mathbb{E}[x_i(s)]$ and U_i .
- ▶ If t_i cannot be replaced by a monotone threshold rule preserving $\mathbb{E}[x_i(s)]$ and U_i , mechanism must involve **high payments when the state is low**.
- ▶ Jointly perturbing (x_i, t_i) increases $\mathbb{E}[x_i(s)]$ while preserving U_i , a contradiction.

FULL SURPLUS EXTRACTION

- ▶ Suppose any pricing schemes, including negative payment (**subsidies**), are feasible.

Proposition 4

There is a feasible mechanism (x_i, t_i) such that $\mathbb{E}[t_i(s)] = 0$ and $x_i(s) = 1$ for all s .

- ▶ Under the interdependent values, any allocation rule is implementable.
(Cr mer and McLean, 1985, 1988; McAfee and Reny, 1992; Lopomo *et al.*, 2022)
- ▶ (Not a corollary of the literature, but a similar construction works.)
- ▶ For example, if there are only two agents, $x_i(s) = 1$ and $t_i(s) = \frac{2s_j - 1}{\text{Cov}(s_i, s_j)}$.

RESOURCE CONSTRAINT

- ▶ Suppose that there is a **resource constraint (RC)** such that $\sum_i x_i(s) \leq 1$.
- ▶ Limitation: in general, we cannot derive the optimal mechanism. (complicated)
 - ▶ Payoff externalities complicate the analysis, and the problem is no longer separable.
- ▶ We restrict attention to **contest allocation rules**: $x_i(s) = 0$ if $s_i < s_j$ for some $j \neq i$.
 - ▶ Allocation occurs only to the agent with the highest report.
 - ▶ This restriction may be natural in practice because of "fairness" and "simplicity".
 - ▶ Technically, this restriction makes the problem separable across agents again.

EXAMPLE: UNIFORM DISTRIBUTION WITH TWO AGENTS

- ▶ Contest and two-sided threshold rules:

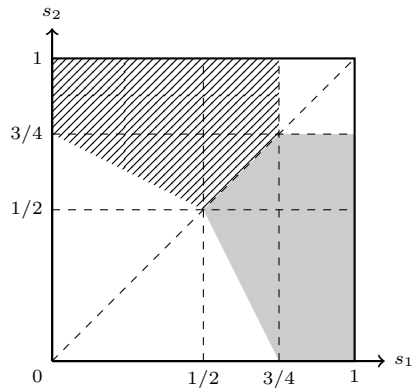
$$x_i(s) = \mathbb{I}\{s_i > \max_{j \neq i} s_j\} \cdot \mathbb{I}\{\text{LR}(s_{-i}) \in [\underline{\tau}(s_i), \bar{\tau}(s_i)]\}$$

- ▶ Scarcity creates **distortions** for top types.

- ▶ This is to deter upward deviations.

- ▶ Contest allocations are **suboptimal**.

- ▶ Dominated by the mechanism:
 - randomly choose an agent i , and
 - allocate the good if $s_i \geq 1/2$.



□ $x_1(s) = x_2(s) = 0$ □ $x_1(s) = 1$ ▨ $x_2(s) = 1$

REFERENCE I

- AUGIAS, V. and UHE, L. (2025). The economics of convex function intervals. *arXiv preprint arXiv:2510.20907*.
- CONDORELLI, D. (2012). What money can't buy: Efficient mechanism design with costly signals. *Games and Economic Behavior*, **75** (2), 613–624.
- CRÉMER, J. and MCLEAN, R. P. (1985). Optimal selling strategies under uncertainty for a discriminating monopolist when demands are interdependent. *Econometrica*, **53** (2), 345–361.
- and MCLEAN, R. P. (1988). Full extraction of the surplus in bayesian and dominant strategy auctions. *Econometrica*, pp. 1247–1257.
- DASKALAKIS, C., DECKELBAUM, A. and TZAMOS, C. (2017). Strong duality for a multiple-good monopolist. *Econometrica*, **85** (3), 735–767.
- DWORCZAK, P. (2026). How to allocate money? *American Economic Journal: Microeconomics*, **18** (1), 312–339.
- HARTLINE, J. D. and ROUGHGARDEN, T. (2008). Optimal mechanism design and money burning. In *Proceedings of the fortieth annual ACM symposium on Theory of computing*, pp. 75–84.
- KATTWINKEL, D. (2020). Allocation with correlated information: Too good to be true. In *Proceedings of the 21st ACM Conference on Economics and Computation*, pp. 109–110.
- and WINTER, A. (2024). Optimal decision mechanisms for committees: Acquitting the guilty. *arXiv preprint arXiv:2407.07293*.
- KLEINER, A. (2022). Optimal delegation in a multidimensional world. *arXiv preprint arXiv:2208.11835*.
- , MOLDOVANU, B. and STRACK, P. (2021). Extreme points and majorization: Economic applications. *Econometrica*, **89** (4), 1557–1593.
- LOPOMO, G., RIGOTTI, L. and SHANNON, C. (2022). Detectability, duality, and surplus extraction. *Journal of Economic Theory*, **204**, 105465.
- MCAFEE, R. P. and MCMILLAN, J. (1992). Bidding rings. *The American Economic Review*, pp. 579–599.

REFERENCE II

- and RENY, P. J. (1992). Correlated information and mechanism design. *Econometrica*, pp. 395–421.
- NIEMEYER, A. and PREUSSER, J. (2022). *Simple allocation with correlated types*. Tech. rep., Technical report, Working Paper.
- NODA, S. and OKADA, G. (2024). No screening is more efficient with multiple objects. *arXiv preprint arXiv:2408.10077*.
- ROCHET, J.-C. (1987). A necessary and sufficient condition for rationalizability in a quasi-linear context. *Journal of mathematical Economics*, **16** (2), 191–200.
- TOKARSKI, F. (2024). Equitable screening. *arXiv preprint arXiv:2402.08781*.
- (2025). Screening with damages and ordeals. *arXiv preprint arXiv:2508.04456*.
- YANG, F. (2025). Costly multidimensional screening. *Review of Economic Studies*, p. rdaf105.
- , DWORCZAK, P. and AKBARPOUR, M. (2025). Comparison of screening devices. *Available at SSRN 4456198*.
- YANG, K. H. and ZENTEFIS, A. K. (2024). Monotone function intervals: Theory and applications. *American Economic Review*, **114** (8), 2239–2270.