

Non-Monetary Mechanism Design without Distributional Information: Using Scarce Audits Wisely

Yan Dai¹



Moïse Blanchard^{2→3}



Patrick Jaillet¹



¹MIT

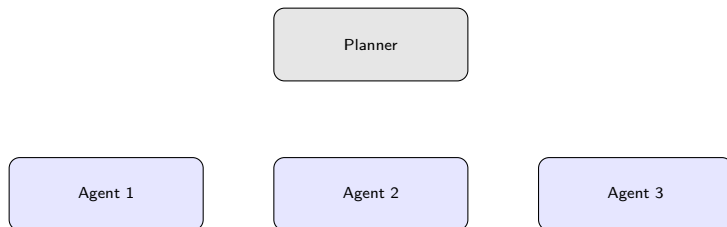
²Columbia

³Georgia Tech

Challenge: Resource Allocation to Strategic Agents

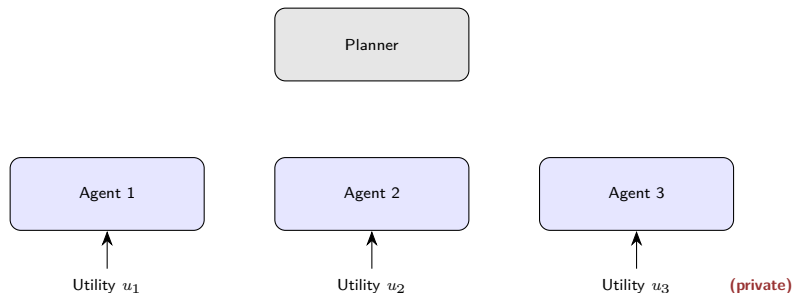
- ① **One central planner** maximize social welfare
- ② **K strategic agents** self-interested (*i.e.*, may lie)
- ③ **1 indivisible item**

Challenge: Resource Allocation to Strategic Agents



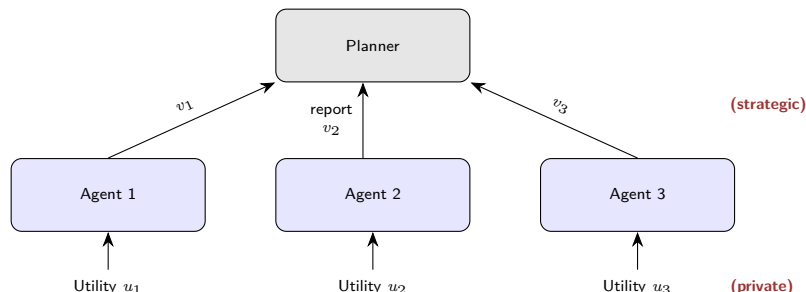
- ❶ **One central planner** maximize social welfare
- ❷ **K strategic agents** self-interested (*i.e.*, may lie)
- ❸ **1 indivisible item**

Challenge: Resource Allocation to Strategic Agents



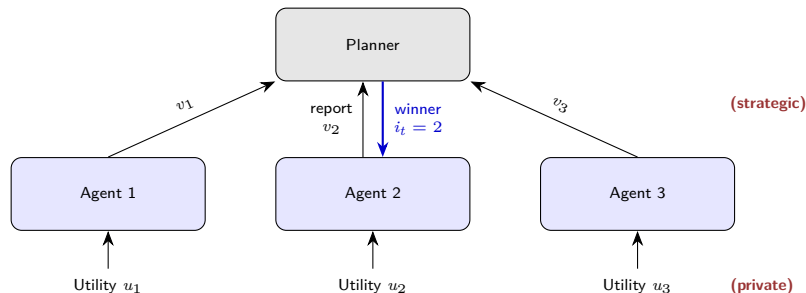
- 1 **One central planner** maximize social welfare
- 2 **K strategic agents** self-interested (*i.e.*, may lie)
- 3 **1 indivisible item**

Challenge: Resource Allocation to Strategic Agents



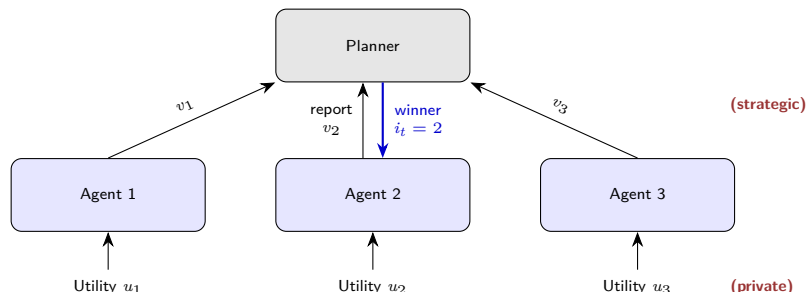
- 1 **One central planner** maximize social welfare
- 2 **K strategic agents** self-interested (*i.e.*, may lie)
- 3 **1 indivisible item**

Challenge: Resource Allocation to Strategic Agents



- 1 **One central planner** maximize social welfare
- 2 **K strategic agents** self-interested (*i.e.*, may lie)
- 3 **1 indivisible item**

Challenge: Resource Allocation to Strategic Agents



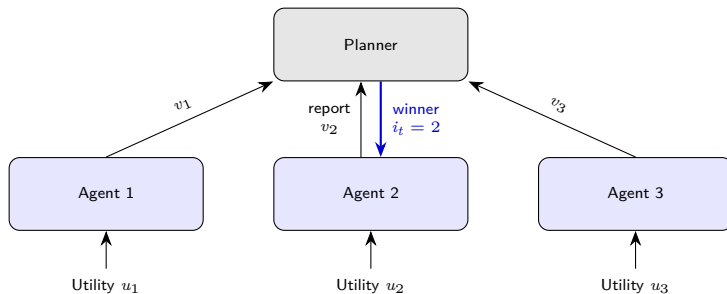
- 1 **One central planner** maximize social welfare
- 2 **K strategic agents** self-interested (*i.e.*, may lie)
- 3 **1 indivisible item**

Two Objectives

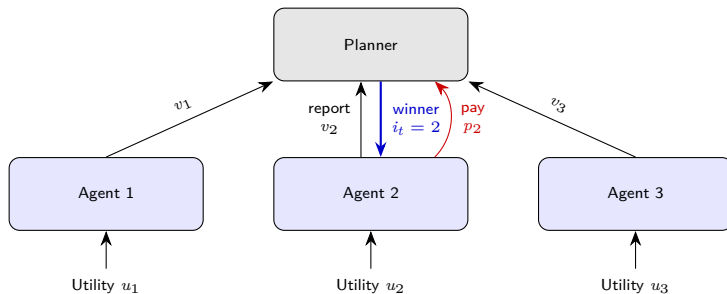
Efficiency. max established utility u_{i_t} (unknown!)

Incentive-Compatibility. truthfully report $v_i \approx u_i$

Classical VCG Mechanisms

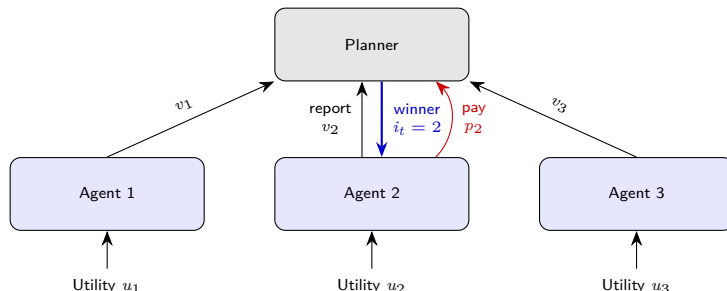


Classical VCG Mechanisms



Monetary mechanisms: VCG family [Vic61; Cla71; Gro73]

Classical VCG Mechanisms

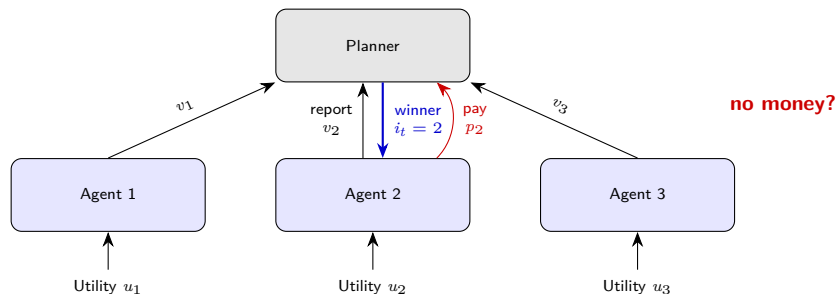


Monetary mechanisms: VCG family [Vic61; Cla71; Gro73]

Money isn't everything!

- 1 Food bank allocations [Pre17; Pre22]
- 2 Healthcare resources [PSÜY24; YBP23]
- 3 GPU in company [ABDVVW22; PSMST22]

Classical VCG Mechanisms

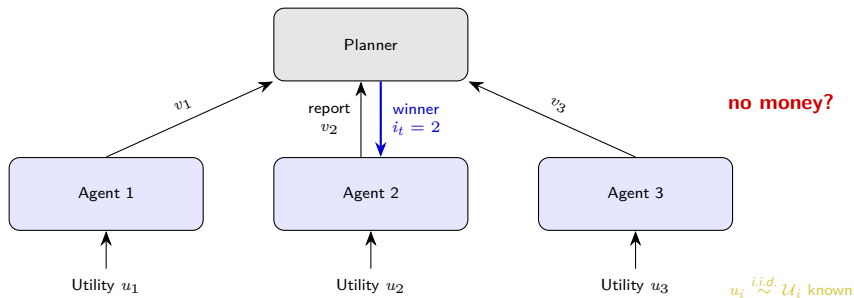


Monetary mechanisms: VCG family [Vic61; Cla71; Gro73]

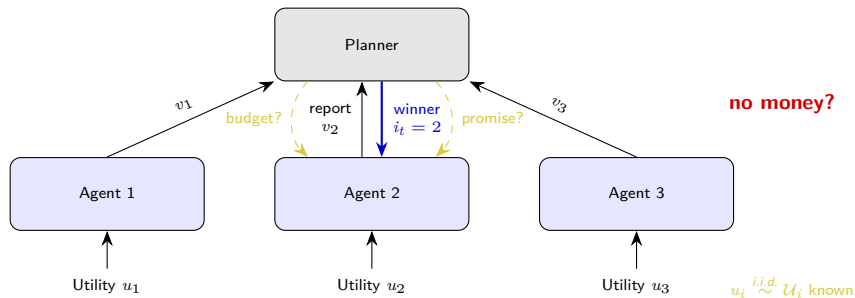
Money isn't everything!

- 1 Food bank allocations [Pre17; Pre22]
- 2 Healthcare resources [PSÜY24; YBP23]
- 3 GPU in company [ABDVVW22; PSMST22]

Non-Monetary Mechanisms?

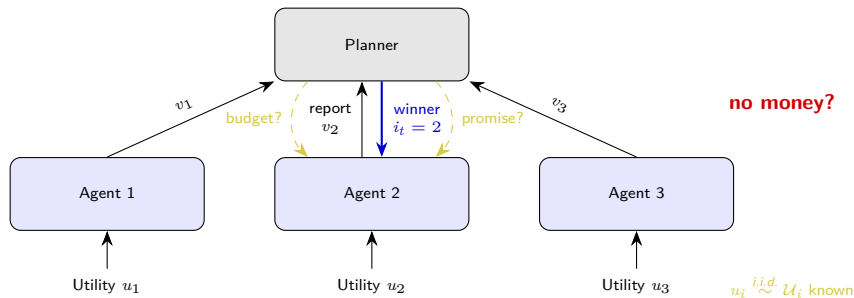


Non-Monetary Mechanisms?



Distribution-Aware approaches [BGS19; GBI21; BJ24]

Non-Monetary Mechanisms?

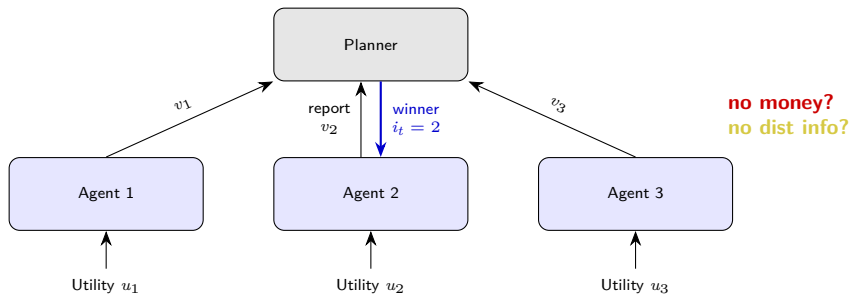


Distribution-Aware approaches [BGS19; GBI21; BJ24]

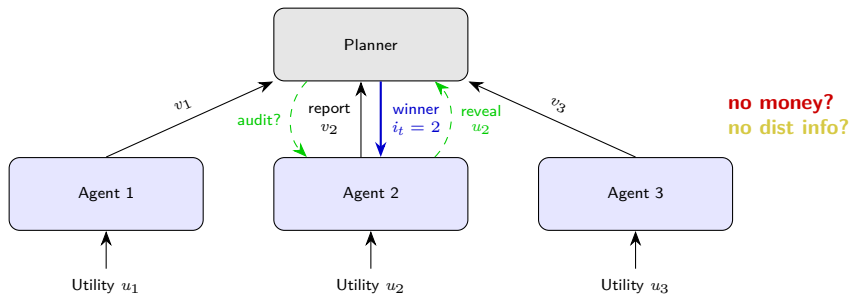
Distributional info *a-priori* is hard!

- 1 Perfect knowledge on agents
- 2 Enormous historical data

Our Prior-Free Mechanism

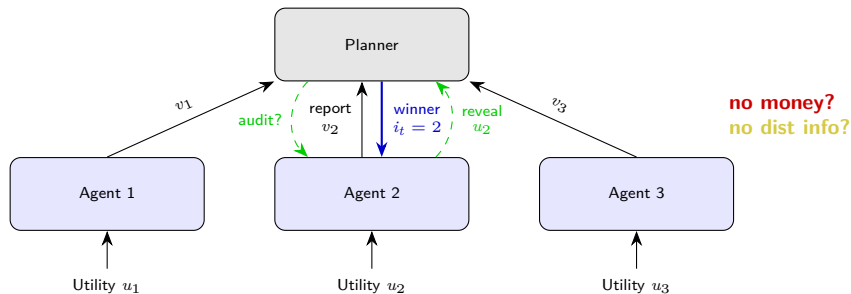


Our Prior-Free Mechanism



Prior-Free mechanism via scarce & powerful “audits”

Our Prior-Free Mechanism



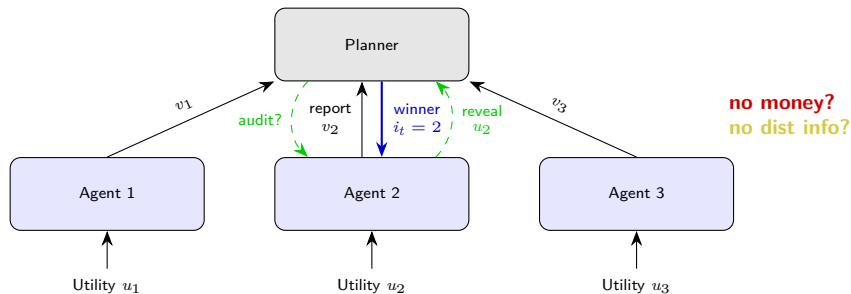
Prior-Free mechanism via scarce & powerful “audits”

Tradeoff between Regret & #Audits

Repeated allocation for T rounds:

❶ **Social Welfare Regret.** $\mathcal{R}_T := \mathbb{E}[\sum_{t=1}^T (\max_i u_{t,i} - u_{t,i_t})]$

Our Prior-Free Mechanism



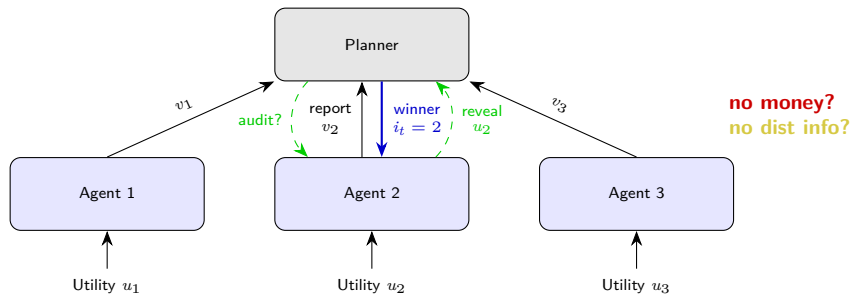
Prior-Free mechanism via scarce & powerful “audits”

Tradeoff between Regret & #Audits

Repeated allocation for T rounds:

- 1 **Social Welfare Regret.** $\mathcal{R}_T := \mathbb{E}[\sum_{t=1}^T (\max_i u_{t,i} - u_{t,i_t})]$
(Different from Online Learning regret!)

Our Prior-Free Mechanism



Prior-Free mechanism via scarce & powerful “audits”

Tradeoff between Regret & #Audits

Repeated allocation for T rounds:

- 1 **Social Welfare Regret.** $\mathcal{R}_T := \mathbb{E}[\sum_{t=1}^T (\max_i u_{t,i} - u_{t,i_t})]$
(Different from Online Learning regret!)
- 2 **Expected Number of Audits.** $\mathcal{B}_T := \mathbb{E}[\sum_{t=1}^T \mathbb{1}[\text{audit}]]$

Tradeoff between Regret & #Audits

∃ Perfect Bayesian Equilibrium (PBE) π^* , s.t.

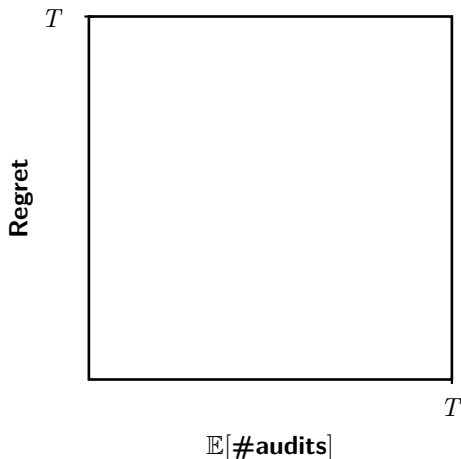


Figure: Regret vs $\mathbb{E}[\text{\#audits}]$
green possible; red impossible

Tradeoff between Regret & #Audits

$\exists$ Perfect Bayesian Equilibrium (PBE) π^* , s.t.

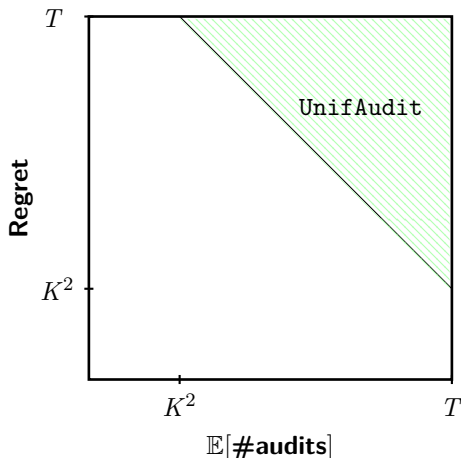


Figure: Regret vs $\mathbb{E}[\text{\#audits}]$

green possible; red impossible

**Simple
Mechanism:
UnifAudit.**

$$\mathcal{R}_T \leq \frac{K^2}{p} \text{ \& } \mathcal{B}_T \leq pT.$$

Tradeoff between Regret & #Audits

$\exists$ Perfect Bayesian Equilibrium (PBE) π^* , s.t.

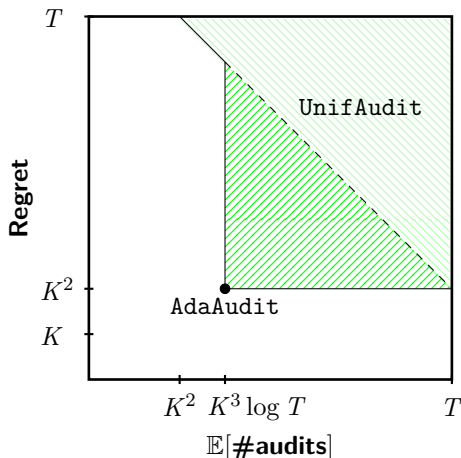


Figure: Regret vs $\mathbb{E}[\text{\#audits}]$

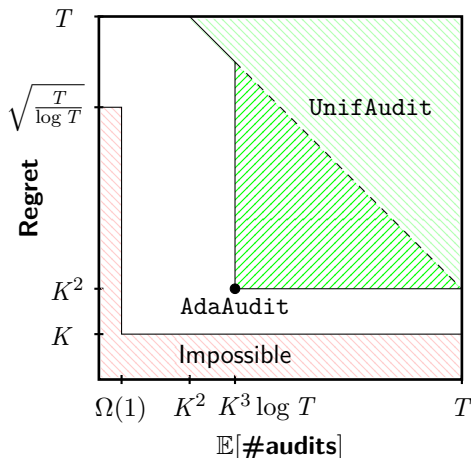
green possible; red impossible

**Main
Mechanism:
AdaAudit.**

$$\mathcal{R}_T \leq K^2 \text{ \& } \mathcal{B}_T = \mathcal{O}(K^3 \log T).$$

Tradeoff between Regret & #Audits

$\exists$ Perfect Bayesian Equilibrium (PBE) π^* , s.t.



Impossible.

$$\mathcal{R}_T = \Omega(K);$$

$$\mathcal{B}_T = \mathcal{O}(1) \implies$$

$$\mathcal{R}_T = \Omega\left(\sqrt{\frac{T}{\log T}}\right).$$

Figure: Regret vs $\mathbb{E}[\text{\#audits}]$

green possible; red impossible

One-Slide Technical Overview

1. **Future Punishment.** Audit reveals $v_{t,i} \neq u_{t,i} \Rightarrow$ never alloc again
2. **Adaptive Audits.**
3. **Learn via Flagging.**
4. **Auxiliary Games.**

One-Slide Technical Overview

1. **Future Punishment.**

Audit reveals $v_{t,i} \neq u_{t,i} \Rightarrow$ never alloc again

2. **Adaptive Audits.**

When i win in round t , audit w.p. $p_{t,i} := 1/V_i^{\text{alive}}$

($V_i^{\text{alive}} := \mathbb{E}_{\text{all agents truthful}} [\sum_{\text{future round}} \text{gain of agent } i]$)

$\Rightarrow$ (almost) always truthful & $\mathcal{B}_T = \tilde{\mathcal{O}}(1)$

(truthful: get $\geq 0 + V_i^{\text{alive}}$; lie: get $\leq 1 + (1 - p_{t,i}) V_i^{\text{alive}}$)

3. **Learn via Flagging.**

4. **Auxiliary Games.**

One-Slide Technical Overview

- Future Punishment.** Audit reveals $v_{t,i} \neq u_{t,i} \Rightarrow$ never alloc again
When i win in round t , audit w.p. $p_{t,i} := 1/V_i^{\text{alive}}$
($V_i^{\text{alive}} := \mathbb{E}_{\text{all agents truthful}}[\sum_{\text{future round}} \text{gain of agent } i]$)
 $\Rightarrow$ (almost) always truthful & $\mathcal{B}_T = \tilde{\mathcal{O}}(1)$
(truthful: get $\geq 0 + V_i^{\text{alive}}$; lie: get $\leq 1 + (1 - p_{t,i}) V_i^{\text{alive}}$)
- Adaptive Audits.** Need to estimate $\mathbb{E}_{\text{all agents truthful}}[\text{gain of agent } i]$ empirically but **can't “condition on”** concentration
(since agents can strategize early; happy to explain offline)
Idea: Let agents “flag” others for biased estimates
(“victims” benefit from truthfully flagging $\Rightarrow$ incentives aligned)
- Learn via Flagging.**
- Auxiliary Games.**

One-Slide Technical Overview

1. Future Punishment.

Audit reveals $v_{t,i} \neq u_{t,i} \Rightarrow$ never alloc again

When i win in round t , audit w.p. $p_{t,i} := 1/V_i^{\text{alive}}$

2. Adaptive Audits.

($V_i^{\text{alive}} := \mathbb{E}_{\text{all agents truthful}}[\sum_{\text{future round}} \text{gain of agent } i]$)

$\Rightarrow$ (almost) always truthful & $\mathcal{B}_T = \tilde{\mathcal{O}}(1)$

(truthful: get $\geq 0 + V_i^{\text{alive}}$; lie: get $\leq 1 + (1 - p_{t,i}) V_i^{\text{alive}}$)

Need to estimate $\mathbb{E}_{\text{all agents truthful}}[\text{gain of agent } i]$ empirically but **can't “condition on”** concentration

3. Learn via Flagging.

(since agents can strategize early; happy to explain offline)

Idea: Let agents “flag” others for biased estimates

(“victims” benefit from truthfully flagging $\Rightarrow$ incentives aligned)

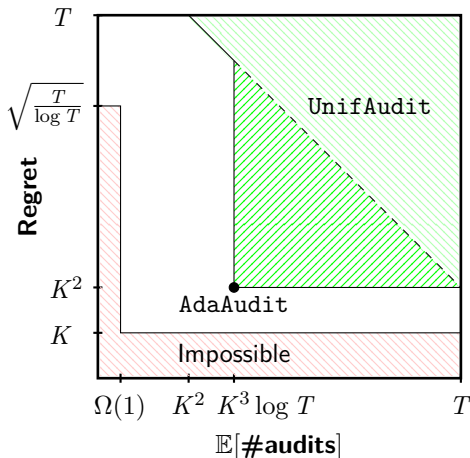
4. Auxiliary Games.

Can't use revelation principle due to unknown / non-unique distributions (happy to explain offline)

How to characterize PBE? Define a “well-behaved” aux game, show aux PBE $\xrightarrow{\text{induce}}$ actual PBE

Main Results & Takeaway

For resource allocation **without money & without dist info...**



Technical Ingredients

- **Future Punishment**
- **Adaptive Audits** ($\mathcal{O}(1)$ regret via $\tilde{\mathcal{O}}(1)$ audits)
- **Learn via Flagging**
("condition on" argument is problematic when strategic)
- **Auxiliary Games**
(revelation principle is inapplicable w/o dist info)

Thank you!

Paper link: <https://arxiv.org/abs/2502.08412>

References

- [ABDVVW22] Andres Alban, Philippe Blaettchen, Harwin De Vries, and Luk N Van Wassenhove. Resource allocation with sigmoidal demands: Mobile healthcare units and service adoption. *Manufacturing & Service Operations Management*, 24(6):2944–2961, 2022.
- [BGS19] Santiago R Balseiro, Huseyin Gurkan, and Peng Sun. Multiagent mechanism design without money. *Operations Research*, 67(5):1417–1436, 2019.
- [BJ24] Moise Blanchard and Patrick Jaillet. Near-optimal mechanisms for resource allocation without monetary transfers. *arXiv preprint arXiv:2408.10066*, 2024.
- [Cla71] Edward H Clarke. Multipart pricing of public goods. *Public choice*, pages 17–33, 1971.
- [GBI21] Artur Gorokh, Siddhartha Banerjee, and Krishnamurthy Iyer. From monetary to nonmonetary mechanism design via artificial currencies. *Mathematics of Operations Research*, 46(3):835–855, 2021.
- [Gro73] Theodore Groves. Incentives in teams. *Econometrica: Journal of the Econometric Society*, pages 617–631, 1973.
- [Pre17] Canice Prendergast. How food banks use markets to feed the poor. *Journal of Economic Perspectives*, 31(4):145–162, 2017.
- [Pre22] Canice Prendergast. The allocation of food to food banks. *Journal of Political Economy*, 130(8):1993–2017, 2022.
- [PSMST22] Sebastian Perez-Salazar, Ishai Menache, Mohit Singh, and Alejandro Toriello. Dynamic resource allocation in the cloud with near-optimal efficiency. *Operations Research*, 70(4):2517–2537, 2022.
- [PSÜY24] Parag A Pathak, Tayfun Sönmez, M Utku Ünver, and M Bumin Yenmez. Fair allocation of vaccines, ventilators and antiviral treatments: Leaving no ethical value behind in healthcare rationing. *Management Science*, 70(6):3999–4036, 2024.
- [Vic61] William Vickrey. Counterspeculation, auctions, and competitive sealed tenders. *The Journal of finance*, 16(1):8–37, 1961.
- [YBP23] Xuecheng Yin, İ Esra Büyüктаhtakin, and Bhumi P Patel. Covid-19: Data-driven optimal allocation of ventilator supply under uncertainty and risk. *European journal of operational research*, 304(1):255–275, 2023.