

Improving Information from Manipulable Data (*JEEA*, 2022)
&
Aligned Recommendations (*teaser*)

Alex Frankel

Navin Kartik



Allocation Problem

Designer uses data generated by an agent to make inference and assign an allocation
prefers to give “higher” types “higher” allocations

E.g., Amazon sees product reviews
determines which products to highlight

Allocation Problem

Designer uses data generated by an agent to make inference and assign an allocation
prefers to give “higher” types “higher” allocations

E.g., Amazon sees product reviews
determines which products to highlight

Given an allocation policy, agent can **manipulate data** to improve allocation

Allocation Problem

Designer uses data generated by an agent to make inference and assign an allocation
prefers to give “higher” types “higher” allocations

E.g., Amazon sees product reviews
determines which products to highlight
→ fake positive reviews



Given an allocation policy, agent can **manipulate data** to improve allocation

Allocation Problem

Designer uses data generated by an agent to make inference and assign an allocation
prefers to give “higher” types “higher” allocations

E.g., Amazon sees product reviews
determines which products to highlight
→ fake positive reviews



Given an allocation policy, agent can **manipulate data** to improve allocation

Manipulation **changes inference** about agent from data



Response to Manipulation

Allocation policy $\rightarrow$ agent manipulation $\rightarrow$ inference from data $\rightarrow$ allocation policy

Response to Manipulation

Allocation policy $\rightarrow$ agent manipulation $\rightarrow$ inference from data $\rightarrow$ allocation policy

- **Fixed point** policy: best response to itself / simultaneous-move Nash eqm
 - Policy is ex post optimal given data it induces
 - Might achieve through adaptive process

Response to Manipulation

Allocation policy $\rightarrow$ agent manipulation $\rightarrow$ inference from data $\rightarrow$ allocation policy

- **Fixed point** policy: best response to itself / simultaneous-move Nash eqm
 - Policy is ex post optimal given data it induces
 - Might achieve through adaptive process
- **Optimal** policy: commitment / Stackelberg solution
 - Maximizes designer's objective accounting for manipulation the policy generates
 - Ex ante but (perhaps) not ex post optimal

Response to Manipulation

Allocation policy $\rightarrow$ agent manipulation $\rightarrow$ inference from data $\rightarrow$ allocation policy

- **Fixed point** policy: best response to itself / simultaneous-move Nash eqm
 - Policy is ex post optimal given data it induces
 - Might achieve through adaptive process
- **Optimal** policy: commitment / Stackelberg solution
 - Maximizes designer's objective accounting for manipulation the policy generates
 - Ex ante but (perhaps) not ex post optimal

This paper: in a simple model,

- ① How does optimal policy compare to fixed point?
- ② What kind of ex post distortions?

Fixed Point vs Optimal (commitment) policy

① How does optimal policy compare to fixed point?

★ Optimal policy is **flatter** than fixed point

Less sensitive to manipulable data

② What ex post distortions are introduced?

★ Commit to **underutilize** data

Best response would put **more weight** on data

Fixed Point vs Optimal (commitment) policy

① How does optimal policy compare to fixed point?

★ **Optimal policy is flatter than fixed point**

Less sensitive to manipulable data

② What ex post distortions are introduced?

★ **Commit to underutilize data**

Best response would put more weight on data

Two interpretations of the result:

■ Designer with commitment power

- Google search, Amazon product rankings, Government targeting
- Positive or prescriptive

Fixed Point vs Optimal (commitment) policy

① How does optimal policy compare to fixed point?

★ **Optimal policy is flatter than fixed point**

Less sensitive to manipulable data

② What ex post distortions are introduced?

★ **Commit to underutilize data**

Best response would put more weight on data

Two interpretations of the result:

■ Designer with commitment power

- Google search, Amazon product rankings, Government targeting
- Positive or prescriptive

■ Allocation determined by competitive market

- Use of credit scores (lending) or test scores (college admissions)
- Market settles on ex post optimal allocations
- Interventions to improve allocation accuracy

Model

Setting

- Agent with type $(\eta, \gamma) \in \mathbb{R}^2$; known joint distribution (finite means and variances)
→ natural action η , gaming ability γ

- Designer wants to match allocation $y \in \mathbb{R}$ to η

$$\text{Utility} \equiv -(y - \eta)^2$$

- Allocation policy $Y(x)$, based agent's data/action $x \in \mathbb{R}$
- Agent chooses x based on (η, γ) and Y (details shortly)

Setting

- Agent with type $(\eta, \gamma) \in \mathbb{R}^2$; known joint distribution (finite means and variances)
→ natural action η , gaming ability γ
- Designer wants to match allocation $y \in \mathbb{R}$ to η
Utility $\equiv -(y - \eta)^2$
- Allocation policy $Y(x)$, based agent's data/action $x \in \mathbb{R}$
- Agent chooses x based on (η, γ) and Y (details shortly)

Expected loss for designer:

$$\text{Loss} \equiv \mathbb{E}[(Y(x) - \eta)^2]$$

Key decomposition:

$$\text{Loss} = \underbrace{\mathbb{E}[(\mathbb{E}[\eta|x] - \eta)^2]}_{\text{Info loss from estimating } \eta \text{ from } x} + \underbrace{\mathbb{E}[(Y(x) - \mathbb{E}[\eta|x])^2]}_{\text{Misallocation loss given estimation}}$$

Linearity assumptions

We assume

- Linear allocation policies for designer:

$$Y(x) = \beta x + \beta_0$$

So β is **allocation sensitivity**, will be strength of incentives

Linearity assumptions

We assume

- Linear allocation policies for designer:

$$Y(x) = \beta x + \beta_0$$

So β is **allocation sensitivity**, will be strength of incentives

- Agent has a linear response function:

Given policy (β, β_0) , agent of type (η, γ) chooses action

$$x = \eta + \beta\gamma$$

- Natural action $x = \eta$ in absence of incentives
- Higher gaming ability γ or stronger incentive $\beta \implies$ more manipulation

Nb: Linear agent response can be microfounded [e.g., agent's payoff $y - (x - \eta)^2 / (2m\gamma)$]

Linearity assumptions

We assume

- Linear allocation policies for designer:

$$Y(x) = \beta x + \beta_0$$

So β is **allocation sensitivity**, will be strength of incentives

- Agent has a linear response function:

Given policy (β, β_0) , agent of type (η, γ) chooses action

$$x = \eta + \beta\gamma$$

- Natural action $x = \eta$ in absence of incentives
- Higher gaming ability γ or stronger incentive $\beta \implies$ more manipulation

Nb: Linear agent response can be microfounded [e.g., agent's payoff $y - (x - \eta)^2 / (2m\gamma)$]

$\rightarrow$ For talk only, restrict to $\beta \geq 0$ and $\rho \equiv \text{Corr}(\eta, \gamma) \geq 0$

Key Idea

Allocation rule $Y(x) = \beta x + \beta_0$; Loss=Info Loss + Misalloc Loss

$\uparrow$ policy sensitivity $\beta \implies \uparrow$ info loss about variable of interest, η

intuitively $\because$ agent's action is given by $x = \eta + \beta\gamma$

(Not Blackwell unless $(\eta, \gamma) \sim$ normal; in general, quadratic loss of best linear estimator)

Key Idea

Allocation rule $Y(x) = \beta x + \beta_0$; Loss=Info Loss + Misalloc Loss

$\uparrow$ policy sensitivity $\beta \implies \uparrow$ info loss about variable of interest, η

intuitively $\because$ agent's action is given by $x = \eta + \beta\gamma$

(Not Blackwell unless $(\eta, \gamma) \sim$ normal; in general, quadratic loss of best linear estimator)

- Fixing agent behavior for policy β , OLS yields
“best response” policy $\hat{\beta}(\beta)$: best linear allocation
given agent's behavior

Key Idea

Allocation rule $Y(x) = \beta x + \beta_0$; Loss=Info Loss + Misalloc Loss

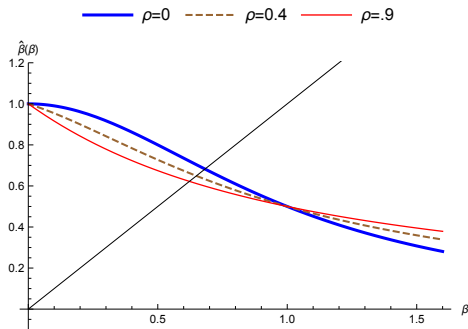
$\uparrow$ policy sensitivity $\beta \implies \uparrow$ info loss about variable of interest, η

intuitively $\because$ agent's action is given by $x = \eta + \beta\gamma$

(Not Blackwell unless $(\eta, \gamma) \sim$ normal; in general, quadratic loss of best linear estimator)

- Fixing agent behavior for policy β , OLS yields “best response” policy $\hat{\beta}(\beta)$: best linear allocation given agent's behavior

- Can show that $\hat{\beta}(\beta)$ is $\downarrow$ in β



Key Idea

Allocation rule $Y(x) = \beta x + \beta_0$; Loss=Info Loss + Misalloc Loss

$\uparrow$ policy sensitivity $\beta \implies \uparrow$ info loss about variable of interest, η

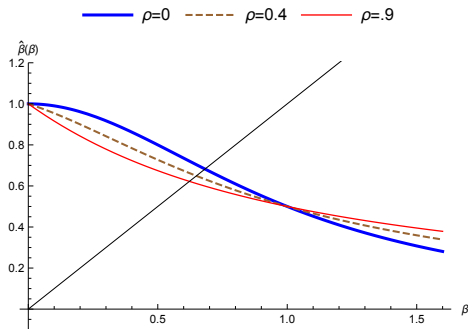
intuitively $\because$ agent's action is given by $x = \eta + \beta\gamma$

(Not Blackwell unless $(\eta, \gamma) \sim$ normal; in general, quadratic loss of best linear estimator)

- Fixing agent behavior for policy β , OLS yields “best response” policy $\hat{\beta}(\beta)$: best linear allocation given agent's behavior

- Can show that $\hat{\beta}(\beta)$ is $\downarrow$ in β

- Fixed-point policy has β^{fp} s.t. $\hat{\beta}(\beta^{\text{fp}}) = \beta^{\text{fp}}$
 - Can view as simultaneous-move Nash Eqm
 - Misallocation loss at fixed point is minimized, 0
 - But info loss can be large



Main Result

Main Result

Designer chooses policy $Y(x) = \beta x + \beta_0$; Agent response $x = \eta + m\beta\gamma$

Nb: Given $\rho \geq 0$, there is a unique fixed point, which is positive: $\beta^{\text{fp}} > 0$

Proposition

For the optimal policy's sensitivity β^* and the fixed point's sensitivity β^{fp} :

- ① (Flattening.) $0 < \beta^* < \beta^{\text{fp}}$.
- ② (Info underutilization.) $\hat{\beta}(\beta^*) > \beta^*$.

Main Result

Designer chooses policy $Y(x) = \beta x + \beta_0$; Agent response $x = \eta + m\beta\gamma$

Nb: Given $\rho \geq 0$, there is a unique fixed point, which is positive: $\beta^{\text{fp}} > 0$

Proposition

For the optimal policy's sensitivity β^* and the fixed point's sensitivity β^{fp} :

- ① (Flattening.) $0 < \beta^* < \beta^{\text{fp}}$.
- ② (Info underutilization.) $\hat{\beta}(\beta^*) > \beta^*$.

Intuition for Flattening:

Loss = Information loss + Misallocation loss.

- First order benefit of $\uparrow \beta$ from 0: constant policy minimizes info loss but not misalloc loss
- First order benefit of $\downarrow \beta$ from β^{fp} : fixed point minimizes misalloc loss, while info loss $\uparrow$ in β

Main Result

Designer chooses policy $Y(x) = \beta x + \beta_0$; Agent response $x = \eta + m\beta\gamma$

Nb: Given $\rho \geq 0$, there is a unique fixed point, which is positive: $\beta^{\text{fp}} > 0$

Proposition

For the optimal policy's sensitivity β^* and the fixed point's sensitivity β^{fp} :

- ① (Flattening.) $0 < \beta^* < \beta^{\text{fp}}$.
- ② (Info underutilization.) $\hat{\beta}(\beta^*) > \beta^*$.

Intuition for Flattening:

Loss = Information loss + Misallocation loss.

- First order benefit of $\uparrow \beta$ from 0: constant policy minimizes info loss but not misalloc loss
- First order benefit of $\downarrow \beta$ from β^{fp} : fixed point minimizes misalloc loss, while info loss $\uparrow$ in β

$\implies$ Local min in $(0, \beta^{\text{fp}})$. Establish it is global min.

Main Result

Designer chooses policy $Y(x) = \beta x + \beta_0$; Agent response $x = \eta + m\beta\gamma$

Nb: Given $\rho \geq 0$, there is a unique fixed point, which is positive: $\beta^{\text{fp}} > 0$

Proposition

For the optimal policy's sensitivity β^* and the fixed point's sensitivity β^{fp} :

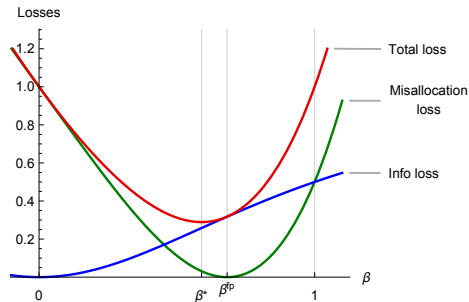
- ① (Flattening.) $0 < \beta^* < \beta^{\text{fp}}$.
- ② (Info underutilization.) $\hat{\beta}(\beta^*) > \beta^*$.

Intuition for Flattening:

Loss = Information loss + Misallocation loss.

- First order benefit of $\uparrow \beta$ from 0: constant policy minimizes info loss but not misalloc loss
- First order benefit of $\downarrow \beta$ from β^{fp} : fixed point minimizes misalloc loss, while info loss $\uparrow$ in β

$\Rightarrow$ Local min in $(0, \beta^{\text{fp}})$. Establish it is global min.



Discussion

- Can nonlinear allocation policies do better? We believe generally yes
 - But linear policies are simple, canonical, practical
 - Straightforward to interpret and discuss sensitivity to and (under/over)utilization of data

Discussion

- Can nonlinear allocation policies do better? We believe generally yes
 - But linear policies are simple, canonical, practical
 - Straightforward to interpret and discuss sensitivity to and (under/over)utilization of data
- If designer wants to reduce manipulation costs, $\downarrow \beta^*$; reinforces punchline
- If manipulation is productive effort, $\uparrow \beta^*$; could alter punchline if sufficiently important

Discussion

- Can nonlinear allocation policies do better? We believe generally yes
 - But linear policies are simple, canonical, practical
 - Straightforward to interpret and discuss sensitivity to and (under/over)utilization of data
- If designer wants to reduce manipulation costs, $\downarrow \beta^*$; reinforces punchline
- If manipulation is productive effort, $\uparrow \beta^*$; could alter punchline if sufficiently important
- Crucial asymmetry in agent behavior $x = \eta + \beta\gamma$
 - If designer instead only wants to match allocation to γ , logic flips $\rightarrow \beta^* > \beta^{\text{fp}}$
 - More generally, $\text{sign}[\beta^* < \beta^{\text{fp}}]$ depends on objective's weight on matching η vs. γ

Discussion

- Can nonlinear allocation policies do better? We believe generally yes
 - But linear policies are simple, canonical, practical
 - Straightforward to interpret and discuss **sensitivity to** and **(under/over)utilization of** data
- If designer wants to reduce manipulation costs, $\downarrow \beta^*$; reinforces punchline
- If manipulation is productive effort, $\uparrow \beta^*$; could alter punchline if sufficiently important
- Crucial asymmetry in agent behavior $x = \eta + \beta\gamma$
 - If designer instead only wants to match allocation to γ , logic flips $\rightarrow \beta^* > \beta^{\text{fp}}$
 - More generally, $\text{sign}[\beta^* < \beta^{\text{fp}}]$ depends on objective's weight on matching η vs. γ
- Hopeful that future research:
counterparts to **flattening** / **underutilizing information** in general allocation problems?

Related Literature

■ Framework of “muddled information” in economics

- Prendergast Topel 1996; Fischer Verrecchia 2000; Benabou Tirole 2006;
Frankel Kartik 2019
- Björkegren Blumenstock Knight 2020
- **Ball 2021**

■ Related flattening to ↓ manipulation in other economic contexts

- Dynamic screening: Bonatti & Cisternas 2019
- Finance: Bond & Goldstein 2015; Boleslavsky, Kelly & Taylor 2017

■ Strategic classification in CS/ML

→ often binary (nonlinear) allocation problems

- Hardt Megiddo Papadimitriou Wooters 2016; Hu Immorlica Vaughan 2019;
Milli Miller Dragan Hardt 2018; Kleinberg Raghavan 2019; Braverman Garg 2019

Aligned Recommendations

(Exceedingly Preliminary)

Motivation

- Algorithms observe user's choices $\rightarrow$ predicts "type" $\rightarrow$ makes recommendations
- Assume that user & algorithm both want accurate predictions/recommendations

Motivation

- Algorithms observe user's choices → predicts “type” → makes recommendations
- Assume that user & algorithm both want accurate predictions/recommendations

Yet users may **distort behavior** to shape future predictions

Motivation

- Algorithms observe user's choices → predicts “type” → makes recommendations
- Assume that user & algorithm both want accurate predictions/recommendations

Yet users may **distort behavior** to shape future predictions

Web advice on incognito mode:

Use It to Not Ruin Your Recommendations

The main reason I use Incognito mode is for YouTube. If I'm doing a one-off search, like for a guide on how to replace the washer on a faucet, I go incognito because I don't want YouTube recommending similar tutorials for the next six months. It's so easy to break the recommendation algorithm if you aren't careful.

The same applies when clicking YouTube video links on Reddit that I don't want to be part of my profile, and when searching for random products on Amazon.

Cen, Ilyas, Allen, Li, Madry (2024):

“Nearly half of participants self-report strategizing in the wild, e.g. ignoring content they actually like to avoid being over-recommended similar content later.”

Motivation

- Algorithms observe user's choices → predicts “type” → makes recommendations
- Assume that user & algorithm both want accurate predictions/recommendations

Yet users may **distort behavior** to shape future predictions

Web advice on incognito mode:

Use It to Not Ruin Your Recommendations

The main reason I use Incognito mode is for YouTube. If I'm doing a one-off search, like for a guide on how to replace the washer on a faucet, I go incognito because I don't want YouTube recommending similar tutorials for the next six months. It's so easy to break the recommendation algorithm if you aren't careful.

The same applies when clicking YouTube video links on Reddit that I don't want to be part of my profile, and when searching for random products on Amazon.

Cen, Ilyas, Allen, Li, Madry (2024):

“Nearly half of participants self-report strategizing in the wild, e.g. ignoring content they actually like to avoid being over-recommended similar content later.”

Algorithm should recognize such distortions

Our Framework

Sender-receiver **common-interest signaling game** with **restricted message space**

- User with private type (θ, η) , where
 - Fundamental component $\theta \sim N(0, \sigma_\theta^2)$
 - Transitory component $\eta = \theta + \varepsilon$, with $\varepsilon \sim N(0, \sigma_\varepsilon^2)$
- User takes action $a \in \mathbb{R}$
- Algorithm observes a and makes inference $T(a) \in \mathbb{R}$ about θ

Our Framework

Sender-receiver **common-interest signaling game** with **restricted message space**

- User with private type (θ, η) , where
 - Fundamental component $\theta \sim N(0, \sigma_\theta^2)$
 - Transitory component $\eta = \theta + \varepsilon$, with $\varepsilon \sim N(0, \sigma_\varepsilon^2)$
- User takes action $a \in \mathbb{R}$
- Algorithm observes a and makes inference $T(a) \in \mathbb{R}$ about θ

Payoffs ($\kappa \in [0, 1]$ is known parameter):

$$\text{User: } -\kappa(a - \eta)^2 - (1 - \kappa)(T - \theta)^2$$

$$\text{Algorithm: } -(T - \theta)^2$$

- So agent wants to take $a = \eta$, but also wants $T(a) = \theta$
- Will distort action from η to help algorithm infer θ

Our Framework

Sender-receiver **common-interest signaling game** with **restricted message space**

- User with private type (θ, η) , where
 - Fundamental component $\theta \sim N(0, \sigma_\theta^2)$
 - Transitory component $\eta = \theta + \varepsilon$, with $\varepsilon \sim N(0, \sigma_\varepsilon^2)$
- User takes action $a \in \mathbb{R}$
- Algorithm observes a and makes inference $T(a) \in \mathbb{R}$ about θ

Payoffs ($\kappa \in [0, 1]$ is known parameter):

$$\text{User: } -\kappa(a - \eta)^2 - (1 - \kappa)(T - \theta)^2$$

$$\text{Algorithm: } -(T - \theta)^2$$

- So agent wants to take $a = \eta$, but also wants $T(a) = \theta$
- Will distort action from η to help algorithm infer θ

We solve for **linear equilibrium**: $A(\theta, \eta) = \alpha\theta + \beta\eta$ and $T(a) = \gamma a$

Observations & What's To Come

- There are distortions, despite alignment over the signaling component (obviously)
- Ex ante user welfare higher with algorithm than without (in best eqm)
 - general point $\therefore$ common-interest signaling

Observations & What's To Come

- There are distortions, despite alignment over the signaling component (obviously)
- Ex ante user welfare higher with algorithm than without (in best eqm)
 - general point $\therefore$ common-interest signaling
- Some limits are clear, e.g.,
 - $\kappa = 0$: no signaling motive, $A^*(\theta, \eta) = \eta$
 - $\kappa = 1$: only signaling motive, $A^*(\theta, \eta) = \theta$ (best eqm)

Observations & What's To Come

- There are distortions, despite alignment over the signaling component (obviously)
- Ex ante user welfare higher with algorithm than without (in best eqm)
 - general point $\therefore$ common-interest signaling
- Some limits are clear, e.g.,
 - $\kappa = 0$: no signaling motive, $A^*(\theta, \eta) = \eta$
 - $\kappa = 1$: only signaling motive, $A^*(\theta, \eta) = \theta$ (best eqm)

In progress:

- Lessons from non-limits, comparative statics
- Also things like dynamic microfoundations