

Rational Inattention via Ignorance Equivalence

Michèle Müller-Ippen (University of St. Gallen)

joint with

Roc Armenter (Federal Reserve Bank of Philadelphia)

Zachary Stangebye (University of Notre Dame)

April 20th, 2024

The views expressed here do not reflect those of the Federal Reserve Board.

Main Contribution

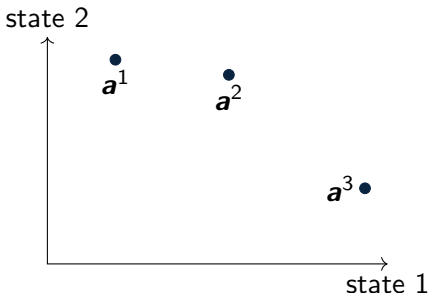
- We propose the novel concept of the **ignorance equivalent** as a parsimonious summary of a decision problem with flexible learning.
 - Ability to learn is '*equivalent*' to access to a larger menu
 - Learning analogue of the certainty equivalent for lotteries

- We propose the novel concept of the **ignorance equivalent** as a parsimonious summary of a decision problem with flexible learning.
 - Ability to learn is '*equivalent*' to access to a larger menu
 - Learning analogue of the certainty equivalent for lotteries
- The equivalence between learning ability and fictional payoffs simplifies analysis, including in multi-player settings.

Model Setup

- An agent is choosing an action from a finite menu $\mathcal{A}$.
- Action $\mathbf{a} \in \mathcal{A} \subseteq \mathbb{R}^I$ yields utility a_i in state $i \in \mathcal{I} = 1, \dots, I$.
- The agent has a prior belief $\pi \in \Delta \mathcal{I}$ over states.

- **Payoffs, states and actions**



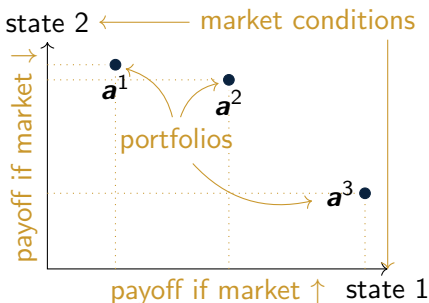
Model Setup

- An agent is choosing an action from a finite menu $\mathcal{A}$.
- Action $\mathbf{a} \in \mathcal{A} \subseteq \mathbb{R}^I$ yields utility a_i in state $i \in \mathcal{I} = 1, \dots, I$.
- The agent has a prior belief $\pi \in \Delta \mathcal{I}$ over states.

- Payoffs, states and actions



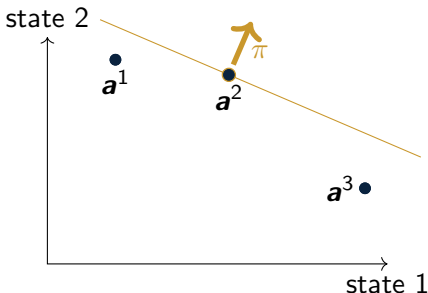
Investor Example



Exogenous Information Baseline

- An agent is choosing an action from a finite menu $\mathcal{A}$.
- Action $\mathbf{a} \in \mathcal{A} \subseteq \mathbb{R}^I$ yields utility a_i in state $i \in \mathcal{I} = 1, \dots, I$.
- The agent has a prior belief $\pi \in \Delta \mathcal{I}$ over states.

- **No information**
 - Agent maximizes expected utility under prior belief π .

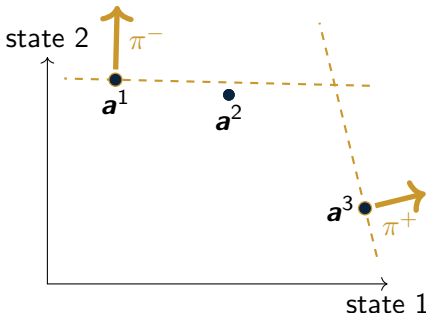


Exogenous Information Baseline

- An agent is choosing an action from a finite menu $\mathcal{A}$.
- Action $\mathbf{a} \in \mathcal{A} \subseteq \mathbb{R}^I$ yields utility a_i in state $i \in \mathcal{I} = 1, \dots, I$.
- The agent has a prior belief $\pi \in \Delta \mathcal{I}$ over states.

- **Incomplete information**

- Agent observes signal draw s ,
- updates belief to π^s ,
- maximizes expected utility under this belief.

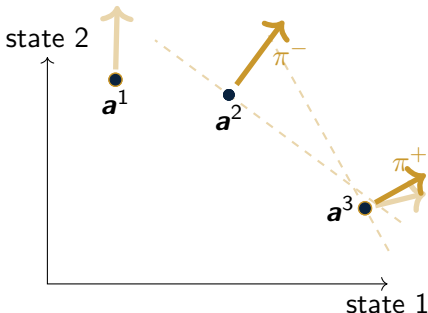


Rational Inattention

- An agent is choosing an action from a finite menu $\mathcal{A}$.
- Action $\mathbf{a} \in \mathcal{A} \subseteq \mathbb{R}^I$ yields utility a_i in state $i \in \mathcal{I} = 1, \dots, I$.
- The agent has a prior belief $\pi \in \Delta \mathcal{I}$ over states.

- **Rational Inattention**

- Agent chooses *which* costly signal to draw,
- ... and proceeds as before.



Welfare: $W(\mathcal{A}, \pi, c) = \mathbb{E}[\text{consumption utility}] - [\text{cost of signal}]$.

- Agent can choose *any* conditional signal $\mathcal{S} = \langle S, \mathbf{q} \rangle$, where $q_i \in \Delta S$ denotes the probability of each realization in state i .

Learning Strategies

- Agent can choose *any* conditional signal $S = \langle S, \mathbf{q} \rangle$, where $q_i \in \Delta S$ denotes the probability of each realization in state i .
 - Obedience principle: Signal recommends actions, $S = \mathcal{A}$, agent obeys. We refer to these as **learning strategies**.

- Agent can choose *any* **learning strategy** $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$, where $q_i(\mathbf{a})$ denotes the likelihood of taking action $\mathbf{a}$ conditional on state i .

- Agent can choose *any* **learning strategy** $\mathcal{S} = \langle \mathcal{A}, \mathbf{q} \rangle$, where $q_i(\mathbf{a})$ denotes the likelihood of taking action $\mathbf{a}$ conditional on state i .
- Admissible costs include all differentiable & prior-concave UPS costs.
Prominent examples:
 - Mutual Information (*Sims JME'03*),
 - some Tsallis costs (*Caplin-Dean-Leahy '19*)
 - Total Information (*Bloedel-Zhong '20*), which subsumes
 - Wald costs (*Morris-Strack '19*), Fisher Information (*Hébert-Woodford '20*).

Ignorance Equivalent

The Ignorance Equivalent

Definition

Payoff vector α is an **ignorance equivalent** of RI problem $(\mathcal{A}, \pi, c)$ if

The Ignorance Equivalent

Definition

Payoff vector α is an **ignorance equivalent** of RI problem $(\mathcal{A}, \pi, c)$ if

- the agent is willing to commit to *always* implement α ,

$$W(\{\alpha\}, \pi, c) \geq W(\mathcal{A}, \pi, c).$$

Agent is willing to forgo learning opportunities that are present in $\mathcal{A}$.

The Ignorance Equivalent

Definition

Payoff vector α is an **ignorance equivalent** of RI problem $(\mathcal{A}, \pi, c)$ if

- the agent is willing to commit to *always* implement α ,

$$W(\{\alpha\}, \pi, c) \geq W(\mathcal{A}, \pi, c).$$

Agent is willing to forgo learning opportunities that are present in $\mathcal{A}$.

- the agent is willing to commit to *never* implement α ,

$$W(\mathcal{A}, \pi, c) \geq W(\mathcal{A} \cup \{\alpha\}, \pi, c).$$

Agent is willing to forgo learning opportunities that arise when α is added to $\mathcal{A}$.

The Ignorance Equivalent

Definition

Payoff vector α is an **ignorance equivalent** of RI problem $(\mathcal{A}, \pi, c)$ if

- the agent is willing to commit to *always* implement α ,

$$W(\{\alpha\}, \pi, c) \geq W(\mathcal{A}, \pi, c).$$

- the agent is willing to commit to *never* implement α ,

$$W(\mathcal{A}, \pi, c) \geq W(\mathcal{A} \cup \{\alpha\}, \pi, c).$$

Since larger menus are always better,

$$\{\alpha\} \sim \mathcal{A} \sim \mathcal{A} \cup \{\alpha\}$$

in terms of the agent's preference over menus given cost c and prior π .

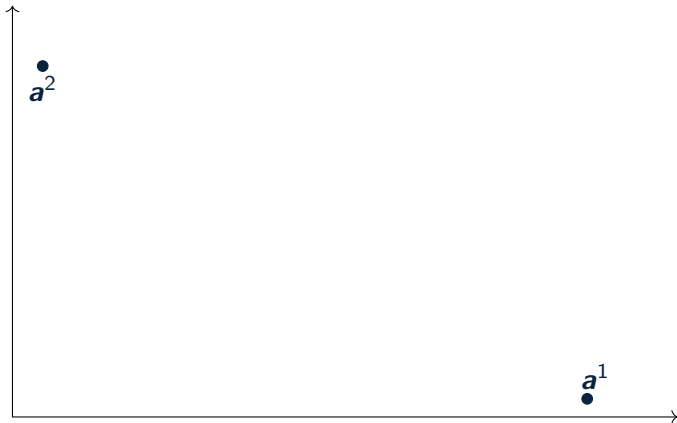
Existence and Uniqueness

Theorem 1

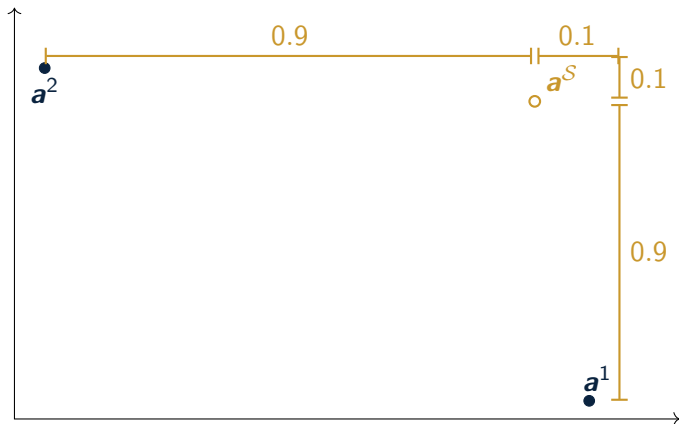
Each RI problem $(\mathcal{A}, \pi, c)$ admits a unique ignorance equivalent α .
The ignorance equivalent can be constructed from any optimal signal $\mathcal{S}$.

Ignorance equivalence: $W(\{\alpha\}, \pi, c) = W(\mathcal{A}, \pi, c) = W(\mathcal{A} \cup \{\alpha\}, \pi, c)$.

Locating the Ignorance Equivalent

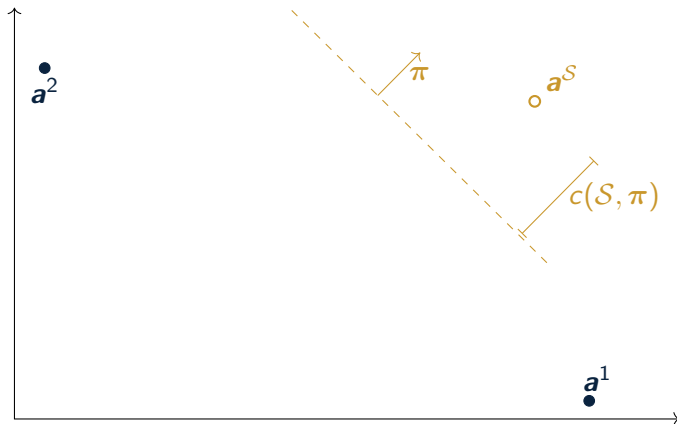


Locating the Ignorance Equivalent



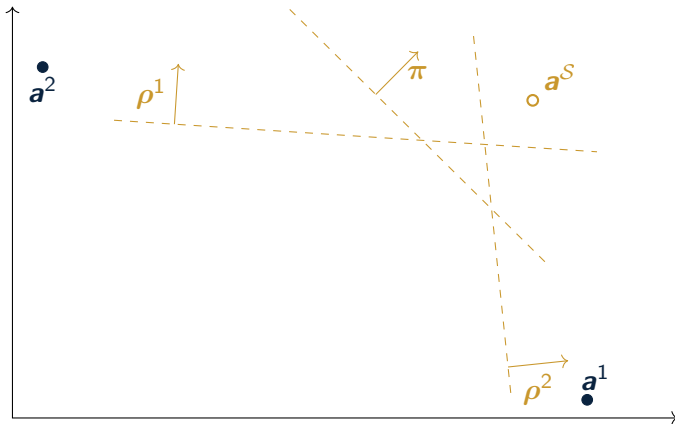
- Strategy $\mathcal{S}$ always implements the high-payoff action with $q_i(\mathbf{a}^i) = 0.9$.
- In each state i , the expected consumption payoff is $a_i^{\mathcal{S}} = \sum_{\mathbf{a} \in \mathcal{A}} q_i(\mathbf{a}) a_i$.

Locating the Ignorance Equivalent



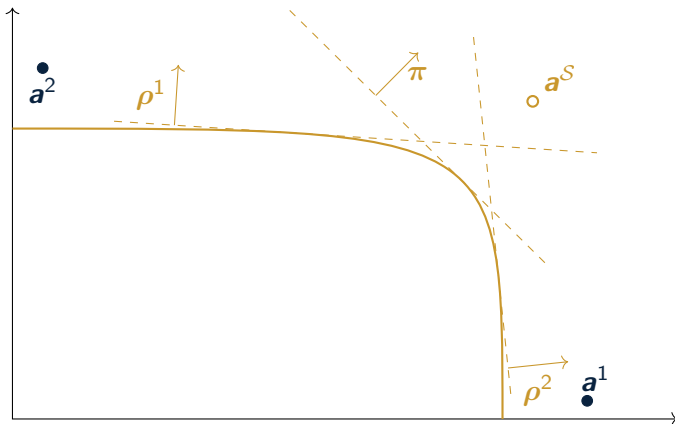
- Following S yields net utility $\pi \cdot a^S - c(S, \pi)$ under prior π .

Locating the Ignorance Equivalent



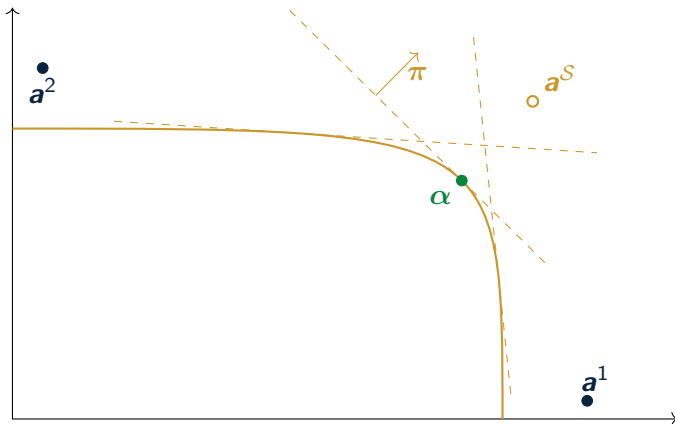
- The same construction shows net utility from S under any belief ρ .

Locating the Ignorance Equivalent



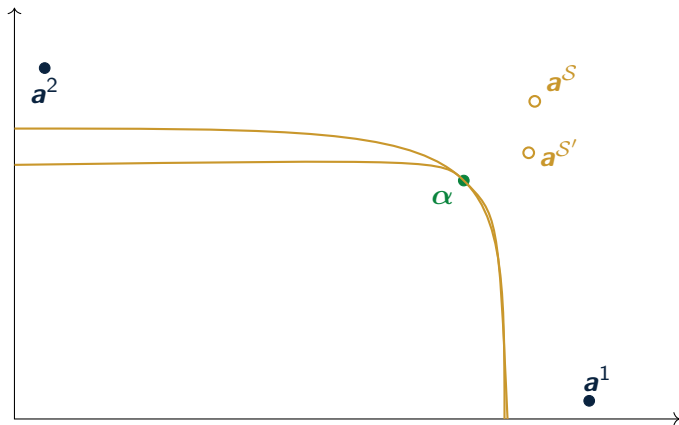
- Dominance: $\mathbf{x} \succsim \mathcal{S} \iff \rho \cdot \mathbf{x} \leq \rho \cdot \mathbf{a}^{\mathcal{S}} - c(\mathcal{S}, \rho) \forall \rho.$

Locating the Ignorance Equivalent



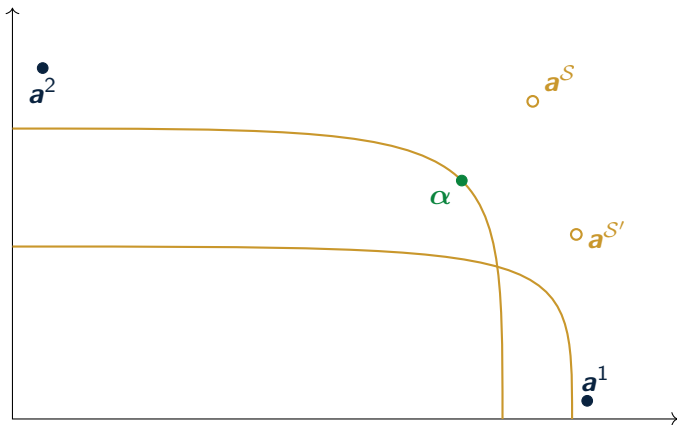
- Unconditional implementation of the ignorance equivalent α
 - is just as good as the optimal S under the prior.
 - is no better than S under *any* belief.

Identifying optimal signals



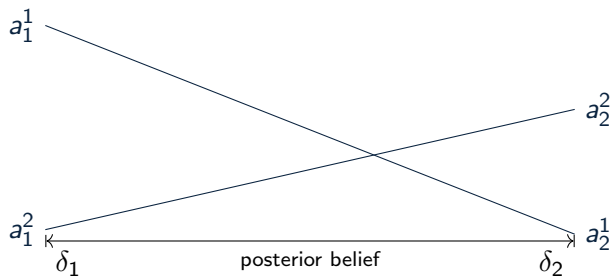
- Optimal signals are exactly those that beat α under any belief.

A suboptimal signal



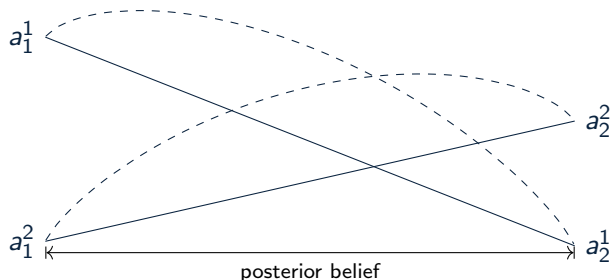
- Optimal signals are exactly those that beat α under any belief.

Posterior-based construction (Kamenica&Gentzkow'11, Caplin&Dean'13)



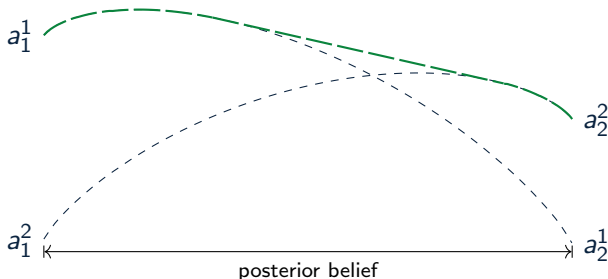
- Posterior belief on axis.
- For each action, draw a line to represent expected utility.

Posterior-based construction (Kamenica&Gentzkow'11, Caplin&Dean'13)



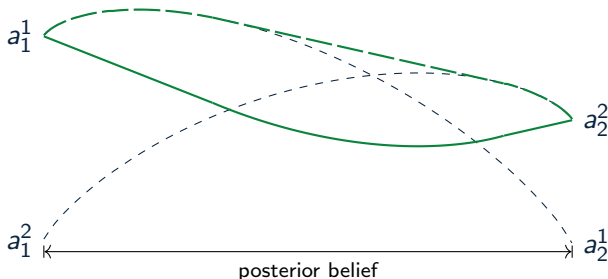
- Posterior belief on axis.
- For each action, draw a line to represent expected utility.
- Subtract the cost potential ϕ to obtain net payoffs.

Posterior-based construction (Kamenica&Gentzkow'11, Caplin&Dean'13)



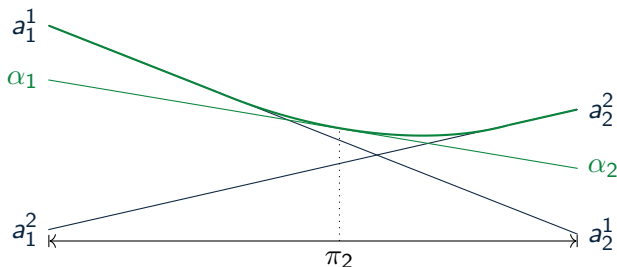
- Posterior belief on axis.
- For each action, draw a line to represent expected utility.
- Subtract the cost potential ϕ to obtain net payoffs.
- Take the concave upper envelope to obtain net utilities with learning.

Posterior-based construction (Kamenica&Gentzkow'11, Caplin&Dean'13)



- Posterior belief on axis.
- For each action, draw a line to represent expected utility.
- Subtract the cost potential ϕ to obtain net payoffs.
- Take the concave upper envelope to obtain net utilities with learning.
- Add back ϕ to obtain equivalent 'no learning' utilities.

Posterior-based construction (Kamenica&Gentzkow'11, Caplin&Dean'13)



- Posterior belief on axis.
- For each action, draw a line to represent expected utility.
- Subtract the cost potential ϕ to obtain net payoffs.
- Take the concave upper envelope to obtain net utilities with learning.
- Add back ϕ to obtain equivalent 'no learning' utilities.
- Supporting hyperplane at prior $\pi \Rightarrow$ ignorance equivalent α .

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Corollary 1: If α is also the ignorance equivalent of $(\mathcal{A}, \pi', c)$, then the two RI problems have the same set of optimal learning strategies.

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Corollary 1: If α is also the ignorance equivalent of $(\mathcal{A}, \pi', c)$, then the two RI problems have the same set of optimal learning strategies.

- ★ is sufficient to identify which menu additions $\mathbf{a}^+ \in \mathbb{R}^I$ are welfare enhancing, $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) > W(\mathcal{A}, \pi, c)$.

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Corollary 1: If α is also the ignorance equivalent of $(\mathcal{A}, \pi', c)$, then the two RI problems have the same set of optimal learning strategies.

- ★ is sufficient to identify which menu additions $\mathbf{a}^+ \in \mathbb{R}^I$ are welfare enhancing, $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) > W(\mathcal{A}, \pi, c)$.

Theorem 3: $\mathbf{a}^+$ adds welfare to $\mathcal{A}$ $\iff$ $\mathbf{a}^+$ adds welfare to $\{\alpha\}$.

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Corollary 1: If α is also the ignorance equivalent of $(\mathcal{A}, \pi', c)$, then the two RI problems have the same set of optimal learning strategies.

- ★ is sufficient to identify which menu additions $\mathbf{a}^+ \in \mathbb{R}^I$ are welfare enhancing, $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) > W(\mathcal{A}, \pi, c)$.

Theorem 3: $\mathbf{a}^+$ adds welfare to $\mathcal{A}$ $\iff$ $\mathbf{a}^+$ adds welfare to $\{\alpha\}$.

- does not distort learning under *any* prior.

Properties of the Ignorance Equivalent

The ignorance equivalent α of RI problem $(\mathcal{A}, \pi, c)$...

- characterizes all optimal learning strategies w/o reference to beliefs.

Corollary 1: If α is also the ignorance equivalent of $(\mathcal{A}, \pi', c)$, then the two RI problems have the same set of optimal learning strategies.

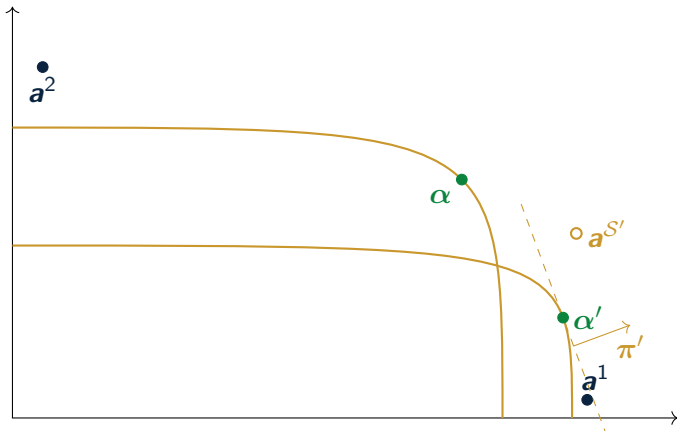
- ★ is sufficient to identify which menu additions $\mathbf{a}^+ \in \mathbb{R}^I$ are welfare enhancing, $W(\mathcal{A} \cup \{\mathbf{a}^+\}, \pi, c) > W(\mathcal{A}, \pi, c)$.

Theorem 3: $\mathbf{a}^+$ adds welfare to $\mathcal{A}$ $\iff$ $\mathbf{a}^+$ adds welfare to $\{\alpha\}$.

- does not distort learning under *any* prior.

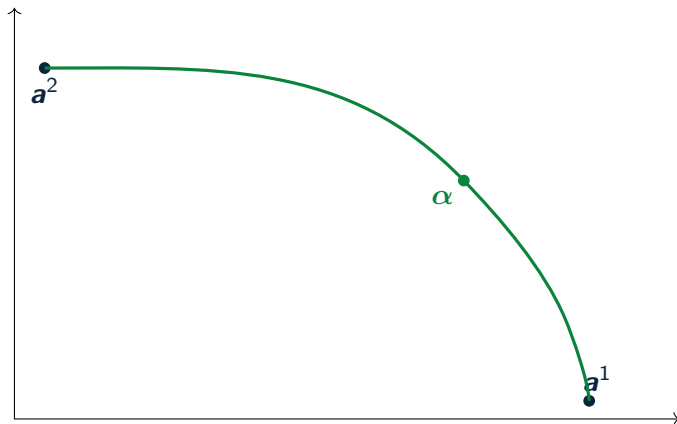
Corollary 3: If learning strategy $\langle \mathcal{A}, \mathbf{q} \rangle$ is optimal in RI problem $(\mathcal{A}, \pi', c)$, then it is also optimal in $(\mathcal{A} \cup \{\alpha\}, \pi', c)$.

Self-selection Property (Corollary 3)



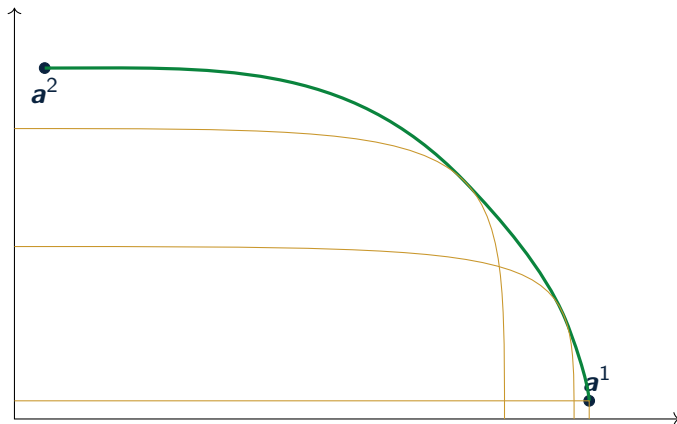
- Stopping at α' is no better than continuing with S' .

Learning-Proof Menu



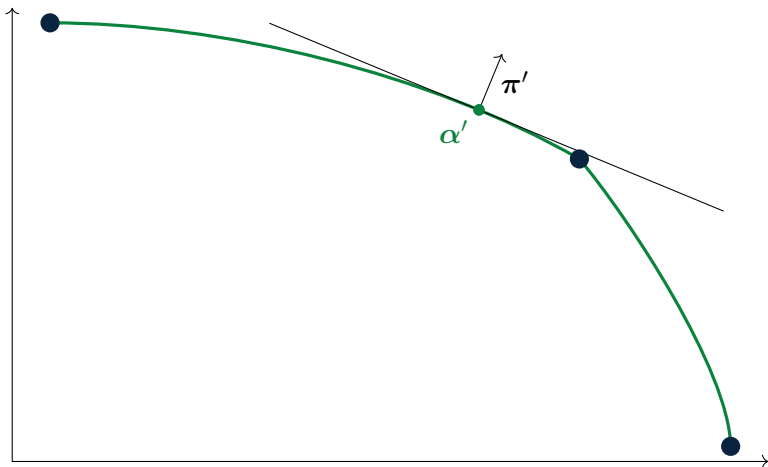
- Together, ignorance equivalents form the **Learning-Proof menu** $\bar{\mathcal{A}}$.

Learning-Proof Menu



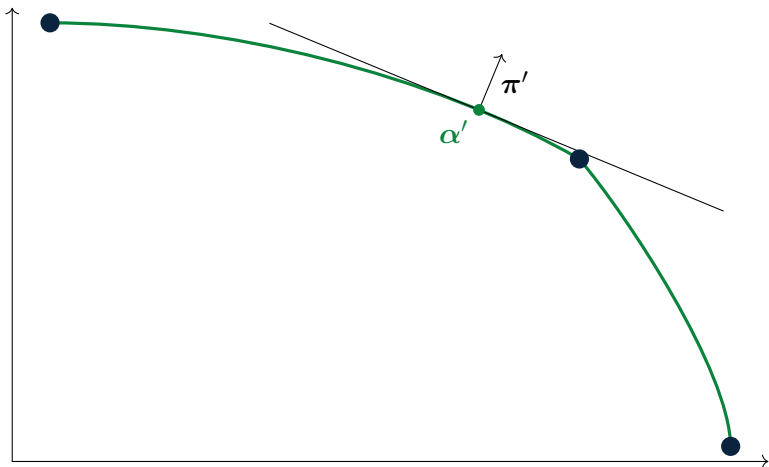
- The learning-proof menu can also be constructed from the signals.

Learning-Proof Menu



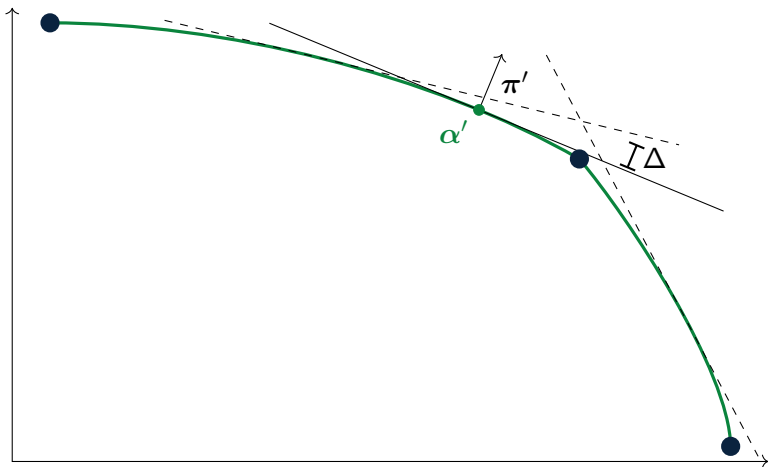
The Learning-Proof Menu is a **EU representation** of the RI problem:
Ability to learn is *as if* agent has access to $\bar{\mathcal{A}}$ rather than $\mathcal{A}$.

Learning-Proof Menu



Offering the learning-proof menu $\bar{\mathcal{A}}$ is helpful for belief elicitation.

Learning-Proof Menu



Welfare impact Δ of exogenous information is intuitive in $\bar{\mathcal{A}}$.

Certainty Equivalent vs Ignorance Equivalent

Certainty equivalent:

- Condenses the appeal of a lottery into a scalar.
- Useful for comparisons across lotteries given a *fixed belief*.

Certainty Equivalent vs Ignorance Equivalent

Certainty equivalent:

- Condenses the appeal of a lottery into a scalar.
- Useful for comparisons across lotteries given a *fixed belief*.
- 'Too' parsimonious for situations with learning:
 - The ignorance equivalent abstracts away from all state dependence.
 - An agent who learns seeks to tailor the choice to the realized state.

Certainty Equivalent vs Ignorance Equivalent

Certainty equivalent:

- Condenses the appeal of a lottery into a scalar.
- Useful for comparisons across lotteries given a *fixed belief*.
- ‘Too’ parsimonious for situations with learning:
 - The ignorance equivalent abstracts away from all state dependence.
 - An agent who learns seeks to tailor the choice to the realized state.

Ignorance equivalent:

- Condenses the entire menu $\mathcal{A}$ into one payoff vector α .
- Retains enough detail to capture learning opportunities.
- Useful for comparisons across signals, menus, and beliefs.

Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability.

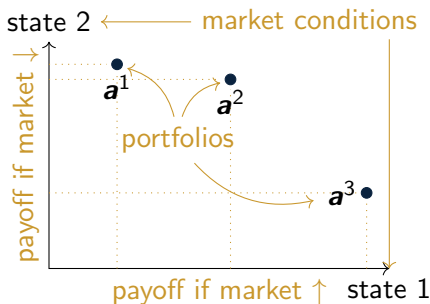
Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

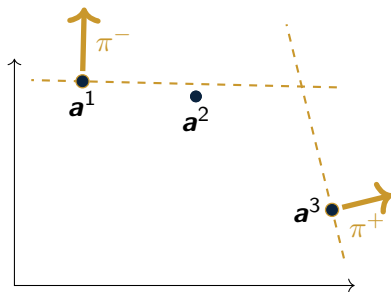


Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

- Investor learns before investing.

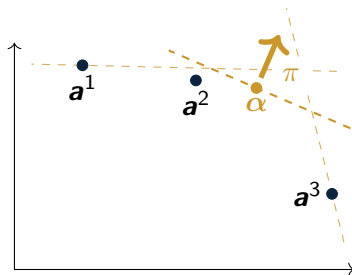


Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

- Investor learns before investing.
- α describes a portfolio that
 - is purchased unconditionally
 - makes her no better off.

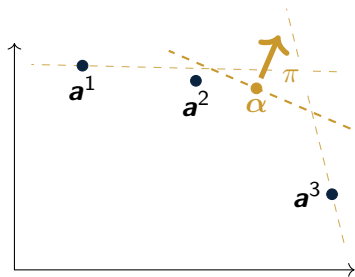


Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

An informed fund manager (he) can offer any payouts $\mathbf{a} \in \mathbb{R}^I$.



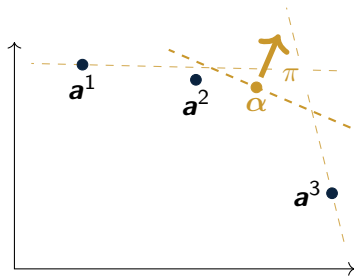
Economic Relevance

The ignorance equivalent and the learning proof menu are not just ‘mental shortcuts’ for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

An informed fund manager (he) can offer any payouts $\mathbf{a} \in \mathbb{R}^I$.

- α is what he wants to offer!
 - Unconditional acceptance.
 - Avoids adverse selection.
 - Reveals no free information.
 - Socially optimal.



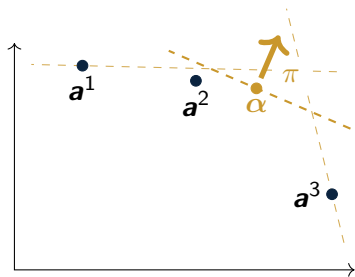
Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

An informed fund manager (he) can offer any payouts $\mathbf{a} \in \mathbb{R}^I$.

- What if prior π is unknown?



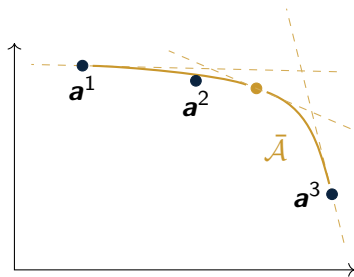
Economic Relevance

The ignorance equivalent and the learning proof menu are not just 'mental shortcuts' for tractability. They arise naturally in economic settings where the menu is designed strategically.

Toy example. A risk-neutral RI investor (she) wants to purchase one of the portfolios in $\mathcal{A}$.

An informed fund manager (he) can offer any payouts $\mathbf{a} \in \mathbb{R}^I$.

- What if prior π is unknown?
 - The manager can offer $\bar{\mathcal{A}}$.



Optimal Allocation

Task: Allocate a single opportunity to one of several RL-agents.

Optimal Allocation

Task: Allocate a single opportunity to one of several RI-agents.

Question: Can agents achieve the first best surplus through trade?

Optimal Allocation

Task: Allocate a single opportunity to one of several RI-agents.

Question: Can agents achieve the first best surplus through trade?

Key complication: Agents can learn and selectively participate only in those contingencies that are most favorable for them.

Optimal Allocation

Task: Allocate a single opportunity to one of several RI-agents.

Question: Can agents achieve the first best surplus through trade?

Key complication: Agents can learn and selectively participate only in those contingencies that are most favorable for them.

- If there is a 'most able' agent, this generalizes the investor example. We construct the full set of terms that implement the first best.
 - ▶ *Crémer&Khalil'92, but with flexible learning*

Optimal Allocation

Task: Allocate a single opportunity to one of several RI-agents.

Question: Can agents achieve the first best surplus through trade?

Key complication: Agents can learn and selectively participate only in those contingencies that are most favorable for them.

- If there is a 'most able' agent, this generalizes the investor example. We construct the full set of terms that implement the first best.
 - ▶ *Crémer&Khalil'92, but with flexible learning*
- What if agents can access different menus and have comparative advantages in learning?
 - For example, teams problem in technology acquisition.

short version

long version

Main Contribution

‘Ignorance Equivalent’ approach to Rational Inattention

- Simplifies intuition

Main Contribution

‘Ignorance Equivalent’ approach to Rational Inattention

- Simplifies intuition
- Yields novel insights

Main Contribution

‘Ignorance Equivalent’ approach to Rational Inattention

- Simplifies intuition
- Yields novel insights
- Economically relevant

Main Contribution

‘Ignorance Equivalent’ approach to Rational Inattention

- Simplifies intuition
- Yields novel insights
- Economically relevant
 - when learning is to be avoided.

‘Ignorance Equivalent’ approach to Rational Inattention

- Simplifies intuition
- Yields novel insights
- Economically relevant
 - *not only* when learning is to be avoided.