

Preferences for Flexibility and Randomization under Uncertainty

Kota Saito

Presenter: Runfa Hu

Micro Theory Reading Group, ACEM, SJTU

April 2026

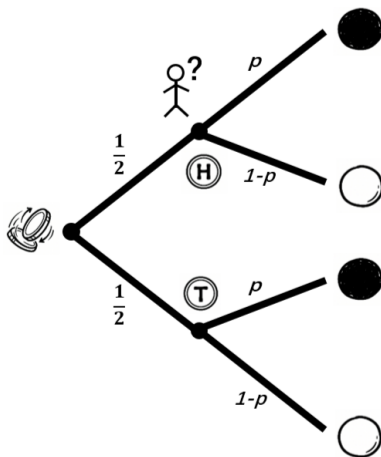
A Hypothetical Experiment

- Bet 1:
 - ▶ A box contains 50 black balls and 50 white balls;
 - ▶ Guess a color, then randomly take one;
 - ▶ Win: \$100, Lose: \$0.
- Bet 2:
 - ▶ A box contains 100 black and white balls;
 - ▶ The ratio is unknown.
- Which one do you prefer?
- Lab experiment: bet 1 $\succ$ bet 2
 - ▶ People seem to have “uncertainty aversion.”
 - ▶ Is it “reasonable?”

Raiffa's Critique

- Consider a coin-toss strategy:
 - ▶ bet on black if heads appears,
 - ▶ bet on white if tails appears.
- Suppose the ratio - [Black : White = $p : 1 - p$]
 - ▶ $\text{Pr}(\text{win}) = \text{Pr}(\text{heads}) \cdot \text{Pr}(\text{black}) + \text{Pr}(\text{tails}) \cdot \text{Pr}(\text{white})$
$$= \frac{1}{2} \cdot p + \frac{1}{2}(1 - p) = \frac{1}{2}$$
 - ▶ bet 1 $\sim$ bet 2
- By tossing a fair coin, you can eliminate the effect of uncertainty!
- Really?

Critique of Raiffa's Critique



- Stage 1: You toss a coin ...
- Stage 2: You stare blankly at the heads-up coin ...
- The coin in your hand tells you nothing about the ratio p !

A Natural Question

- Our trouble:
 - ▶ Experiment: $\text{bet } 1 \succ \text{bet } 2$
 - ▶ Raiffa's Critique: $\text{bet } 1 \sim \text{bet } 2$
 - ▶ Critique of Raiffa's Critique: $\text{bet } 1 \succ \text{bet } 2$
- Two different kinds of people: whether or not you believe the coin-toss strategy eliminates the effects of uncertainty.
- Both beliefs are legitimate: which one is correct depends on you.
- A natural question: what are their differences (of preferences)?

Our Goal

- What are we talking about when we discuss “the differences of people who might believe or disbelieve the uncertainty elimination?”
- How to capture the difference? Does there exist an index that describes the difference (for example, $\delta \in [0, 1]$)?
- How to characterize the index (δ)? What does it mean for different levels of the index?

Back to the Hypothetical Experiment

	p	$1 - p$	$p \in [0, 1]$
	state 1: black	state 2: white	
bet on black f_1	1	0	
bet on white f_2	0	1	
$\frac{1}{2}f_1 + \frac{1}{2}f_2$	$\frac{1}{2}$	$\frac{1}{2}$	

- Bet: state + a set of acts (payoff profiles)
- **believe** in uncertainty elimination: $\{\frac{1}{2}f_1 + \frac{1}{2}f_2\} \sim \{f_1, f_2\}$;
- **disbelieve** uncertainty elimination: $\{\frac{1}{2}f_1 + \frac{1}{2}f_2\} \succ \{f_1, f_2\}$;
- We're talking about **preferences on the sets of payoff profiles!**

Generalization of our Model

	p_1	p_2	p_3	$\cdots$	p_n
	state 1	state 2	state 3	$\cdots$	state n
f_1	r_{11}	r_{12}	r_{13}	$\cdots$	r_{1n}
f_2	r_{21}	r_{22}	r_{23}	$\cdots$	r_{2n}
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$

- $S = \{1, 2, \dots, n\}$ is the set of states
- $p \in C \subseteq \Delta(S)$ is the uncertain probability
- $\mathcal{F} = [a, b]^n$ is all possible payoff profiles ($a \leq r_{ij} \leq b$)
- $\mathcal{A}$ is the set of all nonempty compact subsets of $\mathcal{F}$
- $A \in \mathcal{A}$ is a set of profiles
- We care about binary relation $\succsim$ on $\mathcal{A}$!

Our Goal

- What are we talking about when we discuss “the differences of people who might believe or disbelieve the uncertainty elimination?”
 - ▶ We study the preferences over sets of payoff profiles! ✓
- How to capture the difference? Does there exist an index that describes the difference (for example, $\delta \in [0, 1]$)?
- How to characterize the index (δ)? What does it mean for different levels of the index?

Agent who Believes in Uncertainty Elimination

	p	$1 - p$	$p \in [0, 1]$
	state 1: black	state 2: white	
guess black f_1	1	0	$\times \rho_1$
guess white f_2	0	1	$\times \rho_2$
$\rho_1 f_1 + \rho_2 f_2$	ρ_1	ρ_2	

$$U(A) = \max_{\rho \in \Delta(A)} u \left(\sum_{f \in A} \rho(f) f \right), \text{ where } u(f) := \min_{p \in C} \sum_{s \in S} p(s) f(s)$$

$$\Rightarrow U(\{f_1, f_2\}) = \frac{1}{2} = U \left(\left\{ \frac{1}{2} f_1 + \frac{1}{2} f_2 \right\} \right) \Rightarrow \left\{ \frac{1}{2} f_1 + \frac{1}{2} f_2 \right\} \sim \{f_1, f_2\}$$

Agent who Disbelieves Uncertainty Elimination

	p	$1 - p$	$p \in [0, 1]$
	state 1: black	state 2: white	
guess black f_1	1	0	$\times \rho_1$
guess white f_2	0	1	$\times \rho_2$
$\rho_1 f_1 + \rho_2 f_2$	ρ_1	ρ_2	

$$U(A) = \max_{\rho \in \Delta(A)} \sum_{f \in A} \rho(f) u(f), \text{ where } u(f) := \min_{p \in C} \sum_{s \in S} p(s) f(s)$$

$$\Rightarrow U(\{f_1, f_2\}) = 0 < \frac{1}{2} = U\left(\left\{\frac{1}{2}f_1 + \frac{1}{2}f_2\right\}\right) \Rightarrow \left\{\frac{1}{2}f_1 + \frac{1}{2}f_2\right\} \succ \{f_1, f_2\}$$

Combine Two Kinds of People

- Pessimistic agent **believes** in uncertainty elimination:

$$U(A) = \max_{\rho \in \Delta(A)} u \left(\sum_{f \in A} \rho(f) f \right), \quad \text{where } u(f) = \min_{p \in C} \sum_{s \in S} p(s) f(s)$$

- Pessimistic agent **disbelieves** uncertainty elimination:

$$U(A) = \max_{\rho \in \Delta(A)} \sum_{f \in A} \rho(f) u(f), \quad \text{where } u(f) = \min_{p \in C} \sum_{s \in S} p(s) f(s)$$

- Combining them naively, we get the **random uncertainty-averse (RUA)** utility representation:

$$U(A) = \max_{\rho \in \Delta(A)} \left[\delta u \left(\sum_{f \in A} \rho(f) f \right) + (1 - \delta) \sum_{f \in A} \rho(f) u(f) \right]$$

The Range of RUA

all possible preferences
over sets of payoff profiles

RUA utility representation

?

- We can observe: $(\succsim, \mathcal{A})$
- We have RUA utility $U : \mathcal{A} \rightarrow \mathbb{R}$
- Not all preferences have an RUA representation utility
- We need an if-and-only-if condition!

Some Axioms

Axiom 1 (Weak Order): $\succsim$ is complete and transitive.

Axiom 2 (Continuity): For any $f, g \in \mathcal{F}$ and $A \in \mathcal{A}$, if $\{f\} \succ A \succ \{g\}$ then $\alpha\{f\} + (1 - \alpha)A \succ \beta\{g\} + (1 - \beta)A$ for some α and β in $(0, 1)$.

Axiom 3 (Monotonicity): For any $f, g \in \mathcal{F}$, if $f(s) \geq g(s)$ for all $s \in S$ then $\{f\} \succsim \{g\}$. Moreover, if $f(s) > g(s)$ for all $s \in S$ then $\{f\} \succ \{g\}$.

Axiom 4 (Uncertainty Aversion): For any $f, g \in \mathcal{F}$, if $\{f\} \sim \{g\}$ then $\{\frac{1}{2}f + \frac{1}{2}g\} \succsim \{f\}$.

Axiom 5 (Certainty Set Independence): For any $x \in [a, b]$, $A, B \in \mathcal{A}$, and $\alpha \in [0, 1]$,

$A \succsim B \Leftrightarrow \alpha A + (1 - \alpha)\{(x, \dots, x)\} \succsim \alpha B + (1 - \alpha)\{(x, \dots, x)\}$.

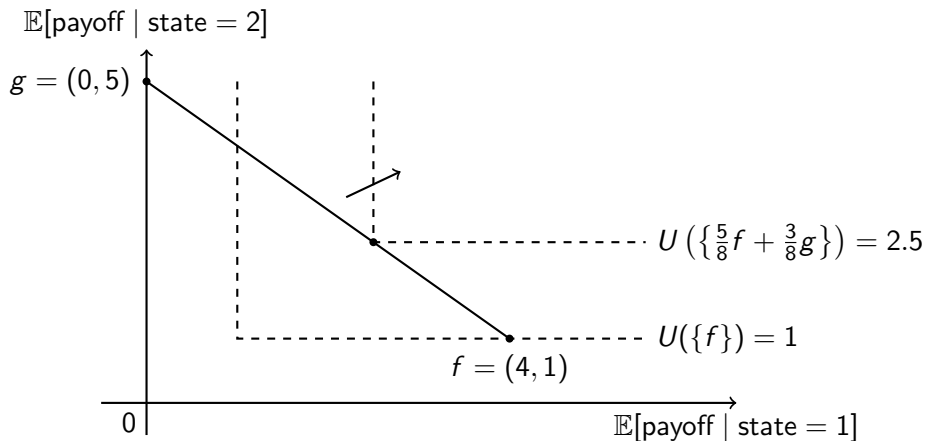
Kreps's (1979) Strategic Rationality Axiom

- $A, B \in \mathcal{A}$ if $A \succsim B$, then $A \sim A \cup B$
- This is **not true** for people who **believe** in the uncertainty elimination!

	p	$1 - p$	$p \in [0, 1]$
	state 1	state 2	
f	4	1	$\times \rho_1$
g	0	5	$\times \rho_2$
$\rho_1 f + \rho_2 g$	$4\rho_1$	$\rho_1 + 5\rho_2$	

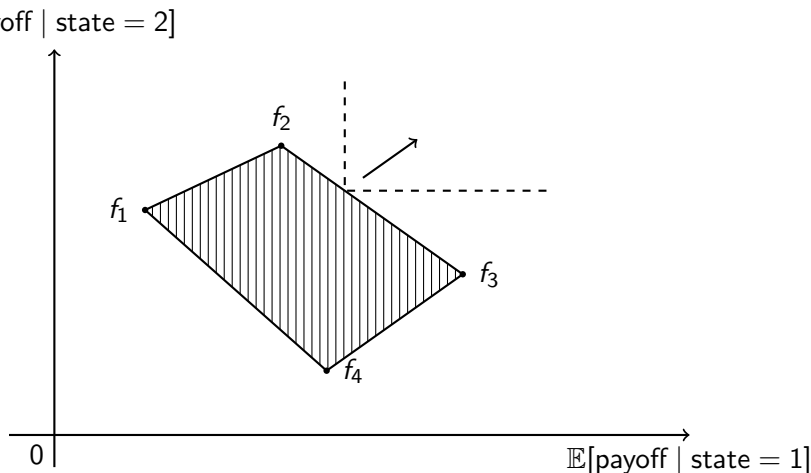
- Claim: $U(\{f\}) > U(\{g\})$ but $U(\{f\}) < U(\{f, g\})$

Kreps's (1979) Strategic Rationality Axiom



- $U(\{f\}) = 1 > 0 = U(\{g\})$
- but $U(\{f\}) = 1 < 2.5 = U(\{\frac{5}{8}f + \frac{3}{8}g\}) = U(\{f, g\})$

Kreps's (1979) Strategic Rationality Axiom



- $A \sim \text{conv}(A)$
- $U(\{f_1, f_2, f_3, f_4\}) = U(\text{conv}(\{f_1, f_2, f_3, f_4\}))$

Kreps's (1979) Strategic Rationality Axiom

- $A, B \in \mathcal{A}$ if $A \succsim B$, then $A \sim A \cup B$
- This is **true** for people who **disbelieve** the uncertainty elimination!

	p	$1 - p$	$p \in [0, 1]$
	state 1	state 2	
f	4	1	$\times \rho_1$
g	0	5	$\times \rho_2$

- $u(\{f\}) = 1, u(\{g\}) = 0$
- $U(\{f, g\}) = \max_{\rho} (\rho_1 \cdot 1 + \rho_2 \cdot 0) = 1 = u(\{f\})$
- If $\{f\} \succ \{g\}$, then the agent will **ignore** g in $\{f, g\}$

Certainty Strategic Rationality

all possible preferences
over sets of payoff profiles

RUA utility representation

$\delta = 0$

$\delta = 1$

- When $\delta = 1$, we need $A \sim \text{conv}(A)$
- When $\delta = 0$, we need
$$A \succsim B \Rightarrow A \sim A \cup B$$
- To capture all $\delta \in [0, 1]$, we need an axiom weaker than both of them!

Axiom 6 Certainty Strategic Rationality

For any $f, g \in \mathcal{F}$ and $x \in [a, b]$,

$$\{(x, \dots, x)\} \succsim \{f, g\} \Rightarrow$$

$$\{(x, \dots, x)\} \sim \{(x, \dots, x), f, g\}$$

Dominance

Definition (Dominance)

For all $\rho, \mu \in \Delta(\mathcal{F})$ such that $\rho = \rho_1 f_1 \oplus \cdots \oplus \rho_m f_m$,

$\mu = \mu_1 g_1 \oplus \cdots \oplus \mu_k g_k$, ρ **dominates** μ if

- ① $x(\rho_1 f_1 + \cdots + \rho_m f_m) \geq x(\mu_1 g_1 + \cdots + \mu_k g_k)$; and
- ② $\rho_1 x(f_1) + \cdots + \rho_m x(f_m) \geq \mu_1 x(g_1) + \cdots + \mu_k x(g_k)$,

where $x(f) := (x, \dots, x)$ such that $f \sim (x, \dots, x)$.

Axiom 7 (Dominance): For any $A, B \in \mathcal{A}$, if for all $\mu \in \Delta(B)$, there exists $\rho \in \Delta(A)$ such that ρ dominates μ , then $A \succsim B$.

The Desired Theorem

Theorem

$\succsim$ satisfies **[Weak Order, Continuity, Monotonicity, Uncertainty Aversion, Certain Set Independence, Certainty Strategic Rationality, and Dominance]** *if and only if* there exists a pair (δ, C) of $\delta \in [0, 1]$ and $C \subset \Delta(S)$ such that $\succsim$ is represented by a function $U : \mathcal{A} \rightarrow \mathbb{R}$ defined by

$$U(A) = \max_{\rho \in \Delta(A)} \left[\delta u \left(\sum_{f \in A} \rho(f) f \right) + (1 - \delta) \sum_{f \in A} \rho(f) u(f) \right]$$

where $u(f) := \min_{p \in C} \sum_{s \in S} p(s) f(s)$ and C is compact and convex.

Our Goal

- What are we talking about when we discuss “the differences of people who might believe or disbelieve the uncertainty elimination?”
 - ▶ We study the preferences over sets of payoff profiles! ✓
- How to capture the difference? Does there exist an index that describes the difference (for example, $\delta \in [0, 1]$)?
 - ▶ We capture the difference by introducing RUA and δ ! ✓
 - ▶ Through axiomatization, we depict all preferences representable by RUA! ✓
- How to characterize the index (δ)? What does it mean for different levels of the index?

$\delta = 1$ or 0

- Indifference to Randomization Axiom: $A \sim \text{conv}(A)$
- Strategic Rationality Axiom: $A \succsim B \Rightarrow A \sim A \cup B$

Proposition

Suppose that $\succsim$ is represented by an RUA representation with (δ, C) .

- $\succsim$ satisfies the indifference to randomization axiom **if and only if** $\delta = 1$.
- $\succsim$ satisfies the strategic rationality axiom **if and only if** $\delta = 0$.

$$\delta_1 \geq \delta_2$$

Definition (Preference for Flexibility)

$\succsim_1$ is said to exhibit a stronger **preference for flexibility** than $\succsim_2$ if, for every $A \in \mathcal{A}$ and $f \in A$,

$$A \succsim_2 \{f\} \Rightarrow A \succsim_1 \{f\}.$$

Definition (Preference for Randomization)

For any RUA preference relations $\succsim_1$ and $\succsim_2$, $\succsim_1$ is said to exhibit a stronger **preference for randomization** than $\succsim_2$ if for any $A \in \mathcal{A}$,

$$(\Delta_{\succsim_2}(A)) \setminus A \neq \emptyset \Rightarrow (\Delta_{\succsim_1}(A)) \setminus A \neq \emptyset,$$

where

$$\Delta_{\succsim}(A) := \arg \max_{\rho \in \Delta(A)} [\delta u(\sum_{f \in A} \rho(f)f) + (1 - \delta) \sum_{f \in A} \rho(f)u(f)].$$

$$\delta_1 \geq \delta_2$$

Proposition

Suppose $\succsim_i$ is represented by an RUA representation with (δ_i, C_i) for each $i \in \{1, 2\}$ and $C_1 = C_2$, where C_1 is not a singleton. Then, the following conditions are **equivalent**:

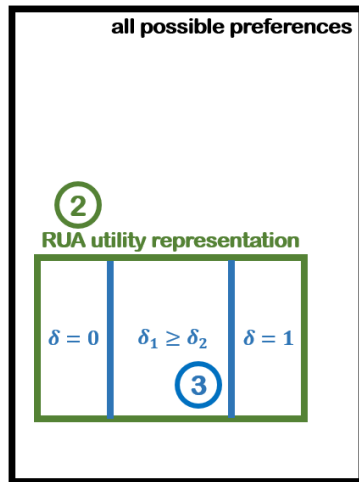
- 1 $\delta_1 \geq \delta_2$;
- 2 $\succsim_1$ exhibits a stronger preference for flexibility than $\succsim_2$;
- 3 $\succsim_1$ exhibits a stronger preference for randomization than $\succsim_2$;

Our Goal

- What are we talking about when we discuss “the differences of people who might believe or disbelieve the uncertainty elimination?”
 - ▶ We study the preferences over sets of payoff profiles! ✓
- How to capture the difference? Does there exist an index that describes the difference (for example, $\delta \in [0, 1]$)?
 - ▶ We capture the difference by introducing RUA and δ ! ✓
 - ▶ Through axiomatization, we depict all preferences representable by RUA! ✓
- How to characterize the index (δ)? What does it mean for different levels of the index?
 - ▶ We characterized δ in each case of $\delta = 0$, $\delta = 1$, $\delta_1 \geq \delta_2$ intuitively! ✓

Conclusion

The differences of people who might **believe or disbelieve** the coin-toss strategy can **eliminate uncertainty**



The differences of **preferences** over **sets of payoff profiles**

Takeaways

What's the difference between people who believe the coin-toss strategy can eliminate uncertainty and those who disbelieve it?

- People who believe it try to make use of the bad choices, while people who disbelieve it just ignore them.
- People who believe it have a stronger preference for flexibility and randomization.