

Lecture 3: Cognitive Imprecision and Choice under Risk

Michael Woodford

Columbia University

A Short Course on Cognitive Imprecision
Northwestern University
November 17, 2025

Risk Attitudes in the Laboratory

- We have already discussed one puzzling aspect of choices observed in the laboratory: behavior not (close to) risk-neutral, even when **stakes are quite small**

— and we have argued that a model of choice based on a **noisy representation** of the gambles offered provides a simple explanation for this

Risk Attitudes in the Laboratory

- We have already discussed one puzzling aspect of choices observed in the laboratory: behavior not (close to) risk-neutral, even when **stakes are quite small**

— and we have argued that a model of choice based on a **noisy representation** of the gambles offered provides a simple explanation for this
- But even more puzzling for EUT: subjects' apparent degree of risk aversion — even whether choices are **risk averse** or **risk seeking** — vary depending on the **nature** of the gambles offered

Kahneman and Tversky (1979)

Problem

In addition to whatever you own, you have been given 1000. You are now asked to choose between (a) winning an additional 500 with certainty, or (b) a gamble with a 50 percent chance of winning 1000 and a 50 percent chance of winning nothing.

Kahneman and Tversky (1979)

Problem

In addition to whatever you own, you have been given 1000. You are now asked to choose between (a) winning an additional 500 with certainty, or (b) a gamble with a 50 percent chance of winning 1000 and a 50 percent chance of winning nothing.

Problem

In addition to whatever you own, you have been given 2000. You are now asked to choose between (a) losing 500 with certainty, and (b) a gamble with a 50 percent chance of losing 1000 and a 50 percent chance of losing nothing.

Kahneman and Tversky (1979)

Problem

In addition to whatever you own, you have been given 1000. You are now asked to choose between (a) winning an additional 500 with certainty, or (b) a gamble with a 50 percent chance of winning 1000 and a 50 percent chance of winning nothing.

84% of subjects choose (a)

Problem

In addition to whatever you own, you have been given 2000. You are now asked to choose between (a) losing 500 with certainty, and (b) a gamble with a 50 percent chance of losing 1000 and a 50 percent chance of losing nothing.

69% of subjects choose (b)

Kahneman and Tversky: “Isolation Effect”

- Problem for EUT: in both cases, subjects are choosing between the **same** probability distributions over **final wealth levels**:

(a) initial wealth + 1500 with certainty

VS

(b) 50 percent chance of initial wealth + 1000,
50 percent chance of initial wealth + 2000

Kahneman and Tversky: “Isolation Effect”

- Problem for EUT: in both cases, subjects are choosing between the **same** probability distributions over **final wealth levels**:

(a) initial wealth + 1500 with certainty

VS

(b) 50 percent chance of initial wealth + 1000,
50 percent chance of initial wealth + 2000

- K-T explanation: people **don't** integrate the initial gain with subsequent gains/losses to evaluate choices in terms of final wealth

— instead, consider second-stage gains/losses only, ignoring the initial gain because it is **common** to both choices: **“isolation effect”**

Kahneman and Tversky: “Reflection Effect”

- This hypothesis renders the two problems no longer **equivalent** — but we need a further hypothesis to explain the result:
 - modal subject is **risk-averse** when choice is framed as between a certain gain and a random **gain**, but instead **risk-seeking** when it's framed as between a certain loss and a random **loss**
 - it isn't the **amount of risk** that explains the degree of penalty for risk, but whether **gains or losses** are involved

Kahneman and Tversky: “Reflection Effect”

- This hypothesis renders the two problems no longer **equivalent** — but we need a further hypothesis to explain the result:
 - modal subject is **risk-averse** when choice is framed as between a certain gain and a random **gain**, but instead **risk-seeking** when it's framed as between a certain loss and a random **loss**
 - it isn't the **amount of risk** that explains the degree of penalty for risk, but whether **gains or losses** are involved
- K-T postulate that changing the **sign** of the prospective payoffs, while preserving their **magnitudes** and probabilities, flips the sign of the typical subject's risk attitude
 - they call this the **“reflection effect”**

Risk Attitude as Adaptation to Cognitive Noise

- These patterns of behavior are not mysterious, if choices are based on a **noisy representation** of the problem

Risk Attitude as Adaptation to Cognitive Noise

- These patterns of behavior are not mysterious, if choices are based on a **noisy representation** of the problem
- Consider a choice between (a) initial transfer Y , plus an additional $C > 0$ with certainty, or (b) initial transfer Y , plus an additional $X > 0$ with probability p (otherwise, no additional amount)
 - If $\mathbf{r}$ is the internal representation of quantities Y, X, p, C , then a decision rule that maximizes DM's expected financial wealth [equivalent to max'ing $E[U(W)]$ if these quantities are all **small relative to the curvature of $U(W)$**] will be to choose (b) if and only if

$$E[Y | \mathbf{r}] + E[pX | \mathbf{r}] > E[Y | \mathbf{r}] + E[C | \mathbf{r}]$$

Risk Attitude as Adaptation to Cognitive Noise

- ... choose (b) if and only if

$$E[Y | \mathbf{r}] + E[pX | \mathbf{r}] > E[Y | \mathbf{r}] + E[C | \mathbf{r}]$$

or equivalently, if and only if

$$E[pX | \mathbf{r}] > E[C | \mathbf{r}]$$

Risk Attitude as Adaptation to Cognitive Noise

- ... choose (b) if and only if

$$E[Y | \mathbf{r}] + E[pX | \mathbf{r}] > E[Y | \mathbf{r}] + E[C | \mathbf{r}]$$

or equivalently, if and only if

$$E[pX | \mathbf{r}] > E[C | \mathbf{r}]$$

- If the parts of $\mathbf{r}$ that convey information about p, X, C are **independent** of the part that depends on Y , then the probability that this holds is **independent of the value of Y** — the “isolation effect”

Risk Attitude as Adaptation to Cognitive Noise

- In addition, if we suppose that the way that the numerical magnitudes X and C are encoded (and the prior over their possible values) are **the same** whether these amounts of money represent **gains or losses**, then the condition for choosing (b) over (a) [the **risk-seeking** (or less risk-averse) choice] in the problem with risky vs. certain **gains**,

$$E[pX | r] > E[C | r],$$

will instead be the condition for choosing (a) over (b) [the **risk-averse** (or less risk-seeking) choice] in the problem with risky vs. certain **losses**

Risk Attitude as Adaptation to Cognitive Noise

- Then if (a) is the **modal** choice in the problem involving gains (indicating **risk aversion** if $C = pX$), we should expect (b) to be the modal choice in the problem involving losses (indicating **risk seeking** if $C = pX$)
— the “reflection effect”

Risk Attitude as Adaptation to Cognitive Noise

- Then if (a) is the **modal** choice in the problem involving gains (indicating **risk aversion** if $C = pX$), we should expect (b) to be the modal choice in the problem involving losses (indicating **risk seeking** if $C = pX$)

— the “reflection effect”

- This will be true **regardless** of whether the inequality

$$E[pX | r] > E[C | r]$$

holds less than 1/2 the time (even though $C = pX$) because the average value of $E[X | r_x]$ is a concave function of X [as implied by the model of logarithmic coding discussed in Lecture 2], or because the average value of $E[p | r_p]$ is smaller than p [as could be true, see below]

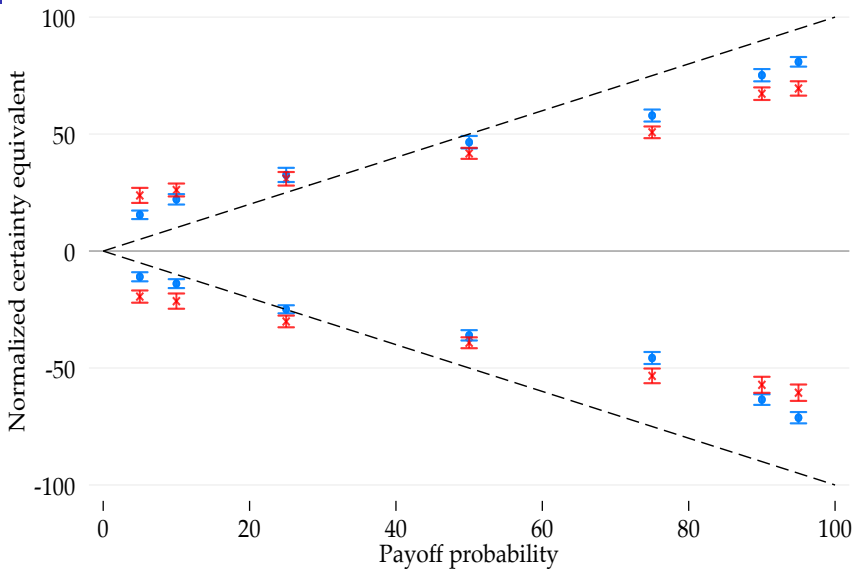
Risk Attitudes in the Laboratory

- K-T further document a more complex pattern of switches between risk-averse and risk-seeking choices

Risk Attitudes in the Laboratory

- K-T further document a more complex pattern of switches between risk-averse and risk-seeking choices
- Tversky and Kahneman (1992): elicit **certainty-equivalent** values for simple lotteries (p, X) , find a **“fourfold pattern”** of risk attitudes:
 - **risk averse** w.r.t. **gains** when p is **substantial**
 - **risk seeking** w.r.t. **gains** when p is **small**
 - **risk averse** w.r.t. **losses** when p is **small**
 - **risk seeking** w.r.t. **losses** when p is **substantial**

“Fourfold Pattern of Risk Attitudes”



[from Enke and Graeber (2023)]

Explaining the Fourfold Pattern of PT

- This pattern can also easily be explained, as an optimal adaptation to the noisy retrieval of **probability information**

Explaining the Fourfold Pattern of PT

- This pattern can also easily be explained, as an optimal adaptation to the noisy retrieval of **probability information**
- Model: experimenter describes lottery (p, X) [note: X may be positive or negative]
 - suppose [for now] that DM's cognitive process can retrieve value of X with **perfect precision**, but value of p only with **noise**:
 - noisy retrieved signal [internal representation of payoff probability]
$$r_p \sim f(r_p | p)$$

with conditional distribution independent of X

- DM's elicited **certainty-equivalent value** for the lottery must be some function $C(X, r_p)$

Explaining the Fourfold Pattern of PT

- Let us again suppose that the decision process is **optimized** to maximize **expected financial wealth** of DM

Explaining the Fourfold Pattern of PT

- Let us again suppose that the decision process is **optimized** to maximize **expected financial wealth** of DM
- Then under standard approaches to incentivizing the choice [**multiple price list**, with one selected at random to implement; or **BDM auction**], **optimal** certainty equivalent will be

$$C = E[pX | X, r_p]$$

where **conditional expectation** is computed using joint distribution for (X, p, r_p) implied by

- **prior distribution** over lotteries (p, X) for which decision process is optimized
- model of noisy coding of probability information

Explaining the Fourfold Pattern of PT

- Suppose further that distribution of p is independent of stake size X , under the prior [true of experiments of Enke and Graeber (2023), or K LW experiment below]; then prediction is simply

$$C = E[p | r_p] \cdot X$$

Explaining the Fourfold Pattern of PT

- Suppose further that distribution of p is independent of stake size X , under the prior [true of experiments of Enke and Graeber (2023), or K LW experiment below]; then prediction is simply

$$C = E[p | r_p] \cdot X$$

- Since r_p is **random** conditional on the true value of p , the model predicts trial-to-trial **variability** of responses

Explaining the Fourfold Pattern of PT

- Suppose further that distribution of p is independent of stake size X , under the prior [true of experiments of Enke and Graeber (2023), or K LW experiment below]; then prediction is simply

$$C = E[p | r_p] \cdot X$$

- Since r_p is **random** conditional on the true value of p , the model predicts trial-to-trial **variability** of responses
- The **median** certainty equivalent across trials, for a given lottery (p, X) , is predicted to be

$$C^{med} = w(p) \cdot X,$$

where

$$w(p) \equiv \text{med}[E[p | r_p] | p]$$

Explaining the Fourfold Pattern of PT

$$w(p) \equiv \text{med}[E[p | r_p] | p]$$

- This provides an interpretation for the “probability weighting function” of prospect theory

Explaining the Fourfold Pattern of PT

$$w(p) \equiv \text{med}[E[p | r_p] | p]$$

- This provides an interpretation for the “probability weighting function” of prospect theory
- The model predicts the **“fourfold pattern”** documented by TK, if there exists an interior probability $\bar{p}$ such that $w(p) > p$ for all $0 < p < \bar{p}$, while $w(p) < p$ for all $\bar{p} < p < 1$
 - and we can easily specify the encoding noise so that Bayesian decision rule has this property
 - intuition: noise in internal representation $\Rightarrow$ **regression to the prior mean** [**“Behavioral Attenuation”**]

A Semi-Analytical Example (Khaw *et al.*, 2025)

- Suppose that the relative odds of the two possible outcomes are encoded by

$$r_p \sim N(\log \frac{p}{1-p}, v^2)$$

— consistent with the finding of Frydman and Jin (2025) that people make **fewer errors in comparisons** of probabilities when either very small or very large; and finding of Enke and Graeber (2023) that **CU is lower** for lotteries with p nearer to 0 or 1

A Semi-Analytical Example (Khaw *et al.*, 2025)

- Suppose that the relative odds of the two possible outcomes are encoded by

$$r_p \sim N\left(\log \frac{p}{1-p}, \nu^2\right)$$

— consistent with the finding of Frydman and Jin (2025) that people make **fewer errors in comparisons** of probabilities when either very small or very large; and finding of Enke and Graeber (2023) that **CU is lower** for lotteries with p nearer to 0 or 1

- And suppose that the **prior** distribution from which true log odds are drawn is also **Gaussian**:

$$\log \frac{p}{1-p} \sim N(\mu, \sigma^2)$$

- Then **posterior** log odds conditional on r_p will also be Gaussian

A Semi-Analytical Example

- Posterior log odds:

$$\log \frac{p}{1-p} \Big| r_p \sim N(\hat{\mu}(r_p), \hat{\sigma}^2)$$

where

$$\begin{aligned} \hat{\mu}(r_p) &\equiv \gamma \cdot r_p + (1 - \gamma) \cdot \mu, & \gamma &\equiv \frac{\sigma^2}{\sigma^2 + \nu^2} < 1 \\ \hat{\sigma}^{-2} &\equiv \sigma^{-2} + \nu^{-2} \end{aligned}$$

A Semi-Analytical Example

- Posterior log odds:

$$\log \frac{p}{1-p} \Big| r_p \sim N(\hat{\mu}(r_p), \hat{\sigma}^2)$$

where

$$\begin{aligned} \hat{\mu}(r_p) &\equiv \gamma \cdot r_p + (1 - \gamma) \cdot \mu, & \gamma &\equiv \frac{\sigma^2}{\sigma^2 + \nu^2} < 1 \\ \hat{\sigma}^{-2} &\equiv \sigma^{-2} + \nu^{-2} \end{aligned}$$

- Using approximation of Daunizeau (2017) for the mean of a logit-normal random variable, this implies

$$\mathbb{E}[p | r_p] \approx \frac{e^{\alpha \hat{\mu}(r_p)}}{1 + e^{\alpha \hat{\mu}(r_p)}}$$

where $0 < \alpha < 1$ is a decreasing function of $\hat{\sigma}$

A Semi-Analytical Example

- Then $E[p | r_p]$ is an increasing function of $r_p \Rightarrow$ its **median** value is the value when r_p takes its median value (conditional on p) $\Rightarrow$ when r_p equals the **true log odds**

A Semi-Analytical Example

- Then $E[p | r_p]$ is an increasing function of $r_p \Rightarrow$ its **median** value is the value when r_p takes its median value (conditional on p) $\Rightarrow$ when r_p equals the **true log odds**
- Hence (in this approximation) $w(p)$ is given by

$$\log \frac{w(p)}{1 - w(p)} \approx \alpha\gamma \log \frac{p}{1 - p} + (1 - \alpha\gamma) \log \frac{\bar{p}}{1 - \bar{p}}$$

where

$$\log \frac{\bar{p}}{1 - \bar{p}} \equiv \frac{\alpha(1 - \gamma)}{1 - \alpha\gamma} \cdot \mu$$

A Semi-Analytical Example

$$\log \frac{w(p)}{1-w(p)} \approx \alpha\gamma \log \frac{p}{1-p} + (1-\alpha\gamma) \log \frac{\bar{p}}{1-\bar{p}}$$

- This has the “inverse-S shape” required to explain the fourfold pattern of PT, with “crossover point” $\bar{p}$ defined above

A Semi-Analytical Example

$$\log \frac{w(p)}{1 - w(p)} \approx \alpha\gamma \log \frac{p}{1 - p} + (1 - \alpha\gamma) \log \frac{\bar{p}}{1 - \bar{p}}$$

- This has the “inverse-S shape” required to explain the fourfold pattern of PT, with “crossover point” $\bar{p}$ defined above
- But note that the exact shape should depend on **degree of cognitive noise** in a given setting, and the **prior** associated with it

A Semi-Analytical Example

$$\log \frac{w(p)}{1 - w(p)} \approx \alpha\gamma \log \frac{p}{1 - p} + (1 - \alpha\gamma) \log \frac{\bar{p}}{1 - \bar{p}}$$

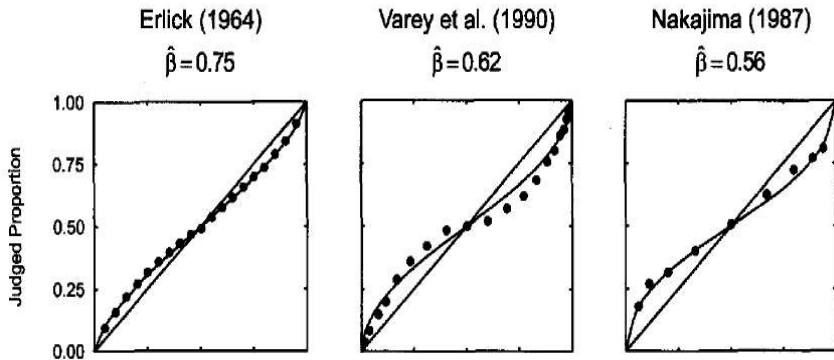
- The predicted functional form — **“linear in log odds”** — is found by Zhang and Maloney (2012) to fit experimental data on **estimation of probabilities, relative frequencies, or relative proportions** in a variety of contexts

A Semi-Analytical Example

$$\log \frac{w(p)}{1 - w(p)} \approx \alpha\gamma \log \frac{p}{1 - p} + (1 - \alpha\gamma) \log \frac{\bar{p}}{1 - \bar{p}}$$

- The predicted functional form — **“linear in log odds”** — is found by Zhang and Maloney (2012) to fit experimental data on **estimation of probabilities, relative frequencies, or relative proportions** in a variety of contexts
 - in experiments where the proportions $(p, 1 - p)$ occur exactly as often as $(1 - p, p)$ (so that the prior should be **symmetric** around $p = 1/2$), the crossover point $\bar{p}$ is found to be **near 1/2**, as above model would predict [recall the figure from Hollands and Dyre (2000), in Lecture 2]

Bias in Judged Proportions



[figure from Hollands and Dyre (2000)]

horizontal axis = true proportion

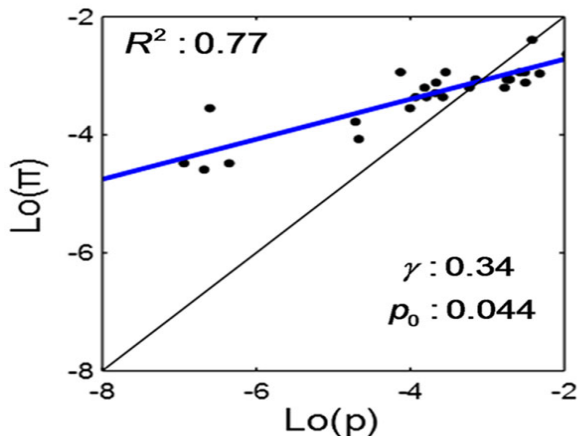
note crossover point near 50% in each case

A Semi-Analytical Example

$$\log \frac{w(p)}{1 - w(p)} \approx \alpha\gamma \log \frac{p}{1 - p} + (1 - \alpha\gamma) \log \frac{\bar{p}}{1 - \bar{p}}$$

- The predicted functional form — **“linear in log odds”** — is found by Zhang and Maloney (2012) to fit experimental data on **estimation of probabilities, relative frequencies, or relative proportions** in a variety of contexts
 - when instead values $p < 1/2$ occur much more often than values $p > 1/2$, the crossover point is found to be much lower [e.g., next slide]
- this may account for the crossover point $\bar{p} < 1/2$ commonly found in empirical estimates of the PT “probability weighting function”

Bias in Judged Frequency of Letter Occurrence



[figure from Zhang and Maloney (2012)]
[plotting data from Attneave (1953)]

actual (p) and estimated (π) log odds for each letter in English text

Remaining Questions

- 1 Is it noise in the retrieval of **numerical information** supplied by experimenter (e.g., value of p) that creates imprecision, as in above exposition, or noise in retrieval of the **outcome of EV calculation?**

— two models equivalent as explanations of Enke-Graeber data, because value of $|X|$ is the **same across all trials**; only p varies

Remaining Questions

- 1 Is it noise in the retrieval of **numerical information** supplied by experimenter (e.g., value of p) that creates imprecision, as in above exposition, or noise in retrieval of the **outcome of EV calculation**?

— two models equivalent as explanations of Enke-Graeber data, because value of $|X|$ is the **same across all trials**; only p varies

- 2 Cognitive noise model can account for the fourfold pattern emphasized by PT; but can it also account for **stake-size effects**?

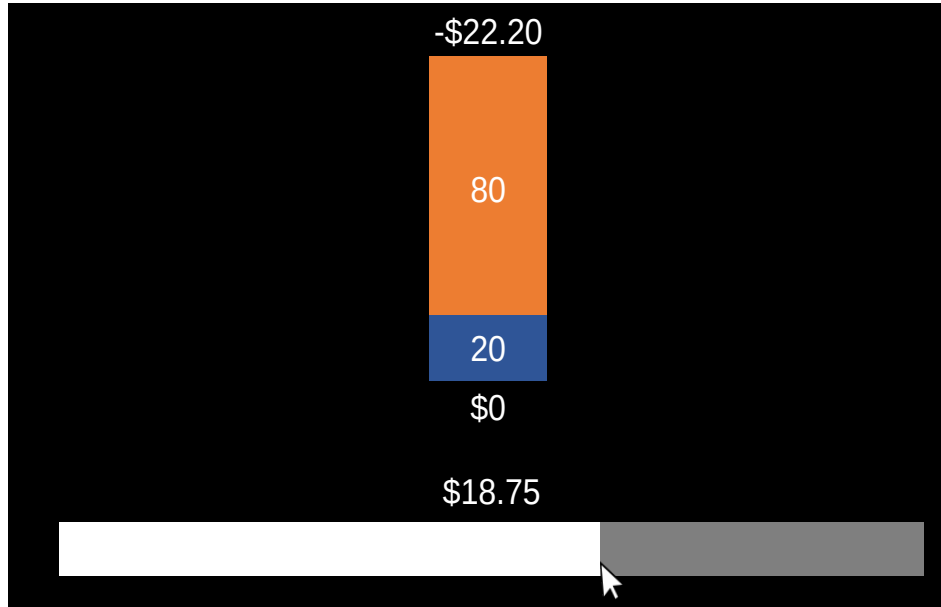
Remaining Questions

- To address these questions, Khaw *et al.* (2025) examine the fit of a model in which risk attitudes result from cognitive imprecision to a dataset in which
 - there is independent variation in lotteries along **multiple** dimensions:
 - gains vs. losses
 - probability of non-zero payoff
 - magnitude of non-zero payoff

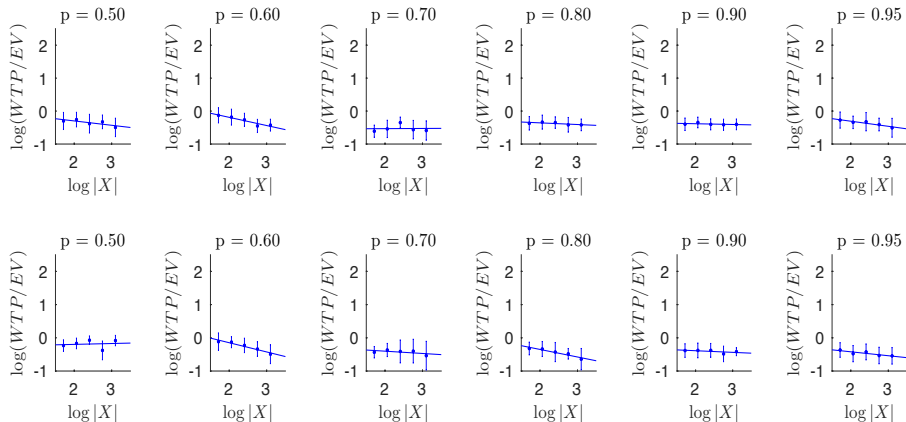
Remaining Questions

- To address these questions, Khaw *et al.* (2025) examine the fit of a model in which risk attitudes result from cognitive imprecision to a dataset in which
 - there is independent variation in lotteries along **multiple** dimensions:
 - gains vs. losses
 - probability of non-zero payoff
 - magnitude of non-zero payoff
 - they collect data not just on a subject's **typical** valuation of some lottery, but on the amount of **trial-to-trial variability** in their judgments

Experimental Interface (Khaw *et al.*, 2025)

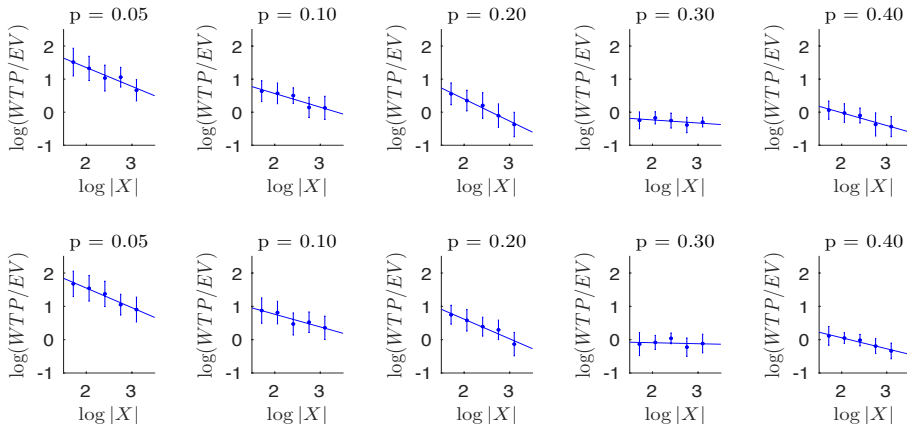


Risk Premium Depends on Both p and X



top line = risky gains; bottom line = risky losses

Risk Premium Depends on Both p and X



top line = risky gains; bottom line = risky losses

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- ② $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- ② $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)
 - also true (for lotteries involving gains) in data of Gonzalez and Wu (2022), which have 15 values of X for each value of p

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- ② $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)
 - also true (for lotteries involving gains) in data of Gonzalez and Wu (2022), which have 15 values of X for each value of p
 - implies change in **sign** of RRA if $|X|$ varies over wide enough range
 - the sign change occurs in KLTW data when $p = 0.20$; in GW data when $p = 0.10$ or 0.25 ; and at lower values of p when wider range of stakes (e.g., Hershey and Schoemaker, 1980)

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- ② $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)
- ③ This function has both a **higher intercept** and **more negative slope**, the smaller is p

Features of Data to Explain

- ① Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- ② $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)
- ③ This function has both a **higher intercept** and **more negative slope**, the smaller is p
 - the way intercept shifts with p confirms the “fourfold pattern” of Tversky and Kahneman (1992)
 - but sign of relative risk premium doesn't depend **only** on $\text{sign}(X)$ and p — stake size also matters

Features of Data to Explain

- 1 Distribution of values of $|WTP|$ conditional on $|X|$ is **similar in gain and loss domains**
- 2 $E[\log WTP / EV]$ roughly an **affine function** of $\log |X|$, with **negative slope** (between 0 and -1)
- 3 This function has both a **higher intercept** and **more negative slope**, the smaller is p
- 4 Variability of $\log |WTP|$ **greater** for smaller values of p

Model with Multiple Kinds of Cognitive Noise

Noisy internal representations:

- information (p, X) defining the lottery encoded by internal states (r_p, r_x) , where

$$r_p \sim N(\log \frac{p}{1-p}, v_z^2)$$

as above

Model with Multiple Kinds of Cognitive Noise

Noisy internal representations:

- the **payoff magnitude** $|X|$ is then encoded by

$$r_x \sim N(\log |X|, v_x^2(r_p))$$

conditional on the draw of r_p

Model with Multiple Kinds of Cognitive Noise

Noisy internal representations:

- the **payoff magnitude** $|X|$ is then encoded by

$$r_x \sim N(\log |X|, v_x^2(r_p))$$

conditional on the draw of r_p

- mean proportional to $\log |X| \Rightarrow$ uniform discriminability of nearby magnitudes in **percentage** terms [“Weber’s Law”]
— consistent with evidence on precision with which monetary amounts (prices) are recalled after time delay (Dehaene and Marques, 2002)
- sign** of X treated as coded without error

Model with Multiple Kinds of Cognitive Noise

Noisy internal representations:

- the **payoff magnitude** $|X|$ is then encoded by

$$r_x \sim N(\log |X|, v_x^2(r_p))$$

conditional on the draw of r_p

- precision** of coding of payoff magnitude allowed to depend on [subjective perception of] likelihood of relevance

— as in Van den Berg and Ma (2018): precision of encoding in working memory depends on probability that a given location will be probed

Model with Multiple Kinds of Cognitive Noise

- Announced WTP: assume

$$\log WTP \sim N(f(r_x, r_p), v_c^2)$$

for some function $f(r_x, r_p)$ [optimized].

Model with Multiple Kinds of Cognitive Noise

- Announced WTP: assume

$$\log WTP \sim N(f(r_x, r_p), v_c^2)$$

for some function $f(r_x, r_p)$ [optimized].

- **Endogenous precision** of coding: both $f(r_x, r_p)$ and $v^2(r_p)$ are chosen to minimize

$$E[L(r_x, r_p) + A(\sigma_x^2/v_x^2(r_p))]$$

where

$$L(r_x, r_p) \equiv E[(WTP - EV)^2 | r_x, r_p]$$

and second term [$A > 0$ a free parameter] represents cost of more precise coding

— if cost is proportional to **number of independent samples** used to code the value of $|X|$, should increase $\sim 1/v_x^2$

Model with Multiple Kinds of Cognitive Noise

- Announced WTP: assume

$$\log WTP \sim N(f(r_x, r_p), v_c^2)$$

for some function $f(r_x, r_p)$ [optimized]

- **Endogenous precision** of coding: both $f(r_x, r_p)$ and $v^2(r_p)$ are chosen to minimize

$$E[L(r_x, r_p) + A(\sigma_x^2 / v_x^2(r_p))]$$

where

$$L(r_x, r_p) \equiv E[(WTP - EV)^2 | r_x, r_p]$$

and second term [$A > 0$ a free parameter] represents cost of more precise coding

— this cost function used by **van den Berg and Ma (2018)** to fit endogenous precision of visual working memory

Model with Multiple Kinds of Cognitive Noise

- Expected loss evaluated under **priors**:

$$\log |X| \sim N(\mu_x, \sigma_x^2) \quad [\text{same for gains and losses}]$$

$$\log(p/1-p) \sim U(\mu_z - \sqrt{3}\sigma_z, \mu_z + \sqrt{3}\sigma_z)$$

- and parameters of prior distributions for $|X|$ and p are chosen to max likelihood of the **values used in experiment** [so $\mu_x, \sigma_x, \mu_z, \sigma_z$ are not additional free parameters]

Model with Multiple Kinds of Cognitive Noise

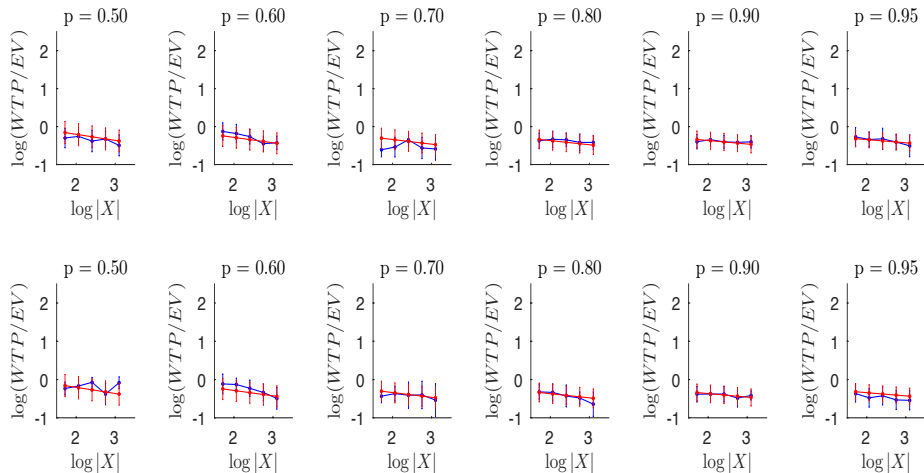
- Expected loss evaluated under **priors**:

$$\log |X| \sim N(\mu_x, \sigma_x^2) \quad [\text{same for gains and losses}]$$

$$\log(p/1-p) \sim U(\mu_z - \sqrt{3}\sigma_z, \mu_z + \sqrt{3}\sigma_z)$$

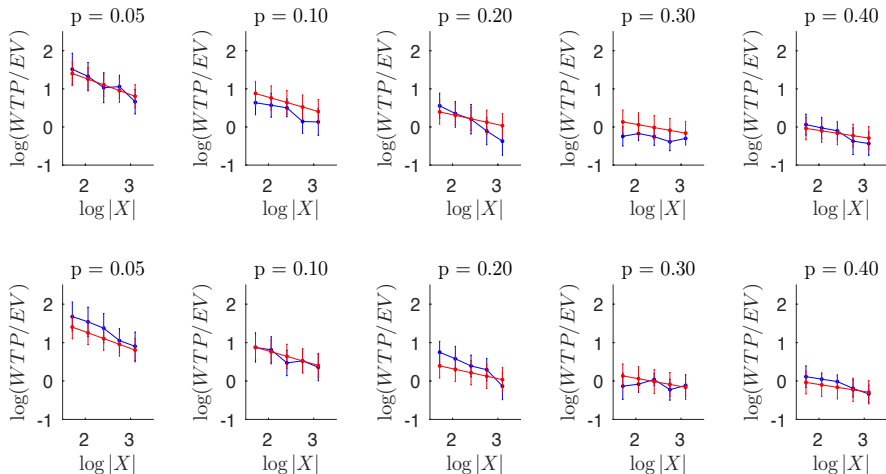
- and parameters of prior distributions for $|X|$ and p are chosen to max likelihood of the **values used in experiment** [so $\mu_x, \sigma_x, \mu_z, \sigma_z$ are not additional free parameters]
- Model thus has **3 free parameters**: ν_z^2, ν_c^2, A , in addition to the parameters of the distribution of values used in the experiment — to explain 220 data moments

Fitted Distributions of WTP



top line = risky gains; bottom line = risky losses
blue = data; red = model predictions

Fitted Distributions of WTP



top line = risky gains; bottom line = risky losses
blue = data; red = model predictions

Which Kind of Noise is Needed?

- Variant models based on cognitive noise:

	LL	BIC
baseline model	-1602.5	3236.3
exogenous noise	-1604.7	3240.7
no payoff noise	-1608.7	3242.4
no probability noise	-1990.0	4005.0
no response noise	-1646.1	3317.2
noisy coding of EV	-1966.4	3957.9

Which Kind of Noise is Needed?

- Variant models based on cognitive noise:

	LL	BIC
baseline model	-1602.5	3236.3
exogenous noise	-1604.7	3240.7
no payoff noise	-1608.7	3242.4
no probability noise	-1990.0	4005.0
no response noise	-1646.1	3317.2
noisy coding of EV	-1966.4	3957.9

- Overall **Bayes factor** in favor of the baseline model, relative to **noisy coding of EV**: greater than 10^{156}

— relative to model with **precise reading of probability**:
greater than 10^{166}

Summary

- A variety of aspects of measured risk attitudes can be explained by a unified model, according to which DM's decision rule is optimally adapted to the presence of **cognitive noise**
- suggesting that risk attitudes (for small gambles) are better viewed as consequences of **imprecise mental calculation**, rather than any actual attitudes toward **risk**

Summary

- A variety of aspects of measured risk attitudes can be explained by a unified model, according to which DM's decision rule is optimally adapted to the presence of **cognitive noise**
 - suggesting that risk attitudes (for small gambles) are better viewed as consequences of **imprecise mental calculation**, rather than any actual attitudes toward **risk**
 - consistent with Oprea (2024) finding of **similar biases** in judgments about the dollar value of obtaining a **fraction** of a monetary amount with **certainty**

Summary

- As in perceptual domains, important to model **variability** of responses and average **bias** in responses jointly

Summary

- As in perceptual domains, important to model **variability** of responses and average **bias** in responses jointly
- A model with **independent noise** in the retrieved values of both **probabilities** and **payoffs** fits better than one with only noise in the retrieved **EV of lottery**
 - and fit is improved by **endogenizing the precision** with which payoffs are encoded/retrieved

Summary

- As in perceptual domains, important to model **variability** of responses and average **bias** in responses jointly
- A model with **independent noise** in the retrieved values of both **probabilities** and **payoffs** fits better than one with only noise in the retrieved **EV of lottery**
 - and fit is improved by **endogenizing the precision** with which payoffs are encoded/retrieved
- Structure of cognitive noise can be disciplined through analogies with what is known about imprecise internal representation of **numbers** and **probabilities** in other contexts — the main distortions are not unique to choice under risk

References

Attneave, F., “Psychological Probability as a Function of Experienced Frequency,” *J. Experimental Psychology* 46: 81-86 (1953).

Daunizeau, J., “Semi-Analytical Approximations to Statistical Moments of Sigmoid and Softmax Mappings of Normal Variables,” posted on *arXiv*, March 3, 2017.

Dehaene, S., and J.F. Marques, “Cognitive Euroscience: Scalar Variability in Price Estimation and the Cognitive Consequences of Switching to the Euro,” *Quarterly J. Experimental Psychology* 55: 705-731 (2002).

Enke, B., and T. Graeber, “Cognitive Uncertainty,” *Quarterly J. Economics* 138: 2021-2067 (2023).

References (p.2)

Erlick, D.E., “Absolute Judgments of Discrete Quantities Randomly Distributed Over Time,” *Journal of Experimental Psychology* 67: 475-482 (1964).

Frydman, C., and L. Jin, “On the Source and Instability of Probability Weighting,” NBER Working Paper no. 31573, April 2025.

Gonzalez, R., and G. Wu, “Composition Rules in Original and Cumulative Prospect Theory,” *Theory and Decision* 92: 647-675 (2022).

Hershey, J.C., and P.J. Schoemaker, “Prospect Theory’s Reflection Hypothesis: A Critical Examination,” *Organizational Behavior and Human Performance* 25: 395-418 (1980).

References (p.3)

Hollands, J.G., and B.P. Dyre, "Bias in Proportion Judgments: The Cyclical Power Model," *Psychological Review* 107: 500-524 (2000).

Kahneman, D., and A. Tversky, "Prospect Theory: An Analysis of Decision Under Risk," *Econometrica* 47: 263-291 (1979).

Khaw, M.W., Z. Li, and M. Woodford, "Cognitive Imprecision and Stake-Dependent Risk Attitudes," NBER Working Paper no. 30417, revised August 2025.

Nakajima, Y., "A Model of Empty Duration Perception," *Perception* 16: 485-520 (1987).

References (p. 4)

Oprea, R., "Decisions Under Risk Are Decisions Under Complexity," *American Econ. Review* 114: 3789-3811 (2024).

Tversky, A., and D. Kahneman, "Advances in Prospect Theory: Cumulative Representation of Uncertainty," *J. Risk and Uncertainty* 5: 297-323 (1992).

Van den Berg, R., and W.J. Ma, "A Resource-Rational Theory of Set-Size Effects in Human Visual Working Memory," *eLife* 7:e34963 (2018).

Varey, C.A., B.A. Mellers, and M.H. Birnbaum, "Judgments of Proportion," *J. Experimental Psych.: Human Perception and Performance* 16: 613-625 (1990).

References (p. 5)

Zhang, H., and L.T. Maloney, "Ubiquitous Log Odds: A Common Representation of Probability and Frequency Distortion in Perception, Action and Cognition," *Frontiers in Neuroscience* 6, art. 1 (2012).