



Behavioural
— COMPUTATIONAL SCIENCE LAB —



AI-Augmented Choice Modeling: From Deep Learning to LLM-Generated Priors

Prateek Bansal

Assistant Professor
National University of Singapore

Summer School, University of Tokyo · 7 September 2026

Agenda

1. Integrating Large Language Models with Human Data to Design Personalized Decoy Nudges
2. A Flexible and Interpretable Reference-Dependent Choice Model with Endogenous Estimation of Reference Points
3. A Theory-Constrained Framework for Data-Driven Learning of Multiple Discrete-Continuous Choices

Integrating Large Language Models with Human Data to Design Personalized Decoy Nudges

Vladimir Maksimenko¹, Qingyao Xin¹, Prateek Gupta², Bin Zhang³, Leonard Lee¹, Prateek Bansal¹

¹National University of Singapore ²Max Planck Institute's Center for Humans and Machines ³Beijing Institute of Technology

Summer School, University of Tokyo · 7 September 2026

Research Context

LLMs can act as simulated human participants, helping reduce the cost of human data collection in behavioural experiments

- Dillion, D., N. Tandon, Y. Gu, and K. Gray (2023), “Can AI Language Models Replace Human Participants?” *Trends in Cognitive Sciences*, 27 (7), 597–600.
- Cui, Z., N. Li, and H. Zhou (2025), “A Large-Scale Replication of Scenario-Based Experiments in Psychology and Management Using Large Language Models,” *Nature Computational Science*, 1–8.

This is especially important for personalization, where data need to be collected from multiple consumer segments

- Sunstein, C. R. (2022), “The Distributional Effects of Nudges,” *Nature Human Behaviour*, 6 (1), 9–10.
- Szaszi, B., D. G. Goldstein, D. Soman, and S. Michie (2025), “Generalizability of Choice Architecture Interventions,” *Nature Reviews Psychology*, 1–12.

While fully replacing human participants with LLMs is not currently feasible, methods that integrate LLM-generated data with small amounts of human data are promising

- Kim, J., M. Kovach, K. M. Lee, E. Shin, and H. Tzavellas (2024), “Learning to Be Homo Economicus: Can an LLM Learn Preferences from Choice?” *arXiv preprint arXiv:2401.07345*.
- Brand, J., A. Israeli, and D. Ngwe (2023), “Using LLMs for Market Research,” Harvard Business School Working Paper No. 23-062, SSRN 4395751.

Research Context

- We integrate LLMs with human data to design **decoy-based nudges**
- The decoy effect occurs when adding a third option to a choice set shifts preferences toward one of the existing options, even though the original options themselves remain unchanged
- Parameters of decoy (here price and offset) **should be personalized** across different consumer segments to maximize the decoy effect at scale.

Ticket 1		Ticket 2			
Price		Price			
\$ 500		\$ 550			
Offset CO ₂		Offset CO ₂			
0 %		100 %			



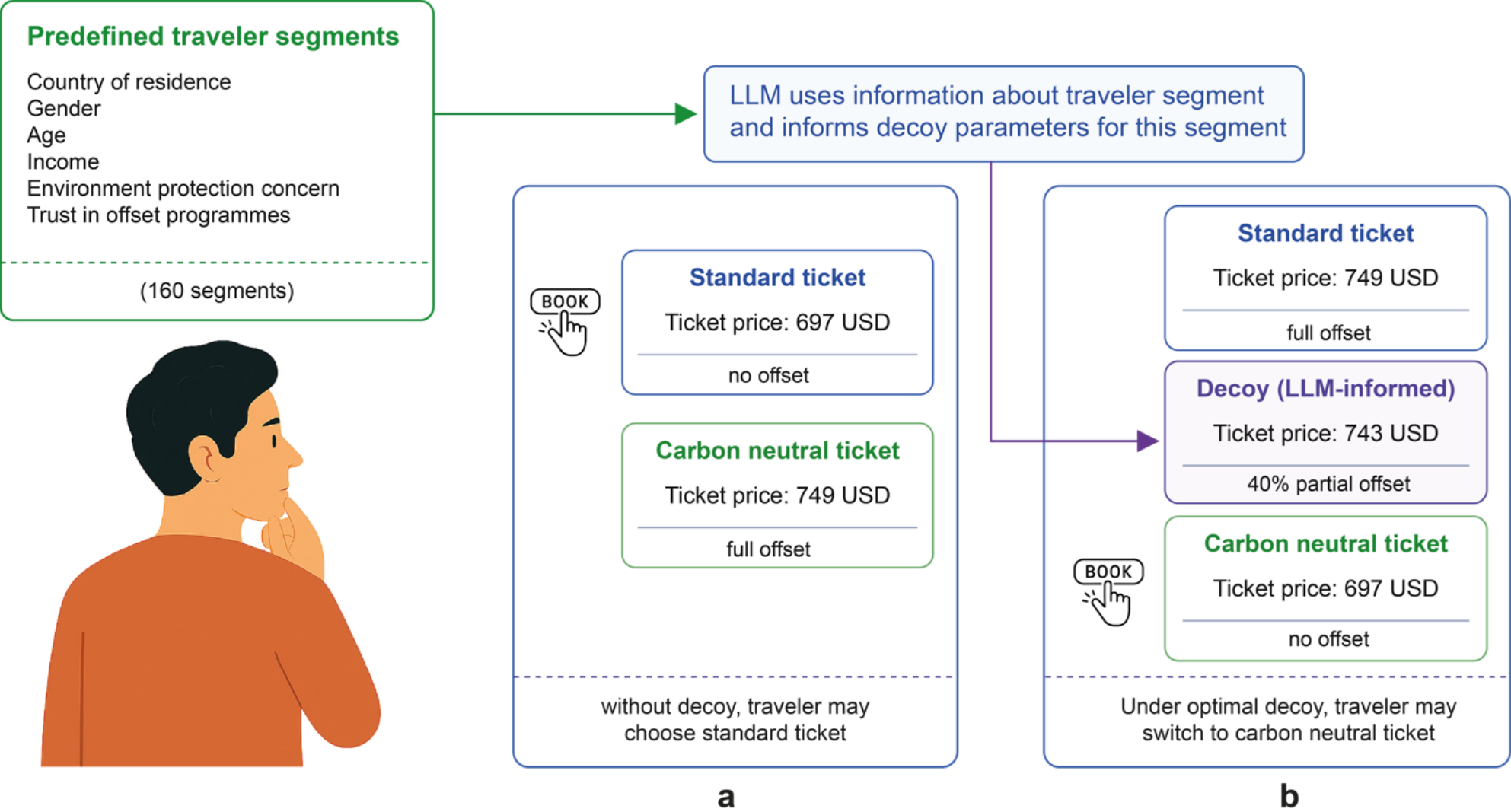
Ticket 1		Ticket 2		Ticket 3	
Price		Price		Price	
\$ 500		\$ 550		\$580	
Offset CO ₂		Offset CO ₂		Offset CO ₂	
0 %		100 %		90 %	

decoy

Research Questions

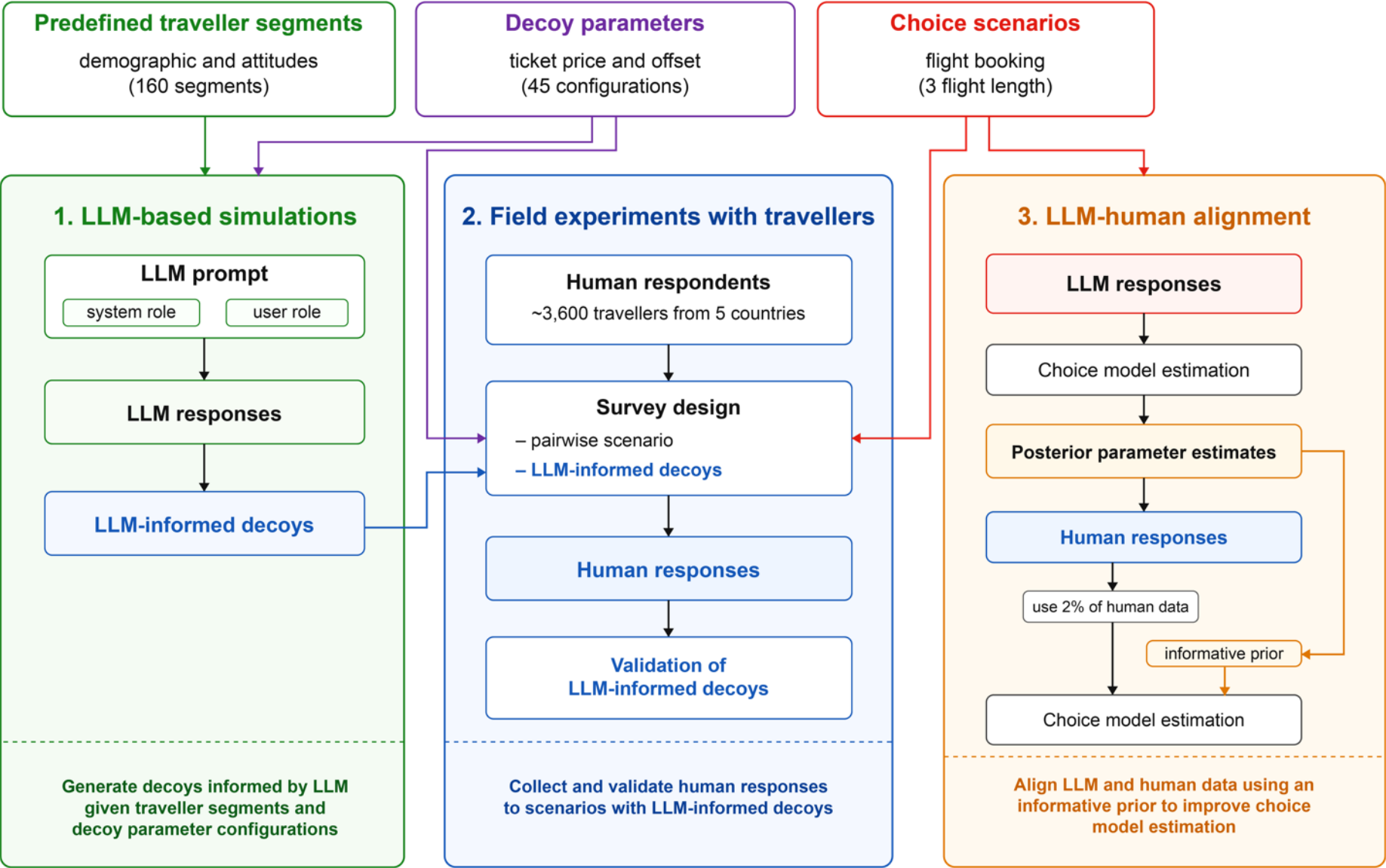
- ? Can LLMs exhibit the decoy effect?
- ? Can LLMs be used to identify segment-specific optimal decoys?
- ? Can LLM-informed priors improve the estimation of choice models when human training data are scarce and certain consumer segments are underrepresented?

Study Design



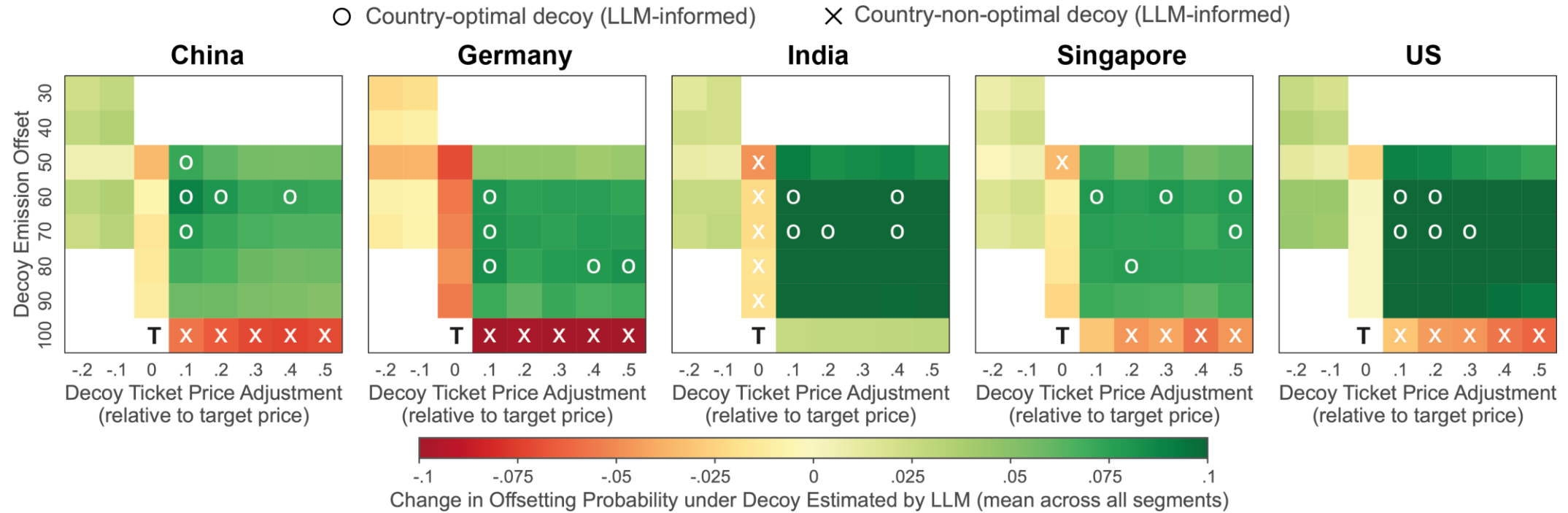
Participant is presented with choice situation without decoy (a) and with decoy (b) where decoy parameters are set by the LLM based on participant's demography

Study Design



Study design includes three subsequent steps

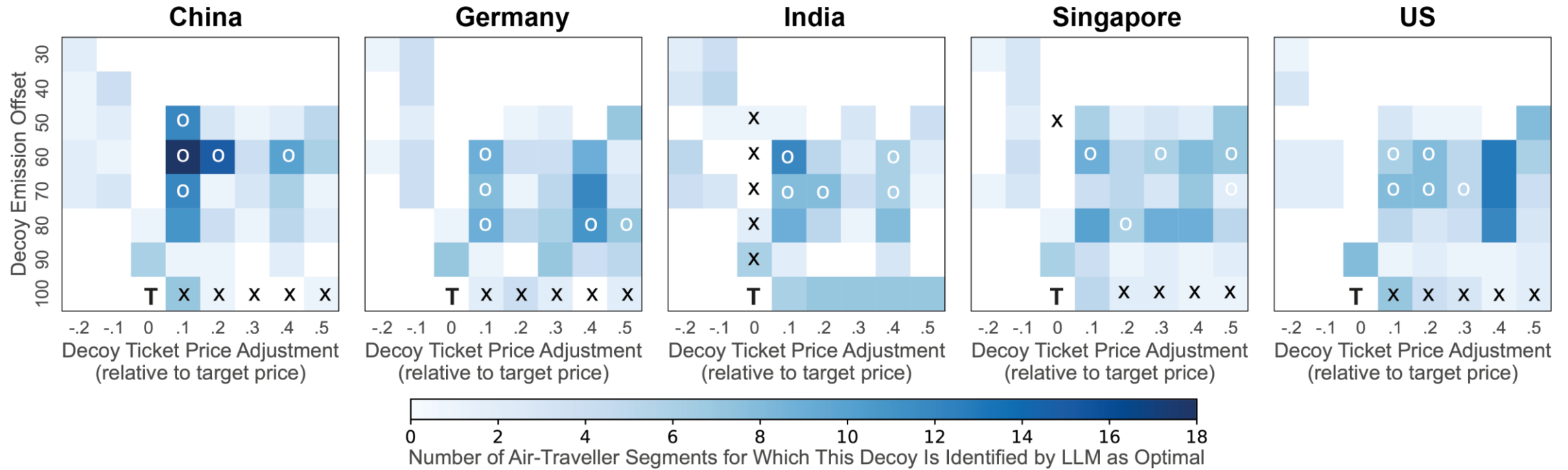
Results of LLM-based Simulations (country-level effect)



- **LLM changes offsetting probability under decoy** (direction of change is reflected by color)
- Magnitude of the decoy **effect depends on the decoy parameters** (different cells have different color)
- We identify LLM-informed **country-optimal decoys** (top 5 with highest magnitude of decoy effect, 'o')
- We identify LLM-informed **country-non-optimal decoys** (top 5 with lowest magnitude of decoy effect, 'x')

Results of LLM-based Simulations (segment-level effect)

○ Country-optimal decoy (LLM-informed) × Country-non-optimal decoy (LLM-informed)



- We identify **individual-optimal decoy** that maximizes magnitude of decoy effect for particular segment
- Some decoys that are not selected as optimal at country level may be optimal for some segments
- Even country-non-optimal decoys may be optimal for some segments

Results: Human validation in Preregistered Experiment across Five Countries

Statistical Model:

$$V_{in} = b_0 + b_1 * Age_i + b_2 * Income_i + b_3 * Gender_i + b_4 * Concern_i + b_5 * Trust_i + b_6 * \Delta Price_{in} + b_7 * \text{CountryOptimal}_{in} + b_8 * \text{CountryNonOptimal}_{in} + b_9 * \text{IndividualOptimal}_{in}$$
$$y_{in} = \begin{cases} 1, & \text{if target is chosen;} \\ 0, & \text{if competitor is chosen.} \end{cases}$$

Hypotheses:

- H1: **country-optimal** decoy generates a stronger decoy effect than the **country non-optimal**, i.e., $b_7 > b_8$
- H2: **individual-optimal** decoy generates a stronger decoy effect than **country-optimal**, i.e., $b_9 > b_7$.

Hypotheses testing:

	Singapore	Germany	US	India	China
Country-optimal (b_7) vs. Country non-optimal (b_8)	z = .633 p = .527	z = 1.913 p = .056	z = 2.509 ✓ p = .012*	z = -1.092 p = .275	z = .898 p = .369
Individual-optimal (b_9) vs Country-optimal (b_7)	z = 1.274 p = .203	z = .322 p = .748	z = -2.581 p = .010*	z = -4.031 p < .001***	z = -2.129 p = .033*
Number of respondents	694	638	736	714	713
Number of choice situations	2776	2552	2944	2856	2852

Results: Human-LLM Alignment. Modeling Framework

Baseline Model:

Attraction Model by Rooderkerk et al (2011) that accounts for decoy effect by adding attraction term to utility specification

$$V_{inj} = \beta_{price} \text{Price}_{inj} + \beta_{offset} \text{Offset}_{inj} + (\beta_{attraction}) \text{ATT}_{inj}$$

Our Modification:

We we allow the attraction coefficient to vary across participants as a function of demographic characteristics and attitudes

$$V_{inj} = \beta_{price} \text{Price}_{inj} + \beta_{offset} \text{Offset}_{inj} + (\beta_{attraction} + \mathbf{D}_i \boldsymbol{\gamma}_{att}) \text{ATT}_{inj}$$

$$\mathbf{D}_i = (\text{Woman}_i, \text{Old}_i, \text{HighIncome}_i, \text{TrustYes}_i, \text{ConcernYes}_i)$$

$$\boldsymbol{\gamma}_{att} = (\gamma_{gender}, \gamma_{age}, \gamma_{income}, \gamma_{concern}, \gamma_{trust})$$

Results: Human-LLM Alignment. Bayesian Model

Model is estimated using Bayesian approach under **two different priors**.

Uninformative Priors:

$$\beta_m \sim \mathcal{N}(0, 100^2), \quad m \in \{offset, price, attraction\}$$

$$\gamma_l \sim \mathcal{N}(0, 100^2), \quad l \in \{gender, age, income, concern, trust\}$$

LLM-informed Priors:

$$\beta_m \sim \mathcal{N}\left(\mu_m^{LLM}, (\xi s_m^{LLM})^2\right), \quad m \in \{offset, price, attraction\}$$

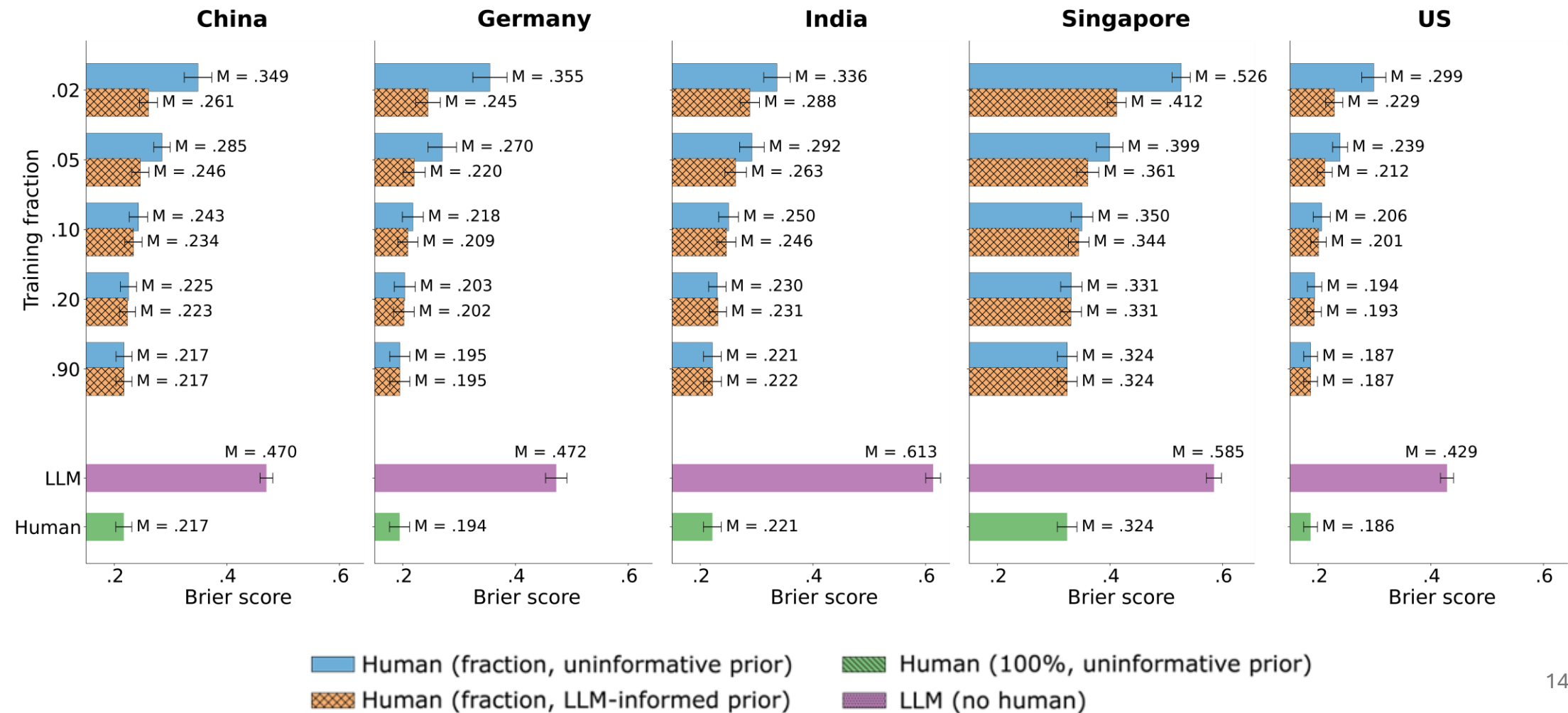
$$\gamma_l \sim \mathcal{N}\left(\mu_l^{LLM}, (\xi s_l^{LLM})^2\right), \quad l \in \{gender, age, income, concern, trust\},$$

μ_m^{LLM} , μ_l^{LLM} , s_m^{LLM} , s_l^{LLM} are the estimates that we obtain by estimating this model on the LLM-based responses

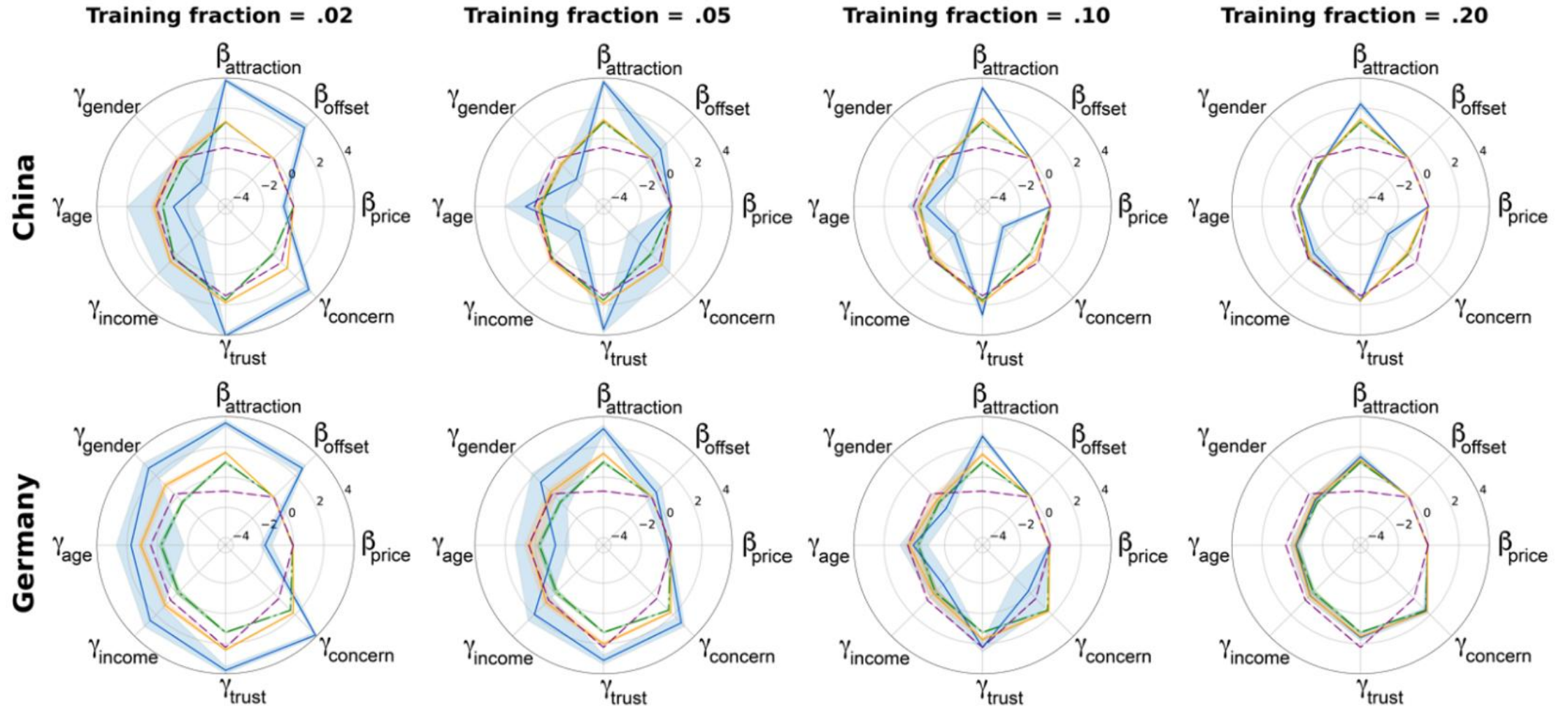
Results: Human-LLM Alignment. Predictive Performance

$$\text{Brier score} = \frac{1}{N} \sum_{n=1}^N \sum_{j=1}^J (P_{nj} - Y_{nj})^2$$

for each choice set n and alternative j , P_{nj} denote the predicted probability of choosing alternative j , and let Y_{nj} denote the observed choice indicator, where $Y_{nj} = 1$ if alternative j was chosen and $Y_{nj} = 0$ otherwise



Results: Human-LLM Alignment. Coefficient Estimates



Conclusions

- LLMs exhibit the decoy effect, but only for specific decoy configurations.
- LLM-informed personalized decoys can enhance nudging effectiveness for some consumer segments but fail for others, resulting in no overall effect at scale.
- LLM-informed priors improve model performance and LLM–human alignment when human data are limited.

A Flexible and Interpretable Reference-Dependent Choice Model with Endogenous Estimation of Reference Points

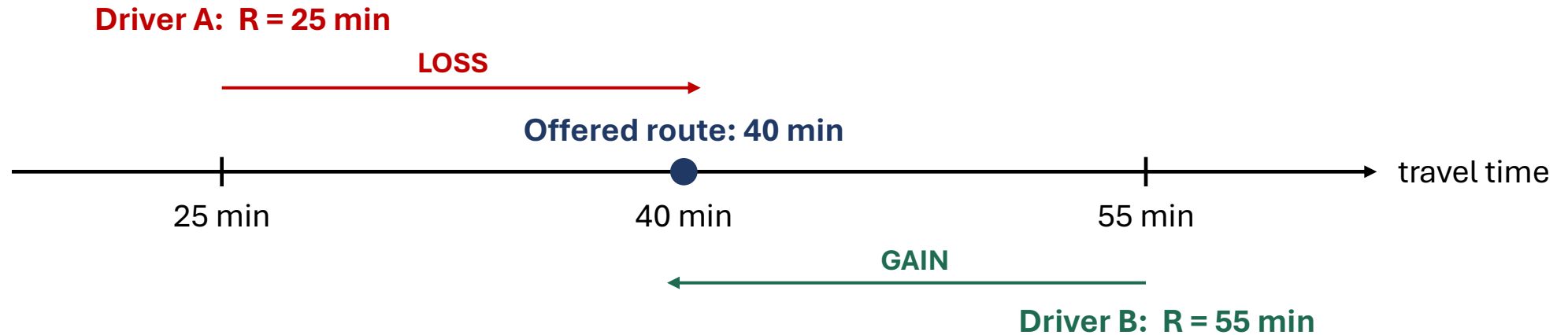
Jinghai Huo¹, Eui-Jin Kim², Irfan Batur¹, Kenan Zhang³, Ram M. Pendyala¹, Prateek Bansal⁴

¹Arizona State University ²Ajou University ³EPFL ⁴National University of Singapore

Summer School, University of Tokyo · 7 September 2026

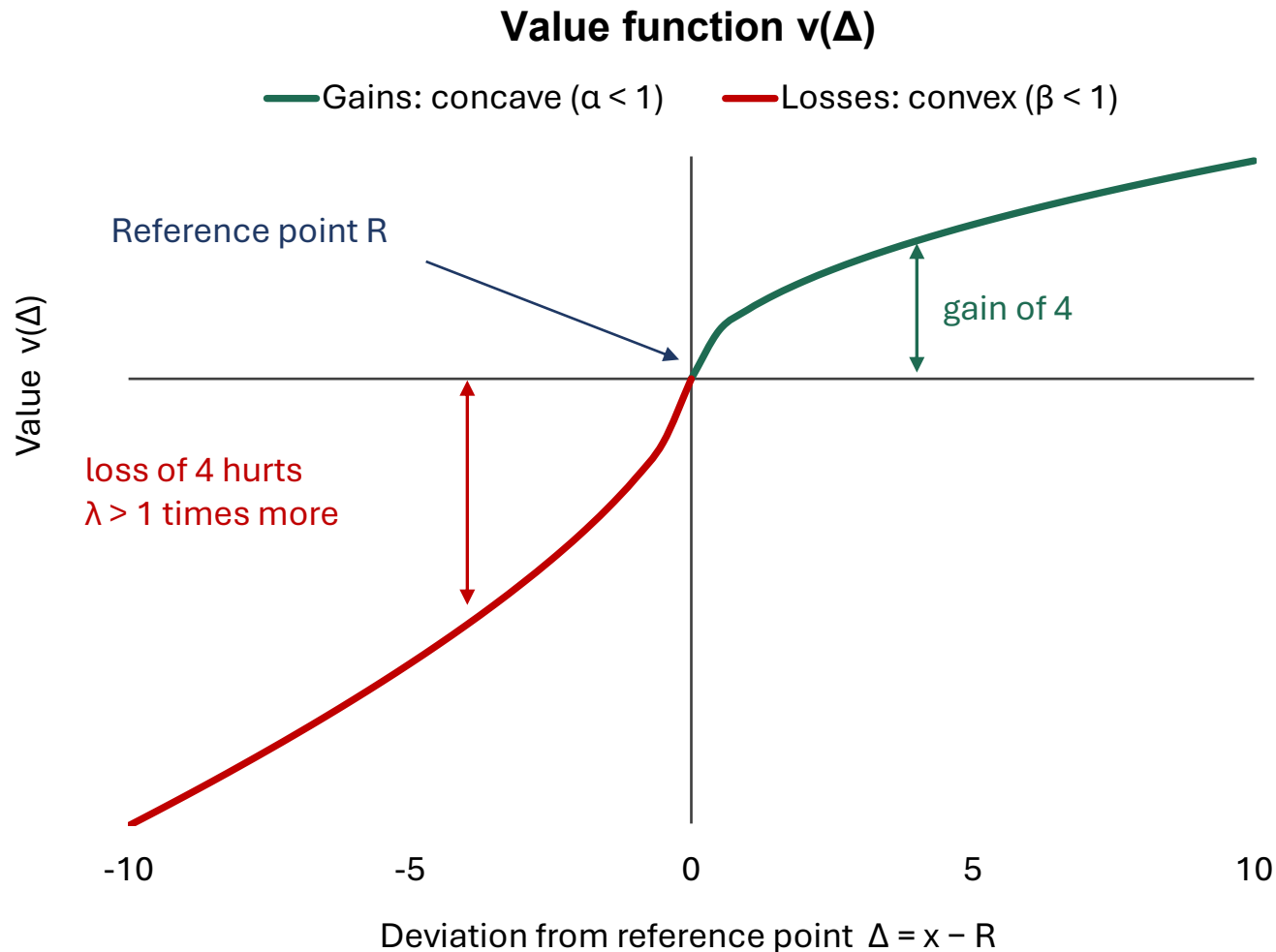
Motivation: people judge outcomes relative to a reference point

- Standard utility theory evaluates attributes in **absolute** terms — 40 minutes is 40 minutes for everyone
- Prospect theory: the same outcome is a **gain** or a **loss** depending on the individual's reference point



- Consequence: ignoring reference points misstates elasticities, willingness-to-pay and policy response
- **But the reference point is rarely observed — and it differs across people**

Prospect theory (1/3): the value function

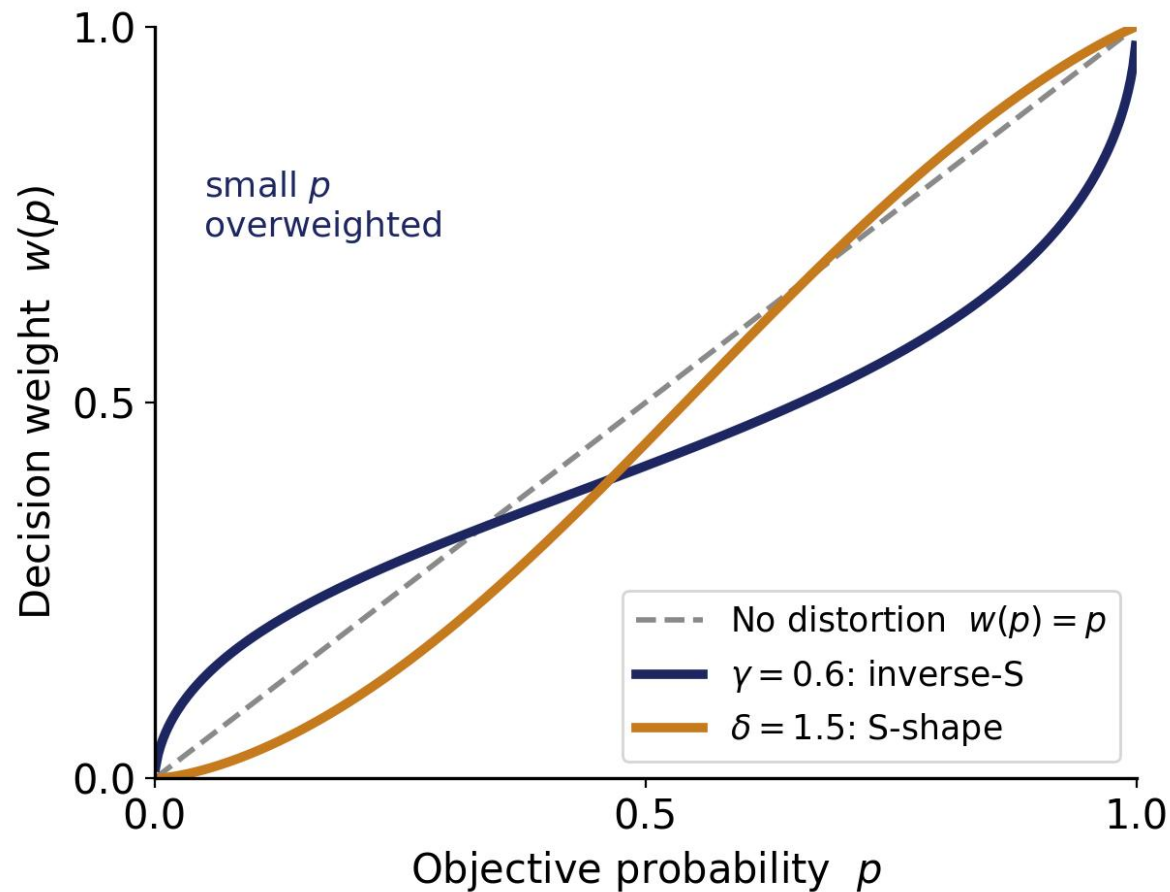


$$v(\Delta) = \Delta^\alpha \quad \text{if } \Delta > 0 \text{ (gain)}$$
$$v(\Delta) = -\lambda(-\Delta)^\beta \quad \text{if } \Delta \leq 0 \text{ (loss)}$$

Notation

- $\Delta = x - R$: deviation of attribute x from reference point R
- $\alpha \in (0,1)$: curvature for gains $\rightarrow$ concave, diminishing sensitivity
- $\beta \in (0,1)$: curvature for losses $\rightarrow$ convex
- $\lambda > 1$: loss aversion — losses loom larger than equal gains

Prospect theory (2/3): probability weighting and prospect value



$$w(p) = \frac{p^\gamma}{[p^\gamma + (1-p)^\gamma]^{1/\gamma}}$$

(gains; replace γ by δ in the loss domain)

- γ, δ : probability distortion in the gain / loss domain
- $\gamma < 1 \rightarrow$ inverse-S: rare events overweighted, likely events underweighted

Prospect value entering utility

$$V^+ = w^+ \cdot v(\Delta) \cdot 1[\Delta > 0] \geq 0$$
$$V^- = w^- \cdot v(\Delta) \cdot 1[\Delta \leq 0] \leq 0$$

$1[\cdot]$ is the indicator function; gains add value, losses subtract it

Prospect theory (3/3): two obstacles to empirical use

Where does **reference point** come from?

- Usual practice: fix R at the status quo for everyone
- Evidence: reference points are heterogeneous and formed from experience, expectations and demographics
- Treating R as fixed pre-processing leaves it inconsistent with the estimated preferences

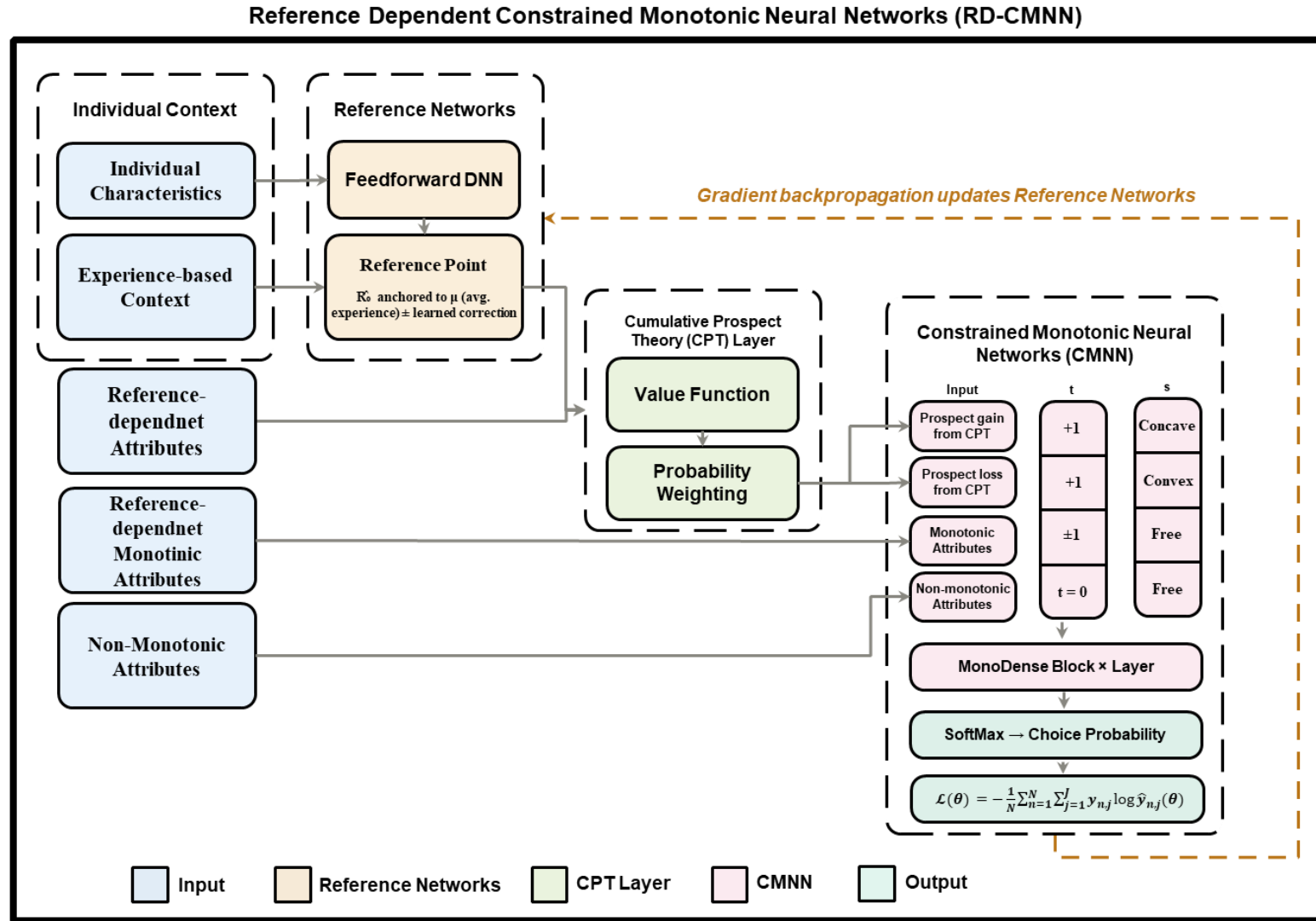
How **flexible** can utility be?

- Parametric forms are interpretable but restrictive, as they impose pre-defined shapes on attribute effects
- Unconstrained DNNs are flexible but may violate monotonicity and curvature
- Violations destroy marginal effects, WTP and elasticities

This study: solve both jointly

- **RD-CMNN**: a reference network learns individual-level R, a CPT layer converts deviations into prospect values, and a constrained monotonic neural network (CMNN) specifies utility
- All three components trained **end-to-end with a single cross-entropy loss** → reference points and behavioural parameters are mutually consistent
- Monotonicity and curvature are guaranteed **by construction**, not by penalty terms

The RD-CMNN architecture



Three components, one differentiable graph: reference networks $\rightarrow$ CPT layer $\rightarrow$ constrained monotonic network; gradients flow all the way back

Component 1: reference networks — learning R for each individual

$$R_{k,n} = x_{k,n} + \tanh \left(f_{\text{ref}}(X_{\text{ind},n}; \varphi) \right) \times \sigma(v_{\text{raw}}) \times \text{RNG}_{k,n}$$

Notation

- $R_{k,n}$: estimated reference point of person n for attribute k
- $x_{k,n}$: experiential anchor (e.g., the reported habitual commute time)
- $X_{\text{ind},n}$: demographics and attitudes of person n
- $f_{\text{ref}}(\cdot; \varphi)$: feedforward network, one per attribute;
- $\tanh(\cdot) \in (-1,1)$: direction of the correction
- $\sigma(v_{\text{raw}}) \in (0,1)$: learned intensity of the correction (trainable)
- $\text{RNG}_{k,n}$: observed range that bounds it: $x_{\text{max}} - x_{\text{min}}$, given, not learned

Key idea

- Start from what the person has experienced
- Let the data learn a bounded, signed correction around that anchor
- No status-quo assumption, yet the estimate stays behaviourally plausible

Empirical instantiation (commute time)

$$\bar{\tau}_{0,n} = \bar{\tau}_n + \tanh \left(f_{\text{ref}}(X_{\text{ind},n}; \varphi_{\text{ref}}) \right) \times \sigma(v_{\text{raw}}) \times (\tau_{\text{max},n} - \tau_{\text{min},n})$$

$\bar{\tau}_n$ anchor $x_{k,n} \rightarrow$ reported average commute time; $\text{RNG}_{k,n} \rightarrow$ individual's experienced commute range

Component 2: CPT layer (from deviations to prospect values)

Deviation

$$\Delta_{k,n,j} = X_{k,n,j} - R_{k,n}$$

sign of Δ decides gain vs. loss

Value function

$$v(\Delta) \text{ with } \alpha, \beta, \lambda$$

diminishing sensitivity + loss aversion

Probability weighting

$$\begin{aligned} w^+ &= w(p^+; \gamma) \\ w^- &= w(p^-; \delta) \end{aligned}$$

subjective distortion of gain/loss frequencies

Prospect values

$$V_{k,n,j}^+ \geq 0 \text{ and } V_{k,n,j}^- \leq 0$$

the inputs handed to the CMNN

- n = individual, j = alternative, k = reference-dependent attribute
- $\hat{p}_k^+$ = empirical share of gain outcomes for attribute k in the training set
- The five CPT parameters $\omega = \{\alpha, \beta, \lambda, \gamma, \delta\}$ are kept strictly positive by log-parameterisation and estimated jointly with the networks

Component 3: CMNN — flexibility with guaranteed shape

- **Monotonicity by construction:** sign constraints on the weight matrix, driven by an indicator t per input

Input group	Monotonicity t	Curvature s
Gain prospect values V^+	$t = +1$ (non-decreasing)	Concave
Loss prospect values V^-	$t = +1$ (non-decreasing)	Convex
Reference-neutral monotonic attributes	$t = \pm 1$ (known direction)	Free
Non-monotonic attributes (e.g. categorical)	$t = 0$ (unconstrained)	Free

$$\text{output} = \rho^s \left(|W^T|_t \cdot x + b \right)$$

W = weights, b = bias, x = layer input

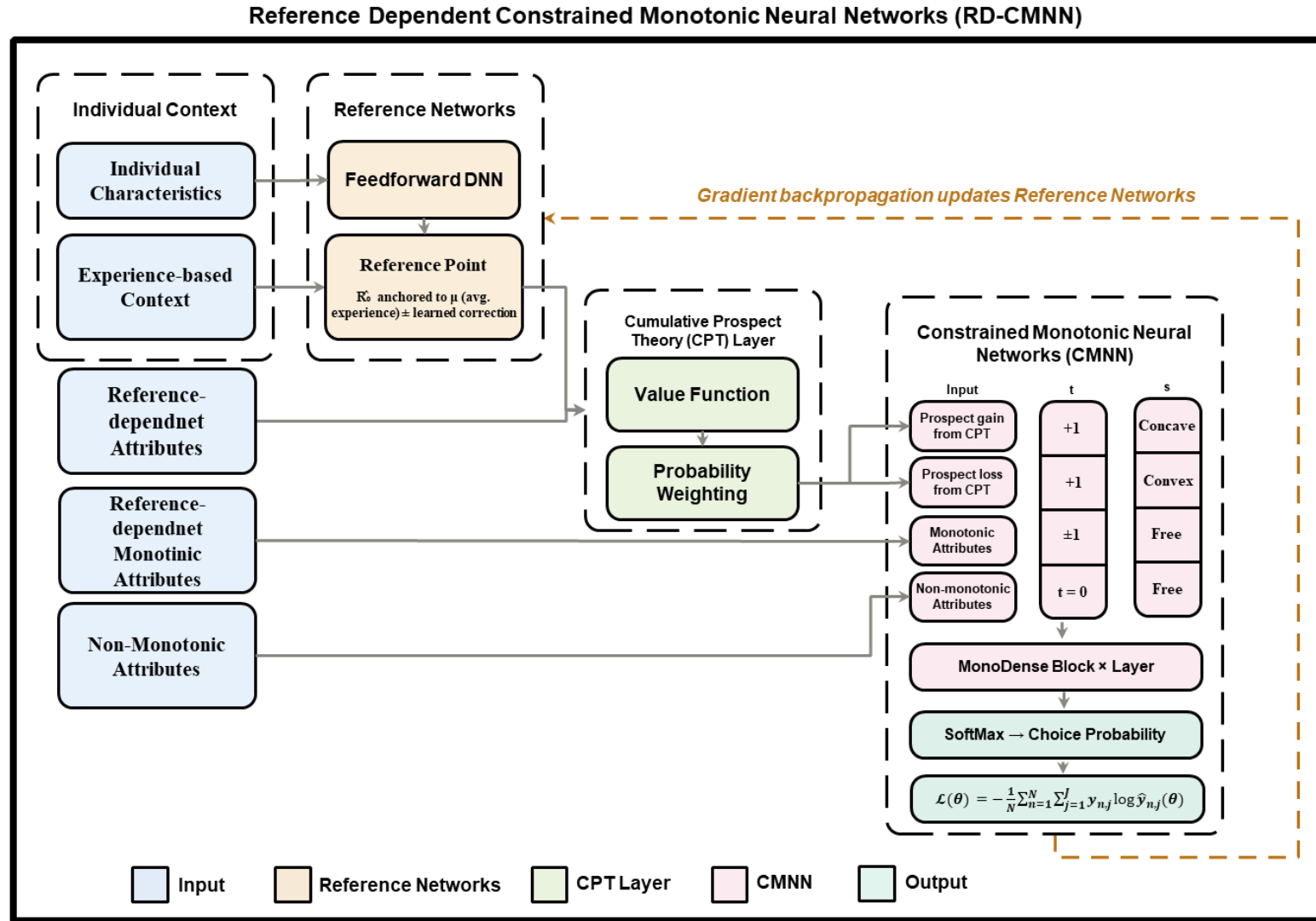
- $|W^T|_t$: sign constraint. $t = +1 / -1 / 0$ per input: utility must rise, fall, or is unconstrained
- ρ^s : curvature partition. $s = (\dot{s}, \hat{s}, \tilde{s})$ splits the hidden neurons into convex, concave and bounded activations

Curvature control

- Weight-constrained networks with a convex activation can only produce convex utility
- CMNN splits the neurons of each layer into three groups: convex ρ , concave $-\rho(-x)$, and bounded $\tilde{\rho}$
- Any monotonic continuous function can then be approximated — concave gains and convex losses included

Utility $U(n,j) = V(x(n,j); \psi) + \varepsilon(n,j)$, with ε i.i.d. Type-I EV $\rightarrow$ logit choice probabilities

The RD-CMNN architecture



Three components, one differentiable graph: reference networks $\rightarrow$ CPT layer $\rightarrow$ constrained monotonic network; gradients flow all the way back

Joint, end-to-end estimation

$$L(\theta) = -\frac{1}{N} \sum_n \sum_j y_{n,j} \log \hat{y}_{n,j}(\theta), \quad \theta = \{\varphi, \omega, \psi\}$$

- φ reference-network weights, $\omega = \{\alpha, \beta, \lambda, \gamma, \delta\}$ CPT parameters, ψ CMNN weights: one loss, one optimisation
- $y_{n,j}$ is the observed choice, $\hat{y}_{n,j}$ the predicted probability

Gradients reach R

Because Δ depends on R, the choice signal itself updates the reference networks, i.e. R is no longer fixed pre-processing

Constraints enforced during training

Projected gradient descent after each update; CPT parameters positive by log-parameterisation; AdamW with exponential LR decay

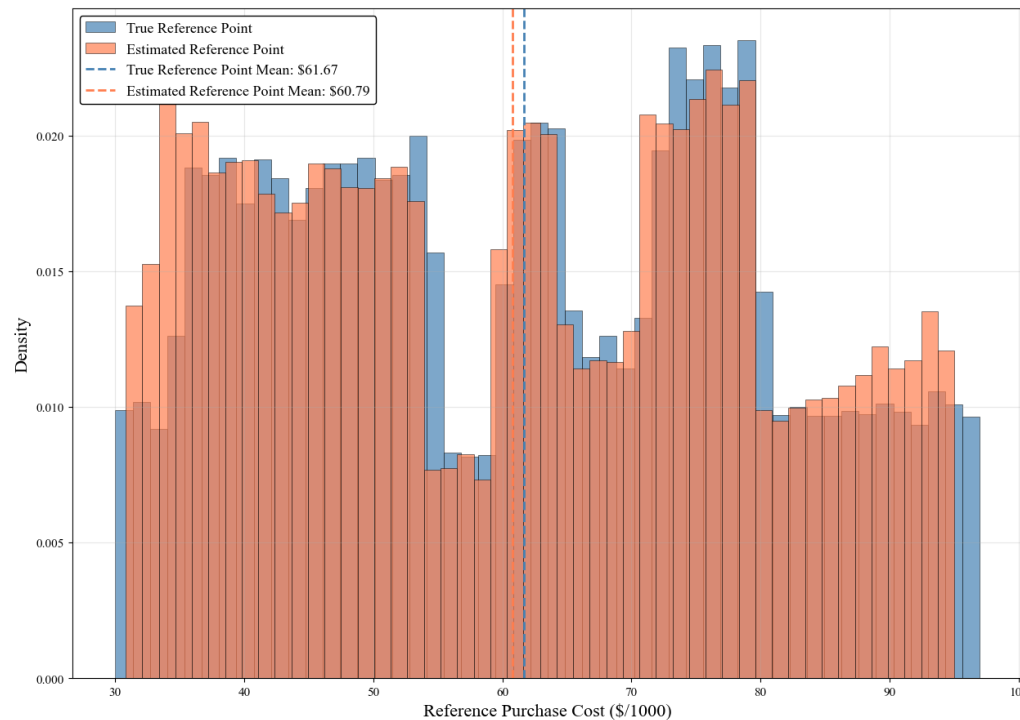
Hyperparameters by Bayesian optimisation

Width, depth, dropout, learning rate and weight decay tuned for both components; grid search is infeasible here

Evaluation: accuracy, AUC and Brier score on held-out individuals

Monte Carlo: can we recover the truth?

- 20,000 synthetic choices from 5,000 individuals; reference points for operating cost and purchase price generated from income, employment and flexible-work status, with interactions
- Three configurations: baseline, stronger loss aversion, and weaker heterogeneity



Estimated vs. true reference purchase price (Config. 1)

Parameter	True	Estimated
Intercept	60.0	59.9
Income	20.0	19.1
Full-time	12.0	11.4
Flexible work	-30.0	-30.4
Full-time × Flexible	-7.0	-7.6
Income × Full-time	5.0	6.1

Reference-point coefficients for purchase price, recovered by regressing the estimated R on the generating variables

Main effects, interactions and the multimodal shape of the reference-point distribution are all recovered, also under weak heterogeneity

Empirical application: cooperative routing enrolment

- Stated-preference survey in **Singapore (400)**, the **UK (405)** and the **US (403)**; 14 choice scenarios each
- Cooperative Routing Scheme (CRS): drivers accept an occasional longer route in exchange for frequent time savings

Section I — driving patterns

- Each respondent reported their **average, minimum and maximum** commute time
- $\bar{\tau}_n$ anchors the reference point; $\tau_{\min,n}$, $\tau_{\max,n}$ bound the learned correction
- Also collected: age, gender, income, education, weekly driving frequency

Sections II–IV — 14 choice tasks

- II: 6 tasks on a two-route network personalised to their own $\bar{\tau}_n$
- III: 4 tasks repeated, now revealing the **enrolment ratio**
- IV: same 4 tasks plus **system-level CO₂ benefits**

Attribute	Current situation		If one enrolls in CRS	
	Route 1	Route 2	Route 1	Route 2
Travel time	59 min	59 min	49 min	68 min
Routing plan			65%	35%
Total travel time over 20 days	1180 min		1113 min	
Time-saving over 20 days			67 min (6% reduction)	
CRS enrolment ratio			80%	
Time savings, all travellers, 20 days			19,080 hours (8% reduction)	
CO ₂ savings over 20 days			56 trees planted	

Model specifications compared

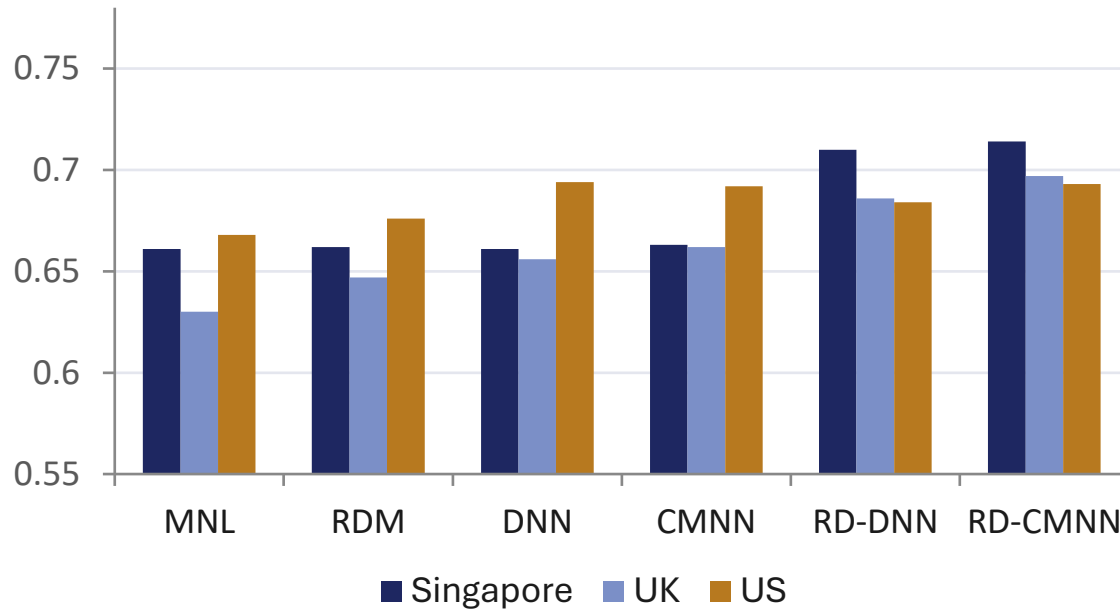
- “Status quo” = the respondent’s **self-reported average commute time** $\bar{\tau}_n$, taken as given and held **fixed** throughout estimation
- “Learned, joint” = the reference point $\hat{\tau}_{0,n}$ is **predicted from individual characteristics** and estimated together with the CPT and network parameters
- Six specifications isolate **one design choice at a time**; all share the same data, folds and metrics

Model	Reference point	Utility function	Shape constraints
MNL	—	Linear	By construction
RDM	Status quo (reported)	Parametric CPT	By construction
DNN	Status quo (reported)	Unconstrained network	None
CMNN	Status quo (reported)	Constrained network	Guaranteed
RD-DNN	Learned, joint	Unconstrained network	None
RD-CMNN (proposed)	Learned, joint	Constrained network	Guaranteed

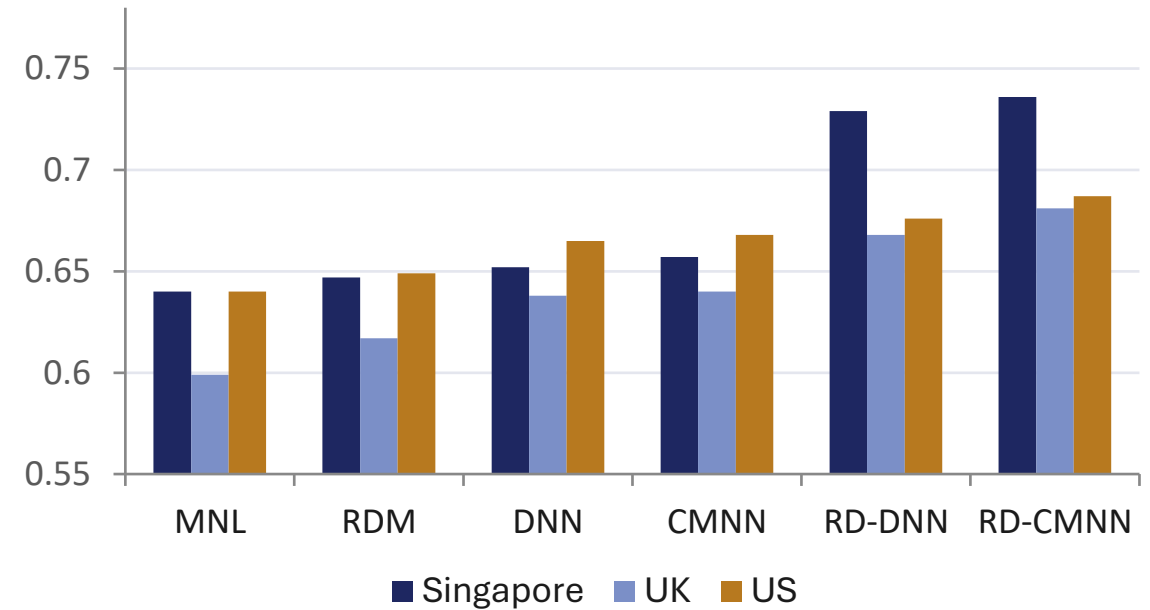
Reference dependence, then flexibility, then a learned reference point (with and without constraints)

Predictive performance: the learned reference point does the work

Accuracy (higher is better)

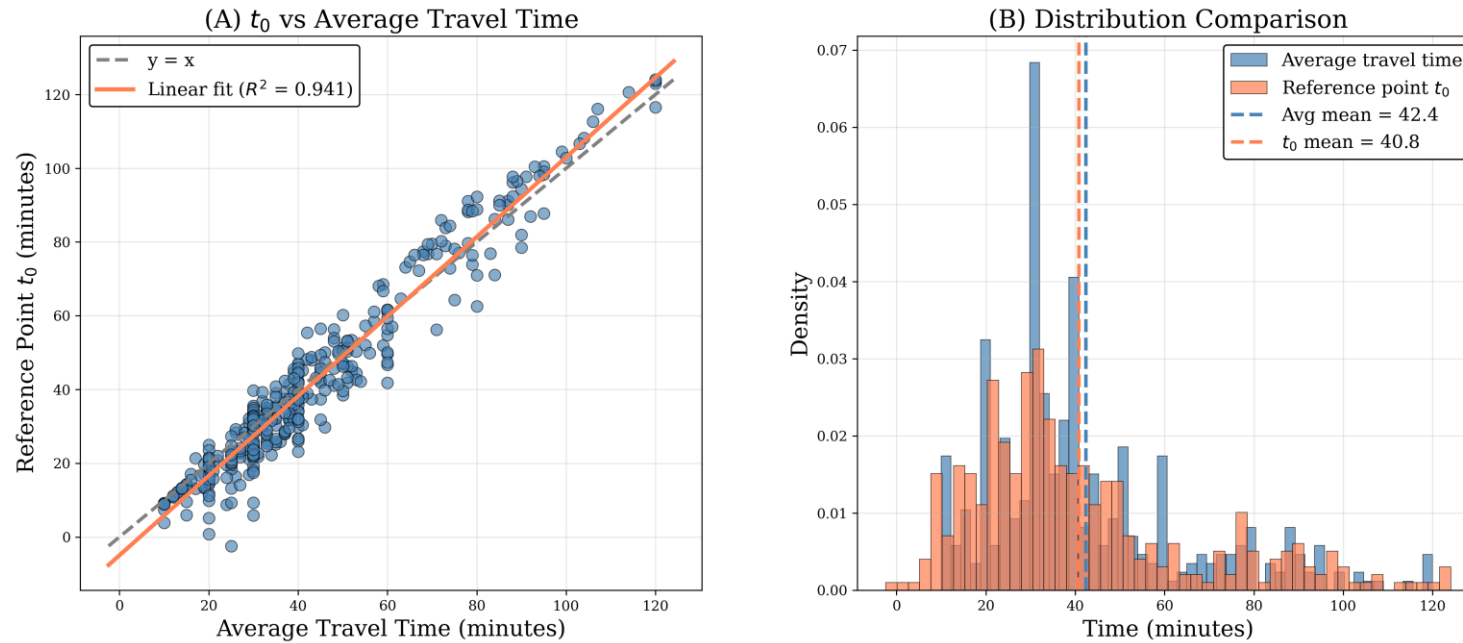


AUC (higher is better)



- Reference-dependent networks (RD-DNN, RD-CMNN) clearly beat models that use the reported status quo; flexibility alone (DNN, CMNN) adds little over RDM
- **RD-CMNN matches RD-DNN — imposing monotonicity costs nothing in predictive accuracy**

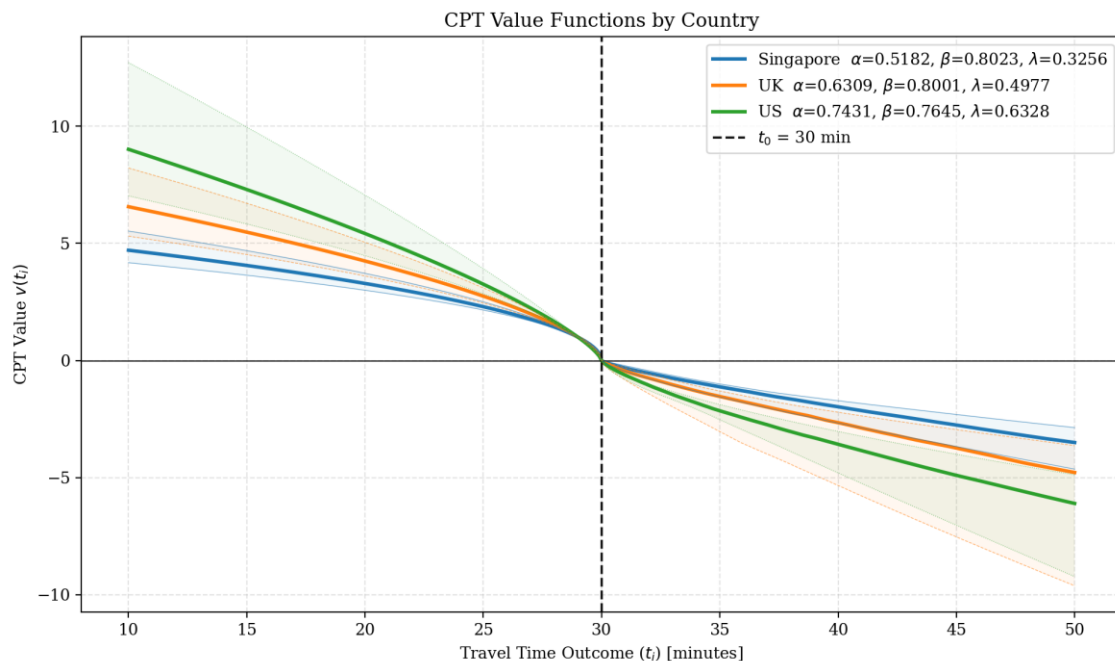
Are the learned reference points behaviourally sensible?



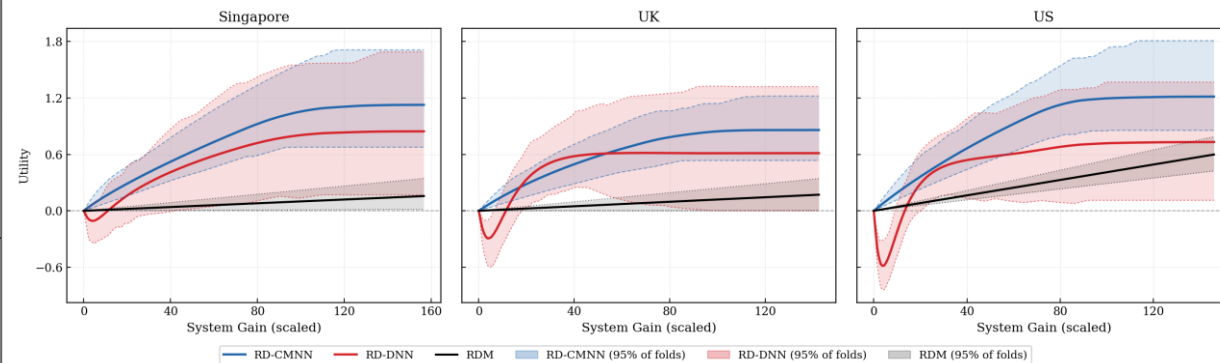
RD-CMNN, Singapore ($n = 400$): estimated reference point vs. reported average commute time

- Nothing forces the model to reproduce the status quo, yet across all three countries R^2 ranges from **0.919 to 0.941** and the means nearly coincide
- The learned reference points stay anchored in experience while allowing systematic individual deviations from it
- Same picture for RD-DNN → the mechanism is robust to the choice of architecture

Interpretability: what constraints buy us



RD-CMNN value functions by country



Partial dependence: utility of system-level gains

- $\alpha < 1$ everywhere: diminishing sensitivity to time savings, strongest in Singapore (0.52) and weakest in the US (0.74) — same ordering as the parametric RDM
- RD-CMNN also recovers convex losses ($\beta < 1$) that the parametric RDM misses, and gives the tightest confidence bands of the three models
- **RD-DNN violates theory:** non-monotonic dips in the partial dependence of system gains, S-shaped gain weighting ($\gamma > 1$), and compressed cross-country differences

Conclusions

- **Reference points can be learned:** A reference network anchored on experience recovers heterogeneous reference points in simulation and yields the largest predictive gain in the empirical study
- **Constraints are free:** RD-CMNN matches the unconstrained RD-DNN in accuracy, AUC and Brier score, while guaranteeing monotonicity and the curvature prospect theory requires
- **Constraints are also necessary:** The unconstrained model violates CPT and utility-theory principles, so its recovered parameters cannot be read behaviourally

A Theory-Constrained Framework for Data-Driven Learning of Multiple Discrete-Continuous Choices

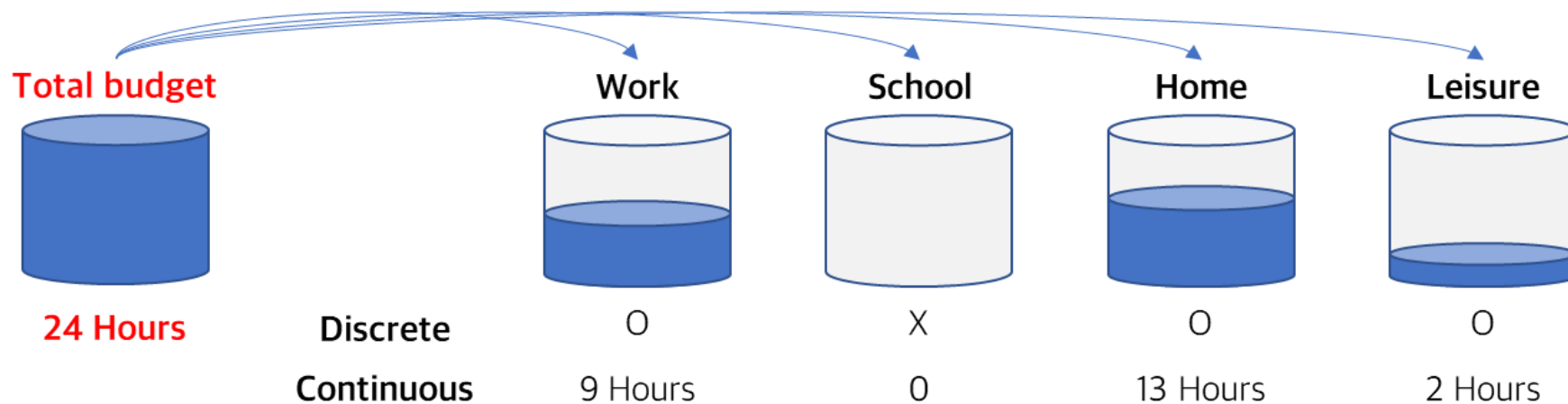
Huichang Lee¹, Jason Hawkins², Dong-Kyu Kim¹, Prateek Bansal³, Eui-Jin Kim⁴

¹Seoul National University, ²University of Calgary, ³National University of Singapore, ⁴Ajou University

Summer School, University of Tokyo · 7 September 2026

Research Context

- Many consumer choice situations are characterized by **simultaneous demand for multiple alternatives** that are imperfect substitutes for one another (e.g., activity time-use, grocery shopping)
- Multiple discrete-continuous choice (MDC) models jointly model **multiple discrete choices and continuous consumption levels**
- Individuals are assumed to **maximize the sum of alternative-specific utilities** by allocating their budget across multiple alternatives



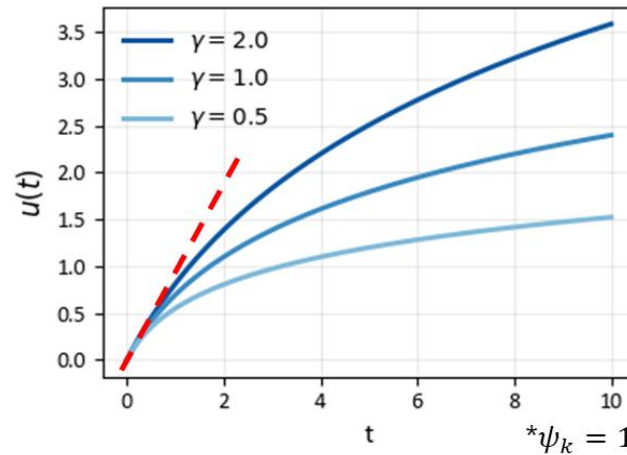
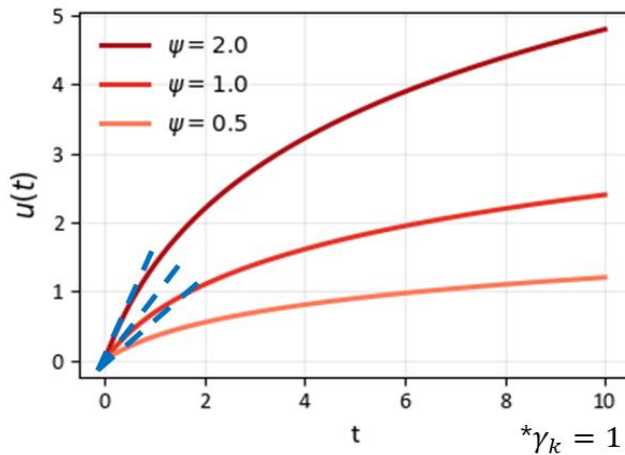
$$\text{Optimal activity time-use } \mathbf{t} = \operatorname{argmax}_{\mathbf{t}} \sum_k U_k(t_k)$$

Research Context

- In MDC models, the satiation effect (i.e., monotonically increasing utility with diminishing marginal utility)
- The multiple discrete-continuous extreme value (MDCEV) model (Bhat, 2008) is most widely used model

The MDCEV model

$$U_k(t_k) = \gamma_k \psi_k \ln\left(\frac{t_k}{\gamma_k} + 1\right)$$

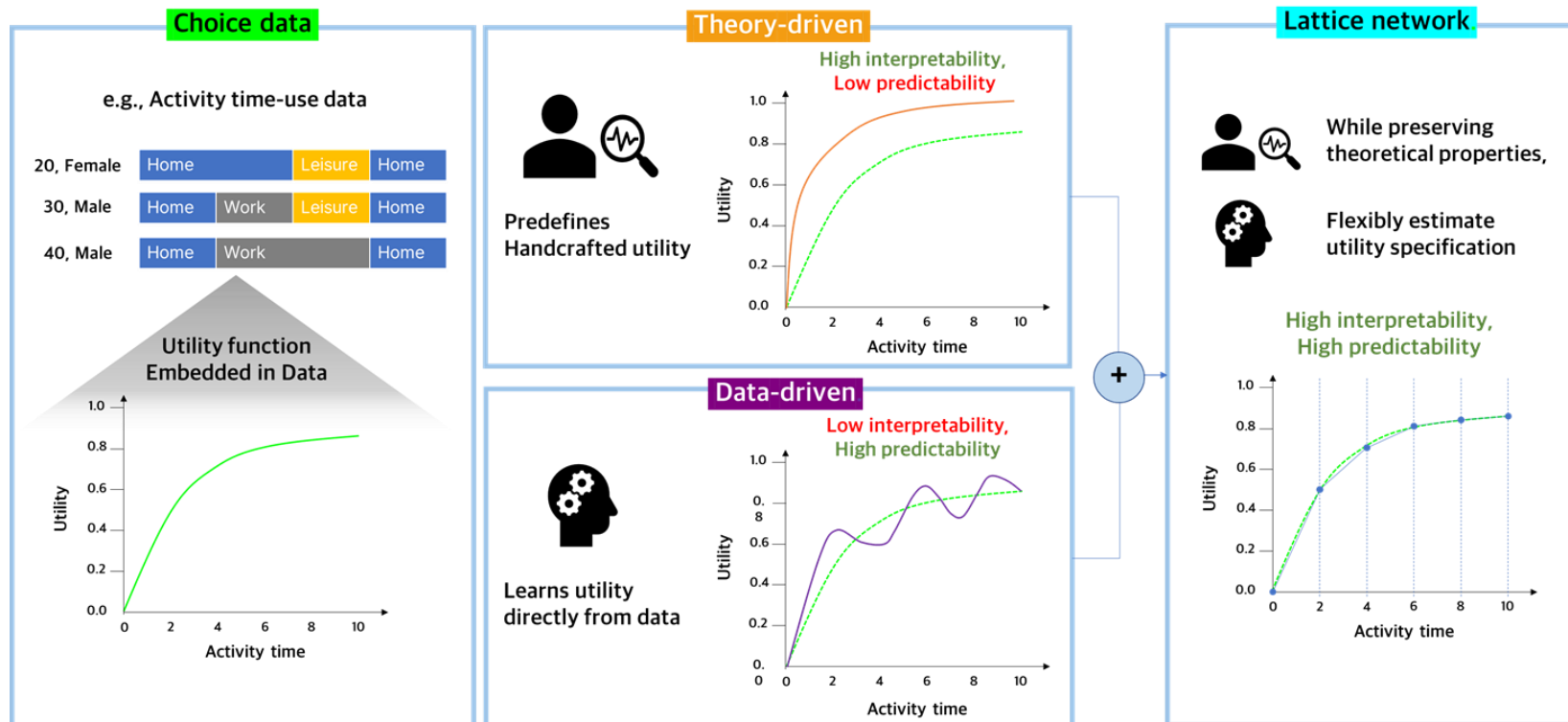


ψ_k : Discrete choice
Baseline marginal utility at zero consumption

γ_k : Continuous choice
Diminishing marginal utility

Motivation

- Theory-driven models (e.g., MDCEV) are highly interpretable with parsimonious utility specifications, but their **limited flexibility** constrains predictive performance
- Data-driven models can achieve higher predictive performance by flexibly learning utility functions from data, but **excessive flexibility may reduce interpretability**
- We develop a **theory-constrained data-driven model** that preserves the KKT structure, balancing predictive performance and interpretability

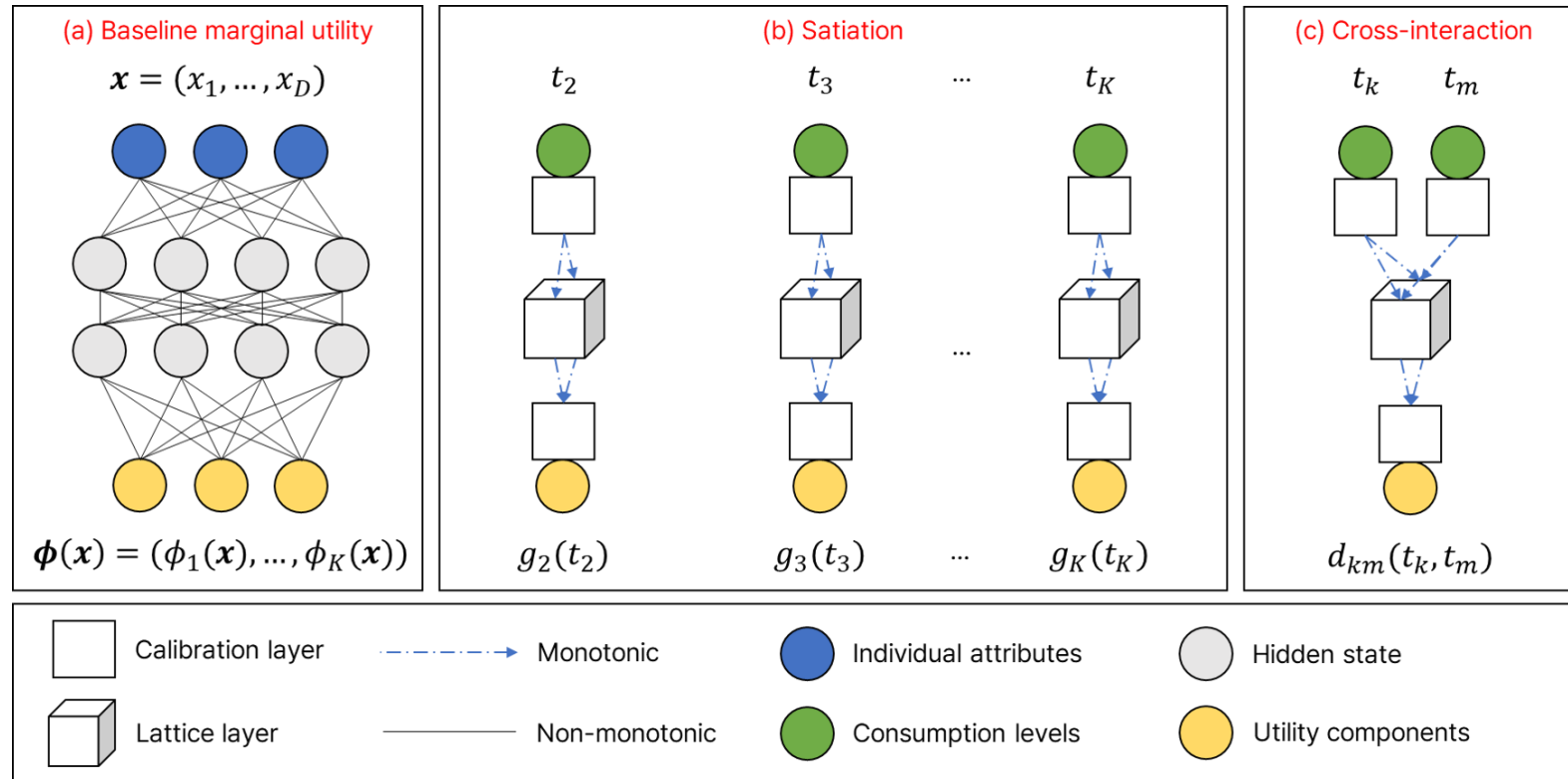


Model formulation

- We generalize the utility structure of MDC models and estimate each utility component directly from data

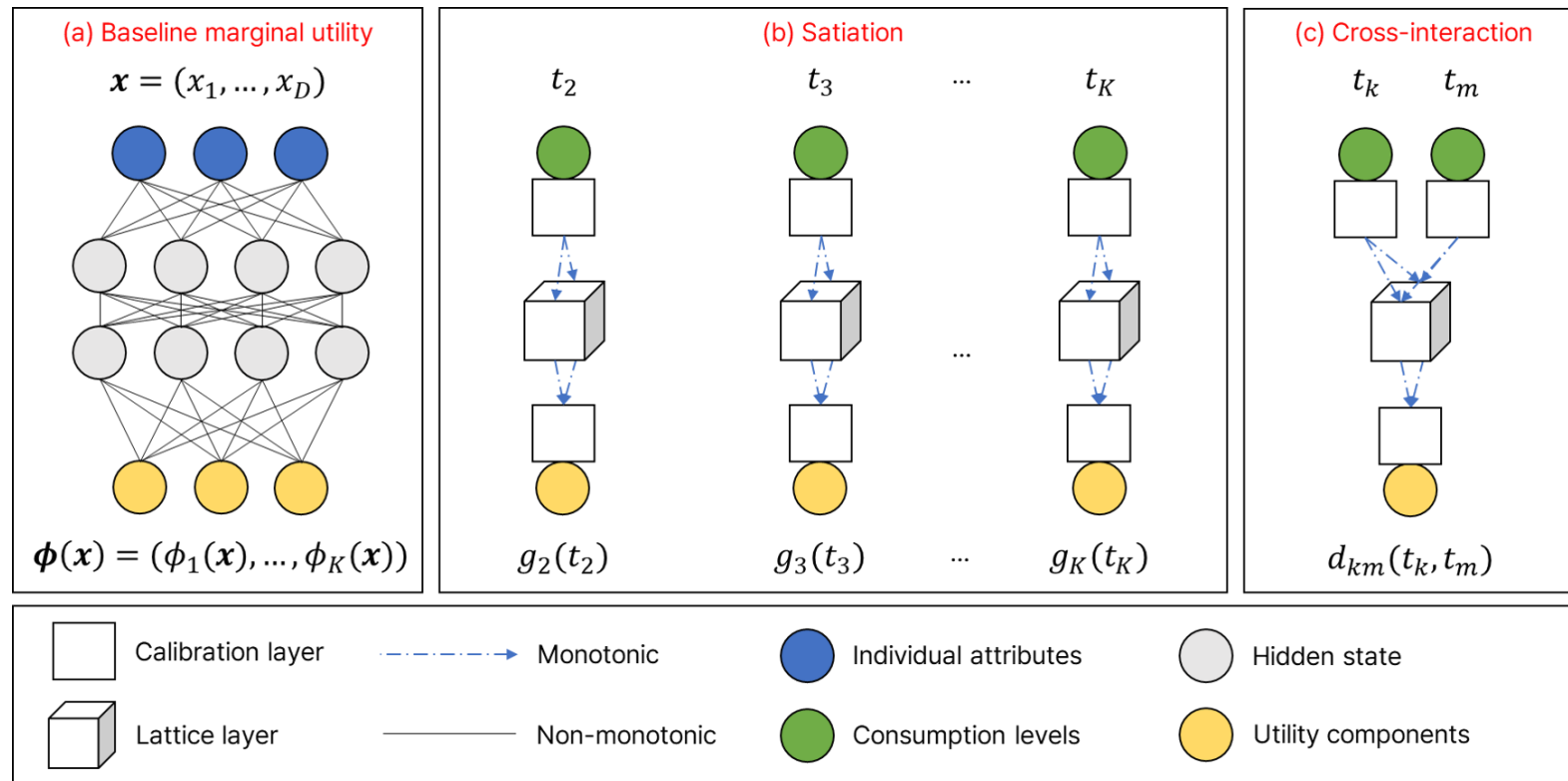
$$U(t) = \phi_1(x)t_1 + \sum_{k=2}^K \phi_k(x)g_k(t_k) + \sum_{k=2}^K \sum_{\substack{m < k \\ m \neq 1}} d_{km}(t_k, t_m)$$

- ϕ_k : baseline marginal utility (parameterized by individual-specific attributes x)
- g_k : satiation effect, d_{km} : cross-alternative interactions



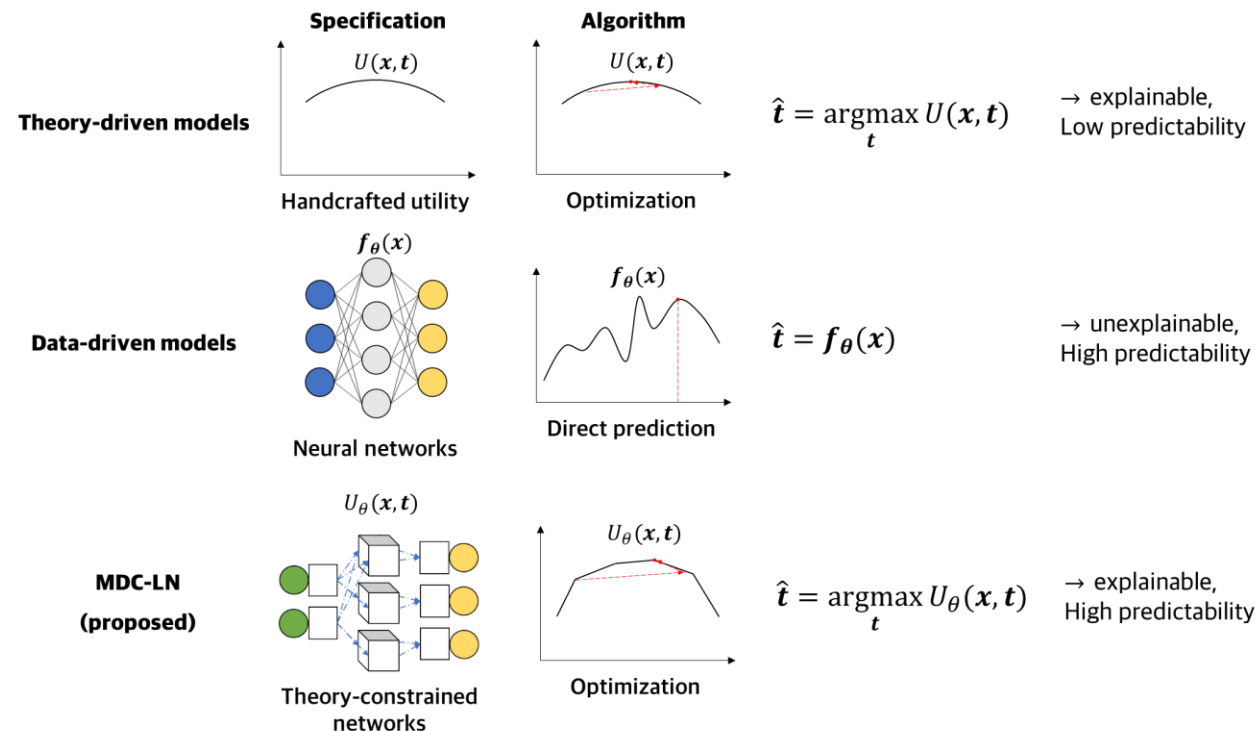
Model formulation

- **Lattice networks (LNs)** (Gupta et al., 2018) preserve **economic shape constraints** (i.e., monotonicity and concavity) through piecewise linear functions and multilinear interpolation
- According to its shape constraints, different network architectures are applied to each utility component:
 - **Baseline marginal utility (a)**: flexible multilayer perceptrons (MLPs)
 - **Satiation (b) and cross-interactions (c)**: shape-constrained LNs



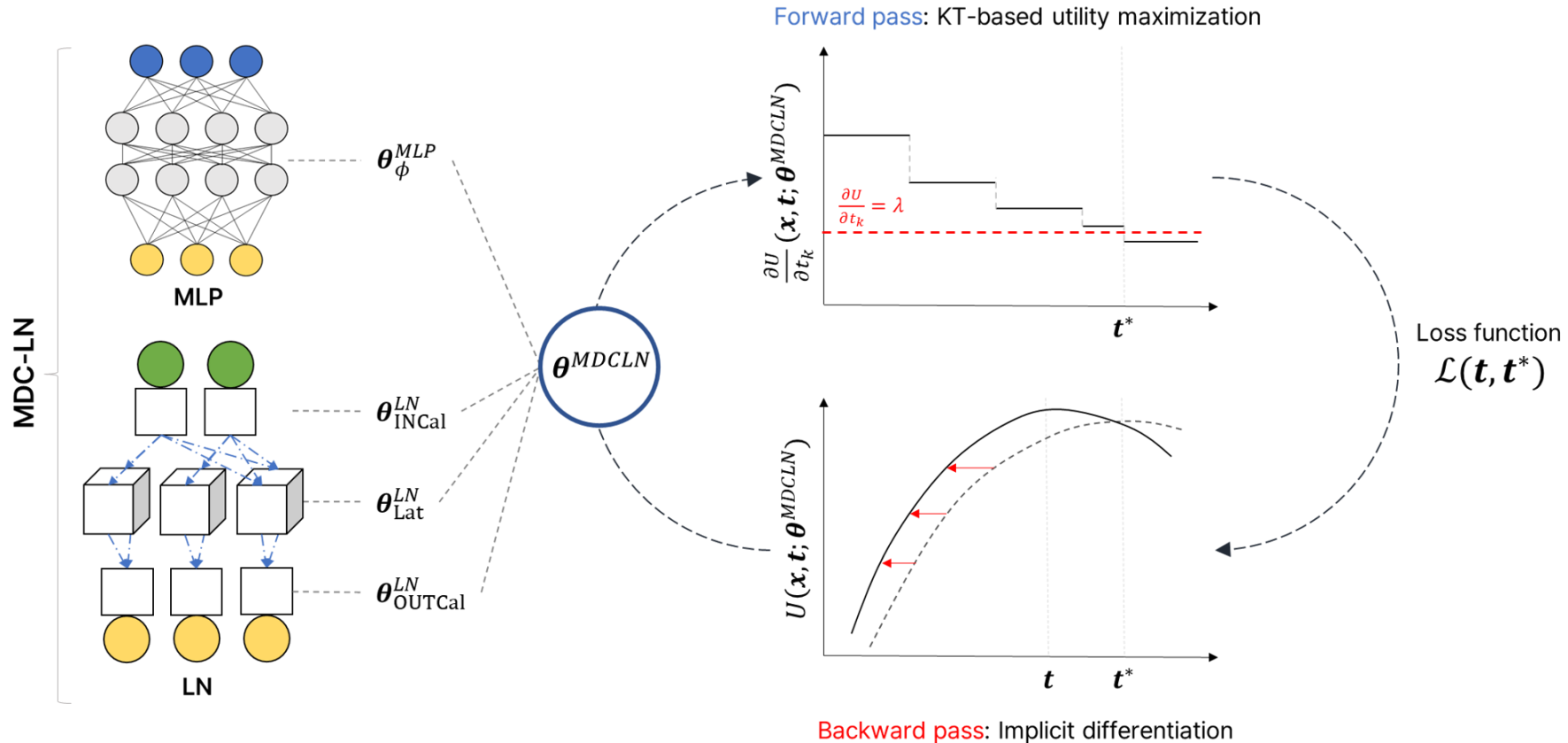
Model training

- Data-driven models are trained by updating a large number of parameters through **backpropagation**
 - Kim and Bansal (2024) specifies the systematic utilities of a discrete choice model using LNs, applies a softmax function to obtain choice probabilities, and updates the model parameters through backpropagation
- **Updating the parameters of MDC-LN is not straightforward**
 - Consumption levels must be obtained by solving a **utility maximization problem**



Model training

- Training MDC-LN involves two alternating optimization steps:
 - **Forward pass:** Given the current model parameters θ^{MDCLN} , solve the utility maximization problem to obtain the optimal consumption vector t^* (using the KT conditions)
 - **Backward pass:** Backpropagate the gradient of the loss between the observed t and predicted t^* to update θ^{MDCLN} (implicit differentiation without unrolling the optimization steps)



Model training: forward pass

- The Lagrangian for the MDC-LN utility maximization problem is

$$Lagr(t) = \phi_1(x)t_1 + \sum_{k=2}^K \phi_k(x)g_k(t_k) + \sum_{k=2}^K \sum_{\substack{m < k \\ m \neq 1}} d_{km}(t_k, t_m) + \lambda(T - \sum_{k=2}^K t_k)$$

KT conditions:

$$\frac{\partial Lagr}{\partial t_1} = 0: \phi_1 - \lambda = 0 \quad (\text{outside good, always consumed})$$

$$\frac{\partial Lagr}{\partial t_k} \leq 0: \phi_k g'_k(t_k) + \sum_{k=2}^K \sum_{\substack{m < k \\ m \neq 1}} \frac{\partial d_{km}}{\partial t_k}(t_k, t_m) - \lambda \leq 0 \quad (\text{equality when consumed, inequality otherwise})$$

- We introduce a normal error term into ϕ_k to capture **unobserved heterogeneity**

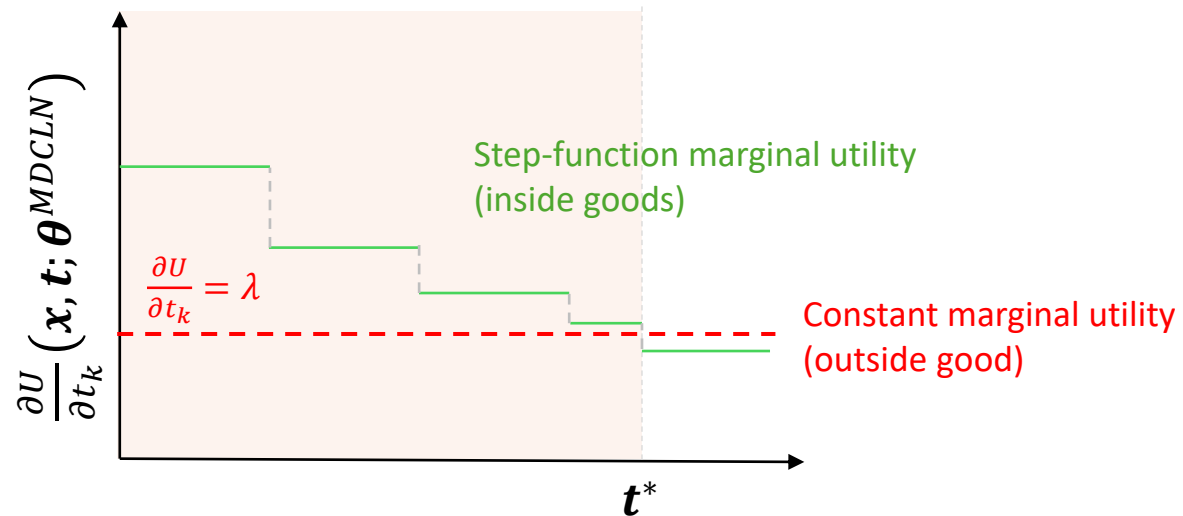
$$\phi_k(x) = \exp\left(h_{\phi,k}\left(x; \theta_{\phi}^{MLP}\right) + \varepsilon_k\right), \quad \varepsilon_k \sim N(0,1)$$

- Rearranging the KT conditions w.r.t. the error term yields **the discrete choice probability** of alternative k

$$\Pr(t_k^* > 0 \mid t_{-k}^*) = \Phi\left(h_{\phi,k}\left(x; \theta_{\phi}^{MLP}\right) + \log g'_k(0) - \log\left(\phi_1 - \sum_{k=2}^K \sum_{\substack{m < k \\ m \neq 1}} \frac{\partial d_{km}}{\partial t_k}(0, t_m^*)\right)\right)$$

Model training: forward pass

- Solving the utility maximization problem is a **major computational bottleneck** in MDC models
- We efficiently compute the forward pass by exploiting the **structural properties** of the utility function
 - **A linear utility** for the outside good
 - **Monotonic, concave, and piecewise-linear utilities** for inside goods



- Further consumption is not beneficial once **marginal utility** falls below the **outside-good marginal utility**
- Consumption is obtained by accumulating segment lengths above this **outside-good marginal utility threshold**
- When cross-interactions are present, consumption is sequentially updated while holding the other alternatives fixed

Model training: backward pass

- The backward pass updates model parameters by backpropagating the gradient of the loss function

$$\arg \min_{\theta^{MDC-LN}} \sum_{n=1}^N [\mathcal{L}_{\text{bce}}(t_n, t_n^*) + \lambda_{\text{rec}} \mathcal{L}_{\text{rec}}(t_n, t_n^*)] + \mathcal{R}_{\text{complex}}(\theta^{MDC-LN})$$

- $\mathcal{L}_{\text{bce}}$ supervises the **discrete-choice component** by penalizing discrepancies between the choice probabilities $\hat{p}_{nk}$ derived from the KT conditions and the observed choice indicator y_{nk}

$$\mathcal{L}_{\text{bce}} = -\frac{1}{K-1} \sum_{k=2}^K [y_{nk} \log \hat{p}_{nk} + (1 - y_{nk}) \log(1 - \hat{p}_{nk})]$$

- $\mathcal{L}_{\text{rec}}$ supervises the **continuous-choice component** by penalizing discrepancies between the observed and predicted continuous consumption levels

$$\mathcal{L}_{\text{rec}} = \frac{\sum_{k=2}^K y_{nk} (t_{nk} - t_{nk}^*)^2}{\sum_{k=2}^K y_{nk}}$$

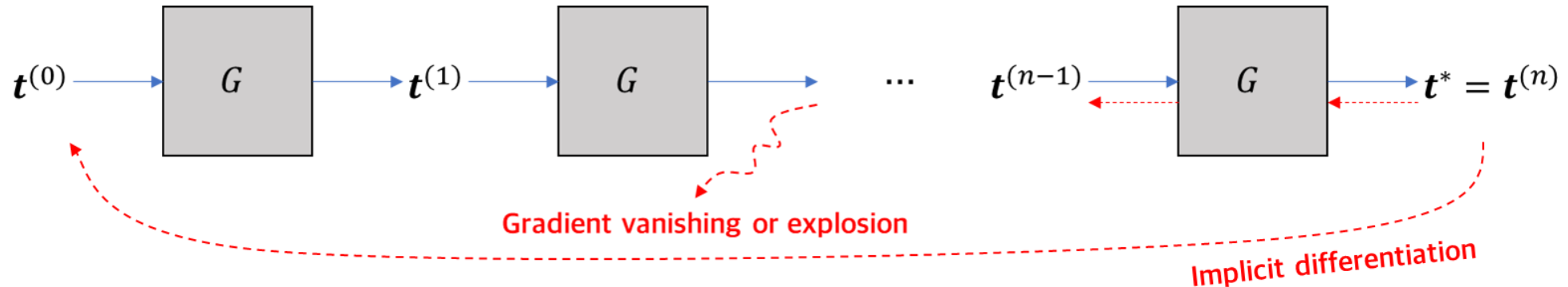
- $\mathcal{R}_{\text{complex}}$ controls **model complexity** by regularizing LN curvature and the Hessian
- We use a **pseudo-likelihood** for training because the full likelihood requires intractable integration and Jacobian

Model training: backward pass

- We decouple the forward pass from differentiation in the backward pass to bypass its non-differentiability
- A gradient-ascent update is used as a **surrogate fixed-point operator** for the backward pass

$$G_k(t_k) = \max \left\{ 0, t_k + \eta \left(\frac{\partial U(t^{(n)})}{\partial t_k} - \frac{\partial U(t^{(n)})}{\partial t_1} \right) \right\}, \quad \eta: \text{step size}$$

- **Implicit differentiation** (Bai et al., 2019) treats the fixed-point solution as an implicit function of the model parameters
- This allows gradients to be computed without tracing intermediate optimization steps



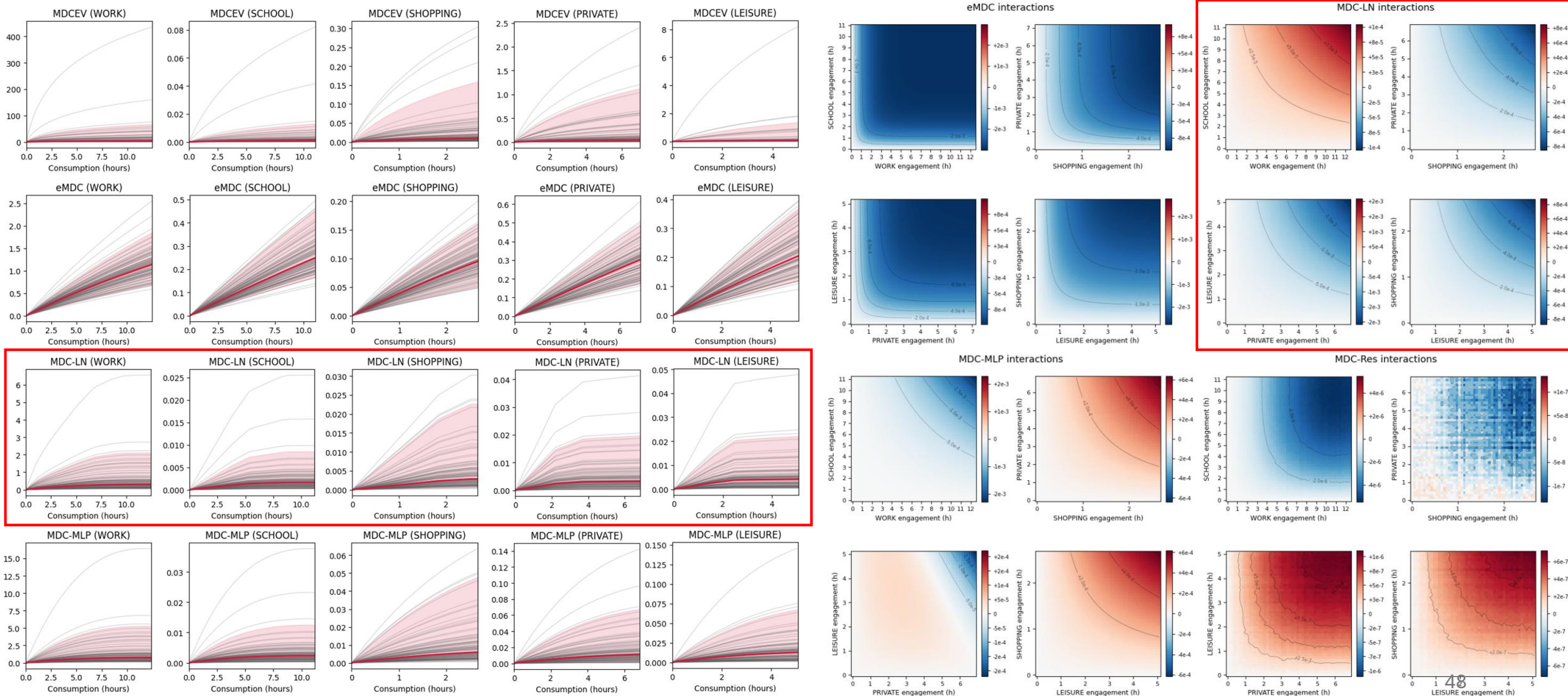
Results: Predictive performance

- We compare the discrete-choice error (Brier score) and continuous-choice error (RMSE) across MDC models
- MDC-LN achieves predictive performance comparable to the unconstrained MDC-MLP benchmark

Activity	Model	True eng. rate	Predicted	Brier score	RMSE
Work	MDCEV	66.0%	63.6%	0.0267	2.76
	MDC-LN		66.4%	0.0237	2.88
	MDC-MLP		66.5%	0.0242	2.67
School	MDCEV	5.0%	4.6%	0.0085	3.74
	MDC-LN		5.2%	0.0078	3.56
	MDC-MLP		5.2%	0.0081	3.91
Shopping	MDCEV	10.6%	7.8%	0.0875	1.72
	MDC-LN		10.7%	0.0726	1.52
	MDC-MLP		10.5%	0.0732	1.57
Private business	MDCEV	14.2%	10.9%	0.1099	3.70
	MDC-LN		16.5%	0.0968	3.48
	MDC-MLP		14.7%	0.0970	3.48
Leisure	MDCEV	15.4%	11.3%	0.1221	2.83
	MDC-LN		16.1%	0.1122	2.73
	MDC-MLP		15.7%	0.1109	2.75

Results: Interpretability

- We estimated the underlying utility functions from Seoul activity time-use data
- MDC-LN flexibly learns utility functions **subject to shape constraints** (i.e., monotonicity and concavity)



Results: Computational burden

- Activity time-use prediction for 11,446 individuals: 7,154s (eMDC; Palma and Hess, 2022) → 159s (MDC-LN)
- MDC-LN substantially reduces prediction time by exploiting its structural properties

Conclusions

- Theory-constrained data-driven MDC models improve the predictive performance and computational efficiency of MDC models.
- MDC-LN achieves predictive performance comparable to the unconstrained benchmark while preserving economic shape constraints.
- The structural properties of MDC-LN alleviate the computational bottlenecks of conventional MDCEV models.



Behavioural
— COMPUTATIONAL SCIENCE LAB —



LinkedIn Page
(Please follow for updates)