

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Regularization more than representational capacity drives heuristic discovery in Bounded Meta-Learned Inference

Permalink

<https://escholarship.org/uc/item/7g07p6qr>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Reining, Lars C.

Radatz, Finn

Cremille, Nora

et al.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Regularization more than representational capacity drives heuristic discovery in Bounded Meta-Learned Inference

Lars C. Reining (lars.reining@tu-darmstadt.de)¹, Finn Radatz¹, Nora Cremille¹ & Constantin A. Rothkopf¹
¹Centre for Cognitive Science, Technical University of Darmstadt

Abstract

Finding computational models that give rise to human-like heuristic decision-making is an open challenge in cognitive science. Meta-learning has been proposed as a promising computational framework for obtaining algorithmic models of cognition by learning from environmental interactions rather than designing normative models. Recent work by Binz et al. (2022) demonstrates that bounded meta-learned inference (BMI) discovers human-like heuristic decision strategies in paired comparison tasks. However, the reasons and mechanisms underlying this discovery remain unclear. Here, we extend research on the BMI framework to better understand its behavior and reverse-engineer its computational principles. We first systematically varied model parameters that influence the network's representational capacity, finding that BMI robustly discovers context-appropriate heuristics across configurations. The model consistently uses single-cue strategies when feature rankings are known and equal weighting strategies when feature directions are known, even when network capacity is strongly reduced compared to the original work. We then reverse-engineer the computational principles underlying this behavior using hierarchical Bayesian modeling. Our analyses reveal that regularization, not network size, drives heuristic discovery. Bayesian models with heuristic priors provide the best fit for regularized network behavior when environmental cues are available. Additionally, we discovered a stronger inductive bias toward equal weighting over single-cue strategies. Thus, we conclude that inductive biases, in the form of priors, have a greater impact on the emergence of heuristics than resource rationality, in the form of networks' representational capacity.

Keywords: meta-learning; representational capacity; heuristic decision-making; Bayesian modeling

Introduction

Paired comparison tasks are used to study human decision-making and its adaptation to different environmental structures (Binz et al., 2022). In such tasks, participants must choose between two options, each represented by a set of features, and learn to weigh these features appropriately based on feedback to their choices. For example, an individual consumer might need to decide between two products based on features such as price, quality, and brand reputation.

For specific paired comparison tasks, an ideal observer can be formulated that describes optimal performance under the assumption of complete knowledge of the underlying generative process (Binz et al., 2022; Geisler, 1989). However, humans often systematically deviate from this ideal (Tversky & Kahneman, 1974), a behavior that has been described

as heuristic decision-making—simple rules or strategies that are conceptualized as enabling quick and efficient decisions (Gigerenzer & Todd, 1999; Martignon & Hoffrage, 2002) while reducing computational effort (Tversky & Kahneman, 1974). Common heuristics include "one-reason" decision-making, where choices are based on a single most important feature (Gigerenzer & Goldstein, 1996), and equal weighting strategies, where all features receive equal importance (Dawes & Corrigan, 1974; Einhorn & Hogarth, 1975).

However, where such heuristic strategies come from remains an open question. An assumption that has been put forward is that they arise naturally from limited computational resources, reflecting a form of resource rationality (Gigerenzer & Todd, 1999; Lieder & Griffiths, 2020).

One instantiation of this assumption that also provided a computational implementation is the work by Binz et al., which proposes that bounded meta-learned inference (BMI) can explain the emergence of heuristic decision-making in paired comparison tasks. BMI extends the meta-learned inference (MI) framework (Binz et al., 2022), where a neural network is meta-trained to learn generalizable inference strategies across a distribution of tasks (Bengio et al., 1991; Thrun & Pratt, 1998). During inference, the network is presented with multiple examples of a specific task and must quickly adapt its strategy based on the feedback received. BMI extends this framework by introducing a regularization term that penalizes complex inference strategies, reducing representational capacity and encouraging simpler strategies (Binz et al., 2022). This framework has successfully captured various cognitive phenomena observed in humans (Binz et al., 2022; Grau-Moya et al., 2022; Jagadish et al., 2024).

Binz et al. demonstrate that BMI can discover human-like heuristic decision-making strategies in paired comparison tasks. More specifically, they meta-train a recurrent neural network on a distribution of paired comparison tasks. When regularization is applied, i.e. BMI instead of MI, and the environment provides specific cues about how features need to be combined, the network learns to use heuristic strategies that resemble human decision-making. For example, when the environment provides information about the ranking of feature importance (e.g., price is more important than quality), the BMI network learns to use a single cue strategy, focusing only on the most important feature. When, instead, the environment provides information about the direction of feature weights (e.g., that a higher price is worse, which might

be obvious to humans but not to a neural network trained on the task), the BMI network learns to use an equal weighting strategy, giving equal importance to all features.

However, the reasons and mechanisms underlying this discovery remain unclear. Binz et al. suggest that resource constraints implemented through regularization drive the emergence of heuristics, but it remains uncertain whether the primary driver is the network’s representational capacity (e.g., number of neurons) or the regularization penalty itself.

To address this question, we conduct two complementary analyses. First, we systematically vary key model parameters, specifically network size and input feature dimensionality, to investigate BMI at the algorithmic level (Marr, 2010) and test whether heuristic discovery by BMI depends on representational capacity. If capacity is the main driver, smaller networks or those with higher input dimensionality should discover heuristics even without regularization. Second, we use hierarchical Bayesian modeling to reverse-engineer the computational principles underlying BMI’s behavior to investigate BMI at the computational level (Marr, 2010). By fitting Bayesian models with different prior structures to the network’s predictions, we can test different hypotheses about the decision-making strategies employed by BMI. This analysis allows us to disentangle the effects of representational capacity and regularization on the emergence of heuristic decision-making in BMI.

Methods

Task description

The task used by Binz et al. (2022) is a sequential paired comparison task where a participant or some model (from now more generally referred to as "agent") must learn to identify better options over a sequence of trials. Each of two options is represented by a set of real-valued (randomly sampled) features x_A and x_B , and the agent must learn to weigh these features to make optimal decisions. The environment is structured into distinct "tasks", where each task consists of a sequence of 10 trials. For every task, a latent weight vector w is drawn from a multivariate normal distribution. The binary target C (correct choice) is generated by comparing weighted feature sums ($y_A = w^T x_A + e_A$ vs. $y_B = w^T x_B + e_B$), with independent Gaussian noise e added to each option.

In each new task, the agent starts with high uncertainty about the current weights w and must reduce this uncertainty sequentially over the trials. After each prediction, the agent receives immediate feedback (the true binary label C), allowing it to update its beliefs and refine its decision strategy for the subsequent trials within that task.

To investigate how environmental structure influences the decision making and learning strategies of the agent, three different cue conditions are implemented. In the *no-cue* condition, the agent has no prior information about the feature weights. In the *direction-cue* condition, the agent is informed about the direction (sign) of each feature weight, indicating whether a feature positively or negatively contributes to the

decision. In the *ranking-cue* condition, the agent is provided with the relative importance of features based on their absolute weight magnitude.

By exposing the agent to a distribution of tasks with different weight vectors w , meta-learning enables it to quickly adapt its decision-making strategy based on feedback within each task.

Training of Meta-Learned Inference Models with varying parameters

Binz et al. (2022) meta-train a Gated Recurrent Unit (GRU) (Cho et al., 2014) architecture to learn the specific weight distribution of the task described above. At every trial t within a task, the GRU receives as input the features of both options and the target label. It combines these inputs with its internal hidden state to produce a posterior distribution over the unknown feature weight vector w . This distribution is modeled as a normal distribution $N(w; \mu_t, \Psi_t)$, where the mean and variance are learned by the GRU.

This distribution is then used to make predictions about which option is superior in the next trial $t + 1$. The GRU’s recurrent structure allows it to maintain and update its internal state across trials. During meta-training, Binz et al. expose the GRU to a variety of tasks with different feature weight vectors w . This exposure should enable the GRU to learn a generalizable strategy for quickly adapting to new tasks based on the feedback it receives.

The original work by Binz et al. proposed that resource constraints, implemented through regularization, drive the emergence of human-like heuristics in BMI. They used a fixed network architecture (128 hidden units, 4 features) and varied only the regularization strength (β). However, this leaves open the question of whether heuristics emerge from the regularization penalty itself or from architectural constraints on the network’s capacity to represent complex strategies. To test this, we varied two key architectural parameters that directly control the model’s capacity. First, the size of the network, specifically the number of hidden units in the GRU. The default value in the original paper was set to 128 units. Second, we varied the number of input features for each option in the decision-making task. The default value in the original paper was set to 4 dimensions. We varied the dimensionality of the input data to 2 and 6.

By examining how these capacity constraints affect heuristic discovery across different cue conditions, we can determine whether bounded capacity alone explains the emergence of heuristics, or whether other mechanisms (such as regularization) play a more decisive role. These variations were investigated in both unregularized pre-trained MI and regularized BMI models as well as in the three cue conditions of (1) no given cue, (2) known directions of features, or (3) known rankings of features. Regularized models include the KL-divergence penalty used by Binz et al. ($\beta = 0.01$). We closely followed the training procedure described by Binz et al. Models were meta-trained using the AMSGrad optimizer (Reddi et al., 2019) ($LR = 3 \times 10^{-4}$). The GRU models have

a single hidden layer, and the number of hidden units is varied as described above. Due to computational constraints, training used 10^5 steps instead of 10^6 . Training loss curves were monitored throughout, and all models showed convergence before the end of meta-training. Validation on the original configuration (128 units, 4 features) confirmed this produced similar results (Figures 2 and 3), making additional training unnecessary. Models are tested on 100 tasks each consisting of 10 trials.

Evaluation and comparing architectures

To evaluate the effects of these parameter variations, we assessed model accuracy (percentage of correct decisions) across trials within a task.

We also examined Gini coefficients, as in the work by Binz et al. (2022). In this context, the Gini coefficient measures how concentrated decision-making is on a single feature:

$$G(w) = \frac{\sum_{i=1}^d \sum_{j=1}^d |w_i - w_j|}{2d \sum_{i=1}^d w_i} \quad (1)$$

where d is the number of features and w are the feature weights sampled at each trial from the learned normal distribution of weights. When decision-making relies entirely on a single feature (one weight equals 1, all others equal 0), the Gini coefficient reaches its maximum at $G_{\max} = 1 - \frac{1}{d}$. This implies that as the number of features increases, the upper bound on the Gini coefficient also increases. In the following, it can be assumed that the Gini coefficients close to 0 indicate an equal weighing strategy and Gini coefficients close to $G_{\max}$ a single cue strategy.

Hierarchical Bayesian Model Specification

To more precisely characterize the computational principles underlying BMI’s behavior, we use hierarchical Bayesian modeling. This allows us to represent different hypotheses about decision strategies as specific prior structures and compare which best explains the network’s predictions.

We design the Bayesian model to infer feature weight distributions based on the model’s predictions and true labels. Our model closely follows the ideal observer defined by Binz et al. (2022) (Figure 1). For each paired comparison task, the decision variable is computed as $Y = w^T x$, where w is the vector of feature weights and x is the vector of feature differences between the two options. Here, the distribution defined by Y represents the ideal observer’s belief about the relative quality of the two options based on the weighted features. The probability of selecting option 1 is then given by:

$$P(\text{choose option 1}) = \Phi\left(\frac{Y}{\sqrt{2}\sigma}\right) \quad (2)$$

where Φ is the standard normal cumulative distribution function, and σ represents decision noise.

This model assumes a perfect mapping from stimuli to choices and probability reports. Probability report here refers

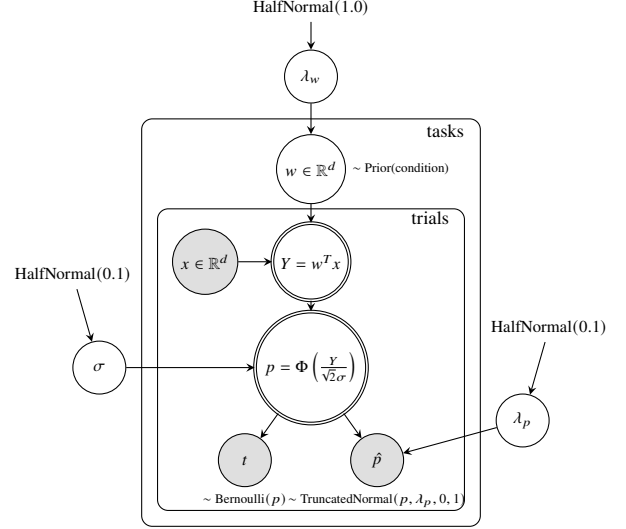


Figure 1: Graphical model representation of the Bayesian model. Shaded nodes represent observed variables, unshaded nodes represent latent variables, and deterministic nodes are depicted with double borders. Plates indicate replication over tasks and trials. Variable definitions: σ = inferred decision noise parameter, λ_w = inferred scale of the weight distribution, λ_p = inferred scale of the prediction noise distribution, $w \in \mathbb{R}^d$ = inferred feature weights vector, $x \in \mathbb{R}^d$ = observed input feature vector, Y = deterministic decision variable, p = deterministic probability of choosing option 1, t = observed ground truth labels, $\hat{p}$ = observed reported probability estimates by the neural network.

to the probability estimates provided by the MI and BMI networks. To account for potential deviations in these reports, we introduce an additional prediction noise. This noise is added to the predicted choice probabilities before generating the reported probabilities. Specifically, the reported probability $\hat{p}$ is modeled as a sample from a truncated normal distribution centered around the predicted choice probability p with a standard deviation λ_p , truncated to the interval $[0, 1]$. This standard deviation captures the variability in the reported probabilities around the true choice probabilities predicted by the model. As shown in Figure 1, the model uses the same unbiased decision variable for both choices and probability estimations. Global parameters include the magnitude of decision noise σ , the overall scale of weights λ_w , and the scale the prediction noise distribution λ_p that determines the precision of probability reports. The prediction noise is meant to allow the Bayesian model to predict responses that are different from the predicted outcome. For each experiment, feature weights are sampled from a prior distribution that depends on the cue condition. For each trial, the model computes a decision variable Y and transforms this into a choice probability as described above, and then generates both binary choices and probability reports.

Model Implementation and Inference We implemented the model using the NumPyro (Phan et al., 2019) probabilistic programming framework. For each of the three cue conditions (ranking, direction, and unrestricted), we generated predictions from both BMI and MI models for 100 tasks with 10 trials each. These predictions served as observed data for the Bayesian model fitting procedure.¹

Prior Structures We implemented five prior structures representing different hypotheses about weight distributions:

- **Unrestricted prior** (*none*): $w_i \sim \mathcal{N}(0, \tau^2)$.
- **Direction prior**: Positive weights only, $w_i \sim \text{HalfNormal}(\tau)$.
- **Equal weighting prior**: All weights share the same value, $w_i = w_j = w$ for all i, j , where $w \sim \mathcal{N}(0, \tau^2)$.
- **Ranking prior**: Weights sorted by absolute magnitude after sampling from $\mathcal{N}(0, \tau^2)$: $|w_1| \leq |w_2| \leq \dots \leq |w_d|$.
- **Single cue prior**: Only the highest-ranked feature has non-zero weight: $w_i = 0$ for $i < d$ and $w_d \sim \mathcal{N}(0, \tau^2)$.

These priors allow us to test specific hypotheses about which decision strategies BMI employs across different cue conditions.

Model comparison For each cue condition (ranking-cue, direction-cue, and no-cue) and model type (MI and BMI), we fitted our Bayesian model with five different prior structures. By comparing the fit of these models, we can determine which prior assumptions best capture the computational strategies employed by the neural networks.

The fitted models were compared by their expected log-predictive density (ELPD) which was estimated by Pareto smoothed importance sampling leave-one-out cross-validation (LOO) of predicting the neural network responses. For this, we relied on the implementation of the ArviZ Python package. ELPD score is an estimate of how well the model generalizes to new data. A higher ELPD score indicates a higher out-of-sample predictive fit and, therefore, a “better” model.

Results

Replication of results of Binz et al. (2022) We successfully replicated the key findings of Binz et al. for the original configuration (4 features, 128 hidden units). Both MI and BMI models showed similar accuracy patterns across the three cue conditions (Figure 2) as reported in the original study: logarithmic accuracy curves in no-cue and ranking-cue conditions, with high initial accuracy in the direction-cue condition. Unregularized models exhibited slightly higher overall accuracy compared to regularized models. Furthermore, the Gini coefficients (Figure 3) of the learned weight distributions closely

¹We performed Bayesian inference using the No-U-Turn Sampler (NUTS). For each model and cue condition, we ran 10 Markov chains with 4,000 warm-up steps and 20,000 sampling steps per chain. Convergence was assessed using the Gelman-Rubin statistic ($\hat{R}$), and chains with divergences exceeding a threshold of 0.5% were excluded from the analysis. Code and data can be found on GitHub.

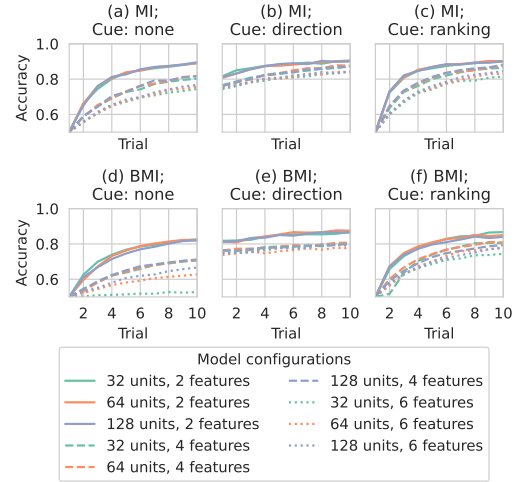


Figure 2: Accuracy across different network sizes (number of hidden units) and input feature counts. The line plots represent the mean accuracy over all tasks for the individual trials within a task. The subplots show results for different regularization (MI or BMI) and cue conditions. 4 features and 128 hidden units corresponds to the original setup by Binz et al. (2022).

matched those reported by Binz et al. In the no-cue condition, Gini coefficients cover the whole range of possible Gini values for both MI and BMI models. When feature directions are known, Gini coefficients tend to be low. When regularization is applied (BMI), Gini coefficients are even closer to zero, indicating a similarity to the equal weighting heuristic. In the ranking-cue condition, Gini coefficients are higher than in the no-cue condition. For BMI models Gini coefficients approach the maximum possible Gini value, indicating a similarity to the single cue heuristic.

Only feature count, not network size influences accuracy Network size (number of hidden units) had minimal impact on accuracy across all conditions (Figure 2). Models with 32 hidden units achieved comparable accuracy to those with 128 units in both MI and BMI configurations.

In contrast, feature count strongly affected accuracy. Higher feature counts led to lower accuracy, particularly in the no-cue condition and especially for BMI models, where networks with 32 hidden units and 6 features achieved only near-chance accuracy.

High initial accuracy in the direction-cue condition Both MI and BMI models consistently demonstrated high accuracy from the first trial in the direction-cue condition across all network sizes and feature counts (Figure 2b, e). In contrast, the no-cue and ranking-cue conditions showed lower initial accuracy that increased over trials (Figure 2a, c, d, f).

No influence of network size on Gini coefficients Network size, defined by the number of hidden units in the GRU, does not seem to have a large influence on the Gini coefficients

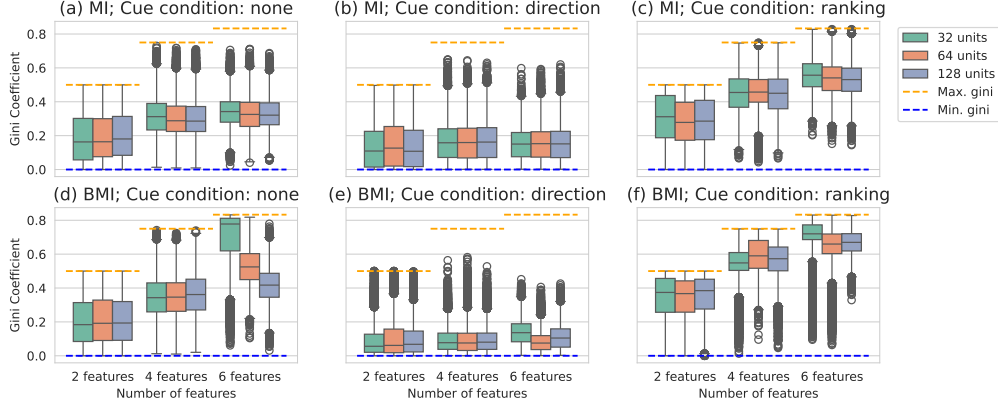


Figure 3: Gini coefficients across different network sizes (number of hidden units) and input feature counts. Each boxplot represents the distribution of Gini coefficients over all tasks and trials. The subplots show results for different regularization (MI or BMI) and cue conditions. The dashed orange line indicates $G_{\max}$ which would be achieved with a single cue heuristic, while the blue dashed line at zero would indicate the equal weighting heuristic. The configuration with 4 features and 128 hidden units corresponds to the original setup by Binz et al. (2022).

of learned weight distributions (Figure 3). Gini coefficients remain relatively stable across different network sizes for both unregularized MI and regularized BMI models in all three cue conditions. Except for 6 features in the no-cue condition, with regularization (BMI) where Gini coefficients tend to be higher for smaller networks, there is no clear trend indicating that larger networks lead to different Gini coefficients. Especially, in the conditions with cues and regularization (direction-cue and ranking-cue), Gini coefficients are consistently low and high, respectively, regardless of network size. Also, in the conditions with cues but without regularization (MI), Gini coefficients are not lower or higher for smaller networks.

No strong influence of feature count on Gini coefficients

While the number of input features has by definition an influence on the possible range of Gini coefficients (see methods), it does not seem to have a large influence on the relative position of Gini coefficients within this range. In Figure 3, Gini coefficients remain relatively stable across different input feature counts for the individual conditions. If anything, the relative position of Gini coefficients seems to shift slightly towards the lower end of the possible Gini range when more features are present. This effect, however, is not very consistent across conditions and model types. Thus, the number of input features does not seem to strongly influence the emergence of heuristic-like weight distributions.

Comparison of model fits with different prior structures

As shown in Figure 4 a and d, the model with the unrestricted prior yields the highest ELPD score in the no-cue condition. This is true for unregularized (MI) and regularized (BMI) networks. Interestingly, for both networks, the model with a ranking prior achieves a comparable fit. In contrast, models with direction, equal weighting, and single cue priors show a poorer fit.

For conditions with environmental cues, the model fit is strongly dependent on network regularization. For MI networks, the model with priors specific to the respective environment achieve optimal fit (Figure 4 b and c). In the environment with a direction-cue, the model with the direction prior results in the highest score, closely followed by the unrestricted prior model. Similarly, when a ranking-cue is introduced, the ranking prior model performs best, again followed closely by the unrestricted prior model.

Models with heuristic priors (equal weighting for direction-cue and single cue for ranking-cue) show intermediate performance, ranking third in their respective conditions. Models implementing priors with constraints incongruent to the environmental cues (for example, equal weighting for a ranking-cue) achieve only low ELPD scores.

When examining conditions with environmental cues but with BMI networks (Figure 4 e and f), a shift can be observed. Models with heuristic priors achieve the best fit, followed closely by models with cue-specific priors (ranking or direction).

Discussion

In this study, we sought to better understand the algorithmic solutions found by bounded meta-learned inference (BMI). After successfully replicating Binz et al. (2022)’s key findings, we approached this task in two distinct ways: we systematically varied the parameters of BMI’s neural network model and the considered tasks, and we used computational cognitive models to reverse engineer the behavior of the BMI models.

The performance of both unregularized MI and regularized BMI models proved generally robust against changes in network size. Networks with as few as 32 hidden units achieved comparable accuracy to those with 128 hidden units across all three cue conditions. This finding suggests that even compar-

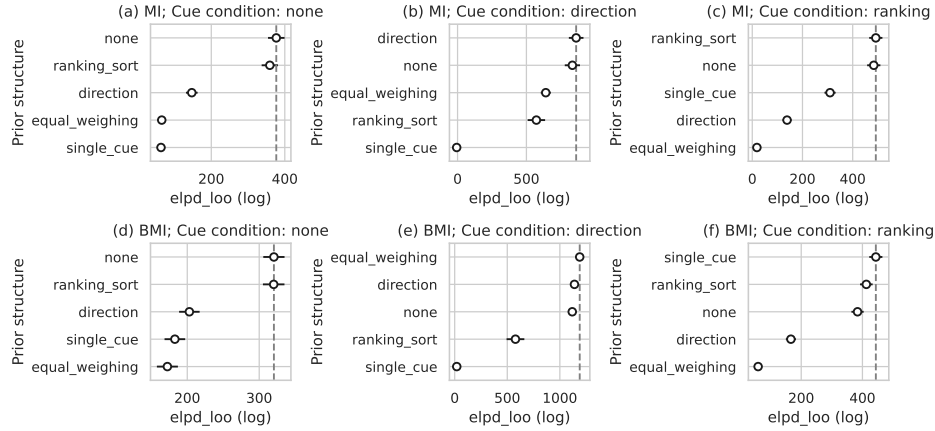


Figure 4: Model comparison between different weight prior structures for each of the three cue and either the neural network with MI or BMI. Shown are the ELPD scores with their standard errors. Larger ELPD values indicate better model fit. The subplots show results for different regularization (MI or BMI) and cue conditions.

atively small networks have sufficient capacity to learn good decision-making strategies in this task. In contrast, performance was sensitive to the input feature count, with models trained on more features exhibiting lower accuracy, particularly in the no-cue condition. This aligns with the intuition that higher dimensionality of the decision problem increases difficulty.

The emergence of heuristic-like decision behavior was largely unaffected by network size or feature count. BMI models consistently discovered the respective heuristics when environmental cues were available, equal weighting for direction-cues and single cue for ranking-cues, regardless of network capacity.

If representational capacity were the main driver of heuristic discovery, smaller networks or those with higher input dimensionality should discover heuristics even without regularization. However, MI models with reduced capacity did not adopt heuristic strategies. Instead, our results suggest that regularization itself, in combination with environmental structure, creates heuristic behavior by introducing inductive biases that favor simpler strategies.

This interpretation is supported by our hierarchical Bayesian modeling results. These demonstrate that heuristic behavior can be modeled by specific prior structures and that such models best explain BMI behavior in the presence of cues. In conditions with environmental cues, BMI models were better described by Bayesian models with heuristic priors than by cue-specific but non-heuristic priors (e.g., direction priors for direction-cue environments).

In contrast, MI models in the same cue conditions were better described by models with cue-specific priors rather than heuristic priors. This difference reveals that regularization alters the computational strategy adopted by the neural networks through meta-learning. Without regularization, networks learn to exploit all available environmental information without simplification. With regularization, networks use

simpler strategies that compress this information into heuristic rules.

Our analysis also revealed an asymmetry in regularization’s inductive biases. The consistently higher initial accuracy in the direction-cue condition and the stronger fit improvement for equal weighting priors suggest a stronger inductive bias toward equal weighting strategies compared to single cue strategies. One possible mechanism is that equal weighting requires only the sign of feature contributions to be encoded and learned, while single-cue strategies require the network to learn to suppress all but the highest-ranked feature. This might be more difficult to encode during meta-training.

The framework proposed by Binz et al. (2022) attributes heuristic discovery to resource constraints that limit the complexity of inference strategies. Our results clarify the mechanism underlying this effect. We find that it is not architectural constraints on capacity, such as network size per se, but rather the regularization penalty, which imposes an inductive bias toward simplicity, that gives rise to heuristics in BMI. Our results suggest heuristics emerge not simply because cognitive systems have limited capacity, but because learning mechanisms favor simple solutions. In humans, preferences for simple strategies could arise from both neural constraints and inductive biases toward efficient decision-making.

Overall, our findings suggest that inductive biases introduced by regularization play a more important role in the emergence of heuristic decision-making than representational capacity alone, at least within the paired comparison setting studied here. Whether this conclusion extends to other decision-making contexts remains an open question. We expect regularization-driven heuristic discovery to generalize when the meta-training environment provides informative cues and the regularization cost of complex strategies is sufficiently high. On a broader view, computational level analyses of intricate algorithmic systems can help better understand their underlying behavior.

Acknowledgments

This work was supported by the German Research Foundation (DFG) under Germany's Excellence Strategy (EXC 3066/1 "The Adaptive Mind" and EXC 3057/1 "Reasonable AI"). C.R. acknowledges support by the European Research Council (ERC; Consolidator Award 'ACTOR', ERC-CoG-101045783). The authors gratefully acknowledge the computing time provided to them on the high-performance computer Lichtenberg II at TU Darmstadt, funded by the German Federal Ministry of Education and Research (BMBF) and the State of Hesse.

References

- Bengio, Y., Bengio, S., & Cloutier, J. (1991). Learning a synaptic learning rule. *IJCNN-91-Seattle International Joint Conference on Neural Networks*, 969 vol.2. <https://doi.org/10.1109/IJCNN.1991.155621>
- Binz, M., Gershman, S. J., Schulz, E., & Endres, D. (2022). Heuristics from bounded meta-learned inference. *Psychological Review*, 129(5), 1042–1077. <https://doi.org/10.1037/rev0000330>
- Cho, K., van Merriënboer, B., Gulcehre, C., Bahdanau, D., Bougares, F., Schwenk, H., & Bengio, Y. (2014, September). Learning Phrase Representations using RNN Encoder-Decoder for Statistical Machine Translation. <https://doi.org/10.48550/arXiv.1406.1078>
- Dawes, R. M., & Corrigan, B. (1974). Linear models in decision making. *Psychological Bulletin*, 81(2), 95–106. <https://doi.org/10.1037/h0037613>
- Einhorn, H. J., & Hogarth, R. M. (1975). Unit weighting schemes for decision making. *Organizational Behavior and Human Performance*, 13(2), 171–192. [https://doi.org/10.1016/0030-5073\(75\)90044-6](https://doi.org/10.1016/0030-5073(75)90044-6)
- Geisler, W. S. (1989). Sequential ideal-observer analysis of visual discriminations. *Psychological review*, 96(2), 267.
- Gigerenzer, G., & Goldstein, D. G. (1996). Reasoning the fast and frugal way: Models of bounded rationality. *Psychological Review*, 103(4), 650–669. <https://doi.org/10.1037/0033-295X.103.4.650>
- Gigerenzer, G., & Todd, P. M. (1999). Fast and frugal heuristics: The adaptive toolbox. In *Simple heuristics that make us smart* (pp. 3–34). Oxford University Press.
- Grau-Moya, J., Delétang, G., Kunesch, M., Genewein, T., Catt, E., Li, K., Ruoss, A., Cundy, C., Veness, J., Wang, J., Hutter, M., Summerfield, C., Legg, S., & Ortega, P. (2022, October). Beyond Bayes-optimality: Meta-learning what you know you don't know. <https://doi.org/10.48550/arXiv.2209.15618>
- Jagadeish, A. K., Coda-Forno, J., Thalmann, M., Schulz, E., & Binz, M. (2024, May). Human-like Category Learning by Injecting Ecological Priors from Large Language Models into Neural Networks. <https://doi.org/10.48550/arXiv.2402.01821>
- Lieder, F., & Griffiths, T. L. (2020). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43, e1. <https://doi.org/10.1017/S0140525X1900061X>
- Marr, D. (2010). *Vision: A computational investigation into the human representation and processing of visual information*. MIT press.
- Martignon, L., & Hoffrage, U. (2002). Fast, frugal, and fit: Simple heuristics for paired comparison. *Theory and Decision*, 52(1), 29–71. <https://doi.org/10.1023/A:1015516217425>
- Phan, D., Pradhan, N., & Jankowiak, M. (2019, December). Composable Effects for Flexible and Accelerated Probabilistic Programming in NumPyro. <https://doi.org/10.48550/arXiv.1912.11554>
- Reddi, S. J., Kale, S., & Kumar, S. (2019, April). On the Convergence of Adam and Beyond. <https://doi.org/10.48550/arXiv.1904.09237>
- Thrun, S., & Pratt, L. (1998). Learning to Learn: Introduction and Overview. In S. Thrun & L. Pratt (Eds.), *Learning to Learn* (pp. 3–17). Springer US. https://doi.org/10.1007/978-1-4615-5529-2_1
- Tversky, A., & Kahneman, D. (1974). Judgment under Uncertainty: Heuristics and Biases. *Science*, 185(4157), 1124–1131. <https://doi.org/10.1126/science.185.4157.1124>