

Dynamic Strategy Selection in Active Function Learning

Nayan Saxena

University of Toronto

Abstract

We test whether people flexibly shift their sampling strategy for learning a functional relationship, based on the strategy’s perceived effectiveness. While general-purpose heuristics such as gathering information evenly across the environment may often approximate optimal sampling policies when opportunities for learning are sparse, these strategies may systematically fail when much of the environment is uninformative. Across several different classes of arbitrary smooth functions, participants ($N = 89$) made sampling choices that were initially consistent with a simple heuristic, but shifted their sampling strategy when this heuristic failed to be informative. People were subsequently more accurate at approximating the true function for smoother functions that can be more easily predicted by sampling evenly, $t(645.3) = 6.803$, $p < .001$; nevertheless, while individuals vary considerably, aggregate judgments across all functions are very accurate, suggesting that people show inductive biases in active learning, but effectively learn when to adjust their policies.

Keywords: Psychology ; Learning ; Reasoning ; Computational Modeling

1. Introduction

Given our limited cognitive resources, a significant portion of the human psyche involves dealing with numerous streams of incomplete information and drawing a sparse number of samples for active decision-making and extrapolation. This process of selectively gathering useful data is crucial, especially for limited-capacity organisms like humans who have a sensorimotor system that senses more data than can be fully processed. Furthermore, in our everyday life, there are several instances where cause-effect functional relationships can only be learned through strategic selection of the points contributing to the learning process. A natural question to then ask is, how do humans choose these samples? and how does human sampling compare to other baselines? The central goal of this paper is to precisely answer this question, through a series of simulation studies and real world experimentation.

Earlier work on function learning has shown that given sufficient evidence humans are capable of learning a wide range of functional relationships (Bott and Heit, 2004; Lucas et al., 2012; Wilson et al., 2015), with inductive biases heavily favouring linear relationships (Brehmer, 1974). There is also research indicating that sampling from regions of high uncertainty is favourable for the efficient learning of linear functions (Jones et al., 2018) while being relatively inaccurate for non-linear functional forms (Kalish et al., 2004; DeLosh et al., 1997). Furthermore, it has been shown that humans tend to employ rather sub-optimal resource-rational heuristics (Gigerenzer, 2008; Lieder and Griffiths, 2020), trading off sampling optimality with the cognitive constraints of employing rather sophisticated sampling heuristics (Gershman et al., 2015).

We build upon earlier research showing that people are able to learn several non-linear functional relationships by employing a relatively simple heuristic strategy, sampling the minimum and maximum x values and evenly spaced points in between, that requires little cognitive effort while providing a comparable outcome on non-linear functions (Gelpi et al., 2021). Through the inclusion of more varied functional forms, we test whether people flexibly shift their sampling strategies, based on the strategy’s perceived effectiveness. While general-purpose heuristics such as gathering information evenly across the environment may

often approximate optimal sampling policies when opportunities for learning are sparse, these strategies may systematically fail when much of the environment is uninformative.

In the following sections, we first present a Gaussian process-based framework for representing the task of function learning, describing how a learner can use their existing knowledge about a functional relationship to interpolate or extrapolate to unknown function values. Then Gaussian mixture modelling clustering is introduced, which is used to group individuals based on their common sampling behavior. Finally, simple candidate strategies are introduced that an active learner might use to select samples and an experimental task to assess humans' sampling behaviour and predictions, and compare these to simulated learning under different sampling policies.

2. Preliminaries

2.1 Gaussian Process Model

We use a Gaussian process (GP) model (Griffiths et al., 2008; Lucas et al., 2015) to provide a general framework for understanding both rule-based and similarity-based function learning. This model uses samples $\mathbf{x}_n = (x_1 \dots x_n)$ to further approximate a learned function f by inducing a Gaussian distribution on the observed $y_i = f(x_i)$ values based on sampled x_i values (Rasmussen and Williams, 2006). For known function outputs f and new unknown points f_* the joint probability distribution is then defined as

$$\begin{pmatrix} f \\ f_* \end{pmatrix} = \mathcal{N}\left(\begin{bmatrix} \mu \\ \mu_* \end{bmatrix}, \begin{bmatrix} K & K_* \\ K_*^T & K_{**} \end{bmatrix}\right) \quad (1)$$

In the above equation, we have $K = k(x, x)$, $K_* = k(x, x_*)$ and $K_{**} = k(x_*, x_*)$, where k denotes the generalized class of Matérn kernel given by:

$$k(x_i, x_j) = \frac{1}{\Gamma(\nu)2^{\nu-1}} \left(\frac{\sqrt{2\nu}}{\ell} d(x_i, x_j)\right)^\nu K_\nu\left(\frac{\sqrt{2\nu}}{\ell} d(x_i, x_j)\right) \quad (2)$$

where d denotes Euclidean distance, K_ν is a modified Bessel function, and Γ is the standard Gamma function. Throughout this paper we use this model with smoothness $\nu = 3$ and varying length-scales for functions with relatively narrow edges and steeper slopes $\ell = 0.1$, as compared to functions which are more continuous $\ell = 1$ with relatively shallow slopes.

2.2 Gaussian Mixture Modelling

Furthermore, to derive sampling policies from experimental data and group individuals based on their common sampling behavior we employ Gaussian mixture modelling clustering (McLachlan et al., 2019; Rasmussen and Williams, 2006; Rasmussen, 1999), where the clusters are fitted through the expectation maximisation (EM) algorithm described below.

Assuming that $\mathbf{x}_n = (x_1 \dots x_n)$ are the samples with the desired number of clusters denoted by k , the cluster means μ_k and mixing coefficients π_k are initially randomly initialised. The algorithm then simply iterates until convergence by alternating between three steps:

1. Compute posterior probabilities $r_k^n = p(z_n = k | \mathbf{x}_n)$, explaining the likelihood of a point originating from cluster k . This is given by $r_k^n = p(z_n = k | \mathbf{x}_n) = \frac{\hat{\pi}_k \mathcal{N}(\mathbf{x}_n | \hat{\mu}_k, \mathbb{I})}{\sum_{i=1}^k \hat{\pi}_i \mathcal{N}(\mathbf{x}_n | \hat{\mu}_i, \mathbb{I})}$.
2. Update the parameters $\hat{\mu}_k = \frac{1}{N_k} \sum_{n=1}^N r_k^{(n)} \mathbf{x}_n$, and $\hat{\pi}_k = \frac{1}{N} \sum_{n=1}^N r_k^{(n)}$ in order to maximise the probability of a point belonging to its cluster. Here, N_k and N denote size of cluster and sample size respectively.
3. Finally the log-likelihood is computed to check for convergence, $\sum_{n=1}^N \log(\sum_{k=1}^K \hat{\pi}_k \mathcal{N}(\mathbf{x}_n | \hat{\mu}_k, \mathbb{I}))$.

3. Sampling Strategies

As we get further from previously-observed points, our uncertainty increases, which could lead to more heterogeneous or inaccurate extrapolation if a function’s minimum and maximum values are not known. Therefore, for a learner with limited opportunities to sample, we expect that the most effective and informative strategies will include sampling the extrema. Throughout this paper we focus on assessing the informativeness of five policies, in increasing order of complexity, uniform random sampling, an equidistant sampling policy that selects

the minimum and maximum values interpolating equally between them, a re-sampling policy that basically samples equidistant points and returns to sample more in certain regions, binary partition sampling and finally uncertainty-based sampling.

Random sampling One simple strategy to learn about a function is to sample randomly from the possible domain of values. While random sampling is straightforward, decision-makers may display substantial sub-optimality in their sampling choices if they happen to sample multiple points in close proximity that are unlikely to be maximally informative. Given human inductive biases towards linearity, this could result in difficulty extrapolating or interpolating non-linear functions where no samples have been drawn. We represent this policy as drawing samples from a distribution where each sample $x_i \sim \text{Uniform}(0, 1)$.

Equidistant sampling As mentioned previously, sampling the minimum and maximum feasible values within the confines of the problem to be solved minimizes the need for extrapolation, and partitioning the remaining space among the remaining samples likewise balances interpolation between the remaining points, making it suitable for relatively accurate prediction across many commonly-encountered functions. While somewhat inflexible, this may represent a highly tractable, efficient heuristic for sampling within a limited domain. We represent this sampling approach as taking the N samples to be drawn and allotting them in such a way that the points sampled roughly reflect $N - 1$ equal partitions of the space to be sampled:

$$x_i \sim \text{Beta}\left(1 + \frac{i\rho}{N-1}, 1 + \frac{(N-1-i)\rho}{N-1}\right) \quad (i \in \{0 \dots N-1\})$$

To reflect a moderate preference for sampling from the peaks of the beta distributions while allowing for some samples to be drawn from nearby values, we chose a free parameter value of $\rho = 10$ that denotes the size of the peaks of the sampling distribution.

Resampling with Equidistant Sampling The two above algorithms reflect sampling procedures that do not dynamically update based on new information. In order to choose

the best samples, however, it may be effective to adapt one’s strategy to account for already sampled points. The simplest approach is to augment the equidistant sampling policy to incorporate re-sampling of previously sampled regions. This reflects the preference of people potentially returning and sampling again in regions where they observe more activity. To achieve this goal, this sampling procedure employs a simple equidistant procedure for the first $N - k$ samples, while keeping track of the maximum observed value of the function denoted by m . Then the remaining k samples are drawn from a randomly generated interval centered around m . Throughout this paper, in-line with our experimental setup, we restrict k to be a randomly generated integer between 1 and 3.

Binary Partition Sampling Another sampling strategy that incorporates subsequent convergence to the function maximum, involves the iterative sampling of points from either ends of the sampling region. Predicated on the previously drawn sample, decision makers might employ this strategy to maximise information gain by sampling in regions of highest uncertainty, which here is the region on the opposite end from the previous sample. This also results in the formation of a narrower sampling interval for subsequent sampling. More formally, after every two samples x_1 and x_2 drawn from either ends of the sampling region, the region gets partitioned into a smaller interval (x_1, x_2) which becomes the new sampling region for the next iteration. Ultimately, the process converges with the final point partitioning the original initial sampling region into almost two equal halves.

Adaptive sampling Motivated by the fact that people preferentially draw samples in regions that resolve current uncertainty, even when this does not maximize information gain (Markant and Gureckis, 2014), we propose another sampling algorithm motivated by a myopically ideal strategy that adapts to new information by identifying the point of highest uncertainty. This sampling procedure uses a GP model that greedily chooses points from the domain-space by iteratively fitting the model on the previously sampled points, and choosing the next sample as the point on the posterior distribution of GP functions with the highest variance.

To test the effectiveness of these algorithms and compare their performance to human sampling strategies as well as the learned functions, we designed an experimental task in which participants must select a small number of samples to try to optimize their performance in a prediction task.

4. Experimental Design

Recruitment and Procedure 89 adult participants were recruited through Prolific and paid £1.25 for completing an online learning task presented via a web-based program.

Familiarization and Exposure In the task, participants were told that they would be playing the role of an oceanographer. The scientist’s job was to learn about a number of possible seas, each being researched for the map of a section of their seafloor. Ultimately, the goal of the participant was to learn the overall relationship between the depth of a seafloor at various given on the seafloor section.

Before participating in the experimental trials, participants were familiarized with a warm-up trial where they were presented with the task of learning about the seafloor of a random sea. They were presented first with an empty canvas and the choice the to sample two points on the seafloor to learn about the depth at those respective points. The x axis of the canvas corresponded to the seafloor and the y axis corresponded to the depth of the seafloor. Participants were told that they could learn about the depth of the sea at a point by clicking on a point. After confirming their choices, participants were asked to draw their predictions of what the seafloor looks like, with the initially chosen points visible on the canvas.

Experimental Task After completing the familiarization trial, participants completed the experimental task. Participants learned about nine total seas, with eight of the seas divided into groups of four each. The two groups corresponded to seas with rough and smooth seafloors respectively, and were counterbalanced with certain individuals randomly being shown the rough group of seas first and vice-versa. Each sea had a predefined relationship

defining the seafloor depth $g(x)$, at any given point x described by the general equation,

$$g(x) = \frac{f_{\text{Beta}}(x|\alpha, \beta) \cdot 2}{f_{\text{Beta}}(\frac{\alpha-1}{\alpha+\beta-2}|\alpha, \beta) \cdot 3} + c \quad (3)$$

Where $c \in \mathbb{R}$ and f_{Beta} is the probability density function for the Beta distribution with shape and scale parameters α and β respectively. An overview of the seas and the respective relationships between the seafloors and depth are provided in Table 1 and the experimental setup is described in Figure 1.

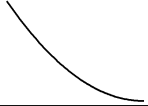
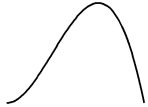
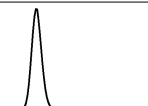
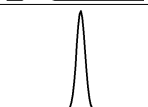
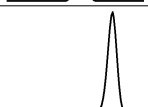
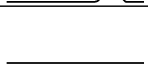
Seafloor Type	Function Name	α	β	c	Graph
Smooth	Smooth 1	1	3	0.083	
	Smooth 2	3	2	0.166	
	Smooth 3	2.5	1	0.243	
	Smooth 4	1.5	2.5	0.333	
Rough	Peaked 1	131	21	0.083	
	Peaked 2	31	111	0.166	
	Peaked 3	141	21	0.243	
	Peaked 4	131	40	0.333	
Flat	Constant	-	-	0.5	

Table 1: Overview of function names, their corresponding parameters and graphs involved in the experiment. The parameters for each function are used in conjunction with Equation 3.

You are **testing** to learn about the depth of this section of the *Hellas Sea*.

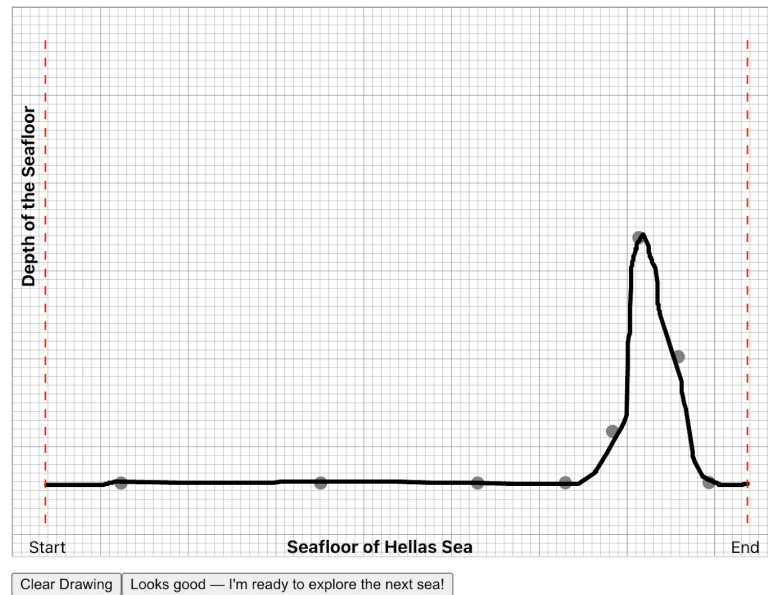
Last experiment!



(a)

Draw the seafloor for Hellas Sea.

To make a drawing, hold down your mouse button and draw from left to right. Your drawing cannot overlap or intersect, and your drawing will be erased if you try to do so. You can stop and start your drawing as you draw, as long as you continue drawing left to right. If you want to start your drawing over, press the clear drawing button.



(b)

Figure 1: Each participant: (a) chooses 8 points to learn about the seafloor depth, and (b) draws their prediction of what the seafloor looks like for each sea.

5. Results

5.1 Sampling Task

Simulation Data To analyze the effectiveness of the identified sampling strategies and establish comparison baselines, we generated synthetic data for 89 participants under each sampling strategy: random, equidistant, resample+equidistant, binary and adaptive sampling, and fitted Gaussian process models for each function using the `sklearn` library in Python. The pooled and average mean squared error (MSE) for the GP fits were then computed, alongside the average minimum absolute distance (Δ) of samples from the horizontal location of the function maxima. A summary of these metrics for each sampling strategy is presented in Table 2, with the important values highlighted for clarity.

Sampling randomly (Row 1) resulted in reasonably good approximations of the true function at the centre but demonstrated wider deviations at the minimum and maximum values. Furthermore, as indicated by the distance from the maximum a large number of the simulated participants failed to effectively sample from the extrema of the function for certain smooth functions, whereas for peaked functions this policy was relatively successful at capturing the extrema. Overall, it was observed that sampling at random works quite well but the pooled and average MSE of simulations of this policy demonstrated a comparative disadvantage at approximating the functions when compared to equidistant sampling across most functions. It should be noted that this policy outperformed the more sophisticated adaptive sampling policy across most functions, and was almost as effective as the equidistant policy at drawing samples closer to the function extrema.

On the contrary, samples drawn using the equidistant sampling policy (Row 2) performed quite well at approximating the true functions based on the pooled and average MSE. Specifically, for smooth functions the equidistant policy generally consistently outperformed other policies like the adaptive sampling policy, binary partition sampling and the resampling + equidistant heuristic. For the peaked functions, the equidistant policy was noted to be the closest in efficacy to the random sampling heuristic while capturing relatively more

Function	Sampling strategy	Pooled MSE	Average MSE	Average Min. Δ from maximum (peak)
Smooth 1	Random	$1.24 \cdot 10^{-5}$	$5.3 \cdot 10^{-5}$	0.1074
	Equidistant	$7.28 \cdot 10^{-8}$	$5.25 \cdot 10^{-7}$	0.0358
	Resample+Equidistant	$4.78 \cdot 10^{-5}$	$6.53 \cdot 10^{-5}$	0.0281
	Binary Partition	$3.56 \cdot 10^{-5}$	0.004	0.0344
	Adaptive	$1.95 \cdot 10^{-7}$	$1.2 \cdot 10^{-6}$	0.0104
	Human	$1.06 \cdot 10^{-6}$	$7.76 \cdot 10^{-5}$	0.0437
Smooth 2	Random	0.000646	0.00208	0.0567
	Equidistant	$8.73 \cdot 10^{-6}$	$3.44 \cdot 10^{-5}$	0.0501
	Resample+Equidistant	0.0065	0.01002	0.0689
	Binary Partition	0.0016	0.0023	0.0576
	Adaptive	$1.45 \cdot 10^{-5}$	0.00014	0.0698
	Human	$4.13 \cdot 10^{-5}$	0.0008	0.0481
Smooth 3	Random	$8.48 \cdot 10^{-6}$	$5.70 \cdot 10^{-5}$	0.1127
	Equidistant	$1.17 \cdot 10^{-7}$	$5.14 \cdot 10^{-7}$	0.0370
	Resample+Equidistant	0.00026	0.0007	0.3105
	Binary Partition	$7.63 \cdot 10^{-5}$	$4.66 \cdot 10^{-5}$	0.1649
	Adaptive	$3.035 \cdot 10^{-7}$	$8.25 \cdot 10^{-7}$	0.0007
	Human	$8.14 \cdot 10^{-7}$	$3.52 \cdot 10^{-5}$	0.0368
Smooth 4	Random	0.003	0.042	0.0659
	Equidistant	0.055	0.179	0.0541
	Resample+Equidistant	0.0083	0.0678	0.0277
	Binary Partition	0.020	0.106	0.0349
	Adaptive	0.215	0.347	0.0451
	Human	0.186	0.332	0.0422
Peaked 1	Random	0.008	0.040	0.0665
	Equidistant	0.007	0.025	0.0465
	Resample+Equidistant	0.014	0.039	0.1995
	Binary Partition	0.019	0.043	0.0563
	Adaptive	0.013	0.025	0.1270
	Human	0.006	0.018	0.0380
Peaked 2	Random	0.005	0.048	0.0649
	Equidistant	0.006	0.032	0.0484
	Resample+Equidistant	0.003	0.016	0.0283
	Binary Partition	0.009	0.027	0.0356
	Adaptive	0.013	0.039	0.1228
	Human	0.005	0.016	0.0323
Peaked 3	Random	0.006	0.036	0.0563
	Equidistant	0.007	0.028	0.0542
	Resample+Equidistant	0.005	0.063	0.0400
	Binary Partition	0.006	0.037	0.0407
	Adaptive	0.010	0.023	0.1142
	Human	0.005	0.016	0.0341
Peaked 4	Random	0.008	0.051	0.0521
	Equidistant	0.007	0.028	0.0471
	Resample+Equidistant	0.009	0.075	0.1168
	Binary Partition	0.028	0.078	0.0698
	Adaptive	0.016	0.025	0.1061
	Human	0.006	0.020	0.0455

Table 2: Pooled MSE, average individual MSE, and average minimum absolute distance (Δ) of samples from the horizontal location of the function maxima across sampling strategies, and human samples. Minimum values for each respective metric are highlighted, with secondary values also highlighted in case of the minimum originating from human samples.

information about the function extrema. Furthermore, for peaked functions, there was also close proximity in terms of performance metrics between equidistant and binary partition or resampling+equidistant policies.

Alternatively, the resampling + equidistant policy (Row 3) performed relatively poorly on the smoother function family as compared to peaked functions. This policy, however, on average, was quite successful at sampling closer to the function extrema across both function families which can be attributed to the resampling of points near regions of highest activity. In terms of the pooled MSE for peaked functions, this policy outperformed the binary partition policy alongside the adaptive sampling policy and was the best performing policy for a several peaked functions. It should be noted that there is a very narrow margin of difference in performance metrics when this policy is compared to adaptive sampling, specifically for peaked functions.

Coming to the binary partition heuristic strategy (Row 4), it was observed that the policy performed best for smoother functions and overall performed relatively poorer as compared to other sampling strategies in terms of the MSE. Furthermore, as indicated by the distance from maxima, it was observed that the policy was comparable to the resample+equidistant heuristic in its tendency to sample relatively closer to the function extrema. The lower MSE for this policy is eclipsed by its capability to sample closer to the maxima which may result in better overall information gain about the function.

Finally, similar to the equidistant heuristic strategy, the adaptive sampling algorithm (Row 5) frequently sampled from regions in the neighbourhood of zero and one. Overall, for smooth functions the adaptive sampling strategy yielded comparable, though slightly higher, MSE values when compared to the equidistant strategy Whereas for peaked functions, the adaptive policy demonstrated significantly inferior MSE values. This policy, for simpler functions, was on average the best at sampling points in close proximity from the function maxima. Although, for more complex functions there was a considerable drop in performance. Given the cognitive costs of adaptive sampling, even when calculated approximately rather than through the representation of the full posterior distribution of candidate functions, we

do not expect people to utilize this strategy to draw samples. Nevertheless, specifically for smoother functions, the similarity of this adaptive strategy to a simple heuristic strategy of sampling evenly spaced points in addition to the minimum and maximum values, suggests that if people use such a strategy, they will select comparably informative points to learn about, improving their subsequent predictions of the true function’s value.

Experimental Data Gaussian process models were also fitted on samples for each of the 89 human participants across all functions. These models were then compared against previously established baselines produced by fitting GP models on generated data from different sampling strategies.

As observed in Table 2 (Row 5), for the smooth functions, the models fitted on human samples in aggregate were able to predict the true function comparable to the equidistant or adaptive sampling strategies in terms of pooled and average approximation error (MSE) as well as distance from maxima. Moreover, for peaked functions the pooled MSE for human samples is closest overall to that of equidistant and resample + equidistant policies. Whereas, average MSE for human samples was most comparable with the resample + equidistant, adaptive and equidistant policies. Overall, for peaked functions, when taking into consideration the average minimum distance from function maxima, human samples were closest to resample + equidistant heuristic followed by the simple equidistant heuristic.

Individual Differences To further distinguish individuals based on their sampling policies, we used Gaussian mixture model clustering which was fitted on data from the 89 participants using the `mclust` library in R programming language. The clustering algorithm, upon data wrangling, partitioned the individuals into three main clusters, each representing a sampling heuristic with a fourth miscellaneous cluster representing around 30% of the participants.

Within each cluster, we then examined the temporal nature of each sampled point to narrow down the sampling order, and in-turn the sampling policy. As demonstrated in Figure 2, each cluster showed behaviour close to three previously outlined sampling policies:

(i) Equidistant sampling, (ii) Binary sampling, and (iii) Resample + Equidistant sampling.

The three sampling strategies were noticed across all functions with the exception of functions Smooth 4 and Smooth 1, as shown in Figure 2. Furthermore, we observed that the equidistant strategy was the most consistently used strategy with around 30% of the participants using it more than 75% of the time. The binary sampling strategy was also consistently employed by a small fraction of participants.

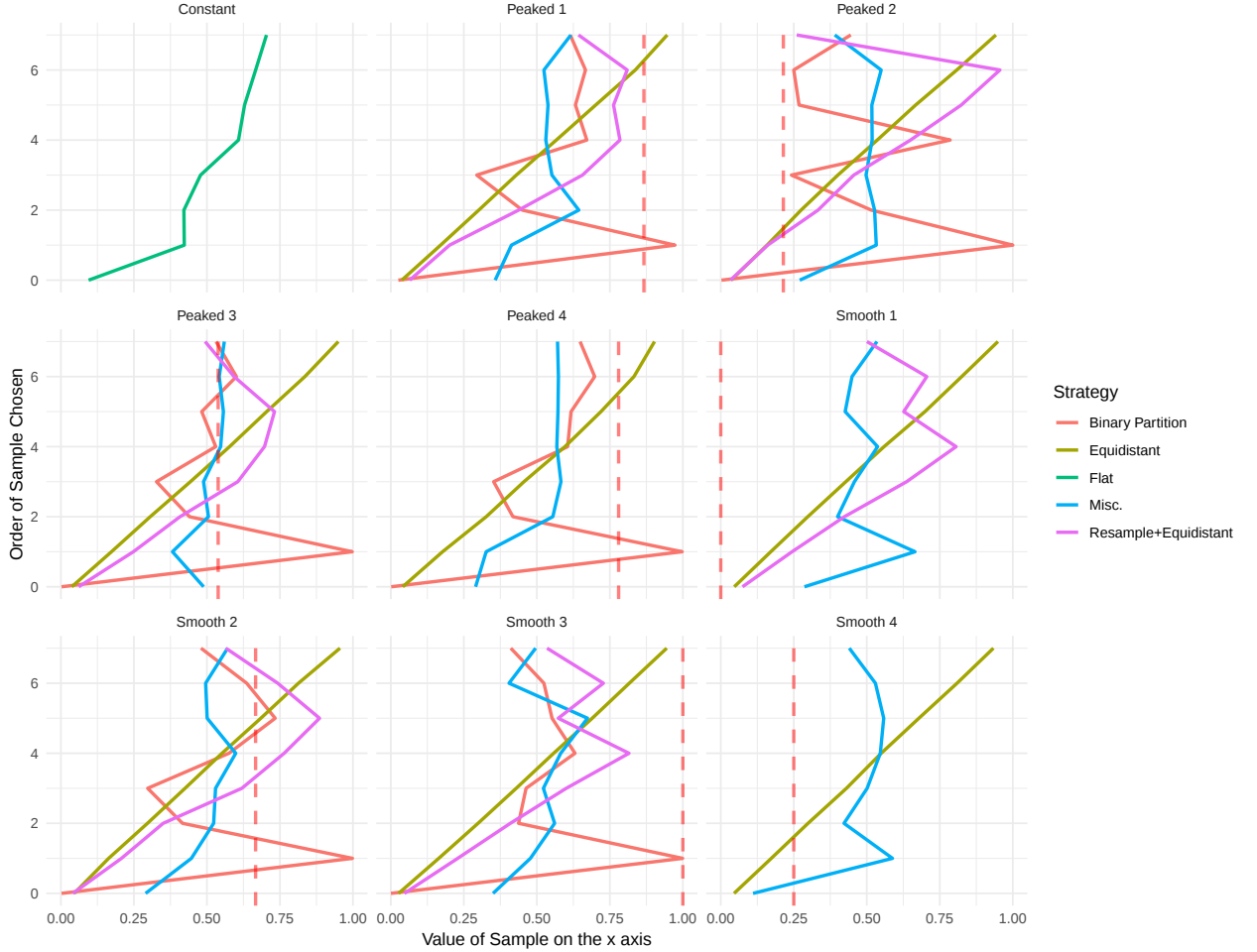


Figure 2: Average sampling patterns observed within each cluster, with the graph showing mean value of the sample drawn with respect to sample order. The dashed red line indicates horizontal location of the true function’s maxima.

The average error for each function and sampling policy was also recorded with certain policies seemingly having lower error rates. Overall, it was observed that binary sampling and resample + equidistant policies have the lowest margin of error as compared to other

policies. To further narrow down on factors that have an influence on the average error we fitted a linear mixed-effects model using the `lme4` package within R programming language. The fitted model, with the response variable being average error, had categorical predictors indicating the counter balancing order, peaked/smooth function and a random effects variable to account for different functions. Reinforcing the fact different functional forms are the key drivers of sampling performance, it was found that across all the variables only the function type was a significant predictor of sampling performance measured by average error ($p < 0.001$). A visual overview of the clustered sampling policies is provided in Figure 3.

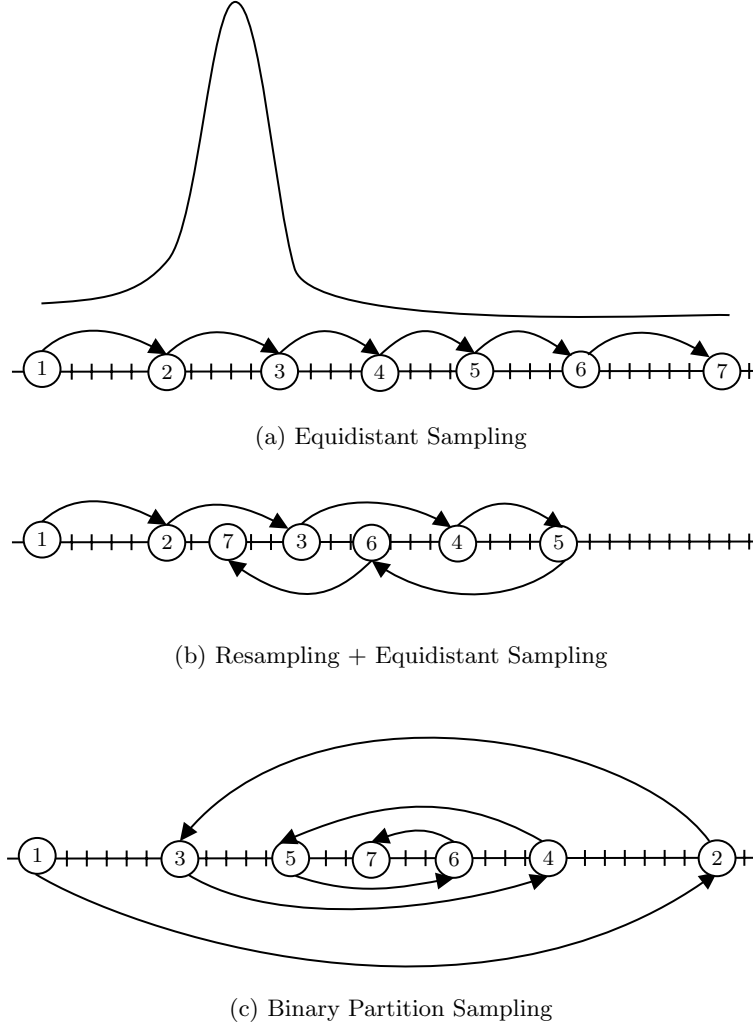


Figure 3: Overview of human sampling policies derived from clustering experimental data.

Finally, the use of sampling policy also seemed to depend on the random counterbalance, which randomly first presents the peaked group of functions to some participants and smooth to others. Overall, the equidistant policy was employed less by participants who encountered peaked functions first. On the contrary, participants who encountered smooth functions first gradually increased their equidistant sampling usage as they progressed throughout the experiment.

5.2 Prediction Task

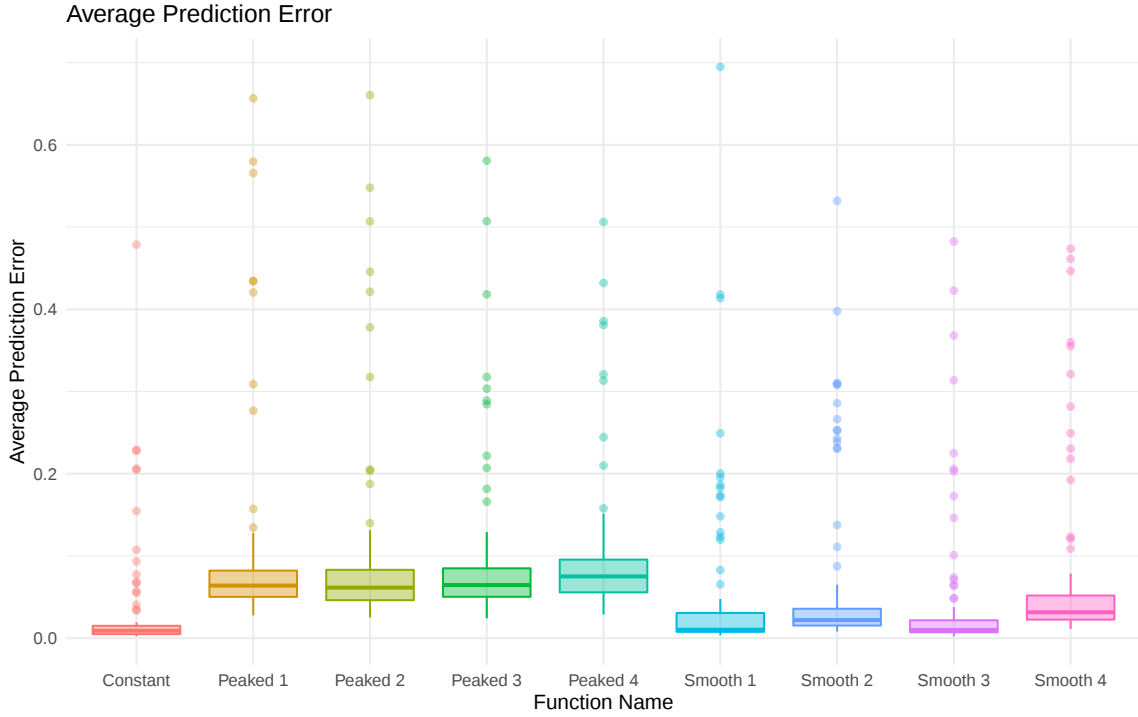


Figure 4: Average prediction error by function type

Participants across all functions performed relatively well on the prediction phase of the experiment. As observed in Figure 5, on average participants converged to a relatively accurate estimation of the true underlying function, especially for smooth functions. Taking a closer look at the average prediction errors, shown in Figure 4, confirmed these findings with peaked functions having higher prediction errors and margins of error, and smooth functions having both narrower margins and lower errors. It was further observed that

participants showed higher error rates for Peaked 4 and Smooth 4 functions, which are the last functions people were presented for both function families.

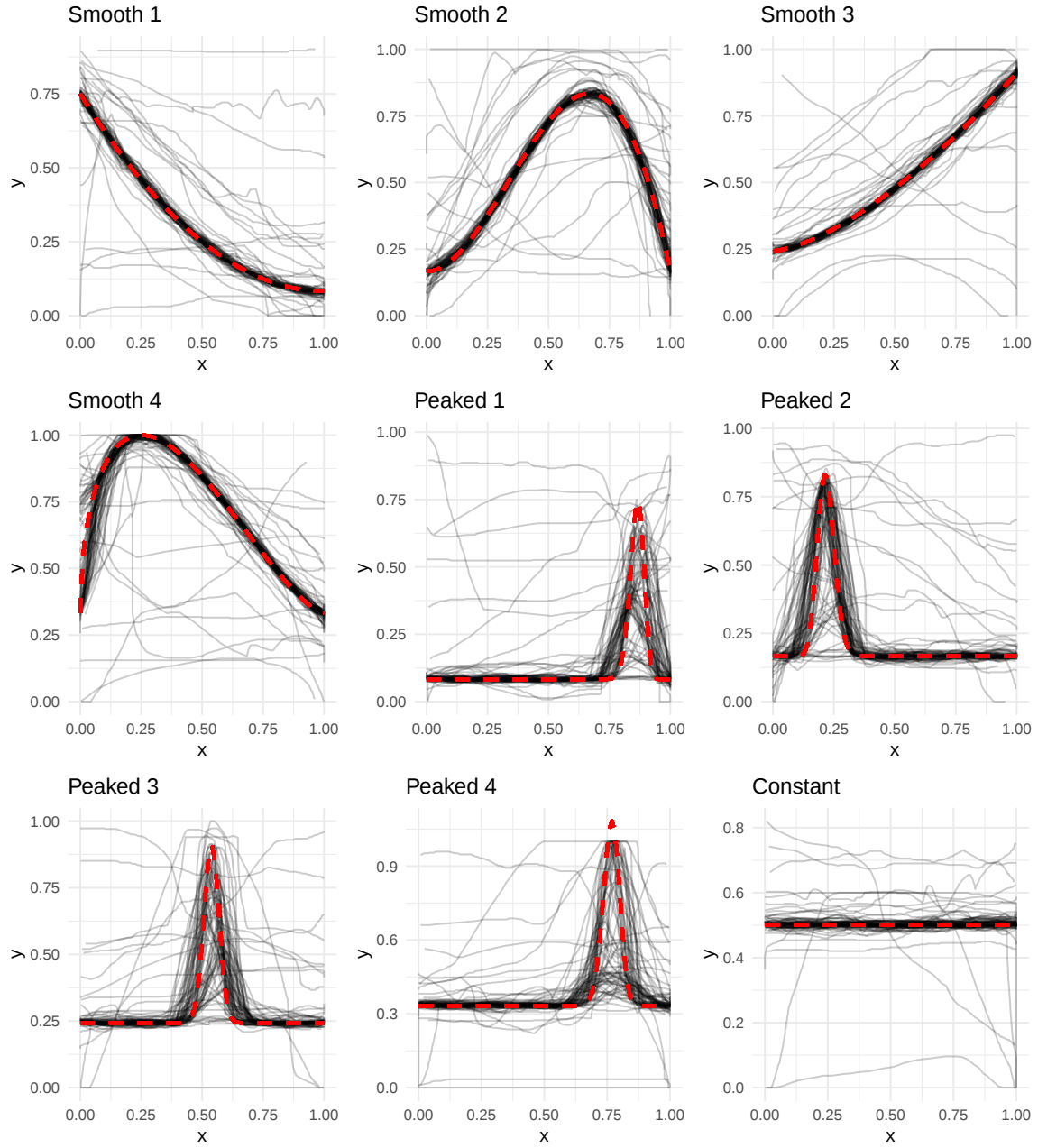


Figure 5: Predicted function fits for each function within the experiment. The red dotted line indicates the true function.

Participants making predictions for smooth functions, at an individual level showed better overall predictive performance for functions closely resembling positive and negative linear relationships such as Smooth 1 and Smooth 3 functions respectively. Performance was slightly worse for functions with more non-linearity in their structure such as Smooth 2 and Smooth 4 with the extrema (maxima) not lying at the edges of the sampling domain. In general, most individuals were successfully able to capture the location of the function maxima and also estimate certain high-level features of the function in cases where the predictions are rather inaccurate.

For peaked functions, participants on average successfully captured location of the peak but at times failed to completely capture the functions shape, often incorrectly predicting the function to be more smoother with relatively shallower slopes. Several individuals on the contrary more accurately identified the sharp sloped nature of the function in their predictions.

6. General Discussion

In this paper, we tested whether people flexibly shift their sampling strategy for learning a functional relationship, based on the strategy’s perceived effectiveness. While general-purpose heuristics such as gathering information evenly across the environment may often approximate optimal sampling policies when opportunities for learning are sparse, these strategies may systematically fail when much of the environment is uninformative. Across several different classes of arbitrary smooth functions, participants made sampling choices that were initially consistent with a simple heuristic, but shifted their sampling strategy when this heuristic failed to be informative. Overall, people were subsequently more accurate at approximating the true function for smoother functions that can be more easily predicted by sampling evenly, nevertheless, while individuals vary considerably, aggregate judgments across all functions were very accurate, suggesting that people show inductive biases in active learning, but effectively learn when to adjust their policies. Respondents showed vast variability in the informativeness of their samples and the accuracy of their predictions;

however, aggregated results showed that a recognizable parametrization of the true function was learned, perhaps the product of a “wisdom of crowds” effect (Steyvers et al., 2009) averaging out individual errors.

Through our experiments, as shown in Figure 3, using model-based clustering we discovered that participants employed three main sampling strategies: (i) Binary Partition Sampling, (ii) Equidistant Sampling, (iii) Resample+ Equidistant Sampling, to learn functional relationships. Consistent with previous work showing adherence of participants to a simpler linear regression policy rather than a generalized GP policy when learning in linear domains (Jones et al., 2018), around 30% of participants employed the equidistant policy, sampling the minimum and maximum x values and evenly spaced points in between, consistently across both smooth and peaked functions, while a small subset of participants consistently employed the binary partition strategy. Overall, the equidistant policy was also employed less by participants who encountered peaked functions first. On the contrary, participants who encountered smooth functions first gradually increased their equidistant sampling usage as they progressed throughout the experiment. As outlined in Table 2, across different functional forms, different sampling heuristics performed closest to human sampling, with the equidistant sampling emerging as the most successful strategy across most functions. The relatively low cognitive cost of the equidistant policy, paired with its efficacy at learning important functional features indicates rational meta reasoning being employed during decision making (Lieder and Griffiths, 2017). Furthermore, specifically for peaked functions, other sampling strategies, also found within our experimental data, like binary partition sampling and resample+equidistant sampling were the best at capturing the region closest to the function maxima.

Our findings also indicate scope for further research, specifically work examining the sampling behaviours of roughly 30% of our participants who, based on clustering, employed a presently unknown sampling heuristic. Furthermore, it is also important to investigate the potential correlation between the predictive performance and the sampling heuristic employed by individuals. Further exploration can also be carried out by delving deeper

into eliminating the human bias of employing a relatively inflexible sampling strategy, while requiring minimum cognitive effort, like the equidistant policy. This can be carried out through the inclusion of adversarial functions into the experimental setup where most important functional features escape the samples generated, as demonstrated in Figure 6.

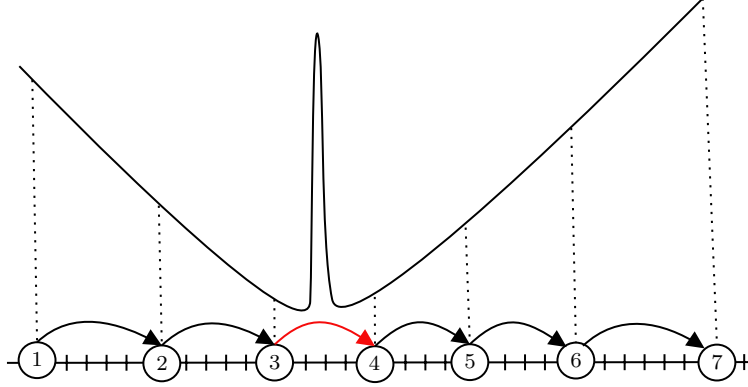


Figure 6: Example of an adversarial function, where the function peak evades samples generated from an equidistant sampling procedure as denoted by the red arrow.

While the true limits of the identified sampling policies are currently unexplored, in the context of active function learning, we have nonetheless pinned down key differences in individual sampling behaviour and explored their varying performance in estimating functions belonging to different families. Another conclusion we draw is that predicated on the complexity of the functional forms, there is some evidence supporting the fact that humans adapt their policies, while maintaining a low cognitive overhead, to accommodate optimal learning. Overall, there is favourable evidence indicating that a more complex and resource-rational human sampling heuristic is at play, manifesting itself as a simpler identifiable sampling policies, depending on the task at hand.

References

- Lewis Bott and Evan Heit. Nonmonotonic Extrapolation in Function Learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 30(1):38–50, 2004.
- Berndt Brehmer. Hypotheses about relations between scaled variables in the learning of probabilistic inference tasks. *Organizational Behavior and Human Performance*, 11(1): 1–27, 1974.
- Edward L DeLosh, Mark A McDaniel, and Jerome R Busemeyer. Extrapolation: The Sine Qua Non for Abstraction in Function Learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23(4):968–986, 1997.
- Rebekah Gelpi, Nayan Saxena, George Lifchits, Daphna Buchsbaum, and Christopher G Lucas. Sampling heuristics for active function learning. 2021.
- Samuel J Gershman, Eric J Horvitz, and Joshua B Tenenbaum. Computational rationality: A converging paradigm for intelligence in brains, minds, and machines. *Science*, 349(6245): 273–278, 2015.
- Gerd Gigerenzer. *Rationality for mortals: How people cope with uncertainty*. Oxford University Press, 2008.
- Thomas L Griffiths, Christopher G. Lucas, Joseph J. Williams, and Michael L Kalish. Modeling human function learning with Gaussian processes. In *Advances in Neural Information Processing Systems*, volume 21, page 553–560, 2008.
- A Jones, E Schulz, B Meder, and A Ruggeri. Active function learning. In *Proceedings of the 40th Annual Meeting of the Cognitive Science Society*, pages 578–583, 2018.
- Michael L. Kalish, Stephan Lewandowsky, and John K. Kruschke. Population of Linear Experts: Knowledge Partitioning and Function Learning. *Psychological Review*, 111(4): 1072–1099, 2004.

- Falk Lieder and Thomas L Griffiths. Strategy selection as rational metareasoning. *Psychological Review*, 124(6):762–794, 2017.
- Falk Lieder and Thomas L Griffiths. Resource-rational analysis: understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43, 2020.
- Christopher G Lucas, Douglas Sterling, and Charles Kemp. Superspace extrapolation reveals inductive biases in function learning. In *Proceedings of the 34th Annual Meeting of the Cognitive Science Society*, pages 713–718, 2012.
- Christopher G. Lucas, Thomas L. Griffiths, Joseph J. Williams, and Michael L. Kalish. A rational model of function learning. *Psychonomic Bulletin & Review*, 22(5):1193–1215, 2015.
- Douglas B. Markant and Todd M Gureckis. A preference for the unpredictable over the informative during self-directed learning. In *Proceedings of the 36th Annual Meeting of the Cognitive Science Society*, pages 958–963, 2014.
- Geoffrey J. McLachlan, Sharon X. Lee, and Suren I. Rathnayake. Finite mixture models. *Annual Review of Statistics and Its Application*, 6(1):355–378, 2019. doi: 10.1146/annurev-statistics-031017-100325. URL <https://doi.org/10.1146/annurev-statistics-031017-100325>.
- Carl Rasmussen. The infinite gaussian mixture model. *Advances in neural information processing systems*, 12, 1999.
- Carl Edward Rasmussen and Christopher K. I. Williams. *Gaussian Processes for Machine Learning*. Adaptive Computation and Machine Learning. MIT Press, Cambridge, Mass, 2006. ISBN 978-0-262-18253-9.

Mark Steyvers, Brent Miller, Pernille Hemmer, and Michael D Lee. The Wisdom of Crowds in the Recollection of Order Information. In *Advances in Neural Information Processing Systems*, pages 1785–1793, 2009.

Andrew Gordon Wilson, Christoph Dann, Christopher G. Lucas, and Eric P. Xing. The Human Kernel. In *Advances in Neural Information Processing Systems*, volume 28, 2015.