

Adaptive Robotic Information Gathering via non-stationary Gaussian processes

The International Journal of
Robotics Research
2024, Vol. 43(4) 405–436
© The Author(s) 2023
Article reuse guidelines:
sagepub.com/journals-permissions
DOI: 10.1177/02783649231184498
journals.sagepub.com/home/ijr



Weizhe Chen , Roni Khardon and Lantao Liu 

Abstract

Robotic Information Gathering (RIG) is a foundational research topic that answers how a robot (team) collects informative data to efficiently build an accurate model of an unknown target function under robot embodiment constraints. RIG has many applications, including but not limited to autonomous exploration and mapping, 3D reconstruction or inspection, search and rescue, and environmental monitoring. A RIG system relies on a probabilistic model's prediction uncertainty to identify critical areas for informative data collection. Gaussian processes (GPs) with stationary kernels have been widely adopted for spatial modeling. However, real-world spatial data is typically non-stationary—different locations do not have the same degree of variability. As a result, the prediction uncertainty does not accurately reveal prediction error, limiting the success of RIG algorithms. We propose a family of non-stationary kernels named Attentive Kernel (AK), which is simple and robust and can extend any existing kernel to a non-stationary one. We evaluate the new kernel in elevation mapping tasks, where AK provides better accuracy and uncertainty quantification over the commonly used stationary kernels and the leading non-stationary kernels. The improved uncertainty quantification guides the downstream informative planner to collect more valuable data around the high-error area, further increasing prediction accuracy. A field experiment demonstrates that the proposed method can guide an Autonomous Surface Vehicle (ASV) to prioritize data collection in locations with significant spatial variations, enabling the model to characterize salient environmental features.

Keywords

Robotic Information Gathering, Informative Planning, non-stationary Gaussian processes, Attentive Kernel

Received 22 December 2022; Revised 1 May 2023; Accepted 29 May 2023

Senior Editor: Jose-Luis Blanco

Associate Editor: Dylan Shell

1. Introduction

Collecting informative data for effective modeling of an unknown physical process or phenomenon has been studied in different domains, for example, Optimal Experimental Design in Statistics (Atkinson 1996), Optimal Sensor Placement in Wireless Sensor Networks (Krause et al., 2008), Active Learning (Settles 2012), and Bayesian Optimization (Snoek et al., 2012) in Machine Learning.

In Robotics, this problem falls within the spectrum of *Robotic Information Gathering (RIG)* (Thrun 2002). RIG has recently received increasing attention due to its wide applicability. Applications include environmental modeling and monitoring (Dunbabin and Marques 2012), 3D reconstruction and inspection (Hollinger et al., 2013; Schmid et al., 2020), search and rescue (Meera et al., 2019), exploration and mapping (Jadidi et al., 2019), as well as active System Identification (Buisson-Fenet et al., 2020).

A RIG system typically relies on a probabilistic model's prediction uncertainty to identify critical areas for informative data collection. Figure 1 illustrates the workflow of a

RIG system, which shows three major forces that drive the progress of RIG: probabilistic models, objective functions, and informative planners.

The defining element distinguishing other *active information acquisition* problems and RIG is the robot embodiment's physical constraints (Taylor et al., 2021). In Active Learning (Biyik et al., 2020) or Optimal Sensor Placement (Krause et al., 2008), an agent can sample arbitrary data in a given space. In RIG, however, a robot must collect data sequentially along the motion trajectories. Consequently, most existing work in RIG is dedicated to a sequential decision-making problem called *Informative (Path) Planning* (Binney et al., 2013; Hollinger and Sukhatme 2014; Lim

Luddy School of Informatics, Computing, and Engineering, Indiana University, Bloomington, IN, USA

Corresponding author:

Lantao Liu, Luddy School of Informatics, Computing, and Engineering, Indiana University, 2401 N Milo B Sampson Ln, Room 080, Bloomington, IN 47408, USA.
Email: lantao@iu.edu

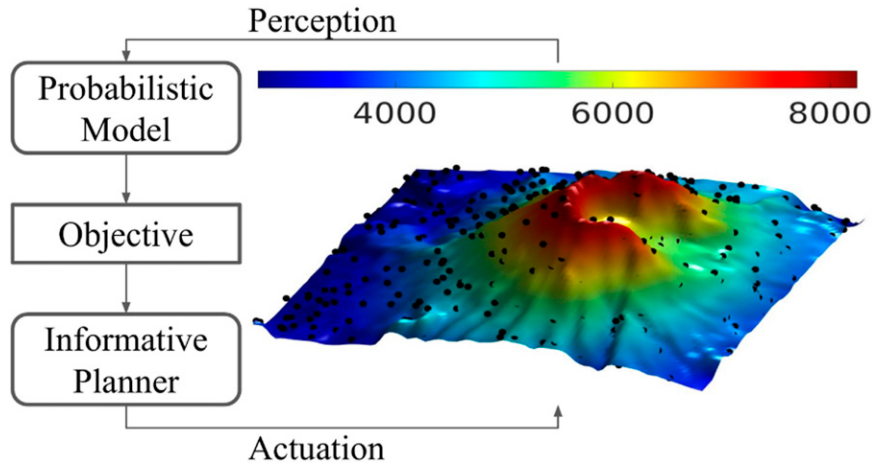


Figure 1. Diagram of a Robotic Information Gathering system. The goal is to autonomously gather informative elevation measurements of Mount St Helens to efficiently build a terrain map unknown a priori. The color indicates elevation, and black dots are collected samples.

et al., 2016; Choudhury et al., 2018; Jadidi et al., 2019; Best et al., 2019). Specifically, Informative Planning seeks an action sequence or a policy by optimizing an objective function that guides the robot to collect informative data, aiming to efficiently build an accurate model of the process under the robot’s motion and sensing cost constraints (Chen and Liu 2019; Popović et al., 2020b). The decisive objective function is derived from the uncertainty of probabilistic models such as Gaussian processes (GPs) (Ghaffari Jadidi et al., 2018), Hilbert maps (Senanayake and Ramos 2017), occupancy grid maps (Charrow et al., 2015a), and Gaussian mixture models (Dhawale and Michael 2020). Since the performance of a RIG system depends on not only planning but also learning, as shown in the feedback loop of Figure 1, a natural question is: how can we further boost the performance by improving the probabilistic models? In this work, we answer this question from the perspective of improving the *modeling flexibility* and *uncertainty quantification* of GPs.

Gaussian process regression (GPR) is one of the most prevalent methods for mapping continuous spatiotemporal phenomena. GPR requires the specification of a kernel, and *stationary* kernels, for example, the radial basis function (RBF) kernel and the Matérn family, are commonly adopted (Rasmussen and Williams 2005). However, real-world spatial data typically does not satisfy stationary models which assume different locations have the same degree of variability. For instance, the environment in Figure 1 shows higher spatial variability around the crater. Due to the mismatch between the assumption and the ground-truth environment, GPR with stationary kernels cannot portray the characteristic environmental features in detail. Figure 2(a) shows the over-smoothed prediction of the elevation map after training a stationary GPR using the collected data shown in Figure 1. The model also assigns low uncertainty to the high-error area, *c.f.*, the circled regions in Figures 2(b) and (c), leading to degraded performance when the model is used in RIG.

Non-stationary GPs, on the other hand, are of interest in many applications, and the past few decades have witnessed great advancement in this research field (Gibbs 1997;

Paciorek and Schervish 2003; Lang et al., 2007; Plagemann et al. 2008a, 2008b; Wilson et al., 2016; Calandra et al., 2016; Heinonen et al., 2016; Remes et al., 2017, 2018). However, prior work leaves room for improvement. The problem is that many non-stationary models learn fine-grained variability at every location, making the model too flexible to be trained without advanced parameter initialization and regularization techniques. We propose a family of non-stationary kernels named *Attentive Kernel* (AK) to mitigate this issue. The main idea of our AK is limiting the non-stationary model to combine a fixed set of correlation scales, that is, primitive length-scales, and mask out data across discontinuous jumps by “soft” selection of relevant data. The correlation-scale composition and data selection mechanisms are learned from data. Figure 2(d) shows the prediction of GPR with the AK on the same dataset used in Figure 2(a). As the arrows highlight, the AK depicts the environment at a finer granularity. Figures 2(e) and (f) show that the AK allocates high uncertainty to the high-error area; thus, sampling the high-uncertainty locations can help the robot collect valuable data to decrease the prediction error further.

1.1. Contributions

The main contribution of this paper is in designing the Attentive Kernel (AK) and evaluating its suitability for Robotic Information Gathering (RIG). We present an extensive evaluation to compare the AK with existing non-stationary kernels and a stationary baseline. The benchmarking task is elevation mapping in several natural environments that exhibit a range of non-stationary features. The results reveal a significant advantage of the AK when it is used in passive learning, active learning, and RIG. We also conduct a field experiment to demonstrate the behavior of the proposed method in a real-world elevation mapping task, where the prediction uncertainty of the AK guides an Autonomous Surface Vehicle (ASV) to identify essential sampling locations and collect valuable data rapidly. Last but not least, we release the code (github.com/weizhe-chen/attentive_kernels) for reproducing all the results.

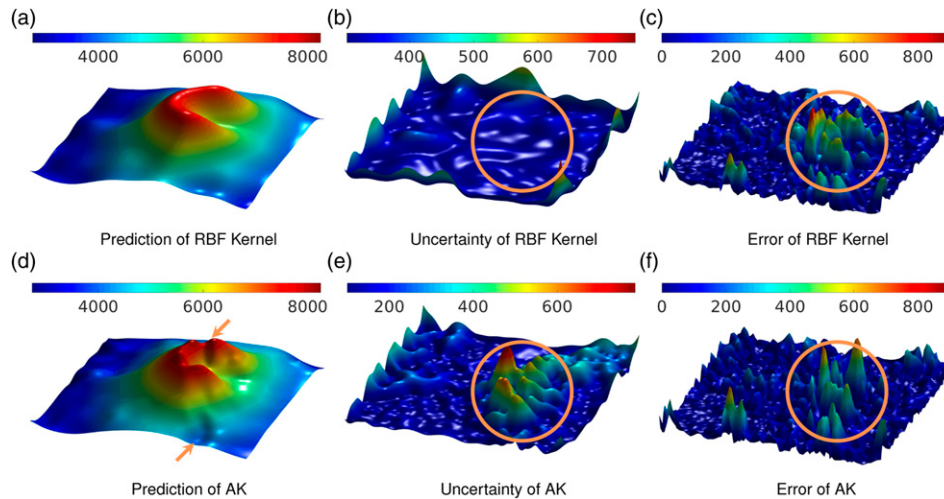


Figure 2. Comparison of Gaussian process regression with radial basis function kernel and Attentive Kernel.

This paper presents an extended and revised version of previous work by [Chen et al. \(2022\)](#). The major modifications include a comprehensive literature review on RIG to contextualize our work, additional evaluation, results, and discussion on the AK, and a substantially improved Python library. Specifically, we provide the following contributions:

- We present a broader and deeper survey on related work to highlight how our work fits into the existing literature on RIG.
- We add more results to the experiments and discuss them in detail to provide further evidence for our conclusions.
- We thoroughly evaluate the AK from various perspectives and discuss its limitations and potential future work.
- We release a new Python library called PyPolo (pypolo.readthedocs.io) for learning, researching, and benchmarking RIG algorithms. This library is a significant improvement and restructure compared to the one presented in [Chen et al. \(2022\)](#).

2. Related work

In this section, we will first survey related work in RIG, which mainly revolves around three pillars: probabilistic models ([Section 2.1.2](#)), objective functions ([Section 2.1.1](#)), and Informative Planning algorithms ([Section 2.1.3](#)). Also, we discuss relevant RIG applications in [Section 2.1.4](#). Then, we categorize prior efforts on non-stationary GPs and how the proposed method relates to the existing solutions ([Section 2.2](#)). Finally, we describe the relationship between RIG and some related research topics to locate our work within the context of existing literature ([Section 2.3](#)).

2.1. Robotic Information Gathering

A RIG system has three essential components:

1. A model to approximate the unknown target function;

2. An objective function that can characterize the model's prediction error;
3. An informative planner that makes *non-myopic* decisions by optimizing the objective function under the robot's embodiment constraints.

We discuss these three aspects in this section.

2.1.1. Objective functions. RIG can be the main goal of some tasks, such as infrastructure inspection ([Bircher et al., 2018](#)), or serve as an auxiliary task for achieving other goals, for example, seeking the biological hotspots in an unknown environment ([McCammon and Hollinger 2018](#)). In the former cases, the objective function is purely “information-driven” ([Ferrari and Wettergren 2021](#); [Bai et al., 2021](#)), while in the latter scenarios, the objective function balances exploration and exploitation ([Marchant and Ramos 2012, 2014](#); [Bai et al., 2016](#)). The objective function can be further extended to multi-objective cases ([Chen and Liu 2019](#); [Ren et al., 2022](#); [Dang 2020](#)).

Many objective functions have been proposed, inspired by Information Theory and Optimal Experimental Design ([Charrow et al., 2015a](#); [Zhang et al., 2020](#); [Carrillo et al., 2015](#)). Information-theoretic objective functions include Shannon's and Rényi's entropy, mutual information, and Kullback–Leibler divergence between the prior and posterior predictive distributions. In the case of multivariate Gaussian distributions, these information measures are all related to the logarithmic determinant of the posterior covariance matrix, which can be intuitively viewed as computing the “size” of the posterior covariance matrix. Optimal design theory directly measures the size by computing the matrix determinant, trace, or eigenvalues. Computing the matrix determinant and eigenvalue is known to be computationally expensive. Therefore, many existing works on objective functions are dedicated to alleviating the computational bottleneck ([Charrow et al. 2015b, 2015a](#); [Zhang et al., 2020](#); [Zhang and Scaramuzza 2020](#); [Gupta et al., 2021](#); [Xu et al., 2021](#)).

Most objective functions are summary statistics of the predictive (co)variance given by a probabilistic model. Only when the predictive (co)variance captures modeling error well, optimizing these objective functions can guide the robot to collect informative data that effectively improve the model's accuracy. From this perspective, improving the uncertainty-quantification capability of probabilistic models can broadly benefit future work based on these objective functions. This aspect is what we strive to improve in this work. As can be seen in the next section, this problem is understudied.

2.1.2. Probabilistic models. Many probabilistic models have been applied to RIG, for example, Gaussian processes (Stachniss et al., 2009; Marchant and Ramos 2012, 2014; Ouyang et al., 2014; Ma et al., 2017; Luo and Sycara 2018; Jang et al., 2020; Popović et al., 2020a; Lee et al., 2022), Hilbert maps (Ramos and Ott 2016; Senanayake and Ramos 2017; Guizilini and Ramos 2019), occupancy grid maps (Popović et al., 2017, 2020b; Saroya et al., 2021), and Gaussian mixture models (O'Meadhra et al., 2018; Tabib et al., 2019). GPs are widely adopted due to their excellent uncertainty quantification feature, which is decisive to RIG. However, the vanilla GP models need to be more computationally efficient to be suitable for real-time applications and multi-robot scenarios. Therefore, related work in RIG mainly discusses GPs in the context of improving computational efficiency and coordinating multiple robots. Jang et al. (2020) apply the distributed GPs (Deisenroth and Ng 2015) to decentralized multi-robot online Active Sensing. Ma et al. (2017) and Stachniss et al. (2009) use sparse GPs to alleviate the computational burden. The mixture of GP experts (Rasmussen and Ghahramani 2001) has been applied to divide the workspace into smaller parts for multiple robots to model an environment simultaneously (Luo and Sycara 2018; Ouyang et al., 2014).

The early work by Krause and Guestrin (2007) is highly related to our work. They use a spatially varying linear combination of localized stationary processes to model the non-stationary pH values in a river. The weight of each local GP is the normalized predictive variance at the test location. This idea is similar to the length-scale selection idea in Section 4.1.1. The main difference is that they manually partition the workspace while our model learns a weighting function from data. To the best of our knowledge, our work is the first to discuss the influence of the probabilistic models' uncertainty quantification on RIG performance.

2.1.3. Informative Planning. The problem of seeking an action sequence or policy that yields informative data is known as Informative Path Planning due to historical reasons (Singh et al., 2007; Meliou et al., 2007). However, the problem is not restricted to path planning. For example, recent work has discussed informative *motion* planning (Teng et al., 2021), informative *view* planning (Lauri et al., 2020), and exploratory *grasping* (Danielczuk et al., 2021). Hence, we adopt the generic term Informative Planning to unify different branches of the same problem.

Early works on Informative Planning propose various *recursive greedy* algorithms that provide performance guarantee by exploiting the *submodularity* property of the objective function (Singh et al., 2007; Meliou et al., 2007; Binney et al., 2013). Note that the performance guarantee is on uncertainty reduction rather than modeling accuracy. Planners based on dynamic programming (Low et al., 2009; Cao et al., 2013) and mixed integer quadratic programming (Yu et al., 2014) lift the assumption on the objective function at the expense of higher computational complexity. These methods solve combinatorial optimization problems in discrete domains, thus scaling poorly in problem size. To develop efficient planners in *continuous* space with motion constraints, Hollinger and Sukhatme (2014) introduce sampling-based informative motion planning, which is further developed to online variants (Schmid et al., 2020; Jadidi et al., 2019). Monte Carlo Tree Search (MCTS) methods are conceptually similar to sampling-based informative planners (Kantaros et al., 2021; Schlotfeldt et al., 2018) and have recently garnered great attention (Arora et al., 2019; Best et al., 2019; Morere et al., 2017; Chen and Liu 2019; Flaspohler et al., 2019). Trajectory optimization is a solid competitor to sampling-based planners. Bayesian Optimization (Marchant and Ramos 2012; Bai et al., 2016; Di Caro and Yousaf 2021) and Evolutionary Strategy (Popović et al., 2017, 2020b; Hitz et al., 2017) are the two dominating methods in this realm. New frameworks of RIG, for example, Imitation Learning (Choudhury et al., 2018), are emerging. Communication constraints (Lauri et al., 2017) and adversarial attacks (Schlotfeldt et al., 2021) have also been discussed.

2.1.4. Relevant applications. Mobile robots can be considered as autonomous data-gathering tools, enabling scientific research in remote and hazardous environments (Li 2020; Bai et al., 2021). RIG has been successfully applied to environmental mapping and monitoring (Dunbabin and Marques 2012). An underwater robot with a profiling sonar can inspect a ship hull autonomously (Hollinger et al., 2013). In Girdhar et al. (2014), the underwater robot performs semantic exploration with online topic modeling, which can group corals belonging to the same species or rocks of similar types. Flaspohler et al. (2019) deploy an ASV for localizing and collecting samples at the most exposed coral head. Hitz et al. (2017) monitor algal bloom using an ASV, which can provide early warning to environmental managers to conduct water treatment in a more appropriate time frame. Manjanna et al. (2018) show that a robot team can help scientists collect plankton-rich water samples via in situ mapping of chlorophyll density. Fernández et al. (2022) propose delineating the sampling locations that correspond to the quantile values of the phenomenon of interest, which helps the scientists to collect valuable data for later analysis. Active lakebed mapping, where the static ground truth is available, can serve as a testbed for ocean bathymetric mapping (Ma et al., 2018). RIG can also be applied to the 3D reconstruction of large

scenes (Kompis et al., 2021) and object surfaces (Zhu et al., 2021). In addition to geometric mapping, semantic mapping is also explored in Atanasov et al. (2014), where a PR2 robot with an RGB-D camera attached to the wrist leverages non-myopic view planning for active object classification and pose estimation. Meera et al. (2019) present a realistic simulation of a search-and-rescue scenario in which Informative Planning maximizes search efficiency under the Unmanned Aerial Vehicle (UAV) flight time constraints. Fixed-wing UAVs use aerodynamics akin to aircraft, so it has a much longer flight time than multi-rotors. Moon et al. (2022) simulate a fixed-wing UAV with a forward-facing camera to search for multiple objects of interest in a large search space.

2.2. Non-stationary Gaussian processes

GPs suffer from two significant limitations (Rasmussen and Ghahramani 2001). The first one is the notorious cubic computational complexity of a vanilla implementation. Recent years have witnessed remarkable progress in solving this problem based on sparse GPs (Quinonero-Candela and Rasmussen 2005; Titsias 2009; Hoang et al., 2015; Sheth et al., 2015; Bui et al., 2017; Wei et al., 2021). The second drawback is that the covariance function is commonly assumed to be stationary, limiting the modeling flexibility. Developing non-stationary GP models that are easy to train is still an active open research problem. Ideas of handling non-stationarity can be roughly grouped into three categories: input-dependent length-scale (Gibbs 1997; Paciorek and Schervish 2003; Lang et al., 2007; Plagemann et al. 2008b, 2008a; Heinonen et al., 2016; Remes et al., 2017), input warping (Sampson and Guttorm 1992; Snoek et al., 2014; Calandra et al., 2016; Wilson et al., 2016; Tompkins et al., 2020a; Salimbeni and Deisenroth 2017), and the mixture of experts (Rasmussen and Ghahramani 2001; Trapp et al., 2020).

Input-dependent length-scale provides excellent flexibility to learn different correlation scales at different input locations. Gibbs (1997) and Paciorek and Schervish (2003) have shown how one can construct a valid kernel with input-dependent length-scales, namely, a *length-scale function*. The standard approach uses another GP to model the length-scale function, which is then used in the kernel of a GP, yielding a hierarchical Bayesian model. Several papers have developed inference techniques for such models and demonstrated their use in some applications (Lang et al., 2007; Plagemann et al. 2008b, 2008a; Heinonen et al., 2016; Remes et al., 2017). Recently, Remes et al. (2018) show that modeling the length-scale function using a neural network improves performance. Note, however, that learning a length-scale function is nontrivial (Wang et al., 2020).

Input warping is more widely applicable because it endows any stationary kernel with the ability to model non-stationarity by mapping the input locations to a distorted space and assuming stationarity holds in the new space. This approach has a tricky requirement: the mapping must be

injective to avoid undesirable folding of the space (Sampson and Guttorm 1992; Snoek et al., 2014; Salimbeni and Deisenroth 2017).

A mixture of GP experts (MoGPE) uses a *gating network* to allocate each data to a local GP that learns its hyper-parameters from the assigned data. It typically requires Gibbs sampling (Rasmussen and Ghahramani 2001), which can be slow. Hence, one might need to develop a faster approximation (Nguyen-Tuong et al., 2008). We view MoGPE as an orthogonal direction to other non-stationary GPs or kernels because any GP model can be treated as the expert so that one can have a mixture of non-stationary GPs.

The AK lies at the intersection of these three categories. Section 4.1.1 presents an input-dependent length-scale idea by weighting base kernels with different fixed length-scales at each location. Composing base kernels reduces the difficulty of learning a length-scale function from scratch and makes our method compatible with any base kernel. In Section 4.1.2, we augment the input with extra dimensions. We can view the augmentation as warping the input space to a higher-dimensional space, ensuring *injectivity* by design. Combining these two ideas gives a conceptually similar model to MoGPE (Rasmussen and Ghahramani 2001) in that they both divide the space into multiple regions and learn localized hyper-parameters. The idea of augmenting the input dimensions has been discussed by Pfingsten et al. (2006). However, they treat the augmented vector as a latent variable and resort to Markov chain Monte Carlo for inference. The AK treats the augmentation vector as the output of a deterministic function of the input, resulting in a more straightforward inference procedure. Also, the AK can be used in MoGPE to build more flexible models.

In robotic mapping, another line of notable work on probabilistic models is the family of Hilbert maps (Ramos and Ott 2016; Senanayake and Ramos 2017; Guizilini and Ramos 2019), which aims to alleviate the computational bottleneck of GPs (O’Callaghan and Ramos 2012) by projecting the data to another feature space and applying a logistic regression classifier in the new space. Since Hilbert maps are typically used for occupancy mapping (Doherty et al., 2016) and reconstruction tasks (Guizilini and Ramos 2017), related work also considers non-stationarity for better prediction (Senanayake et al., 2018; Tompkins et al., 2020b).

2.3. Relationship to other research topics

RIG is a fundamental research problem seeking an answer to the following question:

How does a robot (team) collect informative data to efficiently build an accurate model of an unknown function under robot embodiment constraints?

Depending on how we define *data* and what the unknown *target function* is, RIG appears in the form of Active Dynamics Learning, Active Mapping, Active Localization, and Active Simultaneous Localization and Mapping (SLAM). Figure 3 shows a Venn diagram of these topics.

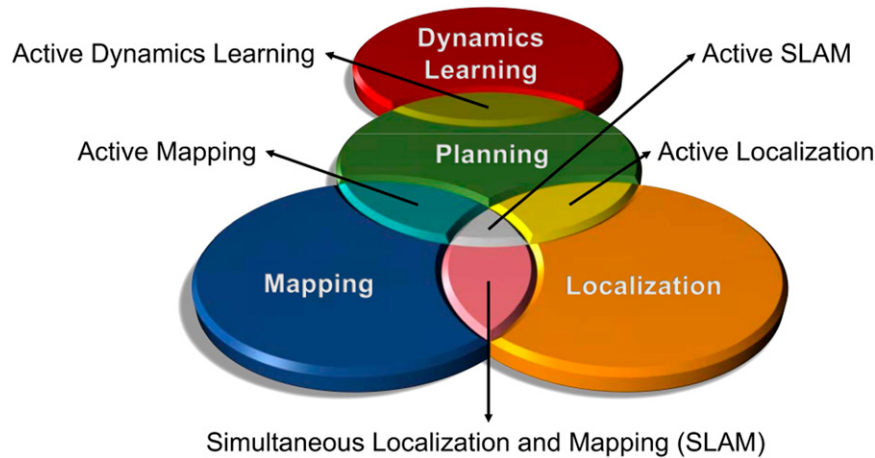


Figure 3. Research topics related to RIG.

Although we evaluate the AK in Active Mapping tasks, other related problems, for example, Active Dynamics Learning, can also benefit from the proposed method if the target function is modeled by a GP. On top of that, guiding the data collection process by minimizing *well-calibrated* uncertainty estimates applies to all these related topics (Rodríguez-Arévalo et al., 2018).

2.3.1. Active dynamics learning. Control synthesis typically depends on the system dynamics. Due to the complex interaction between the robot and the environment, for example, a quadruped running at high speed over rough terrain, mechanical wear and tear, and actuator faults, it may be infeasible to build an accurate dynamics model a priori (Cully et al., 2015). In these cases, the robot must take safe actions and observe its dynamics to explore different behavioral regimes sample-efficiently (Abraham and Murphey 2019). When the robot collects dynamics information to infer the unknown transition function, the RIG problem is known as Active Dynamics Learning or System Identification (Taylor et al., 2021). In this context, informative data refers to the state-action-state pairs or the full state-action trajectories that help efficiently learn an accurate model of the unknown *system dynamics* or *transition function*. The system dynamics can be modeled as fixed-form equations (Jegorova et al., 2020), data-driven models, including *parametric* models (Chua et al., 2018), *non-parametric* models (Calandra et al., 2016), and *semi-parametric* models (Romerres et al., 2019), and the combination of the analytical models and data-driven models (Heiden et al., 2021b). GPs have arguably become the de facto standard in collecting informative data that minimizes the predictive uncertainty of data-driven models to achieve sample-efficient dynamics learning (Rezaei-Shoshtari et al., 2019; Buisson-Fenet et al., 2020; Capone et al., 2020; Lew et al., 2022; Yu et al., 2021). With the rise of Automatic Differentiation (Paszke et al., 2017), a large body of recent work tend to estimate the physical parameters inside *differentiable Rigid-Body Dynamics* models (Sutanto et al.,

2020; Lutter et al., 2021; De Avila Belbute-Peres et al., 2018) or *differentiable robotics simulators* (Hu et al., 2019; Freeman et al., 2021; Werling et al., 2021). The literature emphasizes that calibrating the simulation (Mehta et al., 2021) is essential for both Reinforcement Learning with domain randomization (Ramos et al., 2019; Muratore et al., 2022) or trajectory optimization (Du et al., 2021; Heiden et al., 2021a). In this context, we can consider RIG as Active Simulation Calibration since the robot collects informative trajectories to efficiently learn an accurate model of the unknown simulation parameters under the kinodynamic constraints. Active Simulation Calibration can also directly optimize the task-specific reward. For instance, Muratore et al. (2021) model the policy return as a GP and use Bayesian Optimization to tune the simulation parameters. Liang et al. (2020) learn a task-oriented exploration policy to collect informative data for calibrating task-relevant simulation parameters.

2.3.2. Active perception. When the robot collects data from the *environment* rather than its dynamics, RIG becomes *Active Perception*—an agent (e.g., camera or robot) changes its angle of view or position to perceive the surrounding environment better (Bajcsy 1988; Aloimonos et al., 1988; Bajcsy et al., 2018). If the agent actively perceives the environment to reduce the *localization uncertainty*, the problem is referred to as Active Localization (Fox et al., 1998; Borghi and Caglioti 1998). If the goal is to build the best possible *representation of an environment*, the problem essentially becomes Active Mapping (Lluvia et al., 2021).

2.3.3. Active localization. Localization uncertainty can arise from perceptual degradation (Ebadi et al., 2020), noisy actuation (Thrun 2002), and inaccurate modeling (Roy et al., 1999). Decision-making or planning under uncertainty (LaValle 2006; Bry and Roy 2011; Preston et al., 2022) provides an elegant framework to formulate these problems using partially observable Markov decision processes (POMDPs) (Kaelbling et al., 1998; Cai et al., 2021; Lauri et al., 2022). A principled

approach to address these problems is to plan in the *belief space* (Kaelbling and Lozano-Pérez 2013; Nishimura and Schwager 2021). *Information gathering* is a natural behavior generated by Belief-Space Planning (Platt et al., 2010). Computing optimal policy in belief space is computationally intensive, but useful heuristics enable efficient computation of high-quality solutions (Kim et al., 2019; Prentice and Roy 2009; Zheng et al., 2022). Although the localization uncertainty can come from different sources, in perceptually degraded environments such as subterranean, perception uncertainty outweighs the others. A dedicated topic for this case is perception-aware planning (Zhang 2020). Note that localization is not necessarily positioning a mobile robot on a map (Chaplot et al., 2018); it can also be locating and tracking an object in the workspace of a manipulator with force–torque sensor measurements (Wirnshofer et al., 2020; Schneider et al., 2022).

2.3.4. Active sensing and mapping. Suppose that the data refers to the robot’s observations, for example, camera images or LiDAR point clouds, and the unknown target function is the ground-truth representation of the environment. In that case, RIG can be considered an Active Mapping problem (Placed et al., 2022). Mapping uncertainty can come from *aleatoric* uncertainty inherent in measurement noise and *epistemic* uncertainty due to unknown model parameters and data scarcity (Krause and Guestrin 2007). Active Mapping efficiently builds an accurate model of the environment by minimizing epistemic uncertainty, which is often termed Active Sensing when focusing on the active acquisition of sensor measurements for better *prediction* rather than *model learning* (Cao et al., 2013; MacDonald and Smith 2019; Schlotfeldt et al., 2019; Rückin et al., 2022). When mapping a 3D environment using a sensor with a limited field-of-view, this is known as the Next-Best View problem (Connolly 1985; Bircher et al., 2016; Palomeras et al., 2019; Lauri et al., 2020). Autonomous Exploration is sometimes used interchangeably with Active Mapping (Lluvia et al., 2021). However, the nuances of the assumptions and evaluation metrics of the two domains yield significantly different solutions and robot behaviors. Specifically, Active Mapping typically assumes ideal localization (Popović et al., 2020a) and aims at building an accurate environment map using noisy and sparse observations; thus, the performance is evaluated by reconstruction error against the ground truth. The robot might revisit some complex regions to collect more data if the model prediction is not accurate enough. For example, when performing Active Mapping of a ship hull, the robot should collect more data around the propeller (Hollinger et al., 2013). Autonomous Exploration emphasizes obtaining the global structure of a vast unknown environment, implying that the robot (team) should avoid duplicate coverage; thus, the evaluation criterion is the explored volume (Cao et al., 2021). In contrast to Active Mapping, unreliable localization is one of the major challenges in Autonomous Exploration that should be addressed

(Tranzatto et al., 2022; Papachristos et al., 2019). In this work, our application belongs to the Active Mapping problem, where the better uncertainty quantification of the proposed non-stationary GPR guides the robot to collect more informative data for rapid learning of an accurate map.

2.3.5. Active SLAM. Controlling a robot performing SLAM to reduce both the localization and mapping uncertainty is called active SLAM (Placed et al., 2022). Active Localization and Active Mapping are two conflicting objectives. The former asks the robot to revisit explored areas for potential *loop closure* (Stachniss et al., 2004), while the latter guides the robot to expand *frontiers* for efficient map building (Yamauchi 1997). We refer the interested reader to the corresponding survey papers (Lluvia et al., 2021; Placed et al., 2022).

3. Problem statement

Consider deploying a robot to *efficiently* build a map of an *unknown* environment using only *sparse* sensing measurements of onboard sensors. For instance, when reconstructing a pollution distribution map, the environmental sensors can only measure the pollutant concentration in a *point-wise* sampling manner, yielding sparse measurements along the trajectory. Another scenario is building a large bathymetric map of the seabed. The depth measurements of a multi-beam sonar can be viewed as *point measurements* because the unknown target area is typically vast. Exhaustively sampling the whole environment is prohibitive, if not impossible; thus, one must develop adaptive planning algorithms to collect the most informative data for building an accurate model. Table 1 introduces the notation system used in this paper. We use column vector by default.

3.1. Minimization of error vs. uncertainty

Problem 1. The target environment is an unknown function $f_{\text{env}}(\mathbf{x}) : \mathbb{R}^D \mapsto \mathbb{R}$ defined over spatial locations $\mathbf{x} \in \mathbb{R}^D$. Let $\mathbb{T} \triangleq \{t\}_{t=0}^T$ be the set of decision epochs. A robot at state $\mathbf{s}_{t-1} \in \mathbb{S}$ takes an action $a_{t-1} \in \mathbb{A}$, arrives at the next state $\mathbf{s}_t$ following a transition function $p(\mathbf{s}_t | \mathbf{s}_{t-1}, a_{t-1})$, and collects $N_t \in \mathbb{N}$ noisy measurements $\mathbf{y}_t \in \mathbb{R}^{N_t}$ at sampling locations $\mathbf{X}_t = [\mathbf{x}_1, \dots, \mathbf{x}_{N_t}]^\top \in \mathbb{R}^{N_t \times D}$ when transitioning from $\mathbf{s}_{t-1}$ to $\mathbf{s}_t$. We assume that the transition function is known and deterministic and that the robot state is observable. The robot maintains a probabilistic model built from all the training data collected so far $\mathbb{D}_t = \{(\mathbf{X}_i, \mathbf{y}_i)\}_{i=1}^t$. The model provides predictive mean $\mu_t : \mathbb{R}^D \mapsto \mathbb{R}$ and predictive variance $v_t : \mathbb{R}^D \mapsto \mathbb{R}_{\geq 0}$ functions. Let $\mathbf{x}^\star$ be a test or query location, and $\text{error}(\cdot)$ be an error metric. At each decision epoch $t \in \mathbb{T}$, our goal is to find sampling locations that minimize the *expected error* after updating the model with the collected data

Table 1. Mathematical notations.

Meaning	Example	Remark
Variable	m	Lower-case
Constant	M	Upper-case
Vector	$\mathbf{x}$	Bold, lower-case
Matrix	$\mathbf{X}$	Bold, upper-case
Set/space	$\mathbb{R}$	Blackboard
Cartesian product	$[a, b]^D$	D -dim hypercube
Function	$d(\cdot)$	Typewriter
Special PDF	$\mathcal{N}$	Calligraphy capital
Definition	$\triangleq$	Normal
Transpose	$\mathbf{m}^\top$	Customized command
Euclidean norm	$\ \cdot\ _2$	Customized command

$$\arg \min_{\mathbf{x}_t} E_{\mathbf{x}^*} [\text{error}(f_{\text{env}}(\mathbf{x}^*), \mu_t(\mathbf{x}^*), v_t(\mathbf{x}^*))] \quad (1)$$

The predictive variance is also included in equation (1) because it is required when computing some error metrics, for example, negative log predictive density. Note that the expected error cannot be directly used as the objective function for a planner because the ground-truth function f_{env} is unknown. RIG bypasses this problem by optimizing a surrogate objective.

Problem 2. Assuming the same conditions as Problem 1, find *informative* sampling locations that minimize an uncertainty measure $\text{info}(\cdot)$, for example, entropy:

$$\arg \min_{\mathbf{x}_t} E_{\mathbf{x}^*} [\text{info}(v_t(\mathbf{x}^*))] \quad (2)$$

RIG implicitly assumes that minimizing prediction uncertainty (Problem 2) can also effectively reduce prediction error (Problem 1). This assumption is valid when the model uncertainty is *well-calibrated*. A model with well-calibrated uncertainty gives high uncertainty when the prediction error is significant and low uncertainty otherwise.

3.2. Gaussian process regression

The predictive mean and variance functions are given by a Gaussian process regression (GPR) model in this work. A Gaussian process (GP) is a collection of random variables, any finite number of which have a joint Gaussian distribution (Rasmussen and Williams 2005).

3.2.1. Model specification. We place a Gaussian process *prior* over the unknown target function

$$f_{\text{env}}(\mathbf{x}) \sim \mathcal{GP}(\mathbf{m}(\mathbf{x}), \mathbf{k}(\mathbf{x}, \mathbf{x}')) \quad (3)$$

which is specified by a mean function $\mathbf{m}(\mathbf{x})$ and a covariance function $\mathbf{k}(\mathbf{x}, \mathbf{x}')$, *a.k.a.* *kernel*. After standardizing the training targets y to have a near-zero mean empirically, the mean function is typically simplified to a

zero function, rendering the specification of the covariance function an important choice. Popular choices of the covariance functions are *stationary kernels* such as the RBF kernel and the Matérn family. We refer the interested reader to Rasmussen and Williams (2005) for other commonly used kernels.

This paper uses the RBF kernel to show how we transform a stationary kernel to a non-stationary one using the proposed method. Given two inputs $\mathbf{x}$ and $\mathbf{x}'$, the RBF kernel measures their correlation by computing the following kernel value

$$k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{\|\mathbf{x} - \mathbf{x}'\|_2^2}{2\ell^2}\right) \quad (4)$$

The correlation scale parameter ℓ is called the *length-scale*, which informally indicates the distance one has to move in the input space before the function value can change significantly (Rasmussen and Williams 2005). A given sample should be most correlated to itself; thus, equation (4) gives the largest kernel value when $\mathbf{x} = \mathbf{x}'$. Kernels are typically normalized to ensure that the largest kernel value is 1 and an *amplitude* parameter α can be used to scale the kernel value $\alpha k(\mathbf{x}, \mathbf{x}')$ to a larger range.

GPR assumes a Gaussian likelihood function. The target values y are the function outputs f corrupted by an additive Gaussian white noise

$$p(y|\mathbf{x}) = \mathcal{N}(y|f(\mathbf{x}), \sigma^2) \quad (5)$$

where σ is the observational *noise scale*.

3.2.2. Prediction. Since GP is a *conjugate* prior to the Gaussian likelihood, given N training inputs $\mathbf{X} \in \mathbb{R}^{N \times D}$ and training targets $\mathbf{y} \in \mathbb{R}^N$, the posterior predictive distribution has a closed-form expression:

$$p(f_\star|\mathbf{y}) = \mathcal{N}(f_\star|\mu, v), \quad (6)$$

$$\mu = \mathbf{k}_\star^\top \mathbf{K}_y^{-1} \mathbf{y}, \quad (7)$$

$$v = k_{\star\star} - \mathbf{k}_\star^\top \mathbf{K}_y^{-1} \mathbf{k}_\star \quad (8)$$

where $\mathbf{k}_\star$ is the vector of kernel values between all the training inputs $\mathbf{X}$ and the test input $\mathbf{x}^\star$, $\mathbf{K}_y$ is a shorthand of $\mathbf{K}_\mathbf{x} + \sigma^2 \mathbf{I}$, $\mathbf{K}_\mathbf{x}$ is the covariance matrix given by the kernel function evaluated at each pair of training inputs, and $k_{\star\star} \triangleq k(\mathbf{x}^\star, \mathbf{x}^\star)$.

3.2.3. Learning. The prediction of GPR in equation (6) is readily available with no need to train a model. However, the prediction quality of GPR depends on the setting of *hyper-parameters* $\psi \triangleq [\ell, \alpha, \sigma]$. These are the parameters of the kernel and likelihood function. Hence, optimizing these parameters—a process known as *model selection*—is a common practice to obtain a better prediction. Model selection is typically implemented by maximizing the *model evidence*, *a.k.a.*, *log marginal likelihood*,

$$\ln p(\mathbf{y}|\boldsymbol{\psi}) = \frac{1}{2} \left(\underbrace{-\mathbf{y}^\top \mathbf{K}_y^{-1} \mathbf{y}}_{\text{model fit}} - \underbrace{\ln \det(\mathbf{K}_y)}_{\text{model complexity}} - \underbrace{N \ln(2\pi)}_{\text{constant}} \right)$$

where $\det(\dots)$ is the matrix determinant.

When using GPR with the commonly used stationary kernels to reconstruct a real-world environment, high uncertainty is assigned to less sampled areas, regardless of the prediction error (see Figures 2(b) and 4). However, real-world spatial environments are typically non-stationary, and the high prediction error is more likely to be present in the high-variability region. In other words, the assumption of well-calibrated uncertainty is violated when using stationary kernels. Therefore, we aim to develop non-stationary kernels to improve GPR's uncertainty-quantification capability and prediction accuracy.

4. Methodology

We propose a new kernel called Attentive Kernel to deal with non-stationarity.

Definition 1. Attentive Kernel (AK). Given two inputs $\mathbf{x}, \mathbf{x}' \in \mathbb{R}^D$, vector-valued functions $\mathbf{w}_\theta(\mathbf{x}) : \mathbb{R}^D \mapsto [0, 1]^M$ and $\mathbf{z}_\phi(\mathbf{x}) : \mathbb{R}^D \mapsto [0, 1]^M$ parameterized by θ, ϕ , an amplitude α , and a set of M base kernels $\{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M$, let $\bar{\mathbf{w}} = \mathbf{w}_\theta(\mathbf{X}) / \|\mathbf{w}_\theta(\mathbf{X})\|_2$ and $\bar{\mathbf{z}} = \mathbf{z}_\phi(\mathbf{X}) / \|\mathbf{z}_\phi(\mathbf{X})\|_2$. The AK is defined as

$$\text{ak}(\mathbf{x}, \mathbf{x}') = \alpha \bar{\mathbf{z}}^\top \bar{\mathbf{z}} \sum_{m=1}^M \bar{w}_m \bar{w}'_m k_m(\mathbf{x}, \mathbf{x}') \quad (9)$$

where $\bar{w}_m$ is the m -th element of $\bar{\mathbf{w}}$.

We learn the parametric functions that map each input $\mathbf{x}$ to $\mathbf{w}$ and $\mathbf{z}$. The weight $\bar{w}_m \bar{w}'_m$ gives *similarity attention scores* to combine the set of base kernels $\{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M$. The inner product $\bar{\mathbf{z}}^\top \bar{\mathbf{z}}$ defines a *visibility attention score* to mask the kernel value.

Definition 1 is generic because any existing kernel can be the base kernel. To address non-stationarity, we choose the base kernels to be a set of stationary kernels with the same functional form but different length-scales. Specifically, we use RBF kernels with M length-scales $\{\ell_m\}_{m=1}^M$ that are evenly spaced in the interval $[\ell_{\min}, \ell_{\max}]$:

$$k_m(\mathbf{x}, \mathbf{x}') \triangleq k_{\text{RBF}}(\mathbf{x}, \mathbf{x}' | \ell_m) = \exp\left(-\frac{\|\mathbf{x} - \mathbf{x}'\|_2^2}{2\ell_m^2}\right)$$

Note that the length-scales $\{\ell_m\}_{m=1}^M$ are fixed constants rather than trainable variables. When applying the AK to a GPR, we optimize all the hyper-parameters $\{\alpha, \theta, \phi, \sigma\}$ by maximizing the marginal likelihood and make prediction as in GPR.

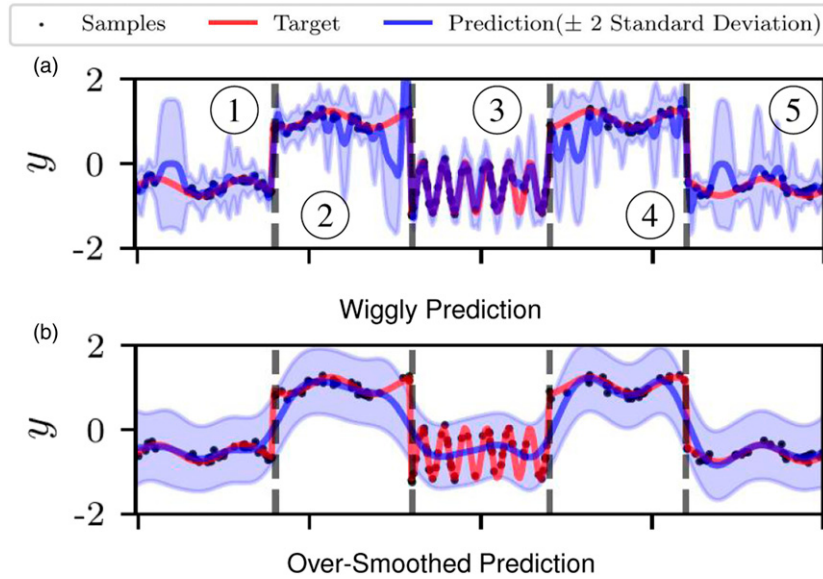


Figure 4. Learning a non-stationary function using GPR with RBF kernel. The target function in red color consists of five partitions separated by vertical dashed lines. The black dots around the function are data points. The function changes drastically in partition#3 and smoothly in the remaining partitions. The transitions between neighboring partitions are sharp. This simple function is challenging for a stationary kernel with a *single* length-scale. GPR with a stationary RBF kernel produces either the wiggly prediction shown in (a) or the over-smoothed prediction in (b). Note that in (a), the prediction in the smooth regions is rugged, and the uncertainty is over-conservative when the training data is sparse. The prediction in (b) only captures the general trend, and every input location seems equally uncertain.

At first glance, the AK looks like a heuristic composite kernel. However, the following sections explain how we design this kernel from the first principles. In short, the kernel is distilled from a generative model called AKGPR that can naturally model non-stationary processes.

4.1. A generative derivation of AK

The example in Figure 4 motivates us to consider using different length-scales at different input locations. Ideally, we need a smaller length-scale for partition#3 and larger length-scales for the others. In addition, we need to break the correlations among data points in different partitions. An ideal non-stationary model should handle these two types of non-stationarity. Many existing works model the input-dependent length-scale as a length-scale function (Lang et al., 2007; Plagemann et al., 2008a; Heinonen et al., 2016). However, parameter optimization of such models is sensitive to data distribution and parameter initialization. We propose a new approach to address this issue that *avoids learning an explicit length-scale function*. Instead, every input location can *select* among a set of GPs with different predefined primitive length-scales and *select* which training samples are used when making a prediction. This idea — selecting instead of inferring an input-independent length-scale — avoids optimization difficulties in prior work. These ideas are developed in the following sections.

4.1.1. Length-scale selection. Consider a set of M independent GPs with a set of base kernels $k_m(\mathbf{x}, \mathbf{x}')$ using predefined primitive length-scales $\{\ell_m\}_{m=1}^M$. Intuitively, if every input location can select a GP with an appropriate length-scale, the non-stationarity can be characterized well. We can achieve this by an *input-dependent weighted sum*

$$f(\mathbf{x}) = \sum_m^M w_m(\mathbf{x}) g_m(\mathbf{x}), \text{ where} \quad (10)$$

$$g_m(\mathbf{x}) \sim \mathcal{GP}(0, k_m(\mathbf{x}, \mathbf{x}')) \quad (11)$$

Here, $w_m(\mathbf{x})$ is the m -th output of a vector-valued weighting function $\mathbf{w}_\theta(\mathbf{x})$ which is parameterized by θ . We denote $\mathbf{w} = [w_1(\mathbf{x}), \dots, w_M(\mathbf{x})]^\top$.

Consider an extreme case where $\mathbf{w}$ is a “one-hot” vector—a binary vector with only one element being one and all other elements being zeros. In this case, $\mathbf{w}$ selects a single appropriate GP depending on the input location. Typically, inference techniques such as Gibbs sampling or Expectation Maximization are required for learning such discrete “assignment” parameters. We lift this requirement by continuous relaxation:

$$\mathbf{w}_\theta(\mathbf{x}) = \text{softmax}(\tilde{\mathbf{w}}_\theta(\mathbf{x})) \quad (12)$$

where $\tilde{\mathbf{w}}_\theta(\mathbf{x})$ is an arbitrary M -dimensional function parameterized by θ . Moreover, using such continuous weights has an advantage in modeling gradually changing non-stationarity, as shown in Figure 5.

Figure 6 shows that length-scale selection gives better prediction after learning from the same dataset as in Figure 4. However, when facing abrupt changes, as shown in the circled area, the model can only select a very small length-scale to accommodate the loose correlations among data. If samples near the abrupt changes are not dense enough, a small length-scale might result in a high prediction error. The following section explains how to handle abrupt changes using instance selection.

4.1.2. Instance selection. Intuitively, an input-dependent length-scale specifies each data point’s neighborhood radius that it can impact. Simply varying the radius cannot handle abrupt changes, for example, in a step function, because data sampled before and after an abrupt change should break their correlations even when

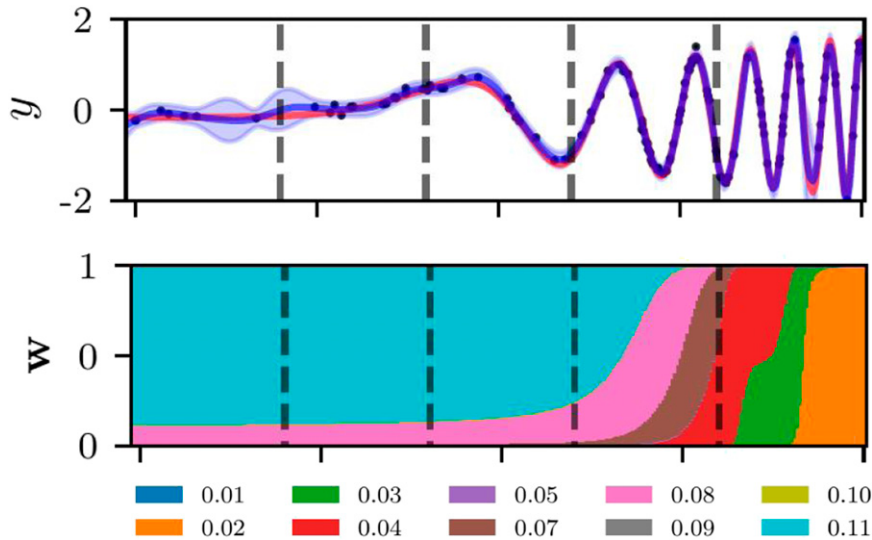


Figure 5. Learning $f(x) = x \sin(40x^4)$ with soft length-scale selection. The $\mathbf{w}$ -plot visualizes the associated weighting vector $\mathbf{w}_\theta(\mathbf{x})$ of each input location. The more vertical length a color occupies, the higher weight we assign to the GP with the corresponding length-scale. The set of predefined length-scales is color-labeled at the bottom. The learned weighting function gradually shifts its weight from smooth GPs to bumpy ones.

they are close in input locations. We need to control the *visibility* among samples: each sample learns only from other samples in the same subgroup. To this end, we associate each input with a *membership vector* $\mathbf{z} \triangleq \mathbf{z}_\phi(\mathbf{x})$ and use a dot product between two membership vectors to control visibility. Two inputs are visible to each other when they hold similar memberships. Otherwise, their correlation will be masked out:

$$\mathbf{k}_m(\mathbf{x}, \mathbf{x}') = \mathbf{z}^\top \mathbf{z}' \mathbf{k}_{\text{RBF}}(\mathbf{x}, \mathbf{x}' | \ell_m) \quad (13)$$

We can view this operation as input *dimension augmentation* where we append $\mathbf{z}$ to $\mathbf{x}$ but use a structured kernel in the joint space of $[\mathbf{x}, \mathbf{z}]$.

Discussing one-hot vectors also helps understand the effect of $\mathbf{z}$. In this case, the dot product is equal to 1 if and only if $\mathbf{z}$ and $\mathbf{z}'$ are the same one-hot vector. Otherwise, the dot product in equation (13) masks out the correlation. This way, we only use the subset of data points in the same group. To make the model more flexible and simplify the parameter optimization, we again use soft memberships:

$$\mathbf{z}_\phi(\mathbf{x}) = \text{softmax}(\tilde{\mathbf{z}}_\phi(\mathbf{x})) \quad (14)$$

Here, $\tilde{\mathbf{z}}_\phi(\mathbf{x})$ is an arbitrary M -dimensional function parameterized by ϕ .

4.1.3. The AKGPR model. Combining the two ideas, we get a new probabilistic generative model developed for non-stationary environments called Attentive Kernel Gaussian Process Regression (AKGPR). Given N inputs $\mathbf{X} \in \mathbb{R}^{N \times D}$ and targets $\mathbf{y} \in \mathbb{R}^N$, the model describes the generative process as follows. We use some shorthands for compact notation:

$$\mathbf{g}_m \triangleq [\mathbf{g}_m(\mathbf{x}_1), \dots, \mathbf{g}_m(\mathbf{x}_N)]^\top,$$

$$\mathbf{f} \triangleq [\mathbf{f}(\mathbf{x}_1), \dots, \mathbf{f}(\mathbf{x}_N)]^\top,$$

$$\mathbf{w}_m \triangleq [\mathbf{w}_m(\mathbf{x}_1), \dots, \mathbf{w}_m(\mathbf{x}_N)]^\top.$$

Here, $\mathbf{w}_m(\mathbf{x})$ is the m -output of equation (12).

- We compute the membership vector $\mathbf{z}_n$ for each input using equation (14). Plugging $\mathbf{z}_n$ and the predefined

length-scales ℓ_m into equation (13), we then compute M covariance matrices $\mathbf{K}_m$ evaluated at every pair of inputs.

- The vector $\mathbf{g}_m$ follows a multivariate Gaussian distribution $\mathcal{N}(\mathbf{0}, \mathbf{K}_m)$ according to the definition of GPs and equation (11). From equation (10), we can see that $\mathbf{f}$ is the summation of M vectors that follows affine-transformed multivariate Gaussian distributions; thus, $\mathbf{f}$ also follows Gaussian distribution:

$$\mathbf{f} = \sum_{m=1}^M \mathbf{W}_m \mathbf{g}_m \sim \mathcal{N}\left(\mathbf{0}, \sum_{m=1}^M \mathbf{W}_m \mathbf{K}_m \mathbf{W}_m^\top\right) \quad (15)$$

where $\mathbf{W}_m$ is a diagonal matrix with $\mathbf{w}_m$ being the N diagonal elements.

- Finally, we can generate the targets $\mathbf{y}$ using the Gaussian likelihood in equation (5).

The plate diagram of this generative process is shown in Figure 7. From equation (15), we observe that the generative process of AKGPR is equivalent to that of a GPR with a new kernel:

$$\mathbf{k}(\mathbf{x}, \mathbf{x}') = \sum_{m=1}^M w_m \underbrace{\mathbf{z}^\top \mathbf{z}' \mathbf{k}_{\text{RBF}}(\mathbf{x}, \mathbf{x}' | \ell_m)}_{\text{hidden in } \mathbf{K}_m} w_m' \quad (16)$$

Since $\mathbf{z}^\top \mathbf{z}'$ is independent of m , we can move it outside the summation to avoid duplicate computation.

Equation (16) is almost the same as the AK in Definition 1, except that this kernel is not normalized yet. When $\mathbf{x} = \mathbf{x}'$, the kernel value $\mathbf{k}(\mathbf{x}, \mathbf{x})$ might be greater than 1. As mentioned in Section 3.2.1, using an amplitude parameter α to adjust the scale of the kernel value is a common practice in GPR. Introducing the amplitude hyper-parameter requires the kernel to be normalized; otherwise, the interplay between the amplitude and the scaling effect of a kernel

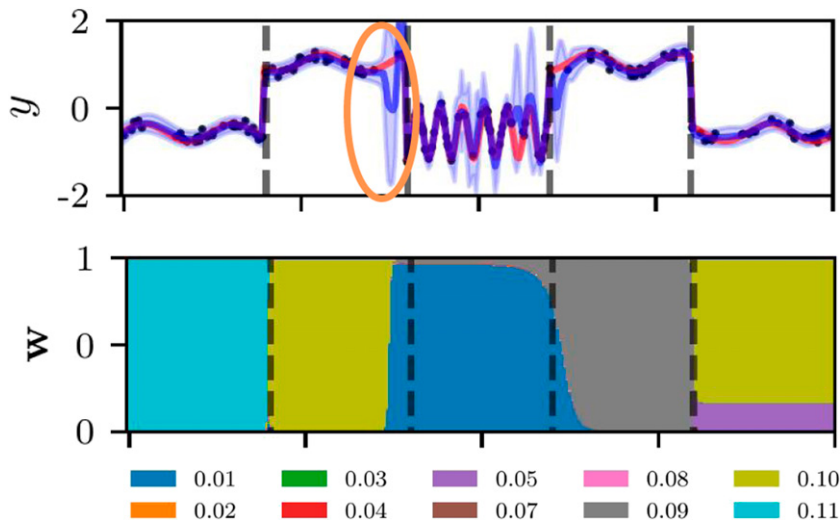


Figure 6. Prediction of length-scale selection.

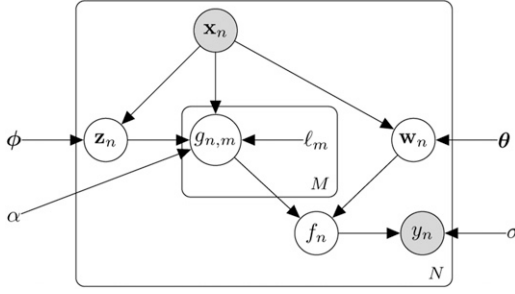


Figure 7. Plate notation of AKGPR.

before normalization makes the optimization difficult because more local optima are introduced due to the symmetries of the parameter space. We normalize $\mathbf{w}$ and $\mathbf{z}$ with ℓ^2 -norm to ensure that the maximum kernel value (when $\mathbf{x} = \mathbf{x}'$) is 1, and α is the only parameter that controls the scale of kernel value. After normalization, we now have the final version of the proposed AK in Definition 1, which can be used in any GP model. From the discussion above, we have:

Proposition 1. The AKGPR generative model is equivalent to a GPR model with the AK defined in Definition 1.

4.2. Applying AK to GPR

We use the AK with a GPR model and optimize all the parameters by maximizing the log marginal likelihood $\ln p(\mathbf{y}|\sigma, \alpha, \theta, \phi)$. Figure 8 shows the prediction results on the example from Figure 4. Now, we can accurately model the highly varying part, the smooth parts, and the abrupt changes. Compared to Figure 4, where the uncertainty mainly depends on the proximity to training samples, the AKGPR assigns higher uncertainty to the high-error locations. The better uncertainty quantification is achieved by putting more weight on the GPs with small length-scales in partition#3 and those with large length-scales in the other partitions. Note that the AKGPR switches the membership vector $\mathbf{z}$ in the circled area to mask the inter-partition correlations, which cannot be realized by length-scale selection in Figure 6. Due to this modeling advantage, the results in Figure 8 are qualitatively better than in Figure 6.

4.3. Remark on the Attentive Kernel

In this section, we discuss how to parameterize the weighting and membership functions in the AK, the computational complexity of the proposed kernel, and some details on hyper-parameter optimization of non-stationary kernels.

4.3.1. Parameterization. To instantiate an AK, we must specify the weighting function $\mathbf{w}_\theta(\mathbf{x})$ and the membership function $\mathbf{z}_\phi(\mathbf{x})$. In the experiments, we find that sharing a single neural network for length-scale selection and instance selection does not affect the performance but reduces the

number of trainable parameters and sometimes helps the training of the instance selection mechanism (see Section 5.2.2). Therefore, we use the same set of parameters $\theta = \phi$ for the two attention mechanisms and choose a simple neural network with two hidden layers (see Section 5.1.3 for more details). Using a simple neural network is an arbitrary choice for simplicity and modeling flexibility. Other parametric functions can also be used, and we leave the study of alternative parameterization to future work.

4.3.2. Computational complexity. Kernel matrix computations are typically performed in a batch manner to take advantage of the parallelism in linear algebra libraries. Figure 9 shows the computational diagram of the self-covariance matrix of an input matrix $\mathbf{X} \in \mathbb{R}^{N \times D}$ for the case where the same function parameterizes $\mathbf{w}_\theta(\mathbf{x})$ and $\mathbf{z}_\phi(\mathbf{x})$.

The computation of a cross-covariance matrix and the case where $\mathbf{w}_\theta(\mathbf{x})$ and $\mathbf{z}_\phi(\mathbf{x})$ are parameterized separately are handled similarly. We first pass $\mathbf{X}$ to a neural network with two hidden layers to get $\mathbf{W} \in \mathbb{R}^{N \times M}$ and $\mathbf{Z} \in \mathbb{R}^{N \times M}$. The computational complexity of this step is $\mathcal{O}(NDH + NH^2 + NHM)$. Then, we compute a visibility masking matrix $\mathbf{O} = \mathbf{Z}\mathbf{Z}^\top$, which takes $\mathcal{O}(N^2M)$. After getting the pairwise distance matrix ($\mathcal{O}(N^2D)$), we can compute the base kernel matrices using different length-scales ($\mathcal{O}(N^2)$). The m -th kernel matrix is scaled by the outer-product matrix of the m -th column of $\mathbf{W}$, which takes $\mathcal{O}(N^2M)$. Finally, we sum up the scaled kernel matrices and multiply the result with the visibility masking matrix to get the AK matrix ($\mathcal{O}(N^2M)$). We defer the discussion of the choices of network size H and number of base kernels M to the sensitivity analysis in Section 5.2.1. In short, these will be relatively small numbers, so the overall computational complexity is still $\mathcal{O}(N^2D)$. In practice, we find that the runtime of the AK experiments is around three times slower than that of the RBF kernel.

4.3.3. Optimization. We note that the model complexity term discussed in Section 3.2.3 is insufficient for preventing over-fitting when training non-stationary kernels for many iterations, a point also mentioned in Tompkins et al. (2020a) in their over-fitting analysis. Although the AK is more robust to over-fitting (see Section 5.2.3), we implement an incremental training scheme to improve the computational efficiency and optimization robustness when using non-stationary kernels in RIG. Specifically, we train the model on all the collected data for N_t iterations after collecting N_t samples at the t -th decision epoch, which corresponds to line 7 to line 10 in Algorithm 1.

This training scheme can be considered a rule-of-thumb early-stopping regularization. We also find that, when using a neural network in a non-stationary kernel, the initial learning rate of the network parameters should be smaller than that of other hyper-parameters. For example, when using the AK, the initial learning rates of θ or ϕ should be smaller than that of $\{\alpha, \sigma\}$.

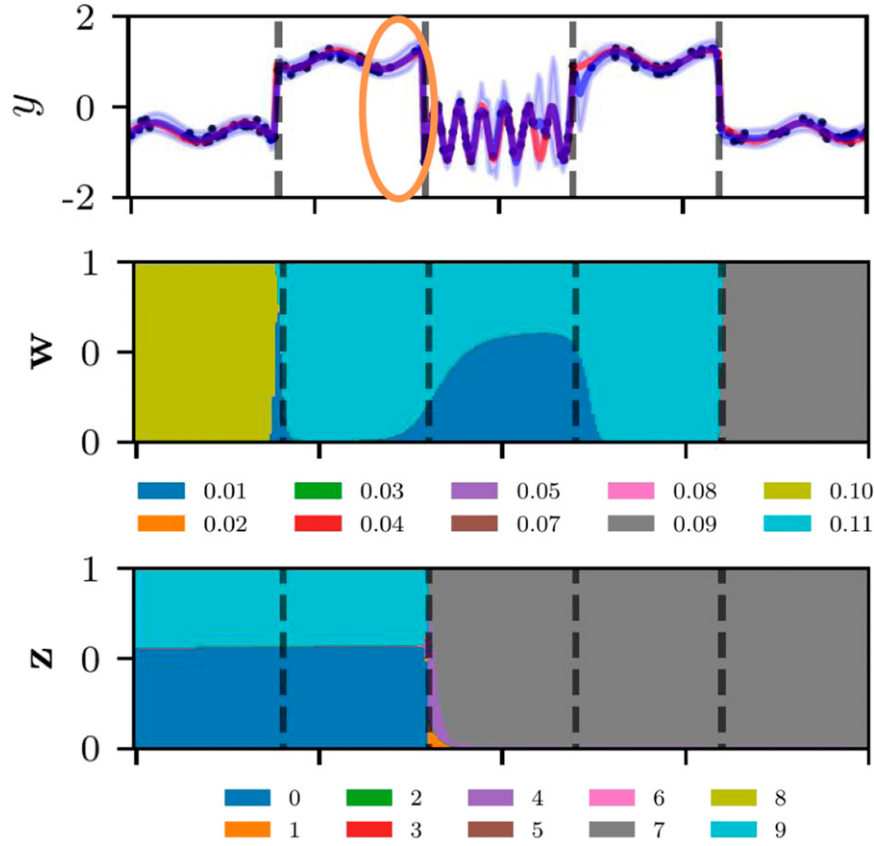


Figure 8. Learning the same function as in Figure 4 using AKGPR. A weight or membership vector is visualized as a stack of bar plots produced by its elements. Different colors represent different length-scales or dimensions of the weight or membership vector.

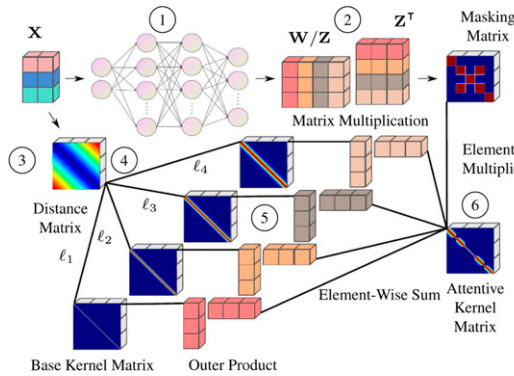


Figure 9. Computational diagram of the AK.

Another important aspect is when to start optimizing the hyper-parameters. Optimizing the parameters when the data is too sparse and not representative can lead to wrong length-scale prediction, which can bias the Informative Planning. In RIG, exploring the environment and sampling data at different locations is necessary before optimizing the hyper-parameters. In our experiments, this is done by following a predefined Bézier curve that explores the environment. An alternative way to achieve this behavior is by fixing the hyper-parameters to some appropriate values and training the model only after collecting a

certain amount of samples. This approach does not require a pilot survey of the environment. However, the user should have some prior knowledge of the target environment in order to set the initial hyper-parameters.

This training setup works well empirically, but we acknowledge that developing more principled ways to learn non-stationary GPs is an essential future direction, which is still an open research problem and has recently received increasing attention (Ober et al. 2021; Van Amersfoort et al. 2021; Lotfi et al. 2022).

Algorithm 1. Active Mapping with the AK.

Arguments: $N_{\max}, \alpha, \sigma, \{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M$
 $w_\theta(\mathbf{x}), z_\phi(\mathbf{x}), \text{strategy}$

- 1: compute normalization and standardization statistics
- 2: $\text{kernel} \leftarrow \text{AK}(\alpha, \{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M, w_\theta(\mathbf{x}), z_\phi(\mathbf{x}))$
- 3: $\text{model} \leftarrow \text{GPR}(\text{kernel}, \sigma)$
- 4: $t \leftarrow 0$
- 5: **while** $\text{model}.N_{\text{train}} < N_{\max}$ **do** ▷ sampling budget
- 6: $\mathbf{x}_{\text{info}} \leftarrow \text{strategy}(\text{model})$ ▷ informative waypoint
- 7: $\mathbf{X}_t, \mathbf{y}_t \leftarrow \text{tracking_and_sampling}(\mathbf{x}_{\text{info}})$ ▷ N_t samples
- 8: $\bar{\mathbf{X}}_t, \bar{\mathbf{y}}_t \leftarrow \text{normalize_and_standardize}(\mathbf{X}_t, \mathbf{y}_t)$
- 9: $\text{model.add_data}(\bar{\mathbf{X}}_t, \bar{\mathbf{y}}_t)$
- 10: $\text{model.optimize}(N_t)$ ▷ maximize marginal likelihood
- 11: $t \leftarrow t + 1$
- 12: **return** model

4.4. Active Mapping with the Attentive Kernel

Algorithm 1 shows how the AK can be used for Active Mapping. The system requires the following input arguments: the maximum number of training data $N_{\max}$, the initial kernel amplitude α , the initial noise scale σ , a set of M base kernels $\{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M$, functions $\mathbf{w}_\theta(\mathbf{x})$ and $\mathbf{z}_\phi(\mathbf{x})$, and a sampling strategy. First, we need to compute the statistics to normalize the inputs $\mathbf{X}$ roughly to the range $[-1, 1]$ and standardize the targets $\mathbf{y}$ to nearly have zero mean and unit variance (line 1). We can get these statistics from prior knowledge of the environment. The workspace extent is typically known, allowing the normalization statistics to be readily calculated. The target-value statistics can be rough estimates or computed from a pilot environment survey (Kemna et al., 2018). Then, we instantiate an AK and a GPR with the given parameters (lines 2–3). At each decision epoch t , the sampling strategy proposes informative waypoints by optimizing an objective function derived from the predictive uncertainty of the GPR (line 6). The robot tracks the informative waypoints and collects N_t samples along the trajectory (line 7). Note that the number of collected samples is typically larger than the number of informative waypoints. The new samples are normalized and standardized and then appended to the model’s training set (lines 8–9). Finally, we maximize the log marginal likelihood for N_t iterations (line 10). The robot repeats predicting (hidden in line 6), planning, sampling, and optimizing until the sampling budget is exceeded (line 5).

5. Experiments

We design our experiments to address the following questions.

- Q1** How does the AK compare to its stationary counterpart and other non-stationary kernels in prediction accuracy and uncertainty quantification?
- Q2** If non-stationary kernels have better uncertainty quantification capability, can we use the uncertainty for active data collection and to further improve the prediction accuracy?
- Q3** Some parameters in the AK need to be determined, that is, the number and range of the primitive length-scales and the network hyper-parameters. Are these parameters hard to tune? Is the performance of AK sensitive to its parameter settings?
- Q4** The AK consists of two ideas: length-scale selection and instance selection. Which one contributes more to the performance in the experiments?
- Q5** How does the AK compare to the other non-stationary kernels in over-fitting?

To answer **Q1**, we use random sampling experiments in Section 5.1.6 to evaluate the AK and the compared kernels. We run the random sampling experiments first because the

performance of a RIG system depends on not only the model’s prediction and uncertainty but also the informative planner. Sampling data uniformly at random (without an informative planner) provides controlled experiments to understand the effects of using different kernels. For **Q2**, we conduct both active learning (Section 5.1.7) and RIG experiments (Section 5.1.8) to disentangle the influence of the model’s uncertainty and the planner. RIG considers the physical constraints of the robot embodiment, while active learning can sample arbitrary locations. We assess the AK via sensitivity analysis, ablation study, and over-fitting analysis to address the remaining questions.

5.1. Simulated experiments

We have conducted extensive simulations in four representative environments that exhibit various non-stationary features. The elevation maps are downloaded from the NASA Shuttle Radar Topography Mission (dwtkns.com/srtm30m). Supplemental materials can be found at weizhechen.github.io/attentive_kernels.

5.1.1. Environments. Figure 10 shows the 3D perspectives of all the environments and the corresponding bird’s-eye views. Note that the 3D plots are rotated for better visualization. When comparing to the model prediction, we use the bird’s-eye map as the ground truth, and we will describe the environmental features in the 3D plots. Looking at environment N17E073 from left to right, it consists of a flat part, a mountainous area, and a rocky region with many ridges. A good non-stationary GP model should use decreasing length-scales from left to right. Also, note that the most complex area (i.e., the red region) occupies roughly one-third of the whole environment. N43W080 presents sharp elevation changes indicated by the arrows while the lakebed is virtually flat. Using a large length-scale can model most of the areas well, albeit better prediction can be achieved by sampling densely around the high-variability spots indicated by the arrows. It is worth noting that better predictions will be more evident in the visualization compared to the evaluation metrics that average across the whole environment since the important area only occupies a small portion of the environment, and the improvements might be negligible in the metrics. In N45W123, the environment has a narrow complex upper part and a smoother lower part. The size of the complex region is smaller than one-third of the environment. There is also a “river” passing through the middle. The right part of N47W124 varies drastically, while its left part is relatively flat. Loosely speaking, N47W124 has the most significant change in spatial variability, followed by N17E073 and then N45W123, so the possible improvement margins of non-stationary models in these environments should also decrease in this order. Only after discovering and sampling the two arrow-indicated spots can non-stationary models show an advantage over a stationary one in predicting environment N43W080.

5.1.2. Robot. We set the extent of the environment to 20×20 meters and simulate a planar robot that has a simple Dubins' car model $[\dot{x}_1, \dot{x}_2, \dot{v}, \dot{\omega}] = [v \cos(\omega), v \sin(\omega), a_1, a_2]$. Here, $\mathbf{x}_b \triangleq [x_1, x_2]^\top$ is the position, $\omega \in [-\pi, \pi)$ is the orientation, and $\mathbf{a} \triangleq [a_1, a_2]^\top$ is the action that represents the change in the linear velocity and angular velocity. The maximum linear velocity is set to 1 m/s, and the control frequency is 10 Hz. Although we assume perfect localization in the simulated experiments, to keep the same interface with the field experiments, we consider that the robot has achieved a goal if it is within a 0.1-meter radius. This radius is an arbitrary choice within the dimension of the robot. The robot has a single-beam range sensor that collects one noisy elevation measurement per second with a unit Gaussian observational noise. In the random and active sampling experiments, the robot can "jump" to an arbitrary sampling location to collect data, so it does not follow Dubins' car model. In the RIG experiment, the robot tracks some informative locations under the Dubins' car kinematic constraint and collects elevation measurements along its trajectory.

5.1.3. Models. The GPR takes two-dimensional sampling locations as inputs and predicts the elevation. We only allow the robot to collect 700 samples, among which the first 50 data points are collected along a pilot survey path pre-computed for the environment. As shown in Figure 11, the path is generated from a Bézier curve with 18 control points. The positions of the control points adapt to the extent of the workspace accordingly. These 50 samples are used to initialize the GPR and compute the statistics to normalize the inputs and standardize the target values. If the statistics are known in advance, the pilot survey is not necessary. One can also use a relatively large length-scale and fix the hyper-parameters of the GPR in the early stage so that the robot can explore the environment and collect diverse data for hyper-parameter optimization and statistics calculation. After normalization and standardization, we initialize the hyper-parameters to $\ell = 0.5$, $\alpha = 1.0$, $\sigma = 1.0$, $\ell_{\min} = 0.01$, $\ell_{\max} = 0.5$. We use the default PyTorch settings for initializing the network parameters. These hyper-parameters and the neural network parameters in the non-stationary kernels are jointly optimized by two Adam optimizers (Kingma and Ba 2014) with initial learning rates 0.01 and 0.001, respectively. We first run an initial optimization of all the parameters for 50 steps. The model's prediction is evaluated on a 50×50 linearly spaced evaluation grid, that is, 2500 query inputs, comparing with the ground-truth elevation values.

We compare the AK with two existing non-stationary kernels: the Gibbs kernel and Deep Kernel Learning (DKL). Since the RBF kernel is widely used in RIG, we also add this kernel as a stationary baseline. The Gibbs kernel extends the length-scale to be any positive function of the input, degenerating to an RBF kernel when using a constant length-scale function. Following Remes et al. (2018), which showed improved results, the length-scale function is modeled using a neural network instead of another Gaussian process. DKL addresses non-stationarity through input

warping. A neural network transforms the inputs to a feature space where the stationary RBF kernel is assumed to be sufficient. We use the same neural network with $2 \times 10 \times 10 \times 10$ neurons and hyperbolic tangent activation function for the AK and DKL and change the output dimension to 1 for the Gibbs kernel because it requires a scalar-valued length-scale function.

Algorithm 2. A Myopic Informative Planning Strategy

Notation

workspace: bounding box of the workspace

N_c : number of candidate locations

model: Gaussian process regression model

$\mathbf{x}_b$: robot's position

- 1: **procedure** rig(workspace, N_c , model, pose)
 - 2: $\mathbf{X}_c \sim \mathcal{U}(\text{workspace}, N_c)$ $\triangleright$ generate candidate locations uniformly at random in the workspace
 - 3: $\mu, \nu \leftarrow \text{model.predict}(\mathbf{X}_c)$ $\triangleright$ predictive mean and variance
 - 4: $\varepsilon \leftarrow \text{entropy}(\nu)$ $\triangleright$ compute predictive entropy
 - 5: $\mathbf{d} \leftarrow \text{distance}(\mathbf{X}_c, \mathbf{x}_b)$ $\triangleright$ pairwise Euclidean distances
 - 6: $\bar{\varepsilon} \leftarrow [\varepsilon - \min(\varepsilon)] / [\max(\varepsilon) - \min(\varepsilon)]$
 - 7: $\bar{\mathbf{d}} \leftarrow [\mathbf{d} - \min(\mathbf{d})] / [\max(\mathbf{d}) - \min(\mathbf{d})]$
 - 8: $\text{score} \leftarrow \bar{\varepsilon} - \bar{\mathbf{d}}$ $\triangleright$ informativeness score
 - 9: $i \leftarrow \arg \max(\text{score})$
 - 10: **return** the i -th candidate in $\mathbf{X}_c$
-

5.1.4. Sampling strategies. We use different sampling strategies in the three sets of experiments. We randomly draw a sample from a uniform distribution at each decision epoch in random sampling experiments. In active sampling experiments, we evaluate the predictive uncertainty on 1000 randomly generated candidate locations and then sample from the location with the highest predictive entropy. While the AK can be plugged into any advanced informative planner for RIG, we use the naive informative planner in Algorithm 2 for simplicity. Specifically, in addition to the predictive entropy, this planner computes the distances from these locations to the robot's position. We normalize the predictive entropy and distance to $[0, 1]$. Each candidate location's informativeness score is defined as the normalized entropy minus the normalized distance. This informativeness score considers the robot's physical constraints and encourages the robot to move to a location with high predictive uncertainty and close to the robot's current position. The planner outputs the informative waypoint with the highest score. A tracking controller is used to move the robot to the waypoint. Note that the number of collected samples N_t varies at different decision epochs depending on the distance from the robot to the informative waypoint.

5.1.5. Evaluation metrics. We care about the prediction performance and whether the predictive uncertainty can effectively reflect the prediction error. Following standard practice in the GP literature, we use *standardized mean squared error (SMSE)* and *mean standardized log loss (MSLL)* to measure these quantities (see Chapter 2.5 in

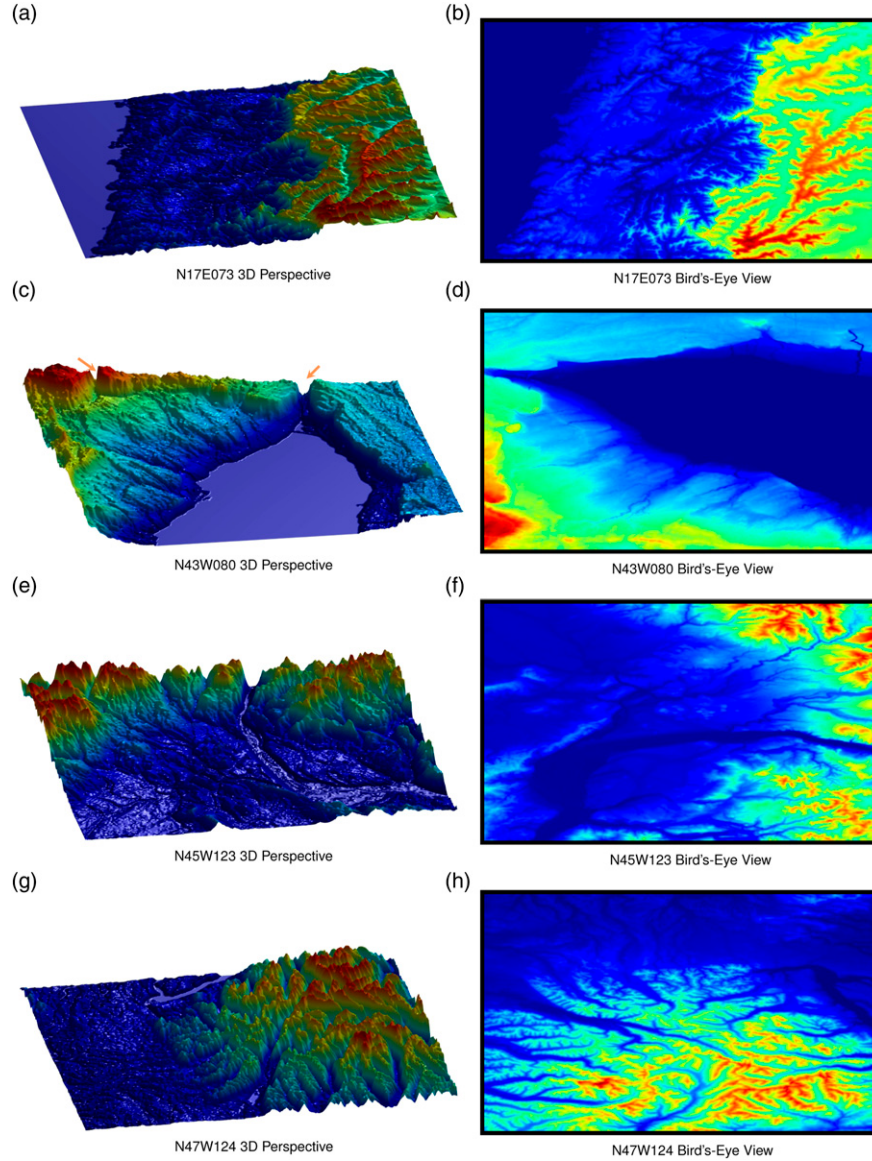


Figure 10. The four environments used in the elevation mapping tasks. Note that the 3D perspectives are rotated and rescaled to highlight the visual features of the environments.

Rasmussen and Williams (2005)). SMSE is the mean squared error divided by the variance of test targets. After this standardization, a trivial method that makes a prediction using the mean of the training targets has an SMSE of approximately 1. To take the predictive uncertainty into account, one can evaluate the negative log predictive density (NLPD), *a.k.a.*, log loss, of a test target,

$$-\ln p(y^* | \mathbf{x}^*) = \frac{\ln(2\pi\nu)}{2} + \frac{(y^* - \mu)^2}{(2\nu)}$$

where μ and ν are the mean and variance in equations (7) and (8). MSLL standardizes the log loss by subtracting the log loss obtained under the trivial model, which predicts using a Gaussian with the mean and variance of the training targets. The MSLL will be approximately zero for naive methods and negative for better methods. In the

experiments, we also measured the root-mean-square error (RMSE) and the mean absolute error (MAE). We report the mean and standard deviation of the metrics over ten runs of the experiments with different random seeds. For a more obvious quantitative comparison, we present all the benchmarking results in Tables 2–4. Each number summarizes a metric curve by averaging the curve across the x-axis, that is, the number of samples, which indicates the averaged *area under the curve*. A smaller area implies a faster drop in the curve. For all the metrics, smaller values indicate better performance.

5.1.6. Random sampling results. Table 2 gives a positive answer to Q1 firmly. The AK consistently outperforms other kernels across all the considered environments and evaluation metrics. To avoid clutter, we only visualize the SMSE and MLSS curves because they are normalized versions of

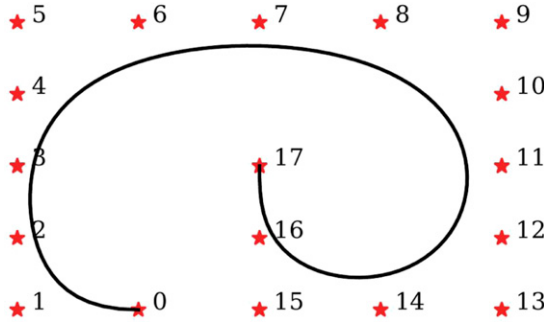


Figure 11. Pilot survey path. The red stars are control points to generate the Bézier curve.

RMSE and NLPD and the results of MAE are consistent with that of RMSE.

Figure 12 shows the metrics versus the number of collected samples of the four kernels in all environments. From the SMSE curves, we can see that the advantage of the AK (i.e., the green line) is most significant in N47W124, followed by N17E073 and then N45W123. This order complies with the changes in the spatial variability of these environments. In environment N43W080, all the lines are overlapped. N43W080 is the environment that has two spots with drastic variations. Too few random samples landed on the two spots to allow the AK to learn better prediction. That said, the MLSS curve of the AK is still outstanding in this environment. The advantage of the AK on uncertainty quantification is significant in all environments. The Gibbs kernel also has better uncertainty quantification than the RBF kernel and DKL.

Figure 13 visually compares kernels' prediction, uncertainty, and absolute error after collecting 570 samples in environment N47W124. Note that the prediction and error maps use the same color scale for easy comparison across different methods. Each uncertainty map uses its color scale—red color only indicates relatively high uncertainty *within* the map. These rules are applied to other heat maps hereafter. The AK learns more detailed environmental features (c.f., Figure 10(h)), hence obtaining better SMSE; the AK also assigns higher uncertainty to the region that is relatively more difficult to model, thus giving better MSLL. As a comparison, the RBF kernel ignores these details and assigns higher uncertainty to the sparsely sampled areas. The Gibbs kernel also has a smooth prediction in the complex region because it learns an incorrect length-scale function. Instead of assigning small length-scales to the complex region, it places them in the lower-right corner, indicated by the high uncertainty. DKL's prediction and uncertainty maps have similar patterns to the Gibbs kernel.

5.1.7. Active sampling results. The objective of active sampling experiments is to investigate whether prediction uncertainty can influence sampling towards significant areas and ultimately enhance accuracy. By comparing the SMSE results of the AK in Table 2 and Table 3, we observe a clear

improvement in accuracy when using the active sampling strategy. Specifically, the relative accuracy improvements are 9%, 15%, 23%, and 6% in N17E073, N43W080, N45W123, and N47W124, respectively, which answers **Q2**. The AK's better uncertainty quantification can further enhance prediction accuracy when the data collection strategy is guided by predictive uncertainty. However, we do not observe consistent improvements when using active sampling with the other kernels. Although they all improve the SMSE in N45W123 and N47W124, they do not improve the accuracy in the other two environments. Note that the relative improvements in N17E073 and N47W124 are smaller because the AK has already achieved good accuracy in these two environments when using random samples, so there is less room to improve than in the other two environments.

The AK still performs the best in the active sampling experiments, as seen in Table 3 and Figure 14. The SMSE curves in Figures 14 and 12 are similar, except that the advantage gap of the AK shrinks in N47W124 and increases in N43W080. We attribute the faster error drop in N43W080 to the better sample distribution. Figures 15(a) and 15(e) show that, when using the AK, more informative samples are collected in the complex regions in N43W080. Figure 15 also shows the prediction, uncertainty, and 570 samples of the three non-stationary kernels in N45W123, where all methods provide better accuracy when using active sampling strategies. The predictions of the AK and the Gibbs kernel are visually similar. The minor difference is located at the lower-right corner, where the AK learns more details (c.f., Figure 10(f)). This difference comes from the different sampling patterns. The AK samples the right part densely while the Gibbs kernel emphasizes the upper-right (c.f., Figures 15(f) and 15(g)). Also, the Gibbs kernel samples the left part of the environment very sparsely. Figure 15(d) shows that DKL is good at depicting the river. However, it connects the two "hotspots" at the upper-right corner, which is an interesting phenomenon: two non-adjacent locations are correlated. This phenomenon can be found in all the DKL predictions (see Figures 13(d) and 17(d)). The cause of this behavior is that the neural network in DKL warps the geometry of the input space, so the correlation of two given data points is no longer proportional to their distance in the original input space. It is nontrivial to explain the prediction uncertainty and sampling distribution of DKL shown in Figure 15(h).

5.1.8. Informative Planning results. The RIG experiments are more challenging than random and active sampling because once the robot decides to visit an informative waypoint, it has to collect the intermediate samples along the trajectory, so the results in Table 4 should not be compared with that of Tables 2 and 3. Given a fixed maximum number of samples, the number of decision epochs of RIG is much smaller than that of active sampling, which makes informed decisions more essential. Table 4 shows that AK consistently leads across all metrics in the four environments with the simple Informative Planning

Table 2. Random sampling benchmarking results.

Environment	Method	SMSE $\downarrow_0^1$	MSLL $\downarrow^0$	NLPD $\downarrow$	RMSE $\downarrow_0$	MAE $\downarrow_0$
N17E073	RBF	$1.33 \pm 0.03 \times 10^{-1}$	$-9.9 \pm 0.1 \times 10^{-1}$	4.59 ± 0.01	$2.33 \pm 0.03 \times 10^1$	$1.69 \pm 0.03 \times 10^1$
	AK	$1.11 \pm 0.04 \times 10^{-1}$	-1.24 ± 0.01	4.34 ± 0.01	$2.13 \pm 0.04 \times 10^1$	$1.50 \pm 0.02 \times 10^1$
	Gibbs	$1.33 \pm 0.01 \times 10^{-1}$	-1.09 ± 0.02	4.50 ± 0.03	$2.33 \pm 0.09 \times 10^1$	$1.66 \pm 0.04 \times 10^1$
	DKL	$1.37 \pm 0.06 \times 10^{-1}$	$-9.7 \pm 0.3 \times 10^{-1}$	4.62 ± 0.03	$2.37 \pm 0.05 \times 10^1$	$1.68 \pm 0.04 \times 10^1$
N43W080	RBF	$7.1 \pm 0.3 \times 10^{-2}$	-1.43 ± 0.02	3.87 ± 0.02	$1.23 \pm 0.03 \times 10^1$	8.13 ± 0.06
	AK	$6.0 \pm 0.5 \times 10^{-2}$	-1.69 ± 0.06	3.62 ± 0.06	$1.11 \pm 0.05 \times 10^1$	7.0 ± 0.2
	Gibbs	$7.2 \pm 0.4 \times 10^{-2}$	-1.48 ± 0.06	3.83 ± 0.06	$1.25 \pm 0.05 \times 10^1$	8.3 ± 0.3
	DKL	$6.6 \pm 0.8 \times 10^{-2}$	-1.49 ± 0.04	3.81 ± 0.04	$1.19 \pm 0.07 \times 10^1$	7.5 ± 0.3
N45W123	RBF	$1.65 \pm 0.07 \times 10^{-1}$	$-9.4 \pm 0.3 \times 10^{-1}$	4.37 ± 0.03	$1.97 \pm 0.04 \times 10^1$	$1.28 \pm 0.03 \times 10^1$
	AK	$1.41 \pm 0.06 \times 10^{-1}$	-1.28 ± 0.02	4.03 ± 0.02	$1.80 \pm 0.04 \times 10^1$	$1.15 \pm 0.02 \times 10^1$
	Gibbs	$1.8 \pm 0.1 \times 10^{-1}$	-1.08 ± 0.01	4.24 ± 0.02	$2.07 \pm 0.07 \times 10^1$	$1.34 \pm 0.02 \times 10^1$
	DKL	$2.0 \pm 0.1 \times 10^{-1}$	$-9.1 \pm 0.1 \times 10^{-1}$	4.41 ± 0.01	$2.18 \pm 0.07 \times 10^1$	$1.42 \pm 0.06 \times 10^1$
N47W124	RBF	$2.26 \pm 0.07 \times 10^{-1}$	$-7.2 \pm 0.1 \times 10^{-1}$	4.77 ± 0.01	$2.77 \pm 0.04 \times 10^1$	$1.97 \pm 0.02 \times 10^1$
	AK	$1.90 \pm 0.05 \times 10^{-1}$	-1.06 ± 0.01	4.43 ± 0.01	$2.53 \pm 0.03 \times 10^1$	$1.77 \pm 0.02 \times 10^1$
	Gibbs	$2.21 \pm 0.08 \times 10^{-1}$	$-7.7 \pm 0.4 \times 10^{-1}$	4.72 ± 0.05	$2.74 \pm 0.05 \times 10^1$	$1.94 \pm 0.03 \times 10^1$
	DKL	$2.34 \pm 0.08 \times 10^{-1}$	$-7.1 \pm 0.2 \times 10^{-1}$	4.78 ± 0.02	$2.82 \pm 0.05 \times 10^1$	$1.98 \pm 0.0 \pm 0.03 \times 10^1$

Table 3. Active sampling benchmarking results.

Environment	Method	SMSE $\downarrow_0^1$	MSLL $\downarrow^0$	NLPD $\downarrow$	RMSE $\downarrow_0$	MAE $\downarrow_0$
N17E073	RBF	$1.41 \pm 0.04 \times 10^{-1}$	$-9.8 \pm 0.02 \times 10^{-1}$	4.61 ± 0.02	$2.38 \pm 0.03 \times 10^1$	$1.70 \pm 0.03 \times 10^1$
	AK	$1.01 \pm 0.02 \times 10^{-1}$	-1.32 ± 0.04	4.36 ± 0.02	$2.00 \pm 0.02 \times 10^1$	$1.43 \pm 0.02 \times 10^1$
	Gibbs	$1.37 \pm 0.06 \times 10^{-1}$	-1.20 ± 0.08	4.59 ± 0.03	$2.35 \pm 0.06 \times 10^1$	$1.72 \pm 0.05 \times 10^1$
	DKL	$1.33 \pm 0.07 \times 10^{-1}$	-1.09 ± 0.05	4.59 ± 0.03	$2.32 \pm 0.06 \times 10^1$	$1.62 \pm 0.05 \times 10^1$
N43W080	RBF	$7.8 \pm 0.2 \times 10^{-2}$	-1.41 ± 0.01	3.96 ± 0.01	$1.28 \pm 0.01 \times 10^1$	9.0 ± 0.1
	AK	$5.1 \pm 0.2 \times 10^{-2}$	-1.72 ± 0.02	3.74 ± 0.03	$1.02 \pm 0.02 \times 10^1$	6.9 ± 0.1
	Gibbs	$8.0 \pm 0.6 \times 10^{-2}$	-1.48 ± 0.05	3.98 ± 0.06	$1.31 \pm 0.06 \times 10^1$	9.8 ± 0.4
	DKL	$7 \pm 1 \times 10^{-2}$	-1.6 ± 0.1	3.9 ± 0.01	$1.2 \pm 0.1 \times 10^1$	8.2 ± 0.06
N45W123	RBF	$1.47 \pm 0.04 \times 10^{-1}$	$-9.7 \pm 0.1 \pm 0.01 \times 10^{-1}$	4.36 ± 0.01	$1.85 \pm 0.02 \times 10^1$	$1.23 \pm 0.02 \times 10^1$
	AK	$1.08 \pm 0.03 \times 10^{-1}$	$-1.55 \pm 0.0 \pm 0.04$	4.16 ± 0.02	$1.57 \pm 0.03 \times 10^1$	$1.14 \pm 0.03 \times 10^1$
	Gibbs	$1.29 \pm 0.06 \times 10^{-1}$	-1.48 ± 0.05	4.30 ± 0.02	$1.73 \pm 0.04 \times 10^1$	$1.28 \pm 0.02 \times 10^1$
	DKL	$1.6 \pm 0.1 \times 10^{-1}$	-1.18 ± 0.04	4.35 ± 0.03	$1.91 \pm 0.07 \times 10^1$	$1.35 \pm 0.04 \times 10^1$
N47W124	RBF	$2.15 \pm 0.05 \times 10^{-1}$	$-7.5 \pm 0.1 \times 10^{-1}$	4.75 ± 0.01	$2.70 \pm 0.03 \times 10^1$	$1.90 \pm 0.03 \times 10^1$
	AK	$1.78 \pm 0.08 \times 10^{-1}$	-1.09 ± 0.07	4.56 ± 0.01	$2.45 \pm 0.06 \times 10^1$	$1.75 \pm 0.03 \times 10^1$
	Gibbs	$2.04 \pm 0.06 \times 10^{-1}$	$-9.9 \pm 0.5 \times 10^{-1}$	4.71 ± 0.02	$2.63 \pm 0.04 \times 10^1$	$1.86 \pm 0.03 \times 10^1$
	DKL	$2.2 \pm 0.1 \times 10^{-1}$	$-8.1 \pm 0.5 \times 10^{-1}$	4.76 ± 0.05	$2.75 \pm 0.09 \times 10^1$	$1.94 \pm 0.05 \times 10^1$

Table 4. Robotic information gathering benchmarking results.

Environment	Method	SMSE $\downarrow_0^1$	MSLL $\downarrow^0$	NLPD $\downarrow$	RMSE $\downarrow_0$	MAE $\downarrow_0$
N17E073	RBF	$1.45 \pm 0.03 \times 10^{-1}$	$-9.7 \pm 0.2 \times 10^{-1}$	4.63 ± 0.02	$2.42 \pm 0.02 \times 10^1$	$1.73 \pm 0.02 \times 10^1$
	AK	$1.14 \pm 0.04 \times 10^{-1}$	-1.27 ± 0.03	4.41 ± 0.04	$2.14 \pm 0.04 \times 10^1$	$1.51 \pm 0.02 \times 10^1$
	Gibbs	$1.43 \pm 0.07 \times 10^{-1}$	-1.16 ± 0.04	4.61 ± 0.04	$2.40 \pm 0.07 \times 10^1$	$1.76 \pm 0.06 \times 10^1$
	DKL	$1.38 \pm 0.09 \times 10^{-1}$	-1.01 ± 0.06	4.61 ± 0.04	$2.38 \pm 0.08 \times 10^1$	$1.67 \pm 0.06 \times 10^1$
N43W080	RBF	$7.7 \pm 0.4 \times 10^{-2}$	-1.40 ± 0.02	3.94 ± 0.02	$1.27 \pm 0.03 \times 10^1$	8.8 ± 0.2
	AK	$6.6 \pm 0.2 \times 10^{-2}$	-1.64 ± 0.04	3.78 ± 0.03	$1.14 \pm 0.02 \times 10^1$	7.69 ± 0.09
	Gibbs	$7.6 \pm 0.9 \times 10^{-2}$	-1.50 ± 0.05	3.91 ± 0.07	$1.25 \pm 0.07 \times 10^1$	9.0 ± 0.6
	DKL	$7.0 \pm 0.1 \times 10^{-2}$	-1.56 ± 0.07	3.85 ± 0.06	$1.19 \pm 0.08 \times 10^1$	8.1 ± 0.6
N45W123	RBF	$1.60 \pm 0.06 \times 10^{-1}$	$-9.3 \pm 0.2 \times 10^{-1}$	4.39 ± 0.02	$1.93 \pm 0.04 \times 10^1$	$1.29 \pm 0.02 \times 10^1$
	AK	$1.32 \pm 0.06 \times 10^{-1}$	-1.43 ± 0.04	4.15 ± 0.03	$1.71 \pm 0.04 \times 10^1$	$1.21 \pm 0.03 \times 10^1$
	Gibbs	$1.38 \pm 0.07 \times 10^{-1}$	-1.34 ± 0.04	4.30 ± 0.03	$1.79 \pm 0.05 \times 10^1$	$1.32 \pm 0.04 \times 10^1$
	DKL	$1.7 \pm 0.2 \times 10^{-1}$	-1.06 ± 0.08	4.41 ± 0.06	$1.99 \pm 0.09 \times 10^1$	$1.40 \pm 0.06 \times 10^1$
N47W124	RBF	$2.23 \pm 0.06 \times 10^{-1}$	$-7.4 \pm 0.1 \times 10^{-1}$	4.76 ± 0.01	$2.75 \pm 0.03 \times 10^1$	$1.94 \pm 0.02 \times 10^1$
	AK	$1.85 \pm 0.04 \times 10^{-1}$	-1.10 ± 0.03	4.48 ± 0.03	$2.50 \pm 0.03 \times 10^1$	$1.79 \pm 0.03 \times 10^1$
	Gibbs	$2.12 \pm 0.08 \times 10^{-1}$	$-9.0 \pm 0.5 \times 10^{-1}$	4.73 ± 0.03	$2.69 \pm 0.05 \times 10^1$	$1.91 \pm 0.02 \times 10^1$
	DKL	$2.36 \pm 0.06 \times 10^{-1}$	$-7.7 \pm 0.4 \times 10^{-1}$	4.78 ± 0.03	$2.83 \pm 0.03 \times 10^1$	$1.99 \pm 0.04 \times 10^1$

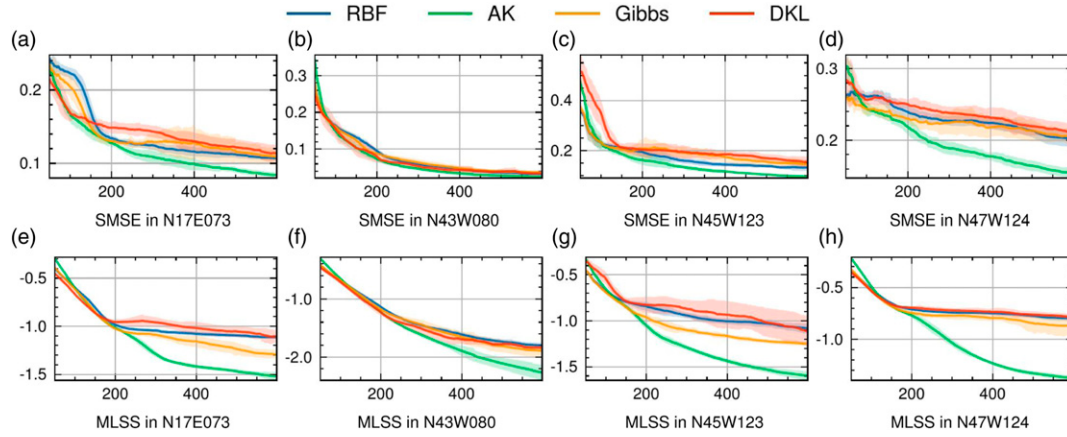


Figure 12. Random sampling metrics versus number of collected samples.

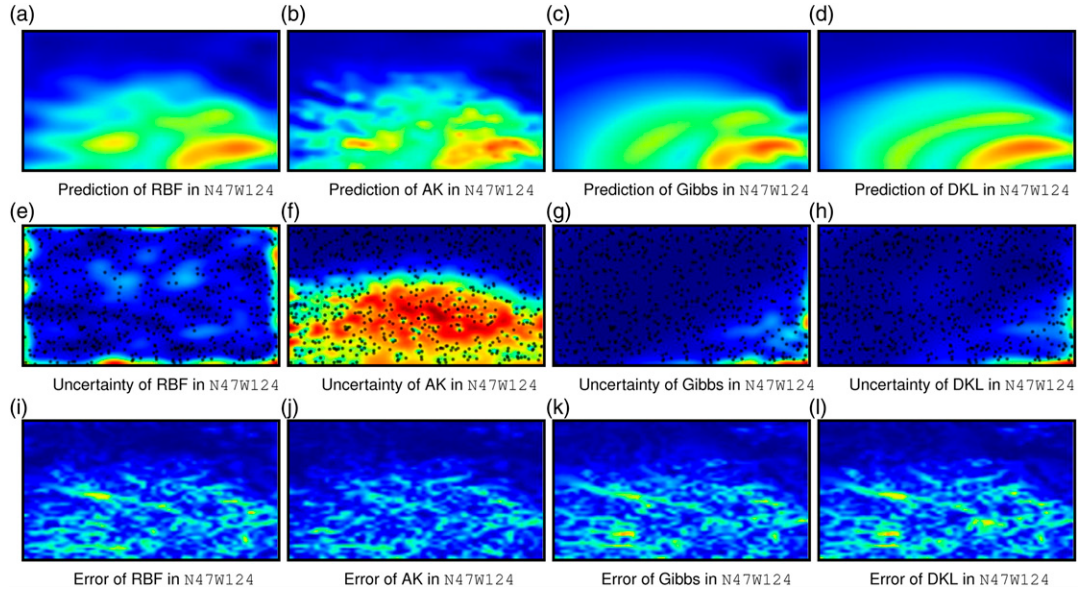


Figure 13. Snapshots of the random sampling experiments with different kernels.

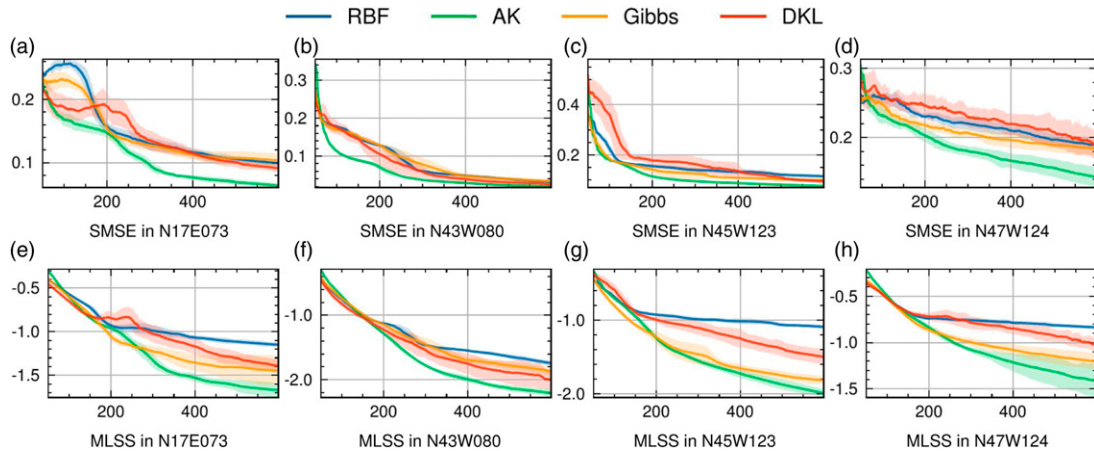


Figure 14. Active sampling metrics versus number of collected samples.

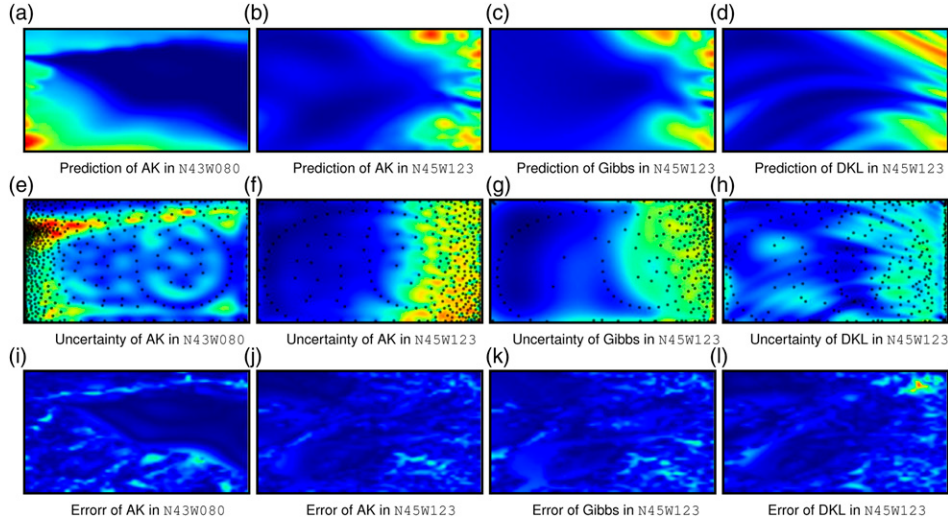


Figure 15. Snapshots of the active sampling experiments with different kernels.

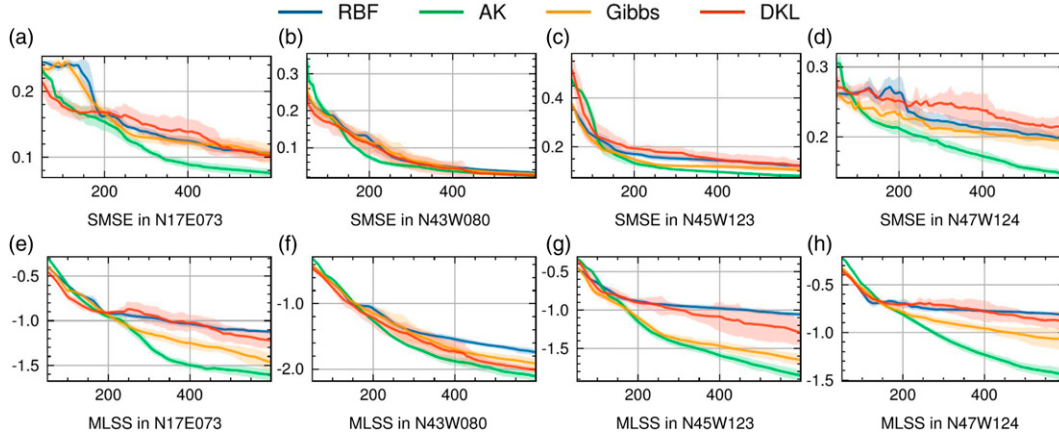


Figure 16. Robotic information gathering metrics versus number of collected samples.

strategy described in Algorithm 2. The conclusions we can draw from Figure 16 are the same as in active sampling experiments. From Figure 16, we can see again that the AK has the fastest error reduction, especially in N47W124. All non-stationary kernels have better MLSS than the stationary baseline. The AK ranks first in MSL, and the Gibbs kernel outperforms DKL.

Figure 17 is a snapshot of different methods' prediction, uncertainty, and absolute error after collecting 400 samples in N17E073. The prediction maps show that the RBF kernel misses many environmental features that non-stationary kernels can capture. We observe the following behaviors by comparing the patterns in the uncertainty maps and error maps.

- Regardless of the prediction errors, the RBF kernel gives the less-sampled area higher uncertainty, so the robot's sampling path uniformly covers the space.
- The AK assigns higher uncertainty in the regions with more significant spatial variation; thus, the sampling path focuses more on the complex region.

- The Gibbs kernel also has higher uncertainty in the rocky region but does not assign high uncertainty to the lower right. Therefore, the sampling path concentrates on the upper-right corner and misses some high-error spots at the bottom.
- When using DKL, the robot also samples the upper-right corner densely, and the prediction error at the bottom of the map is the largest across different methods. However, DKL places high uncertainty in the high-area region, which can guide the robot to visit these spots later.

5.2. Further evaluation and analysis

We evaluate the AK under different parameter settings for sensitivity analysis and compare four variants of the AK for ablation study. The challenges of learning the model are also discussed in this section.

5.2.1. Sensitivity analysis. We use the same experiment configurations as the main experiments in the sensitivity analysis but only run the random sampling strategy. In each analysis, we

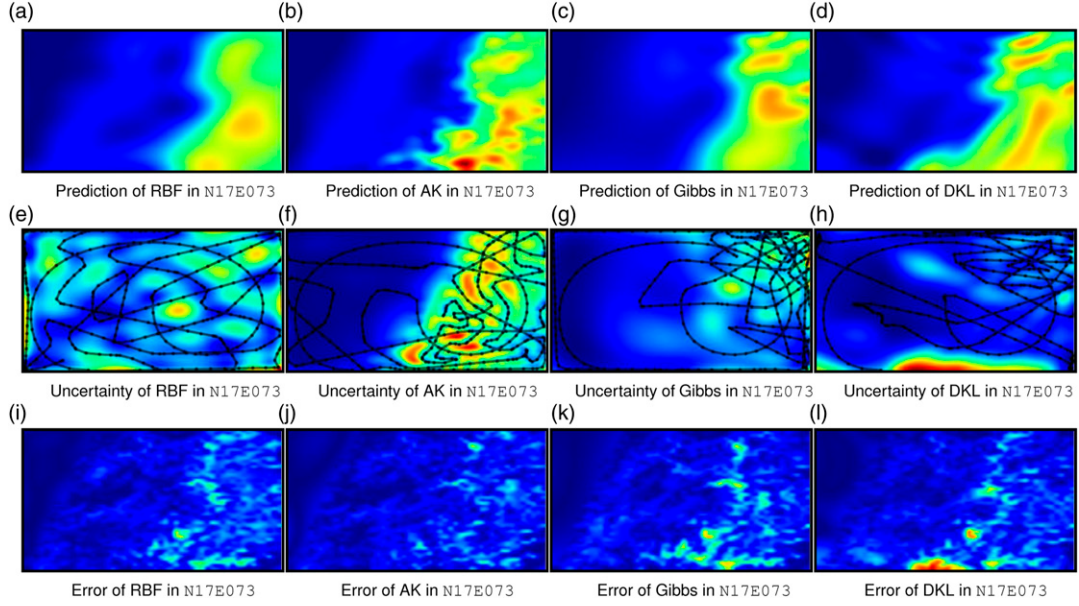


Figure 17. Snapshots of the robotic information gathering experiments with different kernels.

only change one target parameter to different settings and keep all the other parameters fixed. Figure 18 presents the sensitivity analysis results of the number of base kernels M , which should be larger than 2. Increasing M brings better performance, albeit with a diminishing return and higher computational complexity. Choosing a number in the range of [5, 10] is a good trade-off between performance and computational efficiency.

Figure 19 shows that the AK is not sensitive to the number of hidden units in the neural network as long as H is not too small. When $H = 2$, the uncertainty quantification ability decreases, as indicated by the blue MSL curve. In this case, the AK can only blend the minimum and maximum primitive length-scales, and the instance selection mechanism can only use a two-dimensional membership vector.

Smaller $\ell_{\min}$ yields better performance, as shown in Figure 20, albeit with a diminishing improvement. The blue and green lines overlap, meaning that the advantage is negligible when choosing a minimum length-scale smaller than 0.01. If the inputs are normalized to $[-1, 1]$, setting the minimum primitive length-scale to 0.01 is appropriate. It is worth noting that this is the minimum primitive length-scale for the length-scale selection component. It does not mean that the AK can only learn the minimum correlation corresponding to this minimum length-scale because the instance selection component can further decrease the kernel values.

As shown in Figure 21, the AK is robust to the choice of the maximum length-scale as long as it is not too small, for example, 0.2 or 0.3. If the inputs are normalized to $[-1, 1]$, choosing a value in the range [0.5, 1.0] is reasonable.

To conclude, these results are positive indicators of addressing Q3: the AK has robust performance to various parameter settings and does not require laborious parameter tuning.

5.2.2. Ablation study. We compare four AK variants in the ablation study via random sampling experiments. Full means

the AK presented in the paper, Weight represents the AK with only length-scale selection, Mask stands for instance selection alone, and NNx2 uses two separated neural networks to parameterize the similarity attention and visibility attention independently. Figure 22 shows that using only the instance selection component deteriorates the performance significantly, so the length-scale selection component contributes more to the performance, which answers Q4. We do not observe obvious performance change after dropping the instance selection component. Nonetheless, as illustrated in Figure 8, we expect instance selection to provide better modeling of sharp transitions. Since instance selection improves the prediction only in a small region, the improvement might be subtle in the aggregated evaluation metrics. With our current training scheme, using two separate neural networks does not provide a better performance, and one of the MSL curves is surpassed by the one-network version (Figure 22(f)). The two-network implementation might show its strength with a more refined approach to parameter training.

5.2.3. Over-fitting analysis. Non-stationary kernels can enhance the modeling flexibility of GPR, but they are also more susceptible to over-fitting. This can lead to degraded prediction accuracy and uncertainty estimates. To evaluate the robustness of non-stationary kernels, we present an over-fitting analysis in N17E073 and the Mount St Helens environment. The latter is referred to as the volcano environment hereafter. We sample 600 training data from the environment uniformly at random. All the training configurations are the same as in Section 5.1.3, except for the number of optimization iterations. We train all the models for 2000 iterations and evaluate the prediction on the training set and a 100×100 test grid at each optimization step. Figure 23 shows the training and test MSL. In some environments, as

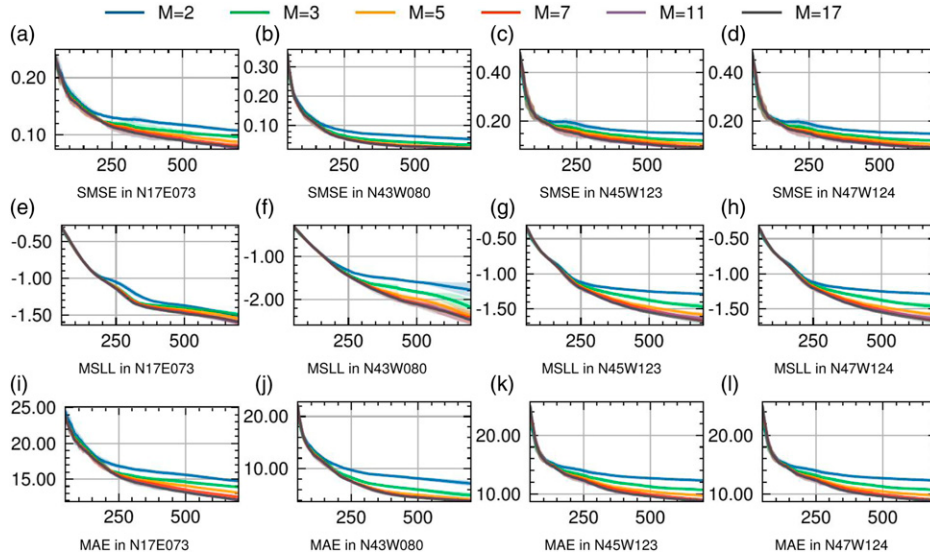


Figure 18. Sensitivity analysis of the number of base kernels M .

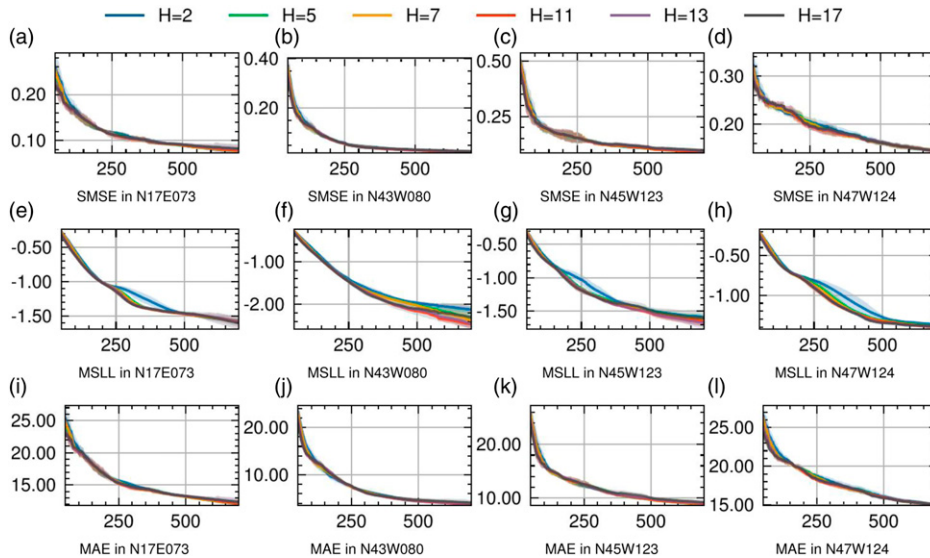


Figure 19. Sensitivity analysis of the number of hidden units H .

shown in Figures 23(a) and 23(b), the AK is fairly robust, while the Gibbs kernel and DKL show a clear over-fitting trend—the training MSL goes down while the test MSL goes up. However, as shown in Figures 23(c) and 23(d), all the non-stationary kernels suffer from over-fitting in some environments, such as N17E073. To mitigate this issue, after collecting one new sample, the optimizer takes only one gradient step on the whole dataset. This heuristic training scheme works well in practice. We have tried to optimize the model for more iterations at each decision epoch. All the non-stationary kernels give poor prediction (the AK is still more robust in this case), and the issue persists even after collecting more data. Overall, the answer to **Q5** is positive: the AK is more robust to over-fitting than other non-stationary kernels, but it can still over-fit in some environments. Developing

more advanced training schemes to mitigate over-fitting is an essential future direction.

5.3. Field experiment

The proposed AK is demonstrated in a RIG task—active elevation, *a.k.a.* bathymetric mapping for underwater terrain. Figure 24(a) shows our robot working in the environment. The goal is to explore an a priori unknown quarry lake and build an elevation map of the underwater terrain. There are two reasons for choosing this task. First, the underwater terrain is static, so the ground-truth environment is available by aggregating the sampled data across different field experiment trials after offsetting the water surface level. Second, the underwater terrain in our target environment has a

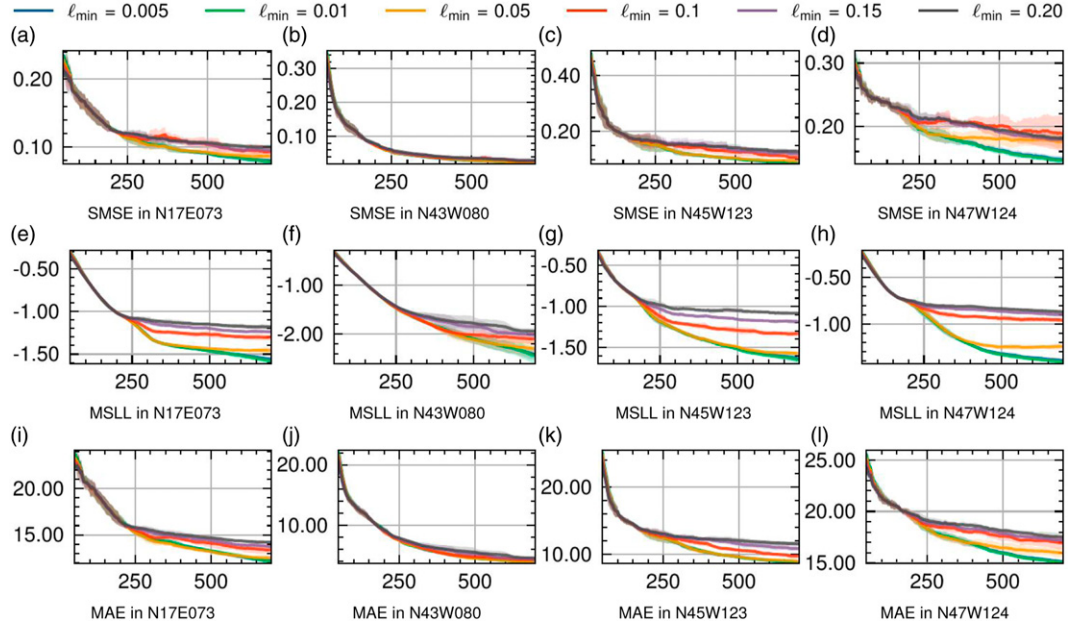


Figure 20. Sensitivity analysis of the minimum primitive length-scale $\ell_{\min}$.

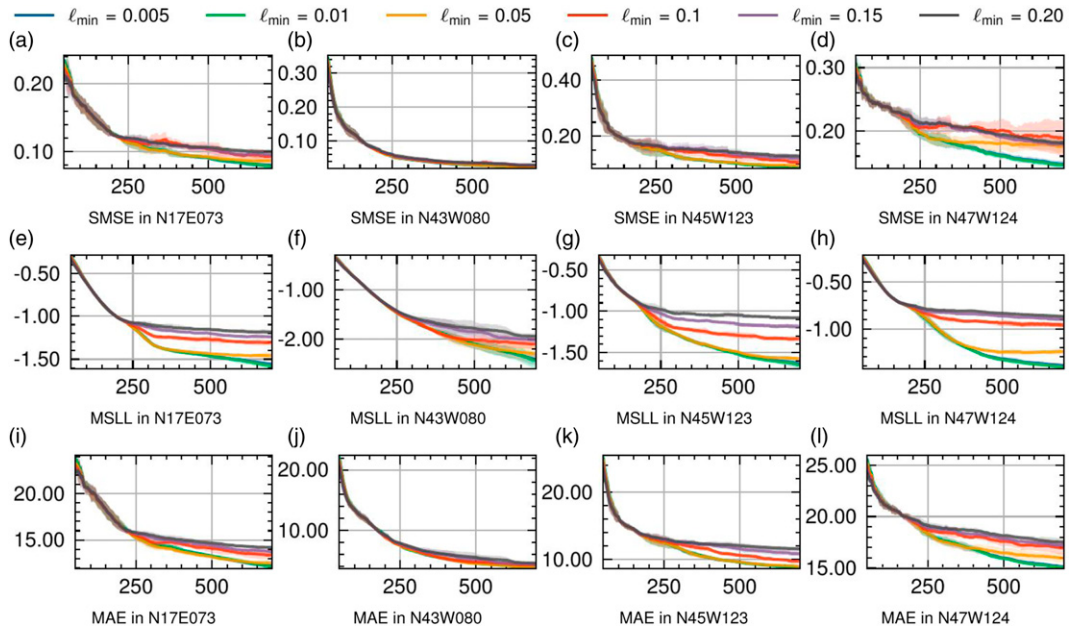


Figure 21. Sensitivity analysis of the maximum primitive length-scale $\ell_{\max}$.

clear separation between “interesting” regions and “boring” areas, which makes it an ideal testbed for RIG with non-stationary GPs.

5.3.1. Target environment. The target environment is a quarry lake formed by seeped-in groundwater and precipitation since mining and quarrying have been suspended for a long time. The floor of the quarry lake is complex in that there are many submerged quarry stones and even abandoned equipment. Our goal is to build an elevation map within the

workspace, that is, the white rectangle shown in Figure 24(c), with a small number of samples. The workspace is 80×88 meters. We chose this workspace because the central part is relatively flat, while the two circled areas have interesting spatial variations. We can vaguely see the environmental features in these circled spots from the satellite imagery.

5.3.2. Hardware setup. We deploy the Autonomous Surface Vehicle (ASV) shown in Figure 24(b). The robot has a single-beam sonar pointing downward to collect depth measurements

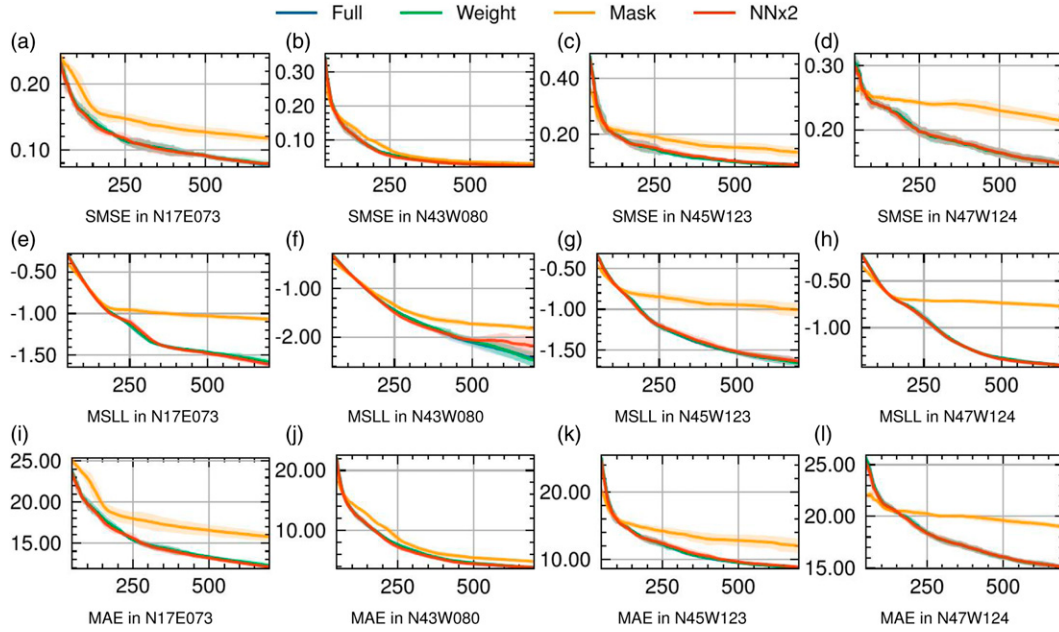


Figure 22. Results of the four variants in the ablation study.

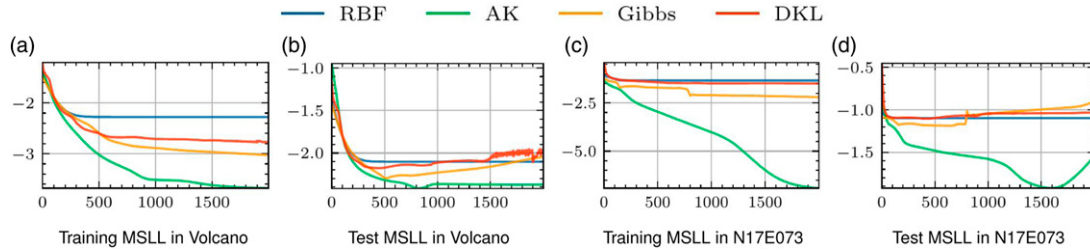


Figure 23. Results of the over-fitting analysis in the *volcano* environment introduced in Figure 1 and N17E073.

and a DJI Manifold 2-C computer for onboard computation. The sonar is the Ping Sonar Altimeter and Echosounder from BlueRobotics. Its maximum measurement distance is 50 meters underwater, and the beam width is 30 degrees. It comes with a Python software interface, and we implemented its ROS driver, which is publicly available at github.com/weizhe-chen/single_beam_sonar. The ASV from Clearpath Robotics has a built-in Extended Kalman Filter (EKF) localization module that fuses the GPS signals and the UM6 Inertial Measurement Unit (IMU) data. The robot also has an embedded WiFi router for communication in the field. The ASV is 1.3, 0.94, and 0.34 meters in length, width, and height, respectively, and is actuated by two thrusters at the rear. It is a differential-drive robot, but its thrusters' maximum forward spinning speed is faster than the backward one. We restrict the maximum linear velocity to 0.7 meters per second and send linear and angular velocities to the robot to track an informative waypoint using a PD controller available at github.com/weizhe-chen/tracking_pid. Since the localization is unreliable, the robot only needs to reach a two-meter-radius circle centered at the waypoint.

5.3.3. Results. Figure 24 shows the snapshots of the model prediction, uncertainty, and sampling path at different

stages. We can see that the prediction uncertainty is effectively reduced after sampling. Most of the samples (i.e., yellow dots) are collected in critical regions with drastic elevation variations. Such a biased sampling pattern allows the robot to model the general trend of smooth regions with a small number of samples while capturing the characteristic environmental features at a fine granularity.

6. Limitations and future work

Although the AK has the same asymptotic computational complexity as the RBF kernel, its empirical runtime is slower than that of the RBF kernel. Thus, one important future work is to speed up the computation. We leverage heuristics to train the non-stationary kernels in our experiments, which can be improved by a more principled training scheme in the future. Using a stationary kernel in non-stationary environments is just one example of *model misspecification*. Investigating the influence of other types of model misspecification on RIG is interesting. For example, the Gaussian likelihood assumes no sensing outliers, and the observational noise scale is the same everywhere. Developing proper ways to handle sensing outliers and

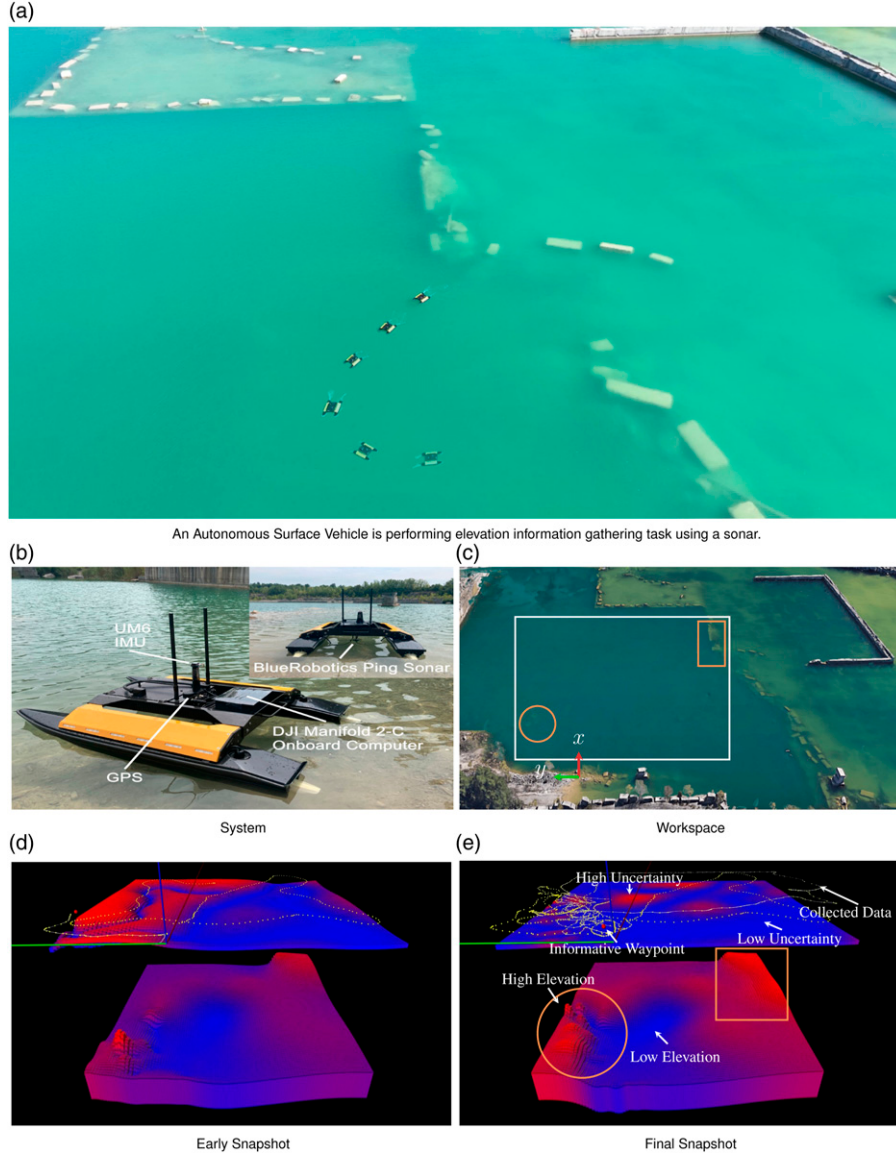


Figure 24. An active elevation mapping field experiment: (a) illustrates the physical space the ASV is mapping and (b) shows the ASV and its components. (c) shows the rectangular workspace for the elevation mapping experiment. We can see two areas with significant elevation features in two highlighted areas, but other regions are opaque. (d) and (f) are two snapshots of the GPR prediction in the rectangular workspace, with the predictive-mean map at the bottom and the uncertainty map (i.e., standard deviation) at the top. In (d), the lower part shows that some features of the highlighted areas have already been detected. The uncertainty is significant on the left side of the workspace and a smaller region in the top right. (e) shows the snapshot at the end of the experiment. The ASV has extensively explored the lower left portion and has a detailed estimate of its elevation map. The smooth portion in the middle shows differences in elevation, which are not visible in the satellite image. The remaining areas of high uncertainty are at boundaries of elevation changes in that region and the top right.

modeling *heteroscedastic noise* can be important future work for RIG. We only tried neural-network parameterization for the weighting function and the membership function. Comparing different parameterization methods for the AK is also valuable. Although we have only showcased the efficacy of AK in elevation mapping tasks, it has potential to benefit other applications such as 3D reconstruction, autonomous exploration and inspection, as well as search and rescue. Exploring its utility in these domains would be interesting. Additionally, while we focused on

non-stationary kernels in the spatial domain, developing spatiotemporal kernels is crucial for RIG in dynamic environments.

7. Conclusion

In this paper, we investigate the uncertainty quantification of probabilistic models, which is decisive for the performance of RIG but has received little attention. We present a family of non-stationary kernels called the Attentive Kernel, which

is simple and robust and can extend any stationary kernel to a non-stationary one. An extensive evaluation of elevation mapping tasks shows that AK provides better accuracy and uncertainty quantification than the two existing non-stationary kernels and the stationary RBF kernel. The improved uncertainty quantification guides the Informative Planning algorithms to collect more valuable samples around the complex area, thus further reducing the prediction error. A field experiment demonstrates that AK enables an ASV to collect more samples in important sampling locations and capture the salient environmental features. The results indicate that misspecified probabilistic models significantly affect RIG performance, and GPR with AK provides a good choice for non-stationary environments.

Acknowledgements

We thank Durgakant Pushp and Mahmoud Ali for their help in conducting the field experiment. We are also grateful for the computational resources provided by the Amazon AWS Machine Learning Research Award. The constructive comments by the anonymous conference reviewers are greatly appreciated.

Declaration of conflicting interests

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding

The author(s) disclosed receipt of the following financial support for the research, authorship, and/or publication of this article: This work was supported by RI: Small: Exploiting Symmetries of Decision-Theoretic Planning for Autonomous Vehicles (grant no 2006886), III: Small: Algorithms and Theoretical Foundations for Approximate Bayesian Inference in Machine Learning (grant no 1906694), and CAREER: Autonomous Live Sketching of Dynamic Environments by Exploiting Spatiotemporal Variations (grant no 2047169).

ORCID iDs

Weizhe Chen  <https://orcid.org/0000-0001-9068-4247>

Lantao Liu  <https://orcid.org/0000-0002-6796-6817>

References

- Abraham I and Murphey TD (2019) Active learning of dynamics for data-driven control using Koopman operators. *IEEE Transactions on Robotics* 35(5): 1071–1083.
- Aloimonos J, Weiss I and Bandyopadhyay A (1988) Active vision. *International Journal of Computer Vision* 1(4): 333–356.
- Arora A, Furlong PM, Fitch R, et al. (2019) Multi-modal active perception for information gathering in science missions. *Autonomous Robots (AURO)* 43(7): 1827–1853.
- Atanasov N, Sankaran B, Ny JL, et al. (2014) Nonmyopic view planning for active object classification and pose estimation. *IEEE Transactions on Robotics (T-RO)* 30(5): 1078–1090.
- Atkinson AC (1996) The usefulness of optimum experimental designs. *Journal of the Royal Statistical Society: Series B (Methodological)* 58(1): 59–76.
- Bai S, Shan T, Chen F, et al. (2021) Information-driven path planning. *Current Robotics Reports* 2: 1–12.
- Bai S, Wang J, Chen F, et al. (2016) Information-theoretic exploration with Bayesian optimization. In: IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), Daejeon, Korea (South), 09–14 October 2016, pp. 1816–1822.
- Bajcsy R (1988) Active perception. *Proceedings of the IEEE* 76(8): 966–1005.
- Bajcsy R, Aloimonos Y and Tsotsos JK (2018) Revisiting active perception. *Autonomous Robots* 42(2): 177–196.
- Best G, Cliff OM, Patten T, et al. (2019) Dec-MCTS: decentralized planning for multi-robot active perception. *The International Journal of Robotics Research (IJRR)* 38(2–3): 316–337.
- Binney J, Krause A and Sukhatme GS (2013) Optimizing waypoints for monitoring spatiotemporal phenomena. *The International Journal of Robotics Research (IJRR)* 32(8): 873–888.
- Bircher A, Kamel M, Alexis K, et al. (2016) Receding horizon” next-best-view” planner for 3d exploration. In: 2016 IEEE International Conference on Robotics and Automation (ICRA). IEEE, pp. 1462–1468.
- Bircher A, Kamel M, Alexis K, et al. (2018) Receding horizon path planning for 3D exploration and surface inspection. *Autonomous Robots* 42(2): 291–306.
- Biyik E, Huynh N, Kochenderfer M, et al. (2020) Active preference-based Gaussian process regression for reward learning. In: Proceedings of Robotics: Science and Systems, Corvallis, Oregon, USA, July 12–16, 2020, pp. 1–10. DOI: [10.15607/RSS.2020.XVI.041](https://doi.org/10.15607/RSS.2020.XVI.041)
- Borghi G and Caglioti V (1998) Minimum uncertainty explorations in the self-localization of mobile robots. *IEEE Transactions on Robotics and Automation* 14(6): 902–911.
- Bry A and Roy N (2011) Rapidly-exploring random belief trees for motion planning under uncertainty. In: 2011 IEEE international conference on robotics and automation. IEEE, pp. 723–730.
- Bui TD, Yan J and Turner RE (2017) A unifying framework for Gaussian process pseudo-point approximations using power expectation propagation. *The Journal of Machine Learning Research (JMLR)* 18(1): 3649–3720.
- Buisson-Fenet M, Solowjow F and Trimpe S (2020) Actively learning Gaussian process dynamics. In: *Learning for Dynamics and Control (L4DC)*. pp. 5–15.
- Cai P, Luo Y, Hsu D, et al. (2021) HyP-DESPOT: a hybrid parallel algorithm for online planning under uncertainty. *The International Journal of Robotics Research* 40(2–3): 558–573.
- Calandra R, Peters J, Rasmussen CE, et al. (2016) Manifold Gaussian processes for regression. In: International Joint Conference on Neural Networks (IJCNN), Vancouver, BC, Canada, 24–29 July 2016, pp. 3338–3345.
- Cao C, Zhu H, Choset H, et al. (2021) TARE: a hierarchical framework for efficiently exploring complex 3D environments. In: Proceedings of Robotics: Science and Systems, Virtual, 12–16 July 2021, pp. 1–9.
- Cao N, Low KH and Dolan JM (2013) Multi-robot informative path planning for active sensing of environmental

- phenomena: a tale of two algorithms. In: International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS), Saint Paul, MN, USA, 06–10 May 2013, pp. 7–14.
- Capone A, Noske G, Umlauf J, et al. (2020) Localized active learning of Gaussian process state space models. In: Conference on Learning for Dynamics and Control (L4DC), Virtual, 10–11 June 2020, 120. pp. 490–499.
- Carrillo H, Latif Y, Rodriguez-Arevalo ML, et al. (2015) On the monotonicity of optimality criteria during exploration in active SLAM. In: 2015 IEEE International Conference on Robotics and Automation (ICRA). IEEE, pp. 1476–1483.
- Chaplot DS, Parisotto E and Salakhutdinov R (2018) Active neural localization. In: International Conference on Learning Representations, Vancouver, BC, Canada, 30 April–03 May 2018, pp. 1–15.
- Charrow B, Kahn G, Patil S, et al. (2015a) Information-theoretic planning with trajectory optimization for dense 3D mapping. In: Proceedings of Robotics: Science and Systems, Rome, Italy, 13–17 July 2015, pp. 1–10.
- Charrow B, Liu S, Kumar V, et al. (2015b) Information-theoretic mapping using Cauchy-Schwarz quadratic mutual information. In: IEEE International Conference on Robotics and Automation (ICRA), Seattle, WA, USA, 26–30 May 2015, pp. 4791–4798.
- Chen W, Khardon R and Liu L (2022) AK: Attentive Kernel for information Gathering. In: Proceedings of Robotics: Science and Systems, New York City, NY, USA, 27 June–1 July 2022, pp. 1–16.
- Chen W and Liu L (2019) Pareto Monte Carlo tree search for multi-objective informative planning. In: Robotics: Science and Systems (RSS), Messe Freiburg, Germany, 22–26 June 2019, pp. 1–10.
- Choudhury S, Bhardwaj M, Arora S, et al. (2018) Data-driven planning via imitation learning. *The International Journal of Robotics Research (IJRR)* 37(13–14): 1632–1672.
- Chua K, Calandra R, McAllister R, et al. (2018) Deep reinforcement learning in a handful of trials using probabilistic dynamics models. *Advances in Neural Information Processing Systems* 31: 1–12.
- Connolly C (1985) The determination of next best views. In: Proceedings. 1985 IEEE International Conference on Robotics and Automation, St. Louis, MO, USA, 25–28 March 1985, volume 2. IEEE, pp. 432–435.
- Cully A, Clune J, Tarapore D, et al. (2015) Robots that can adapt like animals. *Nature* 521(7553): 503–507.
- Dang AT (2020) *Resilient Large-Scale Informative Path Planning for Autonomous Robotic Exploration*. PhD Thesis, University of Nevada, Reno.
- Danielczuk M, Balakrishna A, Brown D, et al (2021) Exploratory grasping: asymptotically optimal algorithms for grasping challenging polyhedral objects. In: Conference on Robot Learning. PMLR, pp. 377–393.
- de Avila Belbute-Peres F, Smith K, Allen K, et al. (2018) End-to-end differentiable physics for learning and control. *Advances in Neural Information Processing Systems* 31: 1–12.
- Deisenroth M and Ng JW (2015) Distributed Gaussian processes. In: International Conference on Machine Learning (ICML), Lille, France, 06–11 July 2015, pp. 1481–1490.
- Dhawale A and Michael N (2020) Efficient parametric multi-fidelity surface mapping. *Robotics: Science and Systems (RSS)* 2: 5.
- Di Caro GA and Yousaf AWZ (2021) Multi-robot informative path planning using a leader-follower architecture. In: 2021 International Conference on Robotics and Automation (ICRA), Xi'an, China, 30 May–05 June 2021, pp. 10045–10051.
- Doherty K, Wang J and Englot B (2016) Probabilistic map fusion for fast, incremental occupancy mapping with 3D Hilbert maps. In: IEEE International Conference on Robotics and Automation (ICRA), Stockholm, Sweden, 16–21 May 2016, pp. 1011–1018.
- Du T, Hughes J, Wah S, et al. (2021) Underwater soft robot modeling and control with differentiable simulation. *IEEE Robotics and Automation Letters* 6(3): 4994–5001.
- Dunbabin M and Marques L (2012) Robots for environmental monitoring: significant advancements and applications. *IEEE Robotics & Automation Magazine (RAM)* 19(1): 24–39.
- Ebadi K, Chang Y, Palieri M, et al. (2020) LAMP: Large-scale autonomous mapping and positioning for exploration of perceptually-degraded subterranean environments. In: 2020 IEEE International Conference on Robotics and Automation (ICRA). IEEE, pp. 80–86.
- Fernández IMR, Denniston CE, Caron DA, et al. (2022) Informative path planning to estimate quantiles for environmental analysis. *IEEE Robotics and Automation Letters* 7(4): 10280–10287.
- Ferrari S and Wettergren TA (2021) *Information-Driven Planning and Control*. Cambridge, Massachusetts, USA: MIT Press.
- Flaspohler G, Preston V, Michel APM, et al. (2019) Information-guided robotic maximum seek-and-sample in partially observable continuous environments. *IEEE Robotics and Automation Letters (RA-L)* 4(4): 3782–3789.
- Fox D, Burgard W and Thrun S (1998) Active markov localization for mobile robots. *Robotics and Autonomous Systems* 25(3–4): 195–207.
- Freeman CD, Frey E, Raichuk A, et al (2021) Brax - A Differentiable Physics Engine for Large Scale Rigid Body Simulation. In: Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 1), Virtual, 06–14 December 2021, pp. 1–13.
- Ghaffari Jadidi M, Valls Miro J and Dissanayake G (2018) Gaussian processes autonomous mapping and exploration for range-sensing mobile robots. *Autonomous Robots (AURO)* 42(2): 273–290.
- Gibbs MN (1997) Bayesian Gaussian processes for regression and classification. *PhD Thesis*, University of Cambridge.
- Girdhar Y, Giguère P and Dudek G (2014) Autonomous adaptive exploration using realtime online spatiotemporal topic modeling. *The International Journal of Robotics Research (IJRR)* 33(4): 645–657.
- Guizilini V and Ramos F (2017) Learning to reconstruct 3D structures for occupancy mapping. In: Robotics: Science and Systems (RSS), Cambridge, MA, USA, 12–16 July 2017, pp. 1–10.

- Guizilini V and Ramos F (2019) Variational hilbert regression for terrain modeling and trajectory optimization. *The International Journal of Robotics Research (IJRR)* 38(12–13): 1375–1387.
- Gupta K, Li PZX, Karaman S, et al. (2021) Efficient computation of map-scale continuous mutual information on chip in real time. In: 2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). IEEE, pp. 6464–6470.
- Heiden E, Macklin M, Narang YS, et al. (2021a) DiSECT: a differentiable simulation engine for autonomous robotic cutting. In: Robotics: Science and Systems, Virtual, 12–16 July 2021, pp. 1–20.
- Heiden E, Millard D, Coumans E, et al. (2021b) NeuralSim: augmenting differentiable simulators with neural networks. In: 2021 IEEE International Conference on Robotics and Automation (ICRA), Xi'an, China, 30 May–05 June 2021, pp. 9474–9481.
- Heinonen M, Mannerström H, Rousu J, et al. (2016) Non-stationary Gaussian process regression with Hamiltonian Monte Carlo. *The Journal of Machine Learning Research (JMLR)* 51: 732–740.
- Hitz G, Galceran E, Garneau MÈ, et al. (2017) Adaptive continuous-space informative path planning for online environmental monitoring. *Journal of Field Robotics (JFR)* 34(8): 1427–1449.
- Hoang TN, Hoang QM and Low BKH (2015) A unifying framework of anytime sparse Gaussian process regression models with stochastic variational inference for big data. In: International Conference on Machine Learning (ICML), Lille, France, 06–11 July 2015, pp. 569–578.
- Hollinger GA, Englot B, Hover FS, et al. (2013) Active planning for underwater inspection and the benefit of adaptivity. *The International Journal of Robotics Research (IJRR)* 32(1): 3–18.
- Hollinger GA and Sukhatme GS (2014) Sampling-based robotic information gathering algorithms. *The International Journal of Robotics Research (IJRR)* 33(9): 1271–1287.
- Hu Y, Liu J, Spielberg A, et al. (2019) Chainqueen: A real-time differentiable physical simulator for soft robotics. In: 2019 International conference on robotics and automation (ICRA). IEEE, pp. 6265–6271.
- Jadidi MG, Miro JV and Dissanayake G (2019) Sampling-based incremental information gathering with applications to robotic exploration and environmental monitoring. *The International Journal of Robotics Research (IJRR)* 38(6): 658–685.
- Jang D, Yoo J, Son CY, et al. (2020) Multi-robot active sensing and environmental model learning with distributed Gaussian process. *IEEE Robotics and Automation Letters (RA-L)* 5(4): 5905–5912.
- Jegorova M, Smith J, Mistry M, et al (2020) Adversarial generation of informative trajectories for dynamics system identification. In: 2020 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). IEEE, pp. 7109–7115.
- Kaelbling LP, Littman ML and Cassandra AR (1998) Planning and acting in partially observable stochastic domains. *Artificial Intelligence* 101(1–2): 99–134.
- Kaelbling LP and Lozano-Pérez T (2013) Integrated task and motion planning in belief space. *The International Journal of Robotics Research* 32(9–10): 1194–1227.
- Kantaros Y, Schlotfeldt B, Atanasov N, et al. (2021) Sampling-based planning for non-myopic multi-robot information gathering. *Autonomous Robots (AURO)* 45(3): 1–18.
- Kemna S, Kroemer O and Sukhatme GS (2018) Pilot surveys for adaptive informative sampling. In: IEEE International Conference on Robotics and Automation (ICRA), Brisbane, QLD, Australia, 21–25 May 2018, pp. 6417–6424.
- Kim SK, Salzman O and Likhachev M (2019) POMHDP: Search-based belief space planning using multiple heuristics. In: Proceedings of the International Conference on Automated Planning and Scheduling, Berkeley, CA, USA, 11–15 July 2019, volume 29, pp. 734–744.
- Kingma DP and Ba J (2014) Adam: a method for stochastic optimization. *arXiv preprint arXiv:1412.6980*.
- Kompis Y, Bartolomei L, Mascaro R, et al. (2021) Informed sampling exploration path planner for 3d reconstruction of large scenes. *IEEE Robotics and Automation Letters (RA-L)* 6(4): 7893–7900.
- Krause A and Guestrin C (2007) Nonmyopic active learning of Gaussian processes: an exploration-exploitation approach. In: International Conference on Machine learning (ICML), Corvallis, OR, USA, 20–24 June 2007, pp. 449–456.
- Krause A, Singh A and Guestrin C (2008) Near-optimal sensor placements in Gaussian processes: theory, efficient algorithms and empirical studies. *The Journal of Machine Learning Research (JMLR)* 9: 235–284.
- Lang T, Plagemann C and Burgard W (2007) Adaptive non-stationary kernel regression for terrain modeling. In: Robotics: Science and Systems (RSS), Atlanta, GA, USA, 27–30 June 2007, pp. 1–8.
- Lauri M, Heinänen E and Frintrop S (2017) Multi-robot active information gathering with periodic communication. In: 2017 IEEE International Conference on Robotics and Automation (ICRA). IEEE, pp. 851–856.
- Lauri M, Hsu D and Pajarinen J (2022) Partially observable markov decision processes in robotics: a survey. *IEEE Transactions on Robotics* 39: 21–40.
- Lauri M, Pajarinen J, Peters J, et al. (2020) Multi-sensor next-best-view planning as matroid-constrained submodular maximization. *IEEE Robotics and Automation Letters* 5(4): 5323–5330.
- LaValle SM (2006) *Planning Algorithms*. Cambridge, United Kingdom: Cambridge University Press.
- Lee J, Feng J, Humt M, et al. (2022) Trust your robots! predictive uncertainty estimation of neural networks with sparse gaussian processes. In: Conference on Robot Learning (CoRL), Auckland, New Zealand, 14–18 December 2022, pp. 1168–1179.
- Lew T, Sharma A, Harrison J, et al. (2022) Safe active dynamics learning and control: a sequential exploration–exploitation framework. *IEEE Transactions on Robotics* 38(5): 1–20.
- Li AQ (2020) Exploration and mapping with groups of robots: recent trends. *Current Robotics Reports* 1: 1–11.

- Liang J, Saxena S and Kroemer O (2020) Learning active task-oriented exploration policies for bridging the sim-to-real gap. In: *Proceedings of Robotics: Science and Systems*. Corvallis, Oregon, USA, July 12–16, 2020, pp. 1–10.
- Lim ZW, Hsu D and Lee WS (2016) Adaptive informative path planning in metric spaces. *The International Journal of Robotics Research* 35(5): 585–598.
- Lluvia I, Lazkano E and Ansuategi A (2021) Active mapping and robot exploration: a survey. *Sensors* 21(7): 2445.
- Lotfi S, Izmailov P, Benton G, et al. (2022) Bayesian model selection, the marginal likelihood, and generalization. In: *Proceedings of the 39th International Conference on Machine Learning, Proceedings of Machine Learning Research*, volume 162. PMLR, pp. 14223–14247.
- Low KH, Dolan J and Khosla P (2009) Information-theoretic approach to efficient adaptive path planning for mobile robotic environmental sensing. In: *International Conference on Automated Planning and Scheduling (ICAPS)*, Thessaloniki, Greece, 19–23 September 2009, volume 19, pp. 233–240.
- Luo W and Sycara K (2018) Adaptive sampling and online learning in multi-robot sensor coverage with mixture of gaussian processes. In: *IEEE International Conference on Robotics and Automation (ICRA)*, Brisbane, QLD, Australia, 21–25 May 2018, pp. 6359–6364.
- Lutter M, Silberbauer J, Watson J, et al. (2021) Differentiable physics models for real-world offline model-based reinforcement learning. In: *2021 IEEE International Conference on Robotics and Automation (ICRA)*, Xi'an, China, 30 May–05 June 2021, pp. 4163–4170.
- Ma KC, Liu L, Heidarsson HK, et al. (2018) Data-driven learning and planning for environmental sampling. *Journal of Field Robotics (JFR)* 35(5): 643–661.
- Ma KC, Liu L and Sukhatme GS (2017) Informative planning and online learning with sparse Gaussian processes. In: *International Conference on Robotics and Automation (ICRA)*, Singapore, 29 May–03 June 2017, pp. 4292–4298.
- MacDonald RA and Smith SL (2019) Active sensing for motion planning in uncertain environments via mutual information policies. *The International Journal of Robotics Research (IJRR)* 38(2–3): 146–161.
- Manjanna SM, Li AQ, Smith RN, et al. (2018) Heterogeneous multi-robot system for exploration and strategic water sampling. In: *IEEE International Conference on Robotics and Automation (ICRA)*, Brisbane, QLD, Australia, 21–25 May 2018, pp. 4873–4880.
- Marchant R and Ramos F (2012) Bayesian optimisation for intelligent environmental monitoring. In: *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, Vilamoura, Portugal, 07–12 October 2012, pp. 2242–2249.
- Marchant R and Ramos F (2014) Bayesian optimisation for informative continuous path planning. In: *IEEE International Conference on Robotics and Automation (ICRA)*, Hong Kong, China, 31 May–07 June 2014, pp. 6136–6143.
- McCammon S and Hollinger GA (2018) Topological hotspot identification for informative path planning with a marine robot. In: *2018 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, pp. 4865–4872.
- Meera AA, Popović M, Millane A, et al. (2019) Obstacle-aware adaptive informative path planning for UAV-based target search. In: *2019 International Conference on Robotics and Automation (ICRA)*, Montreal, QC, Canada, 20–24 May 2019, pp. 718–724.
- Mehta B, Handa A, Fox D, et al. (2021) A User's Guide to calibrating robotic simulators. In: *Conference on Robot Learning*. PMLR, pp. 1326–1340.
- Meliou A, Krause A, Guestrin C, et al. (2007) Nonmyopic informative path planning in spatio-temporal models. In: *The AAAI Conference on Artificial Intelligence (AAAI)*, Vancouver, BC, Canada, 22–26 July 2007, volume 10, pp. 16–17.
- Moon B, Chatterjee S and Scherer S (2022) TIGRIS: an informed sampling-based algorithm for informative path planning. In: *arXiv preprint arXiv:2203.12830*. arXiv, pp. 1–7.
- Morere P, Marchant R and Ramos F (2017) Sequential Bayesian optimization as a POMDP for environment monitoring with UAVs. In: *IEEE International Conference on Robotics and Automation (ICRA)*, Singapore, 29 May–03 June 2017, pp. 6381–6388.
- Muratore F, Eilers C, Gienger M, et al. (2021) Data-efficient domain randomization with bayesian optimization. *IEEE Robotics and Automation Letters* 6(2): 911–918.
- Muratore F, Ramos F, Turk G, et al. (2022) Robot learning from randomized simulations: A review. *Frontiers in Robotics and AI* 9: 799893.
- Nguyen-Tuong D, Peters J and Seeger M (2008) Local Gaussian process regression for real time online model learning. *Advances in Neural Information Processing Systems (NeurIPS)* 21: 1–8.
- Nishimura H and Schwager M (2021) SACBP: belief space planning for continuous-time dynamical systems via stochastic sequential action control. *The International Journal of Robotics Research (IJRR)* 40(10–11): 1167–1195.
- Ober SW, Rasmussen CE and van der Wilk M (2021) The promises and pitfalls of deep kernel learning. In: *Uncertainty in Artificial Intelligence (UAI)*, Virtual, 26–30 July 2021, pp. 1206–1216.
- O'Callaghan ST and Ramos FT (2012) Gaussian process occupancy maps. *The International Journal of Robotics Research (IJRR)* 31(1): 42–62.
- O'Meadhra C, Tabib W and Michael N (2018) Variable resolution occupancy mapping using gaussian mixture models. *IEEE Robotics and Automation Letters (RA-L)* 4(2): 2015–2022.
- Ouyang R, Low KH, Chen J, et al. (2014) Multi-robot active sensing of non-stationary Gaussian process-based environmental phenomena. In: *International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, Paris, France, 05–09 May 2014, pp. 573–580.
- Paciorek CJ and Schervish MJ (2003) Nonstationary covariance functions for Gaussian process regression. In: *Advances on Neural Information Processing Systems (NeurIPS)*, Vancouver, BC, Canada, 08–13 December 2003, pp. 1–8.
- Palomeras N, Hurtós N, Vidal E, et al. (2019) Autonomous exploration of complex underwater environments using a probabilistic next-best-view planner. *IEEE Robotics and Automation Letters* 4(2): 1619–1625.

- Papachristos C, Mascari F, Khattak S, et al. (2019) Localization uncertainty-aware autonomous exploration and mapping with aerial robots using receding horizon path-planning. *Autonomous Robots* 43(8): 2131–2161.
- Paszke A, Gross S, Chintala S, et al. (2017) Automatic Differentiation in PyTorch. In: NIPS 2017 Workshop on Autodiff, Vancouver, BC, Canada, 08–13 December 2019, pp. 1–4.
- Pfingsten T, Kuss M and Rasmussen CE (2006) Nonstationary Gaussian process regression using a latent extension of the input space. In: International Society for Bayesian Analysis (ISBA), Benidorm, Spain, 02–06 June 2006, pp. 1–3.
- Placed JA, Strader J, Carrillo H, et al. (2022) A survey on active simultaneous localization and mapping: state of the art and new frontiers. *arXiv preprint arXiv:2207.00254*.
- Plagemann C, Kersting K and Burgard W (2008a) Nonstationary Gaussian process regression using point estimates of local smoothness. In: Joint European Conference on Machine Learning and Knowledge Discovery in Databases (ECML-KDD). Springer, pp. 204–219.
- Plagemann C, Mischke S, Prentice S, et al. (2008b) Learning predictive terrain models for legged robot locomotion. In: IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), Nice, France, 22–26 September 2008, pp. 3545–3552.
- Platt R, Tedrake R, Kaelbling L, et al. (2010) Belief space planning assuming maximum likelihood observations. In: Proceedings of Robotics: Science and Systems, Zaragoza, Spain, 27–30 June 2010, pp. 1–8.
- Popović M, Hitz G, Nieto J, et al. (2017) Online informative path planning for active classification using UAVs. In: IEEE International Conference on Robotics and Automation (ICRA), Singapore, 29 May–03 June 2017, pp. 5753–5758.
- Popović M, Vidal-Calleja T, Chung JJ, et al. (2020a) Informative path planning for active field mapping under localization uncertainty. In: IEEE International Conference on Robotics and Automation (ICRA), Paris, France, 31 May–04 June 2020, pp. 10751–10757.
- Popović M, Vidal-Calleja T, Hitz G, et al. (2020b) An informative path planning framework for UAV-based terrain monitoring. *Autonomous Robots (AURO)* 44(6): 889–911.
- Prentice S and Roy N (2009) The belief roadmap: Efficient planning in belief space by factoring the covariance. *The International Journal of Robotics Research* 28(11–12): 1448–1465.
- Preston V, Flaspohler G, Michel AP, et al. (2022) Robotic planning under uncertainty in spatiotemporal environments in expeditionary science. *arXiv preprint arXiv:2206.01364*.
- Quinonero-Candela J and Rasmussen CE (2005) A unifying view of sparse approximate Gaussian process regression. *The Journal of Machine Learning Research (JMLR)* 6: 1939–1959.
- Ramos F and Ott L (2016) Hilbert maps: scalable continuous occupancy mapping with stochastic gradient descent. *The International Journal of Robotics Research (IJRR)* 35(14): 1717–1730.
- Ramos F, Possas RC and Fox D (2019) Bayessim: adaptive domain randomization via probabilistic inference for robotics simulators. *arXiv preprint arXiv:1906.01728*.
- Rasmussen C and Ghahramani Z (2001) Infinite mixtures of Gaussian process experts. In: Advances in Neural Information Processing Systems (NeurIPS), Vancouver, BC, Canada, 03–08 December 2001, volume 14, pp. 1–8.
- Rasmussen CE and Williams CKI (2005) *Gaussian Processes for Machine Learning*. Cambridge, Massachusetts, USA: MIT Press.
- Remes S, Heinonen M and Kaski S (2017) Non-stationary spectral kernels. In: Advances on Neural Information Processing Systems (NeurIPS), Long Beach, CA, USA, 04–09 December 2017, pp. 1–10.
- Remes S, Heinonen M and Kaski S (2018) Neural non-stationary spectral Kernel. *arXiv*.
- Ren Z, Srinivasan A, Coffin H, et al. (2022) A local optimization framework for multi-objective ergodic search. In: Proceedings of Robotics: Science and Systems. New York City, NY, USA, June 27–July 1, pp. 1–10. DOI: [10.15607/RSS.2022.XVIII.052](https://doi.org/10.15607/RSS.2022.XVIII.052)
- Rezaei-Shoshtari S, Meger D and Sharf I (2019) Cascaded gaussian processes for data-efficient robot dynamics learning. In: 2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). IEEE, pp. 6871–6877.
- Rodríguez-Arévalo ML, Neira J and Castellanos JA (2018) On the importance of uncertainty representation in active SLAM. *IEEE Transactions on Robotics* 34(3): 829–834.
- Romeres D, Jha DK, DallaLibera A, et al (2019) Semiparametrical gaussian processes learning of forward dynamical models for navigating in a circular maze. In: 2019 International Conference on Robotics and Automation (ICRA). IEEE, pp. 3195–3202.
- Roy N, Burgard W, Fox D, et al (1999) Coastal navigation-mobile robot navigation with uncertainty in dynamic environments. In: Proceedings 1999 IEEE International Conference on Robotics and Automation (Cat. No. 99CH36288C), volume 1. IEEE, pp. 35–40.
- Rückin J, Jin L and Popović M (2022) Adaptive informative path planning using deep reinforcement learning for UAV-based active sensing. In: 2022 International Conference on Robotics and Automation (ICRA). IEEE, pp. 4473–4479.
- Salimbeni H and Deisenroth M (2017) Doubly stochastic variational inference for deep Gaussian processes. In: Advances in Neural Information Processing Systems, Long Beach, CA, USA, 04–09 December 2017, volume 30, pp. 1–12.
- Sampson PD and Guttorp P (1992) Nonparametric estimation of nonstationary spatial covariance structure. *Journal of the American Statistical Association (JASA)* 87(417): 108–119.
- Saroya M, Best G and Hollinger GA (2021) Roadmap learning for probabilistic occupancy maps with topology-informed growing neural gas. *IEEE Robotics and Automation Letters (RA-L)* 6(3): 4805–4812.

- Schlotfeldt B, Atanasov N and Pappas GJ (2019) Maximum information bounds for planning active sensing trajectories. In: 2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). IEEE, pp. 4913–4920.
- Schlotfeldt B, Thakur D, Atanasov N, et al. (2018) Anytime planning for decentralized multi-robot active information gathering. *IEEE Robotics and Automation Letters (RA-L)* 3(2): 1025–1032.
- Schlotfeldt B, Tzoumas V and Pappas GJ (2021) Resilient active information acquisition with teams of robots. *IEEE Transactions on Robotics (T-RO)* 38: 244–261.
- Schmid LM, Pantic M, Khanna R, Ott L, Siegwart R and Nieto J (2020) An efficient sampling-based method for online informative path planning in unknown environments. *IEEE Robotics and Automation Letters (RA-L)* 5(2): 1.
- Schneider T, Belousov B, Chalvatzaki G, et al. (2022) Active exploration for robotic manipulation. *arXiv preprint arXiv:2210.12806*.
- Senanayake R and Ramos F (2017) Bayesian Hilbert maps for dynamic continuous occupancy mapping. In: Conference on Robot Learning (CoRL), Mountain View, CA, USA, 13–15 November 2017, pp. 458–471.
- Senanayake R, Tompkins A and Ramos F (2018) Automorphing kernels for nonstationarity in mapping unstructured environments. In: Conference on Robot Learning (CoRL), Zurich, Switzerland, 29–31 October 2018, pp. 443–455.
- Settles B (2012) Active learning. *Synthesis Lectures on Artificial Intelligence and Machine Learning* 6(1): 1–114.
- Sheth R, Wang Y and Khordon R (2015) Sparse variational inference for generalized GP models. In: International Conference on Machine Learning. PMLR, pp. 1302–1311.
- Singh A, Krause A, Guestrin C, et al. (2007) Efficient planning of informative paths for multiple robots. In: International Joint Conference on Artificial Intelligence (IJCAI), Hyderabad, India, 06–12 January 2007, pp. 2204–2211.
- Snoek J, Larochelle H and Adams RP (2012) Practical bayesian optimization of machine learning algorithms. *Advances in Neural Information Processing Systems* 25: 1–9.
- Snoek J, Swersky K, Zemel R, et al. (2014) Input warping for Bayesian optimization of non-stationary functions. In: International Conference on Machine Learning (ICML), Beijing, China, 21–26 June 2014, volume 32, pp. 1674–1682.
- Stachniss C, Hahnel D and Burgard W (2004) Exploration with active loop-closing for FastSLAM. In: 2004 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)(IEEE Cat. No. 04CH37566), volume 2. IEEE, pp. 1505–1510.
- Stachniss C, Plagemann C and Lilienthal AJ (2009) Learning gas distribution models using sparse Gaussian process mixtures. *Autonomous Robots* 26(2): 187–202.
- Sutanto G, Wang A, Lin Y, et al. (2020) Encoding physical constraints in differentiable newton-euler algorithm. In: *Learning for Dynamics and Control*. PMLR, pp. 804–813.
- Tabib W, Goel K, Yao J, et al. (2019) Real-time information-theoretic exploration with gaussian mixture model maps. In: Robotics: Science and Systems (RSS), Freiburg im Breisgau, Germany, 22–26 June 2019, pp. 1–10.
- Taylor AT, Berrueta TA and Murphey TD (2021) Active learning in robotics: a review of control principles. *Mechatronics* 77: 102576.
- Teng S, Gong Y, Grizzle JW, et al. (2021) Toward safety-aware informative motion planning for legged robots. *arXiv preprint arXiv:2103.14252*.
- Thrun S (2002) Probabilistic robotics. *Communications of the ACM* 45(3): 52–57.
- Titsias M (2009) Variational learning of inducing variables in sparse Gaussian processes. In: International Conference on Artificial Intelligence and Statistics (AISTATS), Clearwater Beach, FL, USA, 16–18 April 2009, pp. 567–574.
- Tompkins A, Oliveira R and Ramos FT (2020a) Sparse spectrum warped input measures for nonstationary Kernel learning. In: Advances in Neural Information Processing Systems (NeurIPS), Vancouver, BC, Canada, 06–12 December 2020, volume 33, pp. 16153–16164.
- Tompkins A, Senanayake R and Ramos F (2020b) Online domain adaptation for occupancy mapping. In: Robotics: Science and Systems (RSS), Virtual, 12–16 July 2020, pp. 1–10.
- Tranzatto M, Dharmadhikari M, Bernreiter L, et al. (2022) Team CERBERUS wins the DARPA subterranean challenge: technical overview and lessons learned. *arXiv preprint arXiv:2207.04914*.
- Trapp M, Peharz R, Pernkopf F, et al. (2020) Deep structured mixtures of gaussian processes. In: International Conference on Artificial Intelligence and Statistics (AISTATS), Palermo, Italy, 03–05 June 2020, pp. 2251–2261.
- van Amersfoort J, Smith L, Jesson A, et al. (2021) On feature collapse and deep kernel learning for single forward pass uncertainty. *arXiv preprint arXiv:2102.11409*.
- Wang K, Hamelijnck O, Damoulas T, et al. (2020) Non-separable non-stationary random fields. In: International Conference on Machine Learning (ICML), Vienna, Austria, 12–18 July 2020, pp. 9887–9897.
- Wei Y, Sheth R and Khordon R (2021) Direct loss minimization for sparse gaussian processes. In: International Conference on Artificial Intelligence and Statistics. PMLR, pp. 2566–2574.
- Werling K, Omens D, Lee J, et al. (2021) Fast and feature-complete differentiable physics engine for articulated rigid bodies with contact constraints. In: Proceedings of Robotics: Science and Systems. Virtual, pp. 1–15. DOI: [10.15607/RSS.2021.XVII.034](https://doi.org/10.15607/RSS.2021.XVII.034)
- Wilson AG, Hu Z, Salakhutdinov R, et al. (2016) Deep Kernel learning. In: International Conference on Artificial Intelligence and Statistics (AISTATS), Cadiz, Spain, 09–11 May 2016, pp. 370–378.
- Wirnshofer F, Schmitt PS, von Wichert G, et al. (2020) Controlling contact-rich manipulation under partial observability. In: Proceedings of Robotics: Science and Systems, Virtual, 12–16 July 2020, pp. 1–10.
- Xu Y, Zheng R, Liu M, et al. (2021) CRMI: Confidence-rich mutual information for information-theoretic mapping. *IEEE Robotics and Automation Letters* 6(4): 6434–6441.
- Yamauchi B (1997) A frontier-based approach for autonomous exploration. In: Proceedings 1997 IEEE International Symposium on Computational Intelligence in Robotics and Automation CIRA'97. 'Towards New Computational Principles for Robotics and Automation'. IEEE, pp. 146–151.

- Yu HSA, Yao D, Zimmer C, et al. (2021) Active learning in Gaussian process state space model. In: Joint European Conference on Machine Learning and Knowledge Discovery in Databases. Springer, pp. 346–361.
- Yu J, Schwager M and Rus D (2014) Correlated orienteering problem and its application to informative path planning for persistent monitoring tasks. In: IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), Chicago, IL, USA, 14–18 September 2014, pp. 342–349.
- Zhang Z (2020) *Active Robot Vision: From State Estimation to Motion Planning*. PhD Thesis, Universität Zürich.
- Zhang Z, Henderson T, Karaman S, et al. (2020) FSMI: fast computation of Shannon mutual information for information-theoretic mapping. *The International Journal of Robotics Research (IJRR)* 39(9): 1155–1177.
- Zhang Z and Scaramuzza D (2020) Fisher information field: an efficient and differentiable map for perception-aware planning. *arXiv preprint arXiv:2008.03324*.
- Zheng D, Ridderhof J, Tsiotras P, et al. (2022) Belief space planning: a covariance steering approach. In: 2022 International Conference on Robotics and Automation (ICRA). IEEE, pp. 11051–11057.
- Zhu H, Chung JJ, Lawrance NR, et al. (2021) Online informative path planning for active information gathering of a 3D surface. In: IEEE International Conference on Robotics and Automation (ICRA), Xi'an, China, 30 May–05 June 2021, pp. 1488–1494.