

State-space Models and Gaussian Processes in Paleoceanography and Paleoclimatology

A DISSERTATION PRESENTED

BY

TAEHEE LEE

TO

THE DIVISION OF APPLIED MATHEMATICS

IN PARTIAL FULFILLMENT OF THE REQUIREMENTS

FOR THE DEGREE OF

DOCTOR OF PHILOSOPHY

IN THE SUBJECT OF

APPLIED MATHEMATICS

BROWN UNIVERSITY

PROVIDENCE, RHODE ISLAND

MAY 2020

Abstract of “State-space Models and Gaussian Processes in Paleoceanography and Paleoclimatology”, by Taehee Lee, Ph.D., Brown University, May 2020.

For the recent decades, statistical (machine) learning and Bayesian modelling have emerged as popular frameworks for various fields of area, regardless of academy or industry. The statistical learning theory has many applicable branches that include state-space models and the Gaussian process. State-space models are designed to deal with latent series under the Markov assumption. The hidden Markov model is applied for the case that each hidden state is defined on a finite set, whereas the Kalman filter, particle filter and Markov-chain Monte Carlo algorithms can treat the continuous cases. The Gaussian process gives a nonparametric prior on the functions. Its flexibility allows various applications including the regression, classification, state-space models and dimensionality reduction.

Paleoceanography and paleoclimatology are the fields of geoscience that discover the past ocean and climate events from the relevant preserved physical characteristics, proxies, as well as their ages. Interdisciplinary research encompassing physics, chemistry and biology has investigated the relationship between ocean and climate events and proxies, and more recently statistical learning is also making progresses. The scarce and agnostic characteristics of data need particular modelling that are more dependent on the data than certain parametric structures.

In this thesis, three applications of the state-space models and Gaussian process to the field of paleoceanography and paleoclimatology are presented. Dual proxy stack construction algorithm addresses the age assignments problem with the radiocarbon (^{14}C) and benthic $\delta^{18}\text{O}$ proxies by adopting a state-space model and the Gaussian process regression. The heteroscedastic Gaussian process regression constructs nonparametric calibration models of alkenone and TEX_{86} proxies for restoring past sea surface temperatures and a novel multimodal Gaussian process regression model is also considered. Two state-space models are applied for reconstructing past atmospheric CO_2 pressures with the boron isotope and benthic $\delta^{18}\text{O}$ proxies: one is a parametric state-space model based on the particle smoother and the other is a nonparametric Gaussian process state-space model. Our study certifies the effectiveness of giving correlation to the hidden states and considering nonparametric data-driven models in paleoceanography and paleoclimatology.

@ Copyright 2020 by Taehee Lee

NOTE: The year of copyright is the same as the year the degree is conferred.

This dissertation by Taehee Lee is accepted in its present form
by the Division of Applied Mathematics as satisfying the
dissertation requirement for the degree of Doctor of Philosophy.

Date _____ (Signature of Advisor)
Charles E. Lawrence, Advisor

Recommended to the Graduate Council

Date _____ (Signature of Reader)
Lorraine E. Lisiecki, Reader

Date _____ (Signature of Reader)
Matthew T. Harrison, Reader

Approved by the Graduate Council

Date _____ (Signature of Dean)
Andrew G. Campbell, Dean of the Graduate School

Curriculum Vitae

Taehee Lee received his Bachelor of Science in Mathematical Sciences with Financial Mathematics as a combined minor and Bachelor of Arts in Economics as a double major from Seoul National University in 2012. During his undergraduate studies he won the Presidential Science Scholarship from the Korea Student Aid Foundation (KOSAF) for 2007 to 2011. Once completing the alternative military service duty by 2014, he started the graduate studies in the Division of Applied Mathematics at Brown University. On September 17th, 2015, He was qualified to become a Ph.D. candidate by passing the preliminary oral exam. He worked as a teaching assistant (TA) for the courses ‘Applied Ordinary Differential Equations’ in Fall 2015 and ‘Essential Statistics’ in Spring 2016. Since September 2016, he has been working as a research assistant (RA) with Professor Charles E. Lawrence as his advisor. During his graduate studies he won the Kwanjeong Educational Foundation Doctoral Program Scholarship from the Kwanjeong Educational Foundation for 2014 to 2019. He also received two transitional master’s degrees from Brown University in the interim: Master of Science in Applied Mathematics in 2017 and Master of Science in Computer Science in 2019. He was elected to SIGMA XI on May 16th, 2016 and is nominated for the Simon Ostrach Fellowship on May 24th, 2020. He successfully defended his Ph.D. thesis on April 16th, 2020.

Preface and Acknowledgments

When I was a high school student good at mathematics, I thought that the world can be explained objectively by a set of well-designed and abstract mathematical models. Such a belief led me to choose mathematics as the major in my undergraduate studies, and economics followed as a double major. Though I could graduate with an excellent academic record in 2012, I became more confused by the following philosophical doubt: can the inevitable subjectivistic viewpoint of the observer be always either ignored, removed, diluted or neutralized completely by a set of objective models? If so, is it okay to consider it to be subordinated to the objectivistic concept? For the next two years of my alternative military service as an assistant librarian in South Korea, I could fortunately read some classical books of the Austrian School and realized the importance of the subjectivistic aspects in theory and methodology. However, the subjectivistic aspects themselves do not explain but just rationalize the given results. How can we find a balance between the subjectivistic and objectivistic aspects in understanding the reality?

The Bayesian framework gives an ideal methodological answer to that question. Suppose that we have something unknown to discover, which is paired with the relevant known. The subjectivistic aspects that are independent from the known are formulated as a prior on the unknown while the objectivistic aspects appear as a set of models that defines the likelihoods of the known given the unknown. The Bayes' theorem derives the posterior of the unknown given the known from the prior and models. That is the reason why I decided to study the Bayesian statistical

learning in the Division of Applied Mathematics at Brown University.

Because most of the coursework in my undergraduate studies was not directly relevant to my current research areas, I could not avoid taking a lot of courses opened in the Division of Applied Mathematics or the Department of Computer Science after starting the Ph.D. program. To dig deep into the details of those courses, reading various textbooks and papers was essential. Weekly lab and personal meeting with my advisor comprise a core part of my studies. Statistical learning is indispensable from actual implementation, thus I needed to be familiar with some popular computer languages. Workshops, conferences, and seminars were good sources for the new ideas and recent research trends. Two transitional master's degrees, poster presentations, published and submitted papers and the award certify my sincere effort in the graduate studies.

My research areas cover a wide range of topics, from designing new statistical models to solving real problems. I picked some of the results from my research to organize this dissertation for proving that my graduate studies of the past six years are sufficient to grant me a Ph.D. degree in Applied Mathematics. To make my dissertation consistent, works in paleoceanography and paleoclimatology that mainly use state-space models and Gaussian processes were chosen. I organized them into three chapters according to their particular applications. Two additional chapters in Appendix were arranged to cover the general background of the models because I hope that my thesis will be also used as a manual for the beginners of the statistical learning theory.

During the graduate studies, I have received a lot of aid and support from many people. My thesis advisor, Prof. Charles Lawrence has contributed much to my studies. He has taught me the Bayesian statistical modelling, given inspiration for resolving difficult parts of the problems, connected me to interdisciplinary data-expert researchers that are relevant to my studies, discussed about new problems and modelling with me, introduced various workshops and conferences, supported expenses of my studies, and been the anchor whom I could rely on. I was also impressed by his gentle and kind manner to me. I indeed owe him the greatest deal of gratitude, next to my family.

I also appreciate the lab members. Prof. William Thompson has been providing good code

implementation. Seonmin Ahn and Taesoo Kim helped me a lot adapt myself to the new surroundings when I first arrived at Providence. Wilson McKerrow and August Guang were good seniors who suggested me proper papers and presented their works. Andrew Kaluzny and Daniel Kunin partially contributed to the works in section 4.3.1 and 2.3.1, respectively. Ashley Conard always actively shows and explains her work in computational biology: I regret not working a lot in research with her. I wish all of them good luck and success in their works.

Most of my research in paleoceanography and paleoclimatology could not be made without the aid of Prof. Lorraine Lisiecki and her graduate student, Devin Rand, at the University of California, Santa Barbara (UCSB). As researchers in paleoceanography, they have consistently collaborated with me by processing the raw data and analyzing my results qualitatively. Prof. Lorraine Lisiecki also participated in my dissertation defense as the second reader of the committee. Devin Rand has followed the rapid transition of my models and steadily given feedbacks to their results. Their efforts contributed a lot to the work in chapter 2. Dr. Geoffrey Gebbie in Woods Hole Oceanographic Institution also contributed to the work.

Prof. Matthew Harrison was willing to be participating in my dissertation defense as the third reader of the committee and be a reference for my job applications. He was also the instructor of the course ‘Essential Statistics’ for which I was a teaching assistant (TA) in Spring 2016. I appreciate his solicitude for me during the overall graduate studies of mine. I am also thankful for the efforts of Prof. Susan Gerbi and her lab members in the project of finding replication origins, though we eventually failed to obtain meaningful results. One of her students, John Urban, particularly spent much time for helping me understand their tools and methodology.

I am grateful to all the other faculty members, staffs, and graduate students in the Division of Applied Mathematics at Brown University. They have indeed contributed to making nice research environment. I also want to express my appreciation of the members of the Korean Graduate Student Association (KGSA) for their small help in my daily lives. Prof. Hyeong-in Choi of Seoul National University recommended me some noteworthy books of the Austrian School that eventually led me to change the majoring fields of studies from pure to applied mathematics.

For the first five years of my graduate studies, the Kwanjeong Educational Foundation had provided me with the Kwanjeong Educational Foundation Doctoral Program Scholarship. I am thankful to Chonghwan Lee, Honorary President of Samyoung Chemical Group who established the foundation. I also appreciate the public education system of Republic of Korea that focuses on not only granting sufficient opportunities to all young Korean students regardless of their families' wealth but also offering good special education for the gifted and talented ones.

Lastly, I need to express my gratitude to my family. My parents have always assisted my learning and studies. My maternal grandmother had sincerely aided me a lot with a keen interest before she passed away in 2012. My maternal grandfather financially supported my parents when I was very young and my father was just a graduate student. My maternal aunt has always been exceedingly supportive to me. I also appreciate all the other family members and relatives for their love.

Contents

I	INTRODUCTION	1
1.1	Paleoceanography and Paleoclimatology	2
1.2	Statistical Learning Theory	6
1.3	Contents	15
1.4	Dictionary	17
2	BENTHIC $\delta^{18}\text{O}$ STACK CONSTRUCTION	21
2.1	Related Work	22
2.2	Research Motive and Data	24
2.3	Core Alignment Algorithm	26
2.4	Stack Construction Algorithm	31
2.5	Results and Discussion	34
2.6	Future Work	42
3	ALKENONE/TEX₈₆ CALIBRATION MODELS	44
3.1	Related Work	46
3.2	Research Motive and Data	47
3.3	Models	50
3.4	Results and Discussion	57
3.5	Future Work	66
4	ATMOSPHERIC CO₂ RECONSTRUCTION USING BORON PROXY	69
4.1	Related Work	70
4.2	Research Motive and Data	71
4.3	Models	72
4.4	Results and Discussion	84
4.5	Future Work	87
5	CONCLUSION	93
	APPENDIX A STATE-SPACE MODELS	97
A.1	Hidden Markov Model	100
A.2	Kalman Filter and Smoother	109
A.3	Particle Filter and Smoother	117
A.4	Markov-chain Monte Carlo Algorithm	122

APPENDIX B	GAUSSIAN PROCESS MODELS	128
B.1	Kernels	130
B.2	Gaussian Process Regression	138
B.3	Gaussian Process Classification	146
B.4	Gaussian Process State-space Models	151
B.5	Gaussian Process Latent Variable Models	155
B.6	Deep Gaussian Process	157
B.7	Variational Inference	159
REFERENCES		185

List of Tables and Illustrations

1.1	An example of how the indirect proxy is utilized for the age inference. Here, five sediment cores are aligned based on their fossil records.	3
1.2	A comparison of the frequentist and Bayesian perspectives.	8
1.3	A diagram of the parametric inference. The relationship between two variables X and Y are explained by the parametric model $f(\cdot \theta)$ (left). The parameter θ is learned from the training data (X, Y) colored green (middle). The unknown Y is explained by the query X colored red and learned parameter θ (right).	11
1.4	Two nonparametric Gaussian process regressions on the same dataset but with different hyperparameters.	12
1.5	A diagram of the nonparametric inference. The relationship between two variables X and Y is not explained by any parametric model but is indirectly controlled by the hyperparameter θ (left). The hyperparameter θ is tuned from the training data (X, Y) colored green (middle). The unknown Y is explained by the query X colored red, tuned hyperparameter θ , and training data colored green (right).	13
1.6	Graphical representations of the generative model (left) and discriminative model (right). Here, θ is a parameter and X is a hidden variable that explains the observation Y	15
2.1	Locations of cores GeoB7920-2, GeoB9508-5, GeoB9526-5, MD95-2042, MD99-2334, ODP658C, GIK13289-2 and SU90-08.	25
2.2	A graphical representation of our core alignment model.	26
2.3	Three steps for the stack construction algorithm.	31
2.4	A simple diagram of the SA-GPR. Each box is iterated until convergence. . . .	35
2.5	Core alignments to stacks. The upper panel (a) shows those to the DNA stack, while the middle panel (b) is our dual proxy local DNEA stack. Stars indicate medians of $\delta^{18}\text{O}$ data classified as inliers and dots represent outliers, after shifting them vertically by the core-specific shift parameters. The darker and brighter gray regions are the 1-sigma and 2-sigma of the stacks, respectively. The dot-dash line indicates their boundary. The lower panel (c) is a portion of the DNEA stack at 10-50 kiloyears.	36
2.6	A part of the DNEA stack constructed with the SE kernel.	37
2.7	Histograms of $\delta^{18}\text{O}$ to our dual proxy DNEA stack. Each $\delta^{18}\text{O}$ is standardized by the mean and standard deviation of the stack at the median of its sampled ages.	37

2.8	The comparison between inferred dual proxy ages and ^{14}C ages of records. Shaded regions indicate 95% confidence regions and the black dashed line is just a diagonal.	38
2.9	Dual proxy age inferences and alignments of GIK13289-2 to DNEA stack. In the upper figure, the darker and brighter areas show the 1-sigma and 2-sigma of the stacks, red points and blue dotted bars indicate medians and 95% confidence intervals of inlier age samples for $\delta^{18}\text{O}$ data, blue points and red dotted bars are the outliers. In the left below figures, cyan bars indicate 95% confidence intervals obtained independently from ^{14}C data, and blue and red areas show the 95% confidence bands of age inferences from ^{14}C only and both proxies, respectively. In the right below figures, medians (dashed lines) and 95% confidence intervals of relative ages to the medians of dual proxy ages are shown.	39
2.10	Dual proxy age inferences and alignments of SU90-08 to DNEA stack, similar to figure 2.9. This figure also shows how multiple observations at the same depths are dealt with in the alignment algorithm.	40
2.11	Histograms of the standardized accumulation rates. The blue histogram comes from the training ^{14}C dataset and the red one from the sampled age paths in constructing the DNEA stack. Thick lines indicate the inferred gamma distributions from the histograms, respectively.	42
3.1	$\text{U}_{37}^{\text{K}'}$ and TEX_{86} data. Each left panel shows the data locations while right plots data values over SSTs.	48
3.2	Some example results from our multimodal GPR.	58
3.3	Homoscedastic Gaussian process regression of $\text{U}_{37}^{\text{K}'}$ on SST. In each panel, red crosses are the outliers and shaded region indicates the 95% confidence band.	60
3.4	Heteroscedastic Gaussian process regression of $\text{U}_{37}^{\text{K}'}$ on SST.	61
3.5	Multimodal Gaussian process regressions of $\text{U}_{37}^{\text{K}'}$ on SST.	62
3.6	Homoscedastic Gaussian process regression of TEX_{86} on SST.	63
3.7	Heteroscedastic Gaussian process regression of TEX_{86} on SST.	64
3.8	Multimodal Gaussian process regressions of TEX_{86} on SST.	65
3.9	Some preliminary results from the site-specific GPR calibration model.	67
4.1	A graphical representation of our particle smoother for the CO_2 reconstruction.	74
4.2	Ice volume and $\delta^{18}\text{O}$ pairs (red) and the chosen model with the 95% confidence band (green).	75
4.3	Atmospheric CO_2 and $\delta^{11}\text{B}$ pairs (red) and the chosen model with the 95% confidence band (green).	76
4.4	(Left) The BIC curve that selects 3 as the number of clusters and (Right) atmospheric CO_2 and ice volume index pairs (green) and the regressions of the clusters (in each color).	80
4.5	Initial and learned values of the parameters in the particle smoother model. The initial values were selected from [Garcia-Olivares & Herrero (2013)].	85

4.6	Results from the particle smoother. (Top) Reconstructed atmospheric CO ₂ with the experimental time series and published lower-upper bands. (Middle) Reconstructed extent of the Antarctic ice sheet indices. (Bottom) Reconstructed ice volume indices with the experimental time series from [Spratt & Lisiecki (2016)].	86
4.7	Results from the Gaussian process state-space model with the SE kernel.	88
4.8	Results from the Gaussian process state-space model with the M25 kernel.	89
4.9	Results from the Gaussian process state-space model with the M15 kernel.	90
4.10	Results from the Gaussian process state-space model with the OU kernel.	91
A.1	A graphical representation of the typical Bayesian inference example. Here the observations Y_n 's are independent and identically distributed samples controlled by the parameter θ	98
A.2	A graphical representation of the typical example of the state-space model. Here the observations Y_n 's are conditionally independent given the hidden states X_n 's.	100
A.3	Another type of the graphical representation of the hidden Markov model for the sequence alignment algorithm. Here, states S, E, I, D and M stand for start, end, insert, delete and match, respectively.	101
A.4	An example of HMM and the inference by the Viterbi and forward-backward algorithm.	106
A.5	An example of Kalman smoother where the 2-dimensional noisy observations are filtered. For each panel, black circles are the observations, the green curve is the true denoised observations (hidden states) and red curve is the inferred hidden states from the noisy observations. The left and right panels use the same observations except for 5% of the contamination in the right panels.	112
A.6	Examples of the particle smoother and mixture Kalman smoother on the same 2-dimensional dataset. Black circles are the noisy observations and green and red curves are the true and inferred denoised hidden states, respectively.	121
B.1	Functions sampled from the Gaussian process with different stationary kernels.	135
B.2	Functions sampled from the Gaussian process with different nonstationary kernels.	137
B.3	A graphical representation of GPR and GPC. Capital X and Y are the training inputs and outputs, respectively, and β indicates the latent state, that is either regression (GPR) or log-odd (GPC). Green x and red y are the query input and the associated output, respectively, and blue b stands for the associated latent variable.	138
B.4	Examples of the Gaussian process regression. Panel (a) and (c) are the homoscedastic cases whereas (b) and (d) are for the heteroscedastic data.	139
B.5	Examples of the GP classification. Each panel consists of four sub-panels. The first sub-panel show the artificial data that are drawn from the distribution depicted in the third sub-panel with 10% contamination. The second and fourth sub-panels show the posterior distributions of the labels with representing their levels with the brightness and the corresponding estimation, respectively.	152

B.6	A graphical representation of GPST. Capital T and Y are the training indices (time) and outputs, respectively, and $X^{(1)}$ and $X^{(2)}$ indicate the latent states, that are hidden states to infer as those of the state-space models in chapter A. Green t and red y are the query index (time) and the associated output, respectively, and blue $x^{(1)}$ and $x^{(2)}$ stand for the associated latent variables.	153
B.7	A graphical representation of DGP. X and Y are the training inputs and outputs, respectively, and β and γ are the latent states (upper-left). Each column vector of β is assumed to follow a GP with X as the input (upper-right). Each column vector of γ is then assumed to follow another GP with β as the input (lower-left). Finally, Y is modelled given γ (lower-right).	158

1

Introduction

Since the dawn of the era of big data, statistical learning (machine learning) has been indispensable to any areas involved with large datasets thus played a crucial role in the innovation and advancement by taking advantage of the fast computer breakthroughs and Bayesian methodologies. As an example, computational biology has shown a remarkable progress in efficiently processing high-throughput data over the past few decades. Computer vision, speech recognition, computational linguistics, robotics, bioinformatics, and marketing are the other good examples.

Geoscience is not an exception. Various subfields including geophysics are actively adopting machine learning tools to analyze the relevant data, such as seismic waves, streamflow, and geomagnetic data. Paleoceanography and paleoclimatology to infer ancient events from proxies are

also depending much on statistical learning a lot in interpreting data. This dissertation aims at broadening such research with state-space models and Gaussian process applications.

I.I PALEOCEANOGRAPHY AND PALEOCLIMATOLOGY

Geoscience is a branch of science that addresses the physical and chemical constitution of the Earth and its atmosphere. It consists of the following four studies: lithosphere, hydrosphere, atmosphere and biosphere. Paleoclimatology studies various types of environmental evidence to understand the past climate of Earth. Scientists utilize the relevant recorders to infer past conditions: such records are called proxies. Such studies have figured out that the Earth's climate system can shift between extremely different states abruptly and examined the causes of past climate changes to unravel how much the recent climate changes can be explained by natural causes and human influences. Paleoceanography is the same sort of studies but related more to the oceans.

I.I.I AGE INFERENCE PROBLEMS

To reconstruct the past oceanic and climate events, one firstly needs to obtain their age information. The problem is that ages are also not measurable as well. In the field of geochronology, age inference is dependent on the proxies, just as the oceanic and climate events.

Age proxies are classified into two categories. Firstly, direct age proxies are allowing to date records for themselves. For example, ^{14}C determinations are assumed to follow the models based on radiocarbon ages, and such radiocarbon ages are translated into calendar ages by calibration curves [Reimer et al. (2013); Hogg et al. (2013)], thus one can access calendar ages via ^{14}C proxies. We call ages inferred by direct age proxies direct ages. Utilizing direct age proxies is dependent on their relationship with ages. If one wants to infer the ages of a sediment core, then the inference is also relevant to the relationship between ages and the associated layer depths.

Direct age proxies are often limited in quantity in each record and only available up to certain age ranges, thus we need complementary methods. Indirect age proxies are the global climate parameters (e.g., benthic $\delta^{18}\text{O}$) that can be utilized as reference for aligning records which have

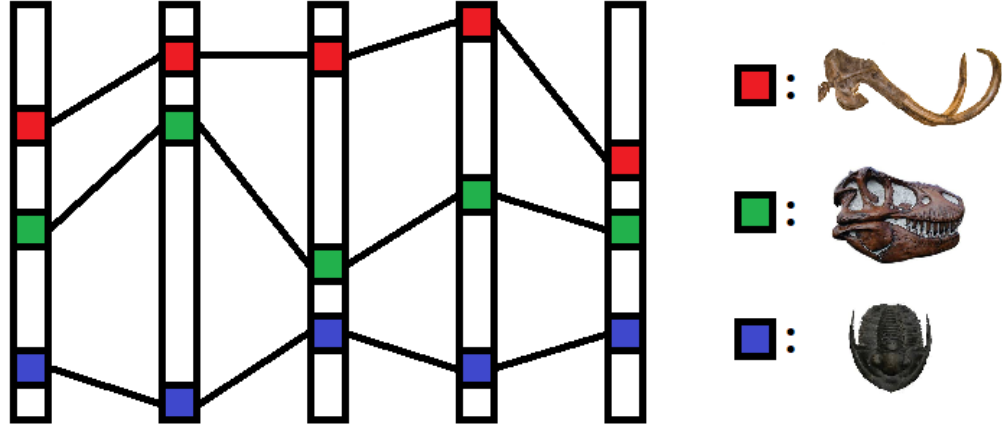


Figure 1.1: An example of how the indirect proxy is utilized for the age inference. Here, five sediment cores are aligned based on their fossil records.

direct proxies. Once records are aligned, direct ages at each record are shared with the others. Figure 1.1 shows an example. Utilizing indirect age proxies can be thus understood as the multiple sequence alignment problem. The dynamic time warping (DTW) [Munich & Perona (1999)] is not sufficient to address this problem because it does not exclude the extreme alignments based on the layer depths and the level of sedimentation rates. Also, pairwise alignment cannot be extended straightforwardly to the multiple sequence alignment because it does not guarantee consistency of the alignments and the time complexity is quadratic of the number of sequences.

Both direct and indirect age proxies can be utilized in the Bayesian framework. The likelihood is consisting of the proxy models given ages, while the prior on ages reflects the prior knowledge from the relationship between ages and secondary information, such as layer depths and sedimentation rates. From the likelihood and prior, one can derive the posterior distribution of ages by the Bayes' theorem. One can simultaneously use the multiple kinds of age proxies straightforwardly from their proxy models if they are conditionally independent given ages. We will further discuss about this topic in chapter 2.

1.1.2 RECONSTRUCTION PROBLEMS

Paleoceanography and paleoclimatology investigate the ocean events and past climate, respectively. We have no way of measuring them directly. Instead, the inference can be done by the

proxy methods that utilize the causality or correlation between some particular preserved physical characteristics and those events. There are various types of proxies.

Mountain glaciers and the polar ice caps and sheets are good sources for the data. The trapped air is a source for direct measurement of the composition of air. Changes in the layering thickness are connected to those in precipitation or temperature. $\delta^{18}\text{O}$ in ice layers represent changes in average sea surface temperature. Sediments contain remnants of preserved vegetation, animals, plankton, or pollen that characterize certain climate zones. Biomarker molecules, chemical signature and isotopic ratios carry further information. Coral rings respond to water temperature, freshwater influx, pH changes and wave action.

Climate and ocean events are highly correlated one another. Proxies could be correlated with multiple events. Statistical inference on them without assuming the underlying theory would be unreliable, especially if only a small amount of data is given. Instead, the theory gives a set of equations that represent the relationships: for example, [Dickson (2010)] characterizes the ocean carbon system with several equilibrium equations. Deriving the explicit calibration model for the climate events to infer given proxies is done by solving the system of those equations. If it is hard to derive such an explicit model, sample-based methods by Markov-chain Monte Carlo algorithm can be considered to be an alternative.

If such a theory is unavailable, then the data-driven methods would be the second-best choice. Here, to avoid overfitting that comes from the curse of dimensionality, either simple parametric models or nonparametric methods applied on only one or two kinds of proxies and events are considered. For example, [Epstein et al. (1953)] formulates the relationship between ice volumes and $\delta^{18}\text{O}$ as a quadratic extrapolation. Though this approach would seem to be limited and ignorant of the other events and proxies at first glance, it gives another way of utilizing multiple proxies by adopting those models as observation distributions in the framework of Bayesian inference.

In either case, the proxy methods depend on the present-day environmental conditions to make statements about past conditions, which is called “uniformitarianism principle” [Scott (1963)]. If the relationships can be explained theoretically based on the fields of physics and/or

chemistry, then it is not too much a stretch to assume it; if the proxies are involved with biological processes, however, then nonstationary issues that stem from the evolution can be problematic. Each proxy has its own span and resolution that characterize its availability: for example, the ^{14}C proxy is limited up to 50 kiloyears and tends to have a low resolution in a record.

This dissertation also contributes to the reconstruction of sea surface temperatures (chapter 3) with alkenone and TEX_{86} proxies and atmospheric CO_2 concentrations (chapter 4) with benthic $\delta^{18}\text{O}$ and boron proxies.

1.1.3 CHARACTERISTICS OF THE DATA

For the last decades, machine learning (or statistical learning) has become very popular in various research areas, as long as analyzing data gives more opportunity. Bayesian analysis and, more recently, deep neural networks [Goodfellow et al. (2016)] have led to innovations. This dissertation also aims at such a purpose in paleoceanography and paleoclimatology. Before applying a model, it is useful to discuss about the characteristics of our data.

Data in paleoceanography and paleoclimatology are consisting of various proxies that have been imprinted for the geological time, from a few thousands to hundreds of millions of years. It means that generating such data is usually infeasible in laboratories: even if one somehow knows how to accelerate the processes, they are also affected by the past environmental conditions (such as salinity, temperatures, acidity, ocean circulations, etc.) that are not known a priori. Thus, data are limited in the observations from the existing records.

It is also difficult to gather such data as much as we want. Firstly, some records are expensive to process. For example, ocean sediment cores are extracted from the thousands of meter deep oceans. Secondly, processing proxies in a record could be either expensive or mostly infeasible due to the tiny amounts thus only a limited number of observations are made in practice. Thirdly, records of high quality are rare: contamination by non-climatic influences could be an issue.

For such reasons, data in paleoceanography and paleoclimatology are often limited in quantity. Also, some proxies have not been fully investigated yet so there is no accurate model to describe

the relationship between them and the associated past climate events. These limitations make the inference more loyal to the data and bring more cautious analysis to the outliers: if data are many, discarding any suspicious data in the preprocessing would not lose information much. However, if data are scarce and likely to include outliers, one needs to classify them more carefully for minimizing information loss.

1.2 STATISTICAL LEARNING THEORY

Statistical learning is a framework for the inference and prediction drawing from statistics and functional analysis [[Hastie et al. \(2009\)](#)]. Nowadays the term ‘machine learning’ is used more widely: it usually aims at accurate predictions so all about the results, while statistical learning tends to focus more on the relationship between variables. The blurred boundary between those two fields comes from the modern trend that machine learning algorithms are often highly depending on the statistical theory and statistical learning is not only limited to understanding the data but also prediction based on it.

However, they are not identical. Suppose that we have a huge amount of high-dimensional data to be exploited for the prediction. A typical approach of machine learning is to search the most predictive and/or robust model to the data, usually at the cost of interpretability. Some data-driven methods such as deep neural network models sometimes outright ignore the underlying data-specific theories that have been studied by experts. This approach could be good if we do not know much about the data and have urgent demands in engineering, while the deviation from the existing theory might make the experts hesitate to take it instead of the original models that could be suboptimal but are highly interpretable in theory.

In contrast, statistical learning is first pursuing the understanding of data probabilistically: it focuses on determining the causality and correlation between variables. Once a probabilistic model on the given data is settled, it is then utilized for deriving the prediction. Though statistical learning can be suboptimal, especially if the model for the data is misspecified, it facilitates the use of previously reported theory germane to the data. Thus, this approach might be better if the

reliability of models is determined by how much they reflect data-specific characteristics.

Because this dissertation research has the purpose in giving more efficient and effective models on the scientific data in paleoceanography and paleoclimatology, the methodology is inevitably more faithful to the theory. For example, it could be fancy to align two sediment records over ages by utilizing the proxy pattern only, but the model will never attract the experts if it does not count on the sedimentation rate models: extreme alignments based on the depths lead to concluding extreme accumulation in past that could be extremely unrealistic. We will discuss more about the methodology in the following subsections.

1.2.1 FREQUENTIST VS. BAYESIAN INFERENCE

Statistics is the academic field involved with the collection, organization, analysis, interpretation, and presentation of data. Furthermore, it deals with the uncertainty that stems from the data. Though the modern statistics is intensively written in the language of mathematics, regarding statistics as a branch of mathematics is deceptive. While mathematics is building a theory with definitions and theorems that can be applied to explain the phenomena, statistics aims at constructing processes for comparing and measuring hypotheses and models to deal with uncertainty. While mathematicians design an ideal setting and wish that it is close enough to explain the reality, statisticians are engrossed in creating models that can encompass its uncertain aspects.

Debates on the frequentist vs. Bayesian reasoning have a long history. Suppose that we have a model that addresses the given data. Frequentists consider that the model is certain but unknown and the associated data can be generated given the model. Thus, they consider that data are random whereas the model is not. Bayesians have a different point of view: they consider that the given data are themselves certain so not random. Instead, they consider models to be random (prior). Data are just explained given models (likelihood). By applying Bayes' theorem [Stuart & Ord (1994)], they measure the models given data (posterior). To be more abstract, the frequentists regard an event's probability as the limit of the relative frequency in multiple trials thus objective while the Bayesian consider it to be the degree of belief so subjective. Figure 1.2

	Frequentists	Bayesians
parameters	unknown constants	random variables
data	randomly generated from models	not random but given
probability	limit of the relative frequency	reasonable expectation or belief
source	observations	observations and prior belief
decision	hypothetical test	model selection via posteriors
goal	explaining the data uncertainty	updating the belief from data

Figure 1.2: A comparison of the frequentist and Bayesian perspectives.

gives a table to compare two reasonings. Note that the modern frequentists use Bayes estimators frequently because they have good frequentist properties: see [Wasserman (2004)]. Here we just discuss the underlying perceptions and philosophies of the frequentists and Bayesians.

One obstacle from the frequentist philosophy in the field of statistical learning is that the inference and prediction become more unrealistic and overfitting to the data as the number of parameters increases. Because they consider parameters to be unknown but certain values, parameter estimation implicitly assumes that the model is true. This approach could be effective if the model has been physically justified but what we do not know are just its parameters. However, it becomes less robust if the model itself is misspecified. To deal with it, frequentists often distinguish the term ‘parameter’ and ‘hidden variable’, where the latter is a random variable that hide in the model as a latent state. Bayesians, however, can treat parameters as random variables so are conceptually free to integrate them out without introducing additional concepts: model uncertainty can be addressed consistently by doing so.

Bayesians have been continuously criticized for their unscientific aspects: if probability is essentially *subjective*, then how can we achieve consensus on the inference and prediction given data? More technically, all results are depending on the choice of priors. If one set a strong prior, then the data will be dominated by our prior belief; if a weak or flat prior is given, then the model is just bothersome without any particular advantages [Gelman (2008)]. Without depending on the conjugate priors, the posteriors are often not possible to be expressed in closed forms. Moreover, posteriors as the conjugate distributions can usually be replicated by the frequentist approach with pseudo-counts on the data, thus not that exclusive to Bayesian, just as it circumvents the

certainty of parameters by introducing hidden variables.

Despite such shortcomings and criticism of Bayesian statistics, it is often both meaningful and practical in the statistical *learning*. Firstly, it is quite common that data are given a lot whereas their characteristics are mostly unknown in practical problems. If so, the models that are applied to the data are either misspecified or involved with too many parameters, thus the overfitting rises as an issue: this can be avoided by assigning particular values to the parameters, and Bayesian inference that justifies the marginalization of parameters is thus practical. In addition, here the models themselves are assumed not to be justified but somehow selected by the solver so “subjective beliefs” in the Bayesian literature would not more be conceptually problematic. Also, data-driven approach is for deriving models from the given data, thus it is more natural to assume that data are not random (as what frequentists say) but given (as Bayesians say) in that case.

Secondly, suppose that we somehow know the characteristics of data, but not sure, then having some candidates instead of a single model is desirable. Frequentist statistics deal with such a problem by the (statistical) hypothesis test. It starts with selecting the most probable and sensible model as the null hypothesis and defines the others to be the alternative hypothesis. Then it either accepts or rejects the null hypothesis based on the relevant test statistic obtained under the null hypothesis and a significant level (or p-value). This approach is quite annoying for the multiple model selection and implies that the tester already has a certain belief on the models: why is one supposed to be null and the other alternative? Bayesian model selection gives a succinct and explicit solution to this problem: it just computes the posteriors of models given their prior beliefs. Moreover, marginalizing models out to get the posterior predictive model is another option for reducing uncertainty.

Thirdly, in some real problems, it is probable that partial information of the parameters and variables to learn is available independent of the data. Bayesian framework gives clear means to deal with this: priors on the parameters and variables are interpreting such information. If so, the priors are not representing the *subjective* beliefs on the values to estimate anymore: they are now *objective* in some sense. Furthermore, Bayesian framework is more intuitive in this case:

researchers are always going forward from the prior knowledge of the objectives to study and supplementing the knowledge with the new findings from the research. This is equivalent to calculating the posterior of the unknown from the prior and likelihood.

Lastly, Bayesian approach could be more robust and stable in practice, especially if the data are scarce. In discrete problems, (a strict form of) frequentist approach might assign zero probabilities to some states: if no observation at a state is made, then it is estimated to never happens. If such possibilities are open, some useful and powerful results in the probability theory cannot be applied. For example, the Hammersley-Clifford theorem [Clifford (1990)] works under the assumption that probability distributions are strictly positive. Giving a nonzero prior can resolve such singularity at once. Also, Bayesian framework prevents the estimation from being sensitive to the initial conditions that might lead estimations to local optima, especially if the algorithm consists of iterations such as Gibbs sampling [Geman & Geman (1984)]. Here, priors are working as “anchors” so that the estimate or sampled parameters or hidden variables are not deviated much from the “probable” ranges.

For the above reasons, I have been working on my research topics as a Bayesian statistician, and this dissertation is also written in the Bayesian language. As discussed in section 1.1.3, data in paleoceanography and paleoclimatology are often lacking an underlying physical dynamics model compared to the other fields that statistical/machine learning are applied to. It is impossible or expensive to generate or gather such data. Thus, Bayesian philosophy that assumes data to be given is practical for this problem.

1.2.2 PARAMETRIC VS. NONPARAMETRIC MODEL

A statistical model is a set of distributions, densities, or regression functions [Wasserman (2004)]. Statistical models can be classified into the following two categories: parametric and nonparametric models. A parametric model can be parametrized by a finite number of parameters that take values in a certain parametric space, which is given independently from the data. A nonparametric model assumes that the data cannot be explained in terms of such a finite set of parameters

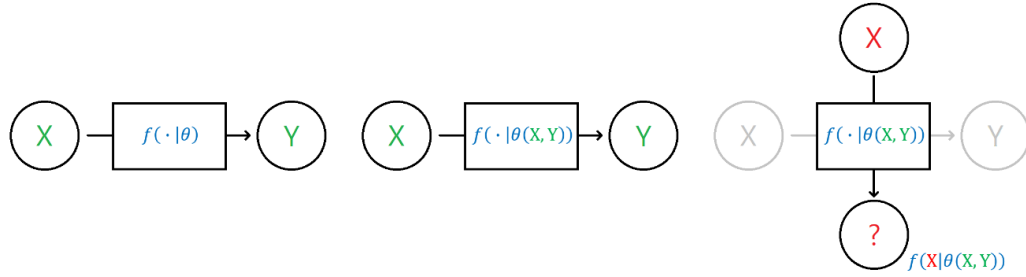


Figure 1.3: A diagram of the parametric inference. The relationship between two variables X and Y are explained by the parametric model $f(\cdot|\theta)$ (left). The parameter θ is learned from the training data (X, Y) colored green (middle). The unknown Y is explained by the query X colored red and learned parameter θ (right).

only. More intuitively, parametric models are described by the functions that are defined by the parameters only, where nonparametric ones are those by not only parameters (often called hyperparameters) but also data themselves to depend on the structural assumptions as few as possible [Wasserman (2006)].

Statistical learning with parametric models can be understood as inferring parameters from the data. Frequentists have some criteria for this procedure, such as the maximum likelihood estimator and (expected) loss, based on the training and validating data. Bayesians consider parameters to be random, so have other criteria for inferring parameters, such as maximum a posteriori probability estimator and Bayesian expected loss. Once a set of parameters are learned, the model is fixed because it is dependent on the parameters only. One more thing that Bayesians can do is to marginalize out those parameters over their posterior distributions to obtain a posterior predictive model: here the model is controlled by the prior hyperparameters. Though fast inference is possible once the parameters are learned, the model complexity is bounded so less flexible in this approach: it is not problematic if the model is investigated to describe the data exactly thus parameters are themselves meaningful, but could be problematic if the model is misspecified. If so, either overfitting or suboptimality could be an issue to deal with. Polynomial regression, logistic regression, finite mixture models, hidden Markov models, K-means, neural networks are examples of the parametric model. Figure 1.3 shows a diagram of the procedure for the parametric inference.

Though the nonparametric models are more data-dependent and free from the structural

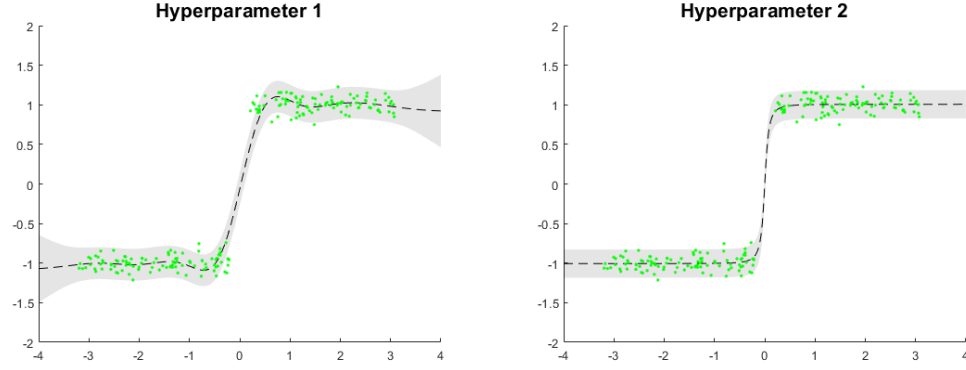


Figure 1.4: Two nonparametric Gaussian process regressions on the same dataset but with different hyperparameters.

models, they are controlled by hyperparameters: see figure 1.4. Thus, nonparametric learning is about tuning hyperparameters, such as bandwidths and number of clusters. One major difference from the parametric learning is that such hyperparameters can often not be estimated by maximizing likelihoods or minimizing losses: such attempts lead the model to be ignorantly loyal to the training data. Instead, they are “tuned” by optimizing objective functions that are regulated for avoiding the overfitting to the training data. Because the tuned model is still a function of data, inference could be done slowly if the time complexity is super-linear and data size is large. A nonparametric model is structural-free so more flexible than the parametric ones, thus useful if any parametric model is hypothetical. Moreover, it is controlled by a small number of hyperparameters thus can avoid overfitting from the lack of data compared to the number of parameters in the parametric models. However, nonparametric models could be more susceptible to suffer from the curse of dimensionality [Geenens (2011)] that means the number of training data to guarantee reliable inference increases tremendously as the dimension of predictors grows. Dirichlet process, Gaussian process, kernel density estimation, K-nearest neighbors and support vector machine with a positive definite kernel are examples of the nonparametric model. Figure 1.5 shows a diagram of the procedure for the nonparametric inference.

This dissertation mainly depends on the following two models: state-space models and the Gaussian process. The formers are mostly parametric whereas the latter is nonparametric. If reliable or justifiable models are given by the data-experts, then adopting parametric models is more

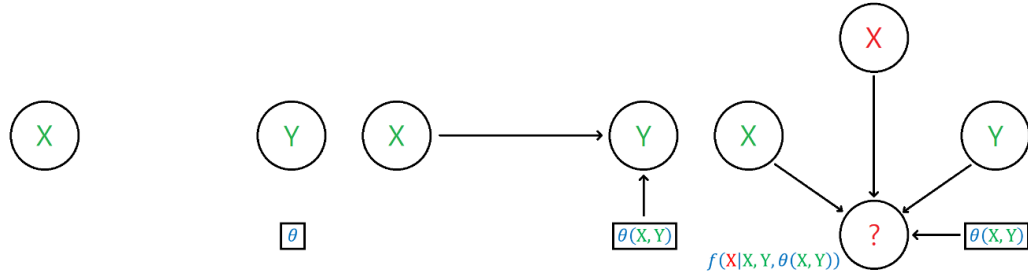


Figure 1.5: A diagram of the nonparametric inference. The relationship between two variables X and Y is not explained by any parametric model but is indirectly controlled by the hyperparameter θ (left). The hyperparameter θ is tuned from the training data (X, Y) colored green (middle). The unknown Y is explained by the query X colored red, tuned hyperparameter θ , and training data colored green (right).

advantageous in the data analysis and learning procedure. Otherwise, data-driven agnostic models could be considered. If the amount of available training data is large, then models based on the deep neural network have been reported to work well in practice. However, due to their excessive number of parameters, it would lead to overfitting to apply such models to the moderate number of data. In such cases, nonparametric models that are depending directly on the training data can be chosen for the data-driven purpose. Note that there is no panacea in statistical learning: the choice of models should be made based on the characteristics of data.

1.2.3 GENERATIVE VS. DISCRIMINATIVE MODEL

Statistical learning is a statistical attempt to find the relationship between variables. There are two types of variables: independent and dependent variables that can be understood as the observable inputs (predictors) and target outcomes resulting from varying inputs, respectively. Practically, what the statistical learning results in are the distribution function of the dependent variable (output) given independent one (input) from the training data. From now on the terms *observations* and *targets* will be used for expressing independent and dependent variables, respectively. Note that this classification is a bit hypothetical and goal-oriented: variables needs not to be the result of other variables.

A generative model starts with investigating the variables without such classification: in mathematical words, it aims at constructing the joint distribution of observations and targets that can

be utilized for deriving the conditional distribution of targets given observations by the Bayes rule. The definition is equivalent to the construction of conditional distribution of observations given targets if a prior on the targets is given in Bayesian terminology: “a target *generates* the associated observation”. This approach is appropriate if 1) the structure of the relationship between the observations and targets are given a priori by experts so only the parameters are needed to be learned or 2) the conditional distribution of observations given targets is easy to deal with – for example, a hidden Markov model assumes that observations are conditionally independent given targets (hidden states). Generative models also give a way of generating pseudo-observations artificially. Mixture models, hidden Markov models, Bayesian networks, latent Dirichlet allocation, Boltzmann machines, variational autoencoders, and energy based-models are the examples of generative models.

A discriminative model is more pragmatic: it directly approximates the conditional distribution (or estimates, if the model is deterministic) of targets given observations. Such distributions are usually hypothetical and independent from the data-specific characteristics because observations are not actually generating the targets. This approach is usually effective if the training data are plentiful and targets are defined in a small finite set or continuous but low-dimensional. However, it either becomes intractable or suffers the curse of dimensionality especially if the targets are defined in a high-dimensional space and expected to be highly correlated one another. K-nearest neighbors, logistic regression, support vector machines, maximum-entropy Markov models and neural networks are the typical examples. Figure 1.6 gives the graphical representations of generative and discriminative models.

In paleoceanography and paleoclimatology, proxies are the raw observations and past climate events or ages are the targets to infer. Here it is reasonable to perceive that past climate events have been imprinting the proxies. Also, it is quite common that multidisciplinary research gives a meaningful prior on the past climate events. Some events, such as ages, are highly correlated one another. Thus, inferences in this dissertation are done in the framework of generative models.

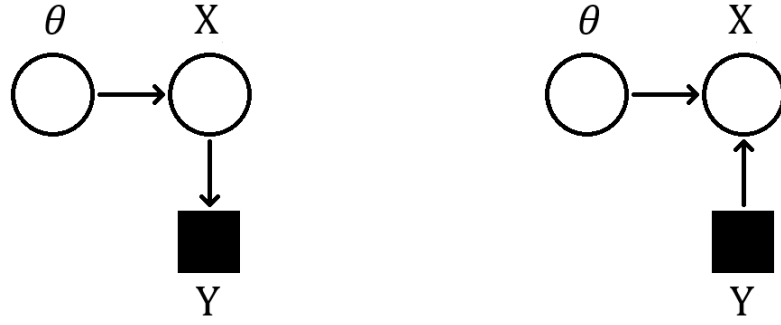


Figure 1.6: Graphical representations of the generative model (left) and discriminative model (right). Here, θ is a parameter and X is a hidden variable that explains the observation Y .

1.3 CONTENTS

As discussed so far, chapter 1 covers the basic concepts and knowledge to give a departure point of the dissertation. It mainly consists of two parts. One is related to the application field, paleoceanography and paleoclimatology, about their issues to be addressed and characteristics distinctive from the other scientific fields. The other briefly introduces the statistical learning theory with some discussions on the methodologies that fit for the application area. Notation and symbols that continuously appear in the rest of the dissertation are covered in the latter part of this chapter.

Chapters 2, 3 and 4 present the applications to paleoceanography and paleoclimatology. Chapter 2 covers how an indirect age proxy, benthic $\delta^{18}\text{O}$, is utilized together with a direct age proxy, ^{14}C , for inferring ages by the multiple sequence alignment algorithm. After briefly introducing the related work, it firstly shows the multiple core alignment algorithm given a reference (or motif/profile), based on the hybrid of particle smoother and Markov-chain Monte Carlo algorithm. Then it presents how to construct a stack that is used as the reference of the multiple core alignment algorithm given core alignments, based on the Gaussian process regression. How these two algorithms form a complete core alignment algorithm and related future work will be covered lastly.

Alkenone and TEX_{86} are proxies for reconstructing sea surface temperatures. Chapter 3 is about the calibration model construction for utilizing alkenone and TEX_{86} based on the heteroscedastic Gaussian process regression. Firstly, an extended version of the heteroscedastic Gaussian process regression from my published work is covered briefly. Then the detailed formulation of novel multimodal Gaussian process regression method is discussed. Lastly, we compare and analyze the constructed calibration models according to the introduced methodologies.

Some climate events are correlated so can be inferred simultaneously from the associated proxies. For example, atmospheric CO_2 pressure and ice volumes are correlated thus can be inferred together from the boron isotopes and benthic $\delta^{18}\text{O}$. Chapter 4 suggests two different models for this work. The first model uses the particle smoother to design the model based on the dynamic system, while the second one is derived by the Gaussian process state-space model. The comparison between the models and relevant future work will be discussed at the end of the chapter.

Finally, the dissertation ends with chapter 5 that covers the conclusion. In Appendix, two additional chapters cover the theoretical background used throughout this dissertation and contain detailed computations that are omitted in the main body to emphasize on our novel works. To be specific, several topics about state-space models are presented in Appendix A. It starts with discrete models such as the hidden Markov model and then deals with continuous but exact models such as the Kalman filter and smoother. Because such exact models are quite limited in some specific cases, approximation methods based on the sampling such as particle filter and smoother and Markov-chain Monte Carlo algorithm follow.

Appendix B discusses Gaussian processes and their applications in statistical learning. After defining some relevant concepts, it classifies kernel functions that determine the performance. Then the applications of regression, classification, state-space models, and dimensionality reduction follow. Lastly, the deep Gaussian process and variational inference are covered to handle the drawbacks of Gaussian processes that stem from the suboptimality and time complexity.

I.4 DICTIONARY

This section lists the notation, symbols and acronyms that are used in the rest of my dissertation. Acronyms are redefined at every chapter.

I.4.1 NOTATION

A Gaussian process needs the mean and kernel covariance functions with data input $\mathbf{X}$ that could be multidimensional, i.e., it has multiple features. We assume that each row of $\mathbf{X}$ represents the input datum. For example, if $\mathbf{X}$ is given as a $N \times D$ matrix, then it has N data and each datum is D -dimensional. To reflect its consistent extensibility, this dissertation abuses some notations. Let μ be the mean function. Then, for any $N \times D$ matrix $\mathbf{X}$, $\mu(\mathbf{X})$ is the $N \times 1$ vector where its n th entry is $\mu(\mathbf{X}_n)$. Similarly, for a kernel covariance function $\mathbb{K}$ and matrices $\mathbf{X}^{(1)}$ and $\mathbf{X}^{(2)}$ that have sizes $N_1 \times D$ and $N_2 \times D$, respectively, $\mathbb{K}(\mathbf{X}^{(1)}, \mathbf{X}^{(2)})$ represents the kernel covariance matrix where its entry at the i th row and j th column is $\mathbb{K}(\mathbf{X}_{i,:}^{(1)}, \mathbf{X}_{j,:}^{(2)})$. To save the space, let $\mu_{\mathbf{X}} = \mu(\mathbf{X})$ and $\mathbb{K}_{\mathbf{X}^{(1)}, \mathbf{X}^{(2)}} = \mathbb{K}(\mathbf{X}^{(1)}, \mathbf{X}^{(2)})$. In the regression models, the observational variance function Λ is additionally required. For any $N \times D$ matrix $\mathbf{X}$, $\Lambda(\mathbf{X})$ is the $N \times N$ diagonal matrix where its n th diagonal entry is $\Lambda(\mathbf{X}_n)$. Similarly, let $\Lambda_{\mathbf{X}} = \Lambda(\mathbf{X})$.

I.4.2 SYMBOLS

$\propto$	proportional to
$\sim$	distributed by
$\mathbb{R}$	the set of real numbers
$\mathbb{R}_+$	the set of nonnegative real numbers
$\vec{0}$	the zero vector
$\mathbb{I}$	the identity matrix
$\mathbf{v} = \mathbf{v}_{1:T}$	a T -dimensional (column) vector
v_i	the i th entry of a vector $\mathbf{v}$

$v_{i:j}$	the truncated vector of a vector v by selecting the entries from i to j
$v^{\mathbb{T}}$	the transpose of a vector v
$diag(v)$	the diagonal matrix where the diagonal entries are from a vector v
$A = [A_{ij}]_{i,j=1}^{M,N}$	a $M \times N$ matrix
A_{ij}	the entry from the i th row and j th column of a matrix A
$A_{i,:}$	the i th row of a matrix A
$A_{:,j}$	the j th column of a matrix A
$A^{\mathbb{T}}$	the transpose of a matrix A
A^{-1}	the inverse of a matrix A
$ A $	the determinant of a matrix A
$diag(A)$	the (column) vector consisting of the diagonal entries of a matrix A
$p(y \Theta)$	the conditional distribution of y given Θ
$\mathcal{N}(y \mu, \Sigma)$	the (multivariate) normal probability density at y with a mean vector μ and covariance matrix Σ
$\mathcal{GP}(\beta \mu, \mathbb{K}_{XX})$	the Gaussian process on β with an input X , mean function μ and covariance kernel function $\mathbb{K}$
$\Gamma(y \alpha, \beta)$	the Gamma probability density at y with a shape parameter α and rate parameter β
$\mathcal{T}(y \mu, \sigma^2, \nu)$	the generalized Student's t-distribution at y with a location parameter μ , scale parameter σ and degree of freedom ν
$\mathbb{E}_{p(X)}[f(X)]$	the expectation of $f(X)$ over the random variable X that follows a probability distribution p
$\mathbb{D}_{KL}(q p)$	the Kullback-Leibler divergence from p to q

I.4.3 ACRONYMS

AR	autoregressive
BIC	Bayesian information criterion
CO ₂	carbon dioxide
$\delta^{11}\text{B}$	a variation derived from the boron isotopes ¹¹ B and ¹⁰ B
$\delta^{18}\text{O}$	a measure of the ratio of stable isotopes ¹⁸ O and ¹⁶ O
DGP	deep Gaussian process
DNN	deep neural network
DTW	dynamic time warping
ELBO	evidence lower bound
EM	expectation-maximization
EP	expectation propagation
FITC	fully independent training conditional
GP	Gaussian process
GPC	Gaussian process classification
GPLVM	Gaussian process latent variable model
GPR	Gaussian process regression
GPST	Gaussian process state-space model
HMC	Hamiltonian Monte Carlo
HMM	hidden Markov model
KL	Kullback-Leibler
LOO-CV	leave-one-out cross-validation
MAP	maximum a posteriori
MCMC	Markov-chain Monte Carlo
ML	maximum likelihood
NN	neural network

NWR	Nadaraya-Watson kernel regression
ODE	ordinary differential equation
OU	Ornstein-Uhlenbeck
PPCA	probabilistic principal component analysis
RNN	recurrent neural network
SE	squared exponential
SSM	state-space model
SST	sea surface temperature
TEX ₈₆	tetraether index of 86 carbon atoms
U ₃₇ ^{K'}	unsaturation index of di- versus tri- unsaturated C ₃₇ alkenones
VFE	variational free energy

2

Benthic $\delta^{18}\text{O}$ Stack Construction

This chapter is based on the paper in preparation [[Lee et al. \(2019\)](#)] (The previous title is 'Dual Proxy Gaussian Process Stack: Integrating Benthic $\delta^{18}\text{O}$ and Radiocarbon Proxies for Inferring Ages on Ocean Sediment Cores'). Thus, some expressions and figures are the same with those in that paper. This work has been done with the collaboration of Dr. Lorraine Lisiecki, Devin Rand, Dr. Geoffrey Gebbie and my advisor, Dr. Charles Lawrence.

$\delta^{18}\text{O}$ is the ratio of stable isotopes ^{18}O and ^{16}O relative to the laboratory standard. It is used as a proxy for inferring the past water temperatures, salinity and global ice volumes in paleoceanography. The $\delta^{18}\text{O}$ of benthic foraminifera (benthic $\delta^{18}\text{O}$) live on the seafloor and it is regarded as a global climate parameter because about half of its variance through time is explained by the

global ice volume [Spratt & Lisiecki (2016)] while deep-water temperature and salinity are relatively similar over space.

The characteristic that benthic $\delta^{18}\text{O}$ is a global parameter has been utilized for the alignment of ocean sediment cores. Once cores are aligned, direct age information obtained by the direct age proxies separately stored in each core is transferred from one to another, thus more complete and comprehensive age inferences to all aligned cores are yielded.

The multiple core alignment problem is not straightforwardly resolved by the pairwise alignment algorithms for the following reasons: 1) The number of pairs is quadratic to the number of cores thus considering all such pairwise alignments becomes intractable as the number of cores increases, and 2) Pairwise alignments do not guarantee the consistency of the multiple alignments, i.e., the pairwise alignment between A and C indirectly obtained by those of core A and B and of B and C would not be consistent to the direct pairwise alignment of A and C.

Inspired by the profile hidden Markov models (HMMs) in computational biology [Durbin et al. (1998)], iterating the step of aligning each core to a profile (or a reference in more common terminology) and that of constructing (updating) the profile from the alignments is an alternative solution for the multiple core alignment problem. For benthic $\delta^{18}\text{O}$, the term ‘stack’ [Lisiecki & Raymo (2005)] is instead used as ‘profile’ that is more prevalent in computational biology.

Thus, the multiple core alignment with stacks is involved with the following two algorithms: the core alignment algorithm and the stack construction algorithm. In this section, we design a continuous core alignment algorithm that simultaneously utilizes ^{14}C as the direct proxy and benthic $\delta^{18}\text{O}$ as the indirect one and a continuous stack construction algorithm based on the homoscedastic Gaussian process regression (GPR) in section B.2.1.

2.1 RELATED WORK

For the direct age inference, radiocarbon (^{14}C) is frequently used as a proxy by allowing to access calendar ages associated to the sampled core depth based on the age models and calibration curves [Reimer et al. (2013); Hogg et al. (2013)]. [Christen & Sergio (2009)] designed an algo-

rithm for constructing sediment core age models from ^{14}C data based on a generalized Student's t-distribution: it addresses the uncertainties from ^{14}C measurements and reservoir effects.

Moreover, various studies have been developing dynamic models that also reflect the varying sediment accumulation rates over time. [Blaauw & Christen (2005)] adopt a piecewise linear approach with automatic section selection and impose constraints on accumulation rates. [Haslett & Parnell (2008)] choose a bivariate monotone Markov process with gamma increments. [Blaauw & Christen (2011)] develop a software which is called Bacon (Bayesian Accumulation) that adopts an autoregressive gamma process for accumulation rates, a generalized Student's t-distribution for radiocarbon data and an adaptive Markov-chain Monte Carlo algorithm for sampling ages.

The benthic $\delta^{18}\text{O}$ utilization for indirect age inferences has the following history. [Lisiecki & Lisiecki (2002)] developed a deterministic alignment algorithm, Match, that uses dynamic programming. [Lisiecki & Raymo (2005)] ran Match on benthic $\delta^{18}\text{O}$ data from 57 globally distributed deep-sea sediment cores for aligning them and calculating a stack, LR04: it is commonly used as a standard reference for benthic $\delta^{18}\text{O}$ change over the past 5.3 Myr.

[Lin et al. (2014)] designed an algorithm HMM-Match based on a probabilistic alignment model to the stack. Its emission model for $\delta^{18}\text{O}$ data is based on Gaussian distribution with time varying mean $\delta^{18}\text{O}$ values from LR04 and a constant core-dependent standard deviation learned by the Baum-Welch algorithm. The transition model explains the probability distribution of accumulation rates based on a log-normal mixture learned by radiocarbon observations from 37 cores. [Ahn (2016)] constructed a new Prob-stack from 180 globally distributed benthic $\delta^{18}\text{O}$ records with an algorithm HMM-Stack that is based on HMM-Match. Prob-stack allows the change of variances over time.

HMM-based models cannot avoid assuming inferred ages on discrete spaces, thus ages must be discretized. This is problematic for the alignments when the cores have higher resolution than the target stack or when a proportional accumulation rate model is chosen as the above-mentioned algorithms. Increasing the resolution of the target stack becomes intractable since the time com-

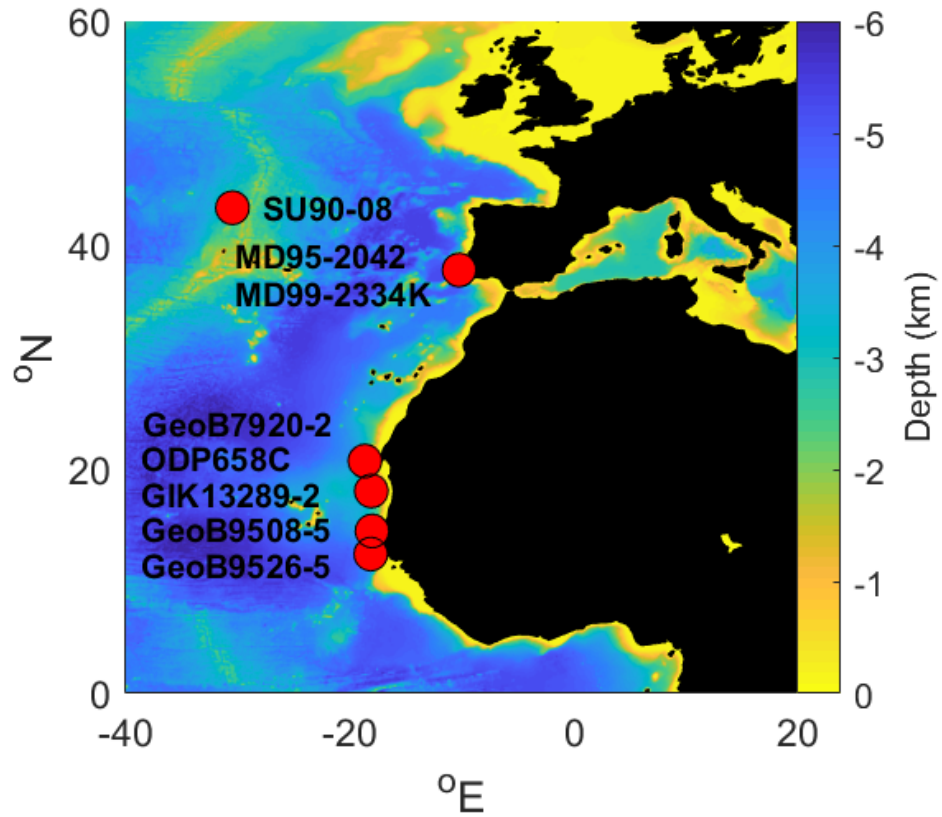
plexity of HMMs is quadratic to the size of hidden spaces. Stack construction algorithms based on the expectation-maximization (EM) algorithm [Dempster et al. (1977)] such as Baum-Welch (section A.1.3) also have a structural problem that stems from the relationship between the resolution and coverage of the cores: stack construction requires a number of cores that have $\delta^{18}\text{O}$ data enough to cover its query ages. Simple interpolation is not a good solution because it does not account for the increasing uncertainty from missing observations.

The following subsections demonstrate the core alignment algorithm and stack construction algorithm that are based on the Signal Alignment algorithm with Gaussian Process Regression profiles (SA-GPR) [Lee et al. (2019)]. Ages are sampled from continuous spaces in the alignment algorithm based on the particle smoother (section A.3.2) and Metropolis-Hastings (section A.4.2), thus the model does not suffer from the resolution-relevant drawbacks as the HMM-based ones. The stack is constructed from the sample-specific GPR models (section B.2.1), so the construction can be done regardless of the resolutions and number of cores.

2.2 RESEARCH MOTIVE AND DATA

The model assumes that benthic $\delta^{18}\text{O}$ data in the input cores show the same pattern over ages: it is roughly true for all ocean sediment cores because benthic $\delta^{18}\text{O}$ is considered as a global parameter, but the local variability between core locations might cause discordance in millennial scale events. [Skinner & Shackleton (2005); Stern & Lisiecki (2014)] report that local effects could result in age model errors up to 4 kiloyears. Such local variability possibly stems from the asynchronous temperature changes and ocean circulation rates by ice melt after the Last Glacial Maximum (19-23 kiloyears ago) [Adkins et al. (2002); McManus et al. (2004); Waelbroeck et al. (2011); Gebbie & Huybers (2012)].

The purpose of our research is for capturing such local variability by constructing a stack for each set of records that share the same local variability: we define such a set to be homogeneous and the associated stack is called a local stack. To validate our stack construction method, we ran six ocean sediment cores, GeoB7920-2, GeoB9508-5, GeoB9526-5, MD95-2042, MD99-2334



Core	Latitude	Longitude	Depth (m)	Citation
GeoB7920-2	20.75	−18.58	2278	Tjallingii et al. (2008) , Collins et al. (2011)
GeoB9508-5	14.5	−17.95	2384	Mulitza et al. (2008)
GeoB9526-5	12.44	−18.06	3223	Zarriess & Mackensen (2011) , Zarriess et al. (2011)
GIK13289-2	18.07	−18.01	2485	Sarnthein et al. (1994)
MD95-2042	37.8	−10.17	3146	Bard et al. (2004) , Shackleton et al. (2000) , Shackleton et al. (2004)
MD99-2334	37.8	−10.17	3146	Skinner et al. (2003) , Skinner & Shackleton (2004)
ODP658C	20.75	−18.58	2273	Demenocal et al. (2000)
SU90-08	43.35	−30.41	3080	Paterne et al. (1999) , Grousset (1993)

Figure 2.1: Locations of cores GeoB7920-2, GeoB9508-5, GeoB9526-5, MD95-2042, MD99-2334, ODP658C, GIK13289-2 and SU90-08.

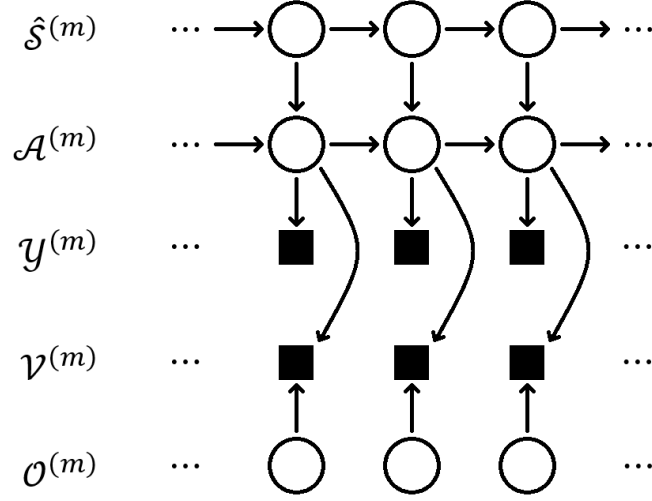


Figure 2.2: A graphical representation of our core alignment model.

and ODP658C, to construct a dual proxy local stack, Deep Northeastern Atlantic (DNEA) stack. We then use this local stack to infer dual proxy ages of two records which are homogeneous but not used in constructing the local stack, GIK13289-2 and SU90-08. Details of the data are given in figure 2.1.

2.3 CORE ALIGNMENT ALGORITHM

The core alignment algorithm aims at inferring ages at the query depths in the input cores: i.e., cores to be aligned are synchronized by their age assignments. Our probabilistic model does not return the ages as their estimates but samples, thus the results are given as a set of sampled age paths for each input core. Before giving the detailed demonstration, the following notations are defined:

- $\mathcal{D}^{(m)} = \{d_n^{(m)}\}_{n=1}^{N_m}$: depths of core m , shallower to deeper, at which the observations are made.
- $\mathcal{A}^{(m)} = \{A_n^{(m)}\}_{n=1}^{B_m}$: hidden ages of core m that are assigned to $\mathcal{D}^{(m)}$.
- $\mathcal{V}^{(m)} = \{v_n^{(m)}\}_{n=1}^{N_m}$: benthic $\delta^{18}\text{O}$ observations at $\mathcal{D}^{(m)}$ of core m . If no value is assigned,

leave it $\emptyset$.

- $\mathcal{Y}^{(m)} = \{\mathcal{Y}_n^{(m)}\}_{n=1}^{N_m}$: ^{14}C measurements at $\mathcal{D}^{(m)}$ of core m . If no value is assigned, leave it $\emptyset$.

The accumulation (sedimentation) rate model controls the ages given depths. Once ages are given to the depths, the models that associated benthic $\delta^{18}\text{O}$ observations or ^{14}C measurements follow are determined by the stack (for benthic $\delta^{18}\text{O}$) or the calibration model (for ^{14}C). Because the vehicles that contains benthic $\delta^{18}\text{O}$ and ^{14}C are different, it does not go too far assuming that they are conditionally independent given ages. Thus, this core alignment problem can be interpreted as a state-space model. A graphical representation of our core alignment model is described in figure 2.2.

2.3.1 TRANSITION MODEL

To define the transition model on the ages given depths, we utilize the dynamics of the accumulation rates. An accumulation rate from $d_{n+1}^{(m)}$ to $d_n^{(m)}$ is defined as follows:

$$S_n^{(m)} \triangleq \frac{A_{n+1}^{(m)} - A_n^{(m)}}{r^{(m)} (d_{n+1}^{(m)} - d_n^{(m)})} \quad (2.1)$$

, where $r^{(m)}$ is a depth-scale constant parameter for standardizing the accumulation rates.

While ages are still defined continuously, the accumulation rates are discretized into the following three states: $\{\mathbb{C}, \mathbb{A}, \mathbb{E}\}$ that stand for contraction, average and expansion, respectively. Let $\hat{S}_n^{(m)}$ be the discretization of $S_n^{(m)}$. Once $\hat{S}_n^{(m)}$ is given, the transition model on the age is determined as follows:

$$p\left(A_n^{(m)} \middle| A_{n+1}^{(m)}, \hat{S}_n^{(m)}\right) \propto \Gamma\left(\frac{A_{n+1}^{(m)} - A_n^{(m)}}{r^{(m)} (d_{n+1}^{(m)} - d_n^{(m)})} \middle| \alpha, \beta\right) \cdot 1\left\{\frac{A_{n+1}^{(m)} - A_n^{(m)}}{r^{(m)} (d_{n+1}^{(m)} - d_n^{(m)})} \in I_{\hat{S}_n^{(m)}}\right\} \quad (2.2)$$

, where $I_{\mathbb{C}} \triangleq (m, 0.9220)$, $I_{\mathbb{A}} \triangleq [0.9220, 1.0850)$ and $I_{\mathbb{E}} \triangleq [1.0850, M)$ for some thresholds

$0.9220 > m \geq 0$ and $M > 1.0850$ to prevent extreme events.

We also give a dependency on the discretized accumulation rates for a state-transition matrix φ as follows:

$$p\left(\hat{S}_n^{(m)} \middle| \hat{S}_{n+1}^{(m)}\right) \triangleq \varphi_{\hat{S}_{n+1}^{(m)}, \hat{S}_n^{(m)}}^{(m)} \quad (2.3)$$

From (2.2) and (2.3), we have the following transition model:

$$p_t\left(A_n^{(m)}, \hat{S}_n^{(m)} \middle| A_{n+1}^{(m)}, \hat{S}_{n+1}^{(m)}\right) \triangleq p\left(A_n^{(m)} \middle| A_{n+1}^{(m)}, \hat{S}_n^{(m)}\right) p\left(\hat{S}_n^{(m)} \middle| \hat{S}_{n+1}^{(m)}\right) \quad (2.4)$$

Thus, our transition model is of two layers and not continuous. Also, note that it is defined backwards: it follows the way how layers are accumulated – from deeper to shallower. By definitions, the transition model guarantees the monotonicity of ages: the deeper the layer is, the older it is. The shape and rate hyperparameters α and β of the gamma distribution in (2.2) are fixed in learning to avoid overfitting: their values are learned a priori from the same ^{14}C dataset for learning the log-normal mixture model in [Lin et al. (2014)]. A gamma distribution fits into the data better than the log-normal mixture and shows thicker tails between the thresholds, thus more robust on the semi-extreme changes in the accumulation rates.

2.3.2 EMISSION MODEL

Here we specify the emission model $p_e(v_n^{(m)}, \gamma_n^{(m)} | A_n^{(m)}) = p(v_n^{(m)} | A_n^{(m)}) p(\gamma_n^{(m)} | A_n^{(m)})$ by defining each observational model $p(v_n^{(m)} | A_n^{(m)})$ and $p(\gamma_n^{(m)} | A_n^{(m)})$.

For benthic $\delta^{18}\text{O}$, we basically accept the assumption that each $\delta^{18}\text{O}$ observation follows a Gaussian distribution with the mean μ_s and variance σ_s^2 defined in the target stack at the given age, as [Ahn (2016)] does. However, this assumption is not robust on the real data that might be contaminated by the noninformative outliers. To reflect it we also add another set of the hidden indicator states, $\mathcal{O}^{(m)} = \{\mathcal{O}_n^{(m)}\}_{n=1}^{N_m}$, that determines whether or not the associated $\gamma_n^{(m)}$ is an

outlier.

$$\begin{aligned}
p(O_n^{(m)}) &\triangleq \text{Bernoulli}(O_n^{(m)} | 0.05) \\
p(v_n^{(m)} | A_n^{(m)}, O_n^{(m)} = 0) &= \mathcal{N}(v_n^{(m)} | \mu_S(A_n^{(m)}) + b^{(m)}, \sigma_S^2(A_n^{(m)})) \\
p(v_n^{(m)} | A_n^{(m)}, O_n^{(m)} = 1) &= \frac{1}{2} \sum_{k=1}^2 \mathcal{N}(v_n^{(m)} | \mu_S(A_n^{(m)}) + 3(-1)^k \sigma_S(A_n^{(m)}) + b^{(m)}, \sigma_S^2(A_n^{(m)}))
\end{aligned} \tag{2.5}$$

, where $b^{(m)}$ is a core-specific shift parameter that synchronizes the $\delta^{18}\text{O}$ observations vertically.

As (2.5), $O_n^{(m)} = 1$ if and only if $y_n^{(m)}$ is an outlier.

From (2.5), $p(v_n^{(m)} | A_n^{(m)})$ is obtained by marginalizing $O_n^{(m)}$ as follows:

$$p(v_n^{(m)} | A_n^{(m)}) = 0.05p(v_n^{(m)} | A_n^{(m)}, O_n^{(m)} = 1) + 0.95p(v_n^{(m)} | A_n^{(m)}, O_n^{(m)} = 0) \tag{2.6}$$

For ^{14}C , we adopt the generalized Student's t-distribution in [Christen & Sergio (2009)] with the calibration curve that consists of the mean μ_C and variance σ_C^2 defined over the calendar ages:

$$p(y_n^{(m)} | A_n^{(m)}) = \mathcal{T}\left(y_n^{(m)} | \mu_C(A_n^{(m)}) + \varrho_n^{(m)}, \frac{b}{a}(\sigma_C^2(A_n^{(m)}) + \varsigma_n^{(m)}), 2a\right) \tag{2.7}$$

, where $\varrho_n^{(m)}$ and $\varsigma_n^{(m)}$ are the input variables (assumed to be fixed) that are given with $y_n^{(m)}$ a priori, and $a = 3$ and $b = 4$ are fixed hyperparameters to control the model variability: we assume the same values that [Christen & Sergio (2009)] suggested.

Note that the above equalities hold if and only if $v_n^{(m)}$ and $y_n^{(m)}$ are assigned by the real values: if $\emptyset$ is assigned instead, then the values are identical to 1, which implies "no information".

The reason why we do not model $p(v_n^{(m)} | A_n^{(m)})$ based on the Student's t-distribution just as $p(y_n^{(m)} | A_n^{(m)})$ is because eventually the outliers in the sampled age paths are discarded before the stack construction algorithm. Note that (2.5) derives the posterior $p(O_n^{(m)} | A_n^{(m)}, v_n^{(m)})$ once both $A_n^{(m)}$ (as a sample) and $v_n^{(m)}$ are given.

2.3.3 DETAILS ON THE INFERENCE

Now we are prepared: a state-space model is fully specified by its transition and emission models. However, model formulation does not tell anything about the tractable algorithm for inference. Here, we adopt a hybrid of the particle smoother (section A.3.2) for initializing sampled age paths and Metropolis-Hastings algorithm (section A.4.2) for refining those paths. Because each $A_n^{(m)}$ is one-dimensional, the proposal distribution can be localized: i.e., after running a few iterations, fix the region at each depth and draw particles in those regions only. Parameters are learned by the EM algorithm from the sampled age paths. Algorithm 1 summarizes the inference procedure of our core alignment algorithm.

Algorithm 1: Core Alignment Algorithm

```

1 initialize parameters
2 for each core do
3   while converging do
4     run particle smoother on the data to draw sampled age paths
5     run Metropolis-Hastings algorithm on the data with the sampled age paths
6     estimate parameters given the sampled and refined age paths
7   end
8   for each sampled and refined age path do
9     estimate outliers
10  end
11 end
12 return sampled and refined age paths with their outliers and parameters

```

A natural question arises as follows: why do we not use (mixture) Kalman smoother against all the odds of the particle smoother described in section A.3.2? Ages are supposed not to be reversible and the accumulation rate model is a bit skewed, thus any Gaussian model cannot replace the gamma distribution in (2.2). The following doubt may stem from our choice of Metropolis-Hastings: section A.4.3 suggests the Hamiltonian Monte Carlo (HMC) as a better MCMC algorithm. Because the transition model is not differentiable, it is not possible to apply HMC here.

Dynamic time warping (DTW) [Munich & Perona (1999); Petitjean et al. (2011, 2014)] is frequently used for the alignment problems, but not adopted here because of the following reasons:

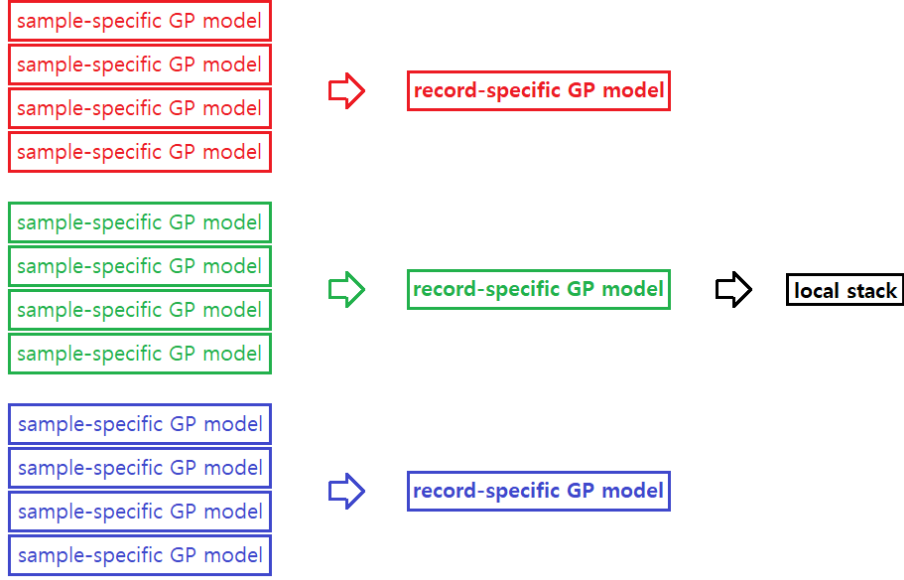


Figure 2.3: Three steps for the stack construction algorithm.

1) DTW has trouble in reflecting the dynamics of the latent series properly if the profile is not of high resolution, 2) results can be strongly dependent on the user-specific pairs and are often forced to match outputs that are apparently asynchronous, 3) the performance is sensitive to the penalties and loss functions in the transition and emission models, 4) the method is deterministic and does not reflect alignment uncertainty (which is beneficial if there are sub-optimal paths) and 5) DTW has a difficulty in addressing signals with heteroscedastic noises by the limitations on the loss function which is defined on the outputs directly.

We also sample outliers *after* drawing sampled and refined age paths that uses the marginalized observational model for benthic $\delta^{18}\text{O}$: this is called the Rao-Blackwellized sampling that efficiently reduces the sampling variance of the estimation.

2.4 STACK CONSTRUCTION ALGORITHM

In this section, we assume to have obtained a set of sampled age paths for each core (record) by running the hybrid algorithm in section 2.3. The stack construction algorithm consists of the following three steps: 1) constructing sample-specific models, 2) averaging sample-specific models to obtain record-specific models and 3) getting the local stack from the record-specific

models. Figure 2.3 visualizes these steps.

2.4.1 SAMPLE-SPECIFIC MODELS

This step focuses on each sampled and refined age path (the k th one) of the core m with the associated benthic $\delta^{18}\text{O}$ observations $(\tilde{\mathcal{A}}^{(m,k)}, \tilde{\mathcal{V}}^{(m,k)})$ after classifying and discarding outliers and synchronizing the benthic $\delta^{18}\text{O}$ observations vertically by the core-specific shift parameter $b^{(m)}$. Firstly, we applied homoscedastic Gaussian process regression (GPR) in section B.2.1 to the pair by regarding $\tilde{\mathcal{A}}^{(m,k)}$ and $\tilde{\mathcal{V}}^{(m,k)}$ as inputs and output, respectively, and derived the following posterior predictive model as the sample-specific one:

$$p(v|a, \tilde{\mathcal{A}}^{(m,k)}, \mathcal{V}^{(m)}) = p(v|a, \tilde{\mathcal{A}}^{(m,k)}, \tilde{\mathcal{V}}^{(m,k)}) \triangleq \mathcal{N}(v|\mu_{m,k}(a), \sigma_{m,k}^2(a)) \quad (2.8)$$

The Ornstein-Uhlenbeck (OU) kernel was selected, and kernel hyperparameters are tuned and shared by the sampled age paths of the same core. We adopted the exact GPR model without variational extension for dealing with large dataset because each core has less than 1000 benthic $\delta^{18}\text{O}$ observations thus such extension was not needed.

2.4.2 RECORD-SPECIFIC MODELS

A record-specific model is designed to represent the corresponding record. Because ages are latent thus random, we define the record-specific model as the marginalized posterior predictive model that does not depend on the choice of ages. Because each $\tilde{\mathcal{A}}^{(m,k)}$ is an independent and identically distributed samples drawn from $p(\mathcal{A}|\mathcal{V}^{(m)})$ by the core alignment algorithm in section 2.3, the following approximation is obtained:

$$\begin{aligned} p(v|a, \mathcal{V}^{(m)}) &= \int p(v|a, \mathcal{A}, \mathcal{V}^{(m)}) p(\mathcal{A}|\mathcal{V}^{(m)}) d\mathcal{A} \\ &\approx \frac{1}{K_m} \sum_{k=1}^{K_m} p(v|a, \tilde{\mathcal{A}}^{(m,k)}, \mathcal{V}^{(m)}) = \frac{1}{K_m} \sum_{k=1}^{K_m} \mathcal{N}(v|\mu_{m,k}(a), \sigma_{m,k}^2(a)) \end{aligned} \quad (2.9)$$

If there is only one record m in the homogeneous set, $p(v|a, \mathcal{V}^{(m)})$ is then supposed to be a local stack that represents the set. Because we want to define the local stack as a Gaussian model, by the model reduction method based on the moment-matching [Murphy (2012)], we have the following Gaussian approximation (2.10) for the Gaussian mixture (2.9) as the record-specific model:

$$\begin{aligned}
p(v|a, \mathcal{V}^{(m)}) &\approx \frac{1}{K_m} \sum_{k=1}^{K_m} \mathcal{N}(v|\mu_{m,k}(a), \sigma_{m,k}^2(a)) \approx \mathcal{N}(v|\mu_m(a), \sigma_m^2(a)) \\
\mu_m(a) &\triangleq \frac{1}{K_m} \sum_{k=1}^{K_m} \mu_{m,k}(a) \\
\sigma_m^2(a) &\triangleq \frac{1}{K_m} \sum_{k=1}^{K_m} \left(\sigma_{m,k}^2(a) + (\mu_{m,k}(a) - \mu_m(a))^2 \right)
\end{aligned} \tag{2.10}$$

2.4.3 LOCAL STACK

Now we have a set of record-specific models obtained in the previous section. Let us recall the assumption that the input cores are homogeneous, i.e., they share the same local variability of the global parameter benthic $\delta^{18}\text{O}$ thus are supposed to follow the same Gaussian model, a local stack:

$$p(v|a, \mathcal{V}) = \mathcal{N}(v|\mu_S(a), \sigma_S^2(a)) \tag{2.11}$$

Let us assume that we have drawn M independent and identically distributed observations $\{v_m\}_{m=1}^M$ from $p(v|a, \mathcal{V})$. Here, our variational method assumes that the joint distribution of $(v_1, v_2, \dots, v_M)$ is approximated by the product of the record-specific models $p(v|a, \mathcal{V}^{(m)})$, i.e., $\{v_m\}_{m=1}^M$ is also considered as a set of samples where each v_m was drawn from the corresponding

$p(v|a, \mathcal{V}^{(m)})$, respectively. The objective function to minimize is given as follows:

$$\begin{aligned}
\mathbb{D}_{\text{KL}} \left(p(\cdot|a, \mathcal{V})^M \middle| \middle| \prod_{m=1}^M p(\cdot|a, \mathcal{V}^{(m)}) \right) &= \sum_{m=1}^M \mathbb{D}_{\text{KL}} (p(\cdot|a, \mathcal{V}) || p(\cdot|a, \mathcal{V}^{(m)})) \\
&= \sum_{m=1}^M \mathbb{D}_{\text{KL}} (\mathcal{N}(\cdot|\mu_S(a), \sigma_S^2(a)) || \mathcal{N}(\cdot|\mu_m(a), \sigma_m^2(a))) \\
&= \frac{1}{2} \sum_{m=1}^M \left(\frac{\sigma_S^2(a)}{\sigma_m^2(a)} + \frac{(\mu_S(a) - \mu_m(a))^2}{\sigma_m^2(a)} - 1 + \log \frac{\sigma_m^2(a)}{\sigma_S^2(a)} \right)
\end{aligned} \tag{2.12}$$

The minimizer $(\mu_S(a), \sigma_S^2(a))$ of the KL divergence (2.12) are given by the following closed forms:

$$\begin{aligned}
\mu_S(a) &= \sum_{m=1}^M \frac{\mu_m(a)}{\sigma_m^2(a)} / \sum_{m=1}^M \frac{1}{\sigma_m^2(a)} \\
\sigma_S^2(a) &= M / \sum_{m=1}^M \frac{1}{\sigma_m^2(a)}
\end{aligned} \tag{2.13}$$

Theoretically, one can analytically compute $\mu_S(a)$ and $\sigma_S^2(a)$ for an arbitrary query age a by computing the sample-specific models in (2.8), the record-specific models in (2.10) and finally the local stack in (2.13), thus the stack is continuously defined. In practice, however, computing the mean and variance in (2.13) for a fine discrete set of ages and interpolating them for arbitrary query inputs is more efficient than the theoretic method: note that (2.13) implies that the constructed stack has continuous mean and variance functions.

2.5 RESULTS AND DISCUSSION

The steps described in section 2.3 and 2.4 present the core idea of SA-GPR [Lee et al. (2019)]: note that the same framework is working for any transition and emission models. Figure 2.4 shows a diagram of SA-GPR: here we use the words ‘motif’ instead of ‘profile’ or ‘stack’ and ‘signal’ instead of ‘sequence’ or ‘core’, as the more general terms, i.e., both core alignments and stacks are obtained by iterating the core alignment algorithm in section 2.3 and stack construction algorithm in section 2.4 until convergence.

We constructed a dual proxy (that means both direct and indirect proxies were utilized) local

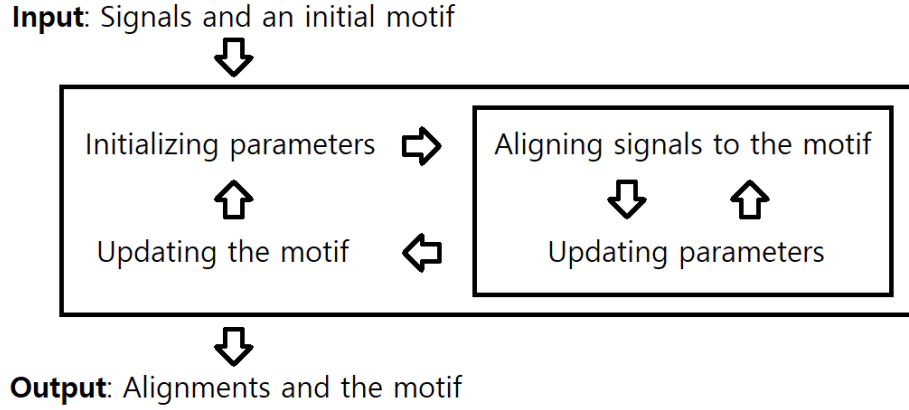


Figure 2.4: A simple diagram of the SA-GPR. Each box is iterated until convergence.

DNEA stack for the deep northeast Atlantic by assuming that two cores from the Iberian margin (MD95-2042 and MD99-2334) and four cores from the northwest African continental slope (GeoB7920-2, GeoB9508-5, GeoB9526-5 and ODP658C) are synchronous enough to be considered homogeneous. The availability of ^{14}C allows direct access to the calendar ages to enhance the age inferences: this is especially useful in the Holocene (after 11.7 ka) where the $\delta^{18}\text{O}$ signal-to-noise is low. We used the regional Deep North Atlantic (DNA) stack from [Lisiecki & Stern (2016)] to initialize the iterative algorithm. The open-source software for the stack construction and core alignment algorithms is available at <http://github.com/eilion/DPGP-Stack>, runs on MATLAB.

The upper and middle panels of figure 2.5 shows the alignments of the six cores to DNA stack and to DNEA stack. Variability in benthic $\delta^{18}\text{O}$ in the DNA stack is considerably larger than that in the DNEA stack: note that most of the medians are inside the 1-sigma of the DNA stack, which is only supposed to contain 68% of them. Higher variances in the DNA stack compared to the DNEA stack may stem from benthic $\delta^{18}\text{O}$ differences within the broader north Atlantic region, record-specific mean shifts applied to the DNEA stack but not the DNA stack, and/or the discrete nature of the algorithm used to construct the DNA stack. Another contributing factor to the tighter DNEA stack variance is the automatic detection and removal of outliers. The tighter variance will contribute to less uncertainty in ages inferred from this stack.

The smoothness of the local stack stems from the fact that the GP model captures correlations

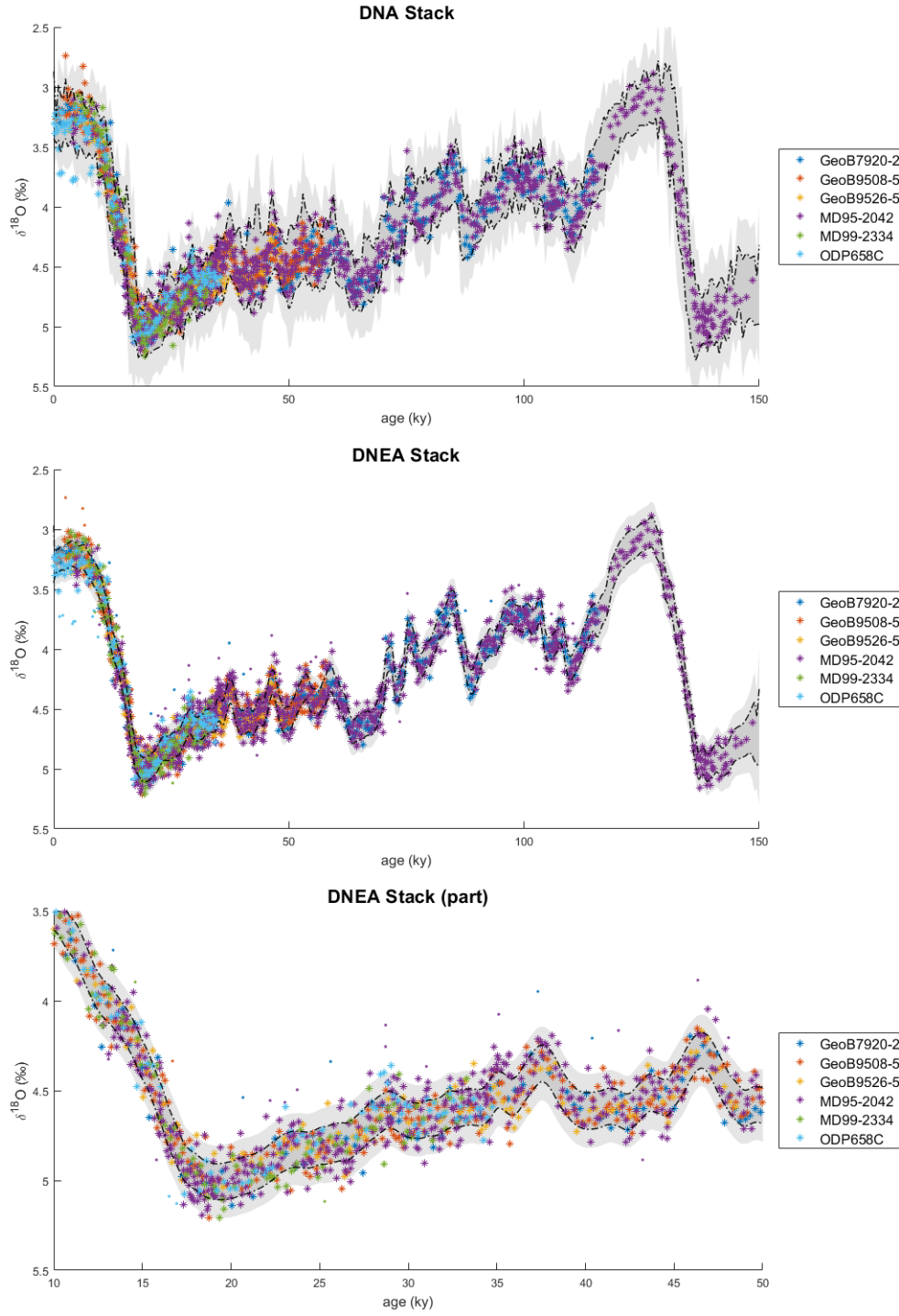


Figure 2.5: Core alignments to stacks. The upper panel (a) shows those to the DNA stack, while the middle panel (b) is our dual proxy local DNEA stack. Stars indicate medians of $\delta^{18}\text{O}$ data classified as inliers and dots represent outliers, after shifting them vertically by the core-specific shift parameters. The darker and brighter gray regions are the 1-sigma and 2-sigma of the stacks, respectively. The dot-dash line indicates their boundary. The lower panel (c) is a portion of the DNEA stack at 10-50 kiloyears.

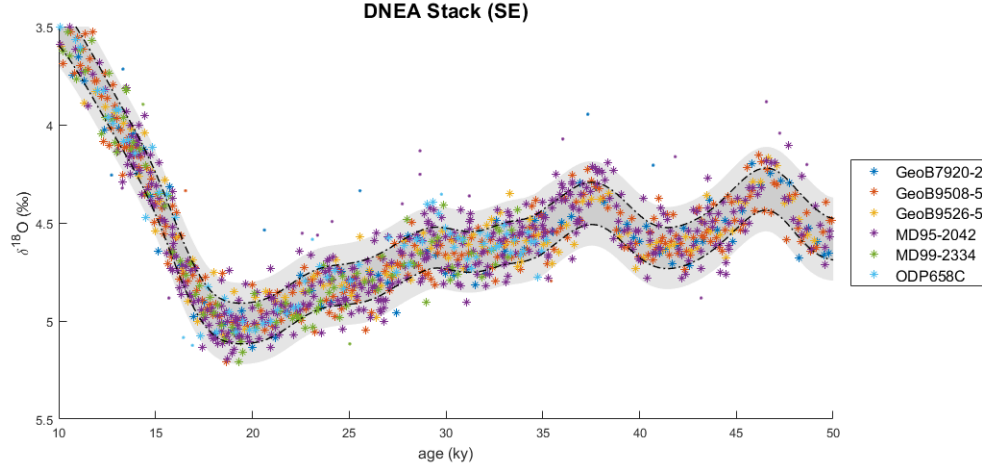


Figure 2.6: A part of the DNEA stack constructed with the SE kernel.

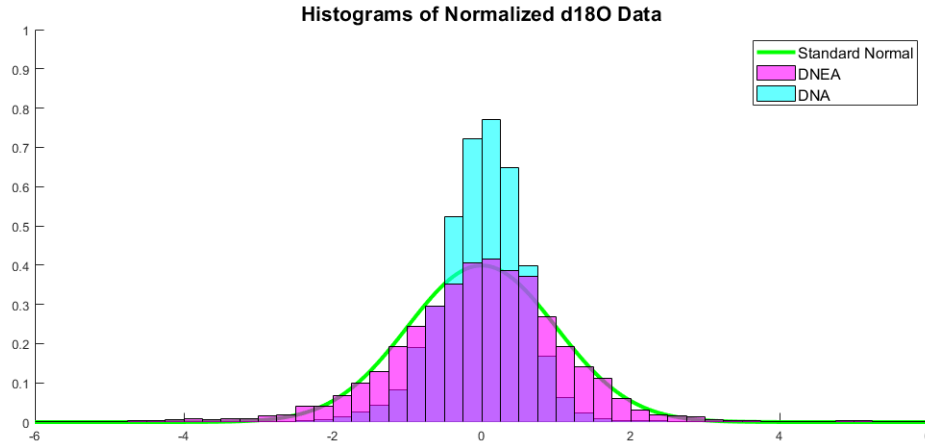


Figure 2.7: Histograms of $\delta^{18}\text{O}$ to our dual proxy DNEA stack. Each $\delta^{18}\text{O}$ is standardized by the mean and standard deviation of the stack at the median of its sampled ages.

between all the data points with a heavier weight placed on near neighbors, thus limiting sudden large changes. Although our dual proxy stack is smoother than the DNA stack, it still captures well-known millennial-scale climate events. For example, the bottom panel of figure 2.5 shows four peaks at 24, 29, 38 and 46 kiloyears, which correspond to the Heinrich events H2 to H5 [Skinner et al. (2013); Lisiecki & Stern (2016)]. The ability to resolve such short-lived features will improve the accuracy of age estimates for cores with high-resolution $\delta^{18}\text{O}$ records. We also tried the squared-exponential (SE) kernel but the results failed to reflect H2 and H3 at 24 and 29 kiloyears respectively, as figure 2.6.

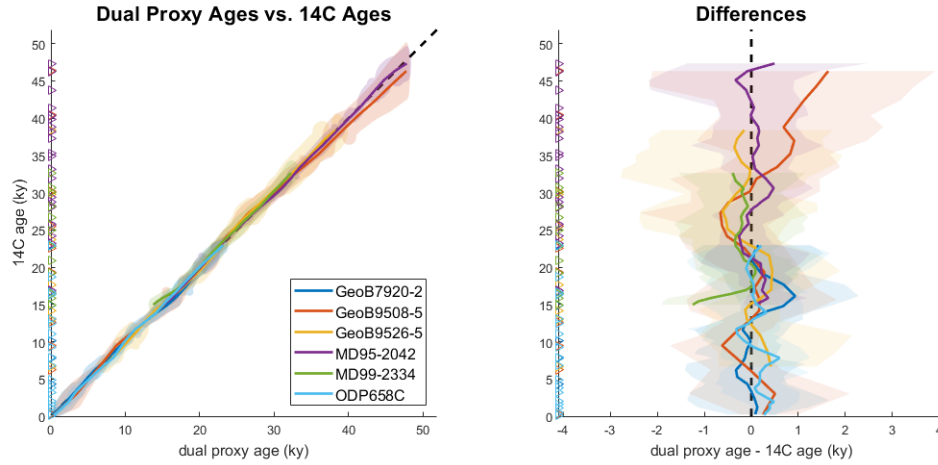


Figure 2.8: The comparison between inferred dual proxy ages and ^{14}C ages of records. Shaded regions indicate 95% confidence regions and the black dashed line is just a diagonal.

Figure 2.7 shows the histograms of standardized $\delta^{18}\text{O}$ from the six cores to the DNA and DNEA stacks, respectively. The fact that there is only a small departure from the standard normal distribution to the DNEA stack supports the validity of the Gaussian assumption. Figure 2.8 compares inferred ages by using both benthic $\delta^{18}\text{O}$ and ^{14}C (dual proxy) to those using ^{14}C only. Overall agreement between cores (to within uncertainty) supports our assumption that benthic $\delta^{18}\text{O}$ is synchronous and homogeneous among sites included in constructing DNEA stack. Some departures from the diagonal are expected, considering influences from their $\delta^{18}\text{O}$ data.

We also present the alignments of two ocean sediment cores from the same region, GIK13289-2 and SU90-08. We infer ages for these cores from their proxies and alignment to the DNEA stack based on the assumption that the aligned cores share the same local $\delta^{18}\text{O}$ effect as the stack. Figure 2.9 and 2.10 show the alignment results. The upper panels show that the shifted and aligned $\delta^{18}\text{O}$ data mainly fall within the stack's confidence intervals. The lower left panels show that the inferred ages mainly pass through the confidence intervals from individual ^{14}C proxies (SU90-08 has ^{14}C only between 10-40 kiloyears and beyond that ages are inferred only based on $\delta^{18}\text{O}$ alignment to the stack). Because stacks with dual proxy ages have narrower confidence intervals than single proxy stacks, they are more informative for stack alignment and produce smaller age uncertainties.

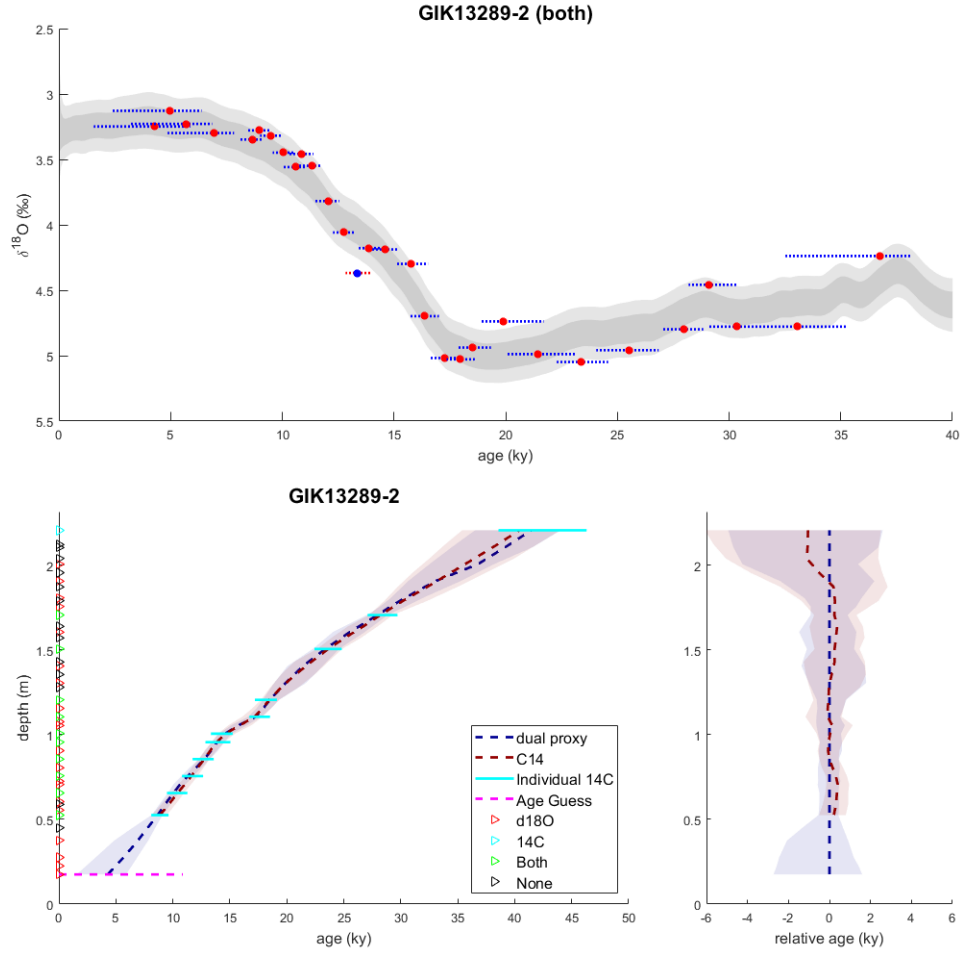


Figure 2.9: Dual proxy age inferences and alignments of GIK13289-2 to DNEA stack. In the upper figure, the darker and brighter areas show the 1-sigma and 2-sigma of the stacks, red points and blue dotted bars indicate medians and 95% confidence intervals of inlier age samples for $\delta^{18}\text{O}$ data, blue points and red dotted bars are the outliers. In the left below figures, cyan bars indicate 95% confidence intervals obtained independently from ^{14}C data, and blue and red areas show the 95% confidence bands of age inferences from ^{14}C only and both proxies, respectively. In the right below figures, medians (dashed lines) and 95% confidence intervals of relative ages to the medians of dual proxy ages are shown.

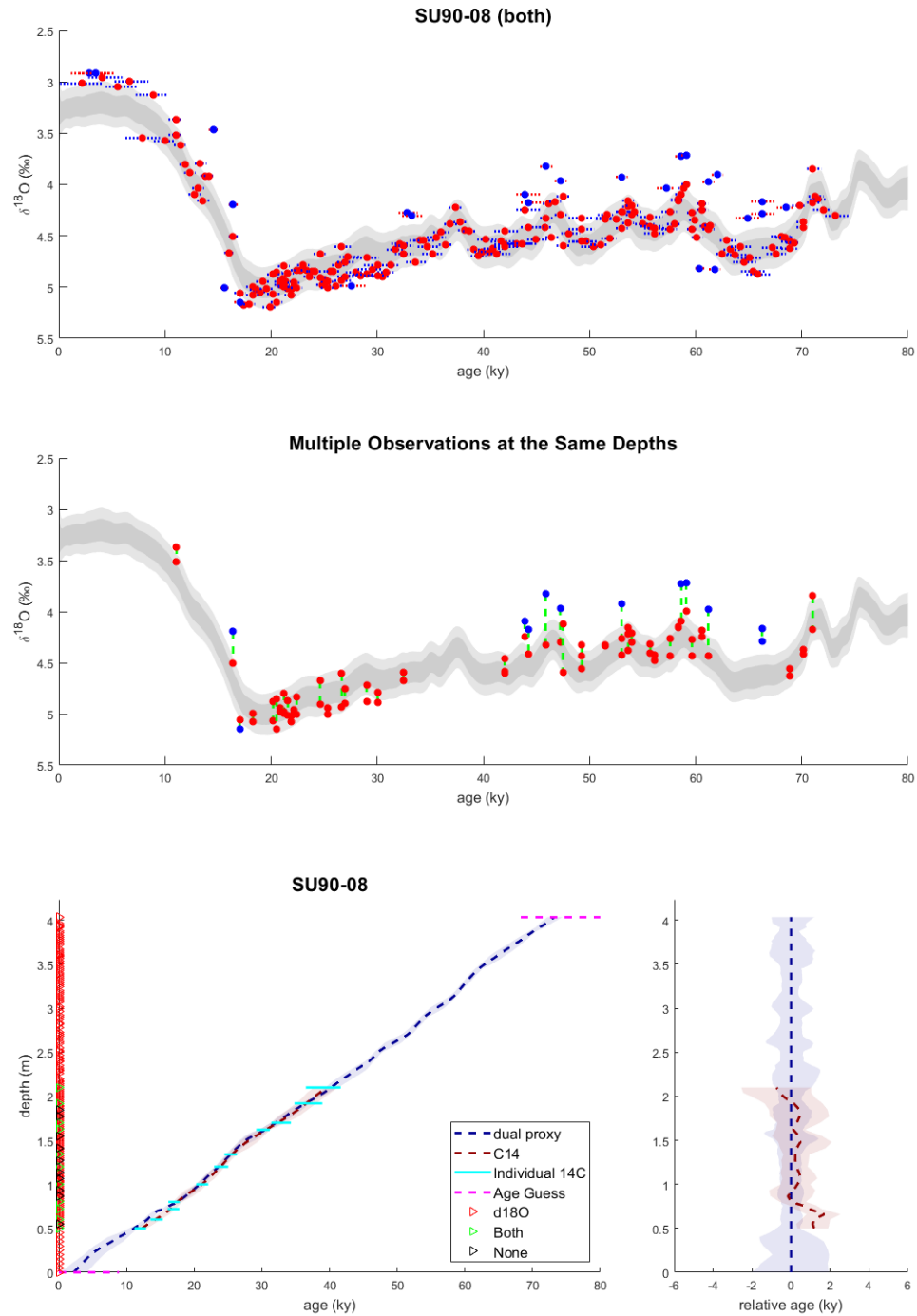


Figure 2.10: Dual proxy age inferences and alignments of SU90-08 to DNEA stack, similar to figure 2.9. This figure also shows how multiple observations at the same depths are dealt with in the alignment algorithm.

Our core alignment and stack construction algorithms share the advantages of the other stack-dependent alignment algorithms. Aligning cores does not only focus on finding the best matches between cores but also integrates the fragmented information stored in each core into the stack. Once the stack is constructed from a set of cores, it can be used as an alignment target for new cores.

Even if continuous sampling is available, it soon becomes insignificant without a continuously defined stack. Our stack construction algorithm depending on GPR is superior to the interpolation algorithms because it is theoretically clearer to exploit data to fill in unobserved outputs: the closer the query age is to the data, the higher the associated benthic $\delta^{18}\text{O}$ is correlated with them, based on the kernel covariance functions. Moreover, our method is nonparametric thus more robust on the general data.

There are several advantages of the new dual proxy method compared to using $\delta^{18}\text{O}$ or ^{14}C only. First, it gives a direct method for age inferences in a $\delta^{18}\text{O}$ stack. The dual proxy method can also be applied for other types of direct age constraints that can be expressed as distributions given their ages, such as dated magnetic reversals or volcanic ash layers. Second, ^{14}C determinations act as “anchor” which stabilizes the learning of transition parameters. Radiocarbon ages effectively protect the parameter learning procedure during the first EM iteration from ill-posed initial values. Third, the integration of both proxies should give more accurate age estimates by using more data. High resolution $\delta^{18}\text{O}$ measurements complement low-resolution ^{14}C data and substantial uncertainty bounds on ^{14}C -only inferences where ^{14}C data are sparse.

In addition, our stack construction algorithm can construct a stack with a limited number of cores because it uses all of the data from every core for inference of every age in the stack: previous methods such as HMM-Stack [Ahn (2016)] require a sufficient number of high-resolution records to estimate the mean and variance separately for each point in the stack. In addition, this dual proxy method returns means and variances in continuous time at arbitrary ages. Previous methods rely on less reliable interpolations between discrete ages.

While the state-transition matrix in section 2.3.1 can be learned either core-specifically or set-

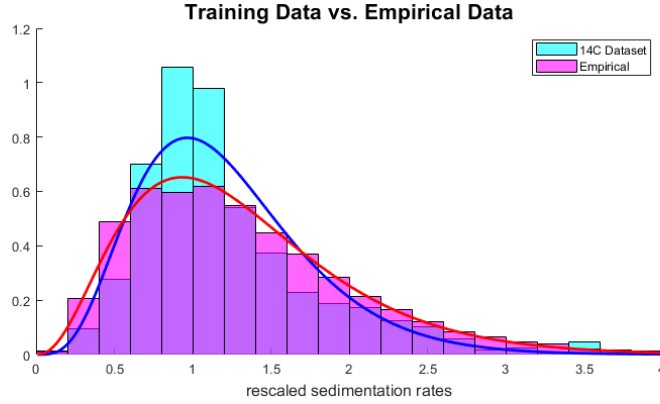


Figure 2.11: Histograms of the standardized accumulation rates. The blue histogram comes from the training ^{14}C dataset and the red one from the sampled age paths in constructing the DNEA stack. Thick lines indicate the inferred gamma distributions from the histograms, respectively.

specifically, the underlying accumulation model based on the gamma distribution is pre-trained and fixed. Gamma distributions tend to have thicker tails than log-normal mixtures learned from the same data, so alignments are more dependent on the emission model than the transition model, which makes the model more faithful to the observations. The comparison of our gamma distribution inferred from the training data and that from the empirical accumulation rates of cores in constructing the DNEA stack is in figure 2.11.

Here we avoid using the heteroscedastic GPR in section B.2.2 (instead of the homoscedastic GPR that we used) for regularization. However, if $\delta^{18}\text{O}$ observations of the cores have higher resolution, this method will model more properly.

2.6 FUTURE WORK

As mentioned at the end of section 2.3.3, the reason why we did not consider HMC instead of Metropolis-Hastings is because the transition model is not differentiable. By giving two-layer hidden states, we could assume the second-order dependency of the ages given depth. If it is possible to alter this transition model to be differentiable, then the full model becomes differentiable with respect to the hidden states thus the core alignment algorithm will run faster by replacing the Metropolis-Hastings with HMC: note that the main bottleneck is the Metropolis-Hastings step. However, this cannot be done without the experts because sedimentation rate model itself

is geological thus must be consistent with the real data and theories certified by the community of paleoceanography.

The next step of local stack construction for each of the homogeneous sets one-by-one is, clearly, designing a more comprehensive construction algorithm that deal with multiple homogeneous sets simultaneously, i.e., considering how to allow homogeneous sets to communicate one another. With interdisciplinary collaborators, I am currently working on the hierarchical stack construction algorithm that automatically differentiate the dependency of $\delta^{18}\text{O}$ observations over cores with respect to their similarity. Here we cannot avoid considering variational extension because all data must be processed together while our stack construction algorithm in section 2.4 deviates the drawback of the GP models relevant to the size of data by constructing record-specific models and then merging them into a local stack, which depends on the assumption that the input cores are homogeneous.

Utilizing more age-relevant proxies, such as magnetic reversals or volcanic ash layers, is also an interesting topic. As discussed in section 2.5, the comprehensive framework of our method is flexible in processing age proxies, either direct or indirect. It could be utilized for constructing multiple stacks for indirect proxies other than $\delta^{18}\text{O}$ and utilizing more proxies at once gives an opportunity to exploit the more complete age information.

3

Alkenone/TEX₈₆ Calibration Models

This section shows some extra works from the conference paper [[Lee & Lawrence \(2019\)](#)] that constructs a nonparametric alkenone calibration model based on the heteroscedastic Gaussian process regression.

The alkenone is a widely used proxy for inferring sea surface temperatures (SSTs) in paleoceanography and paleoclimatology. It is a long-chain unsaturated methyl and ethyl n-ketones produced by some phytoplankton species. Species that produce alkenone respond to environmental changes such as water temperatures by modifying the relative proportions of the saturated and unsaturated alkenones: more saturated alkenones are produced at higher temperatures. The

realization depends on the following indices U_{37}^K and $U_{37}^{K'}$ [Tierney & Tingley (2018)]:

$$\begin{aligned} U_{37}^K &= \frac{C_{37:2} - C_{37:4}}{C_{37:2} + C_{37:3} + C_{37:4}} \\ U_{37}^{K'} &= \frac{C_{37:2}}{C_{37:2} + C_{37:3}} \end{aligned} \quad (3.1)$$

, where $C_{37:2}$, $C_{37:3}$ and $C_{37:4}$ are the di-, tri- and tetra-unsaturated C_{37} alkenones. Though $C_{37:4}$ might be beneficial in prediction [Bard (2001); Bendle & Rosell-Mele (2004)], it does not exist in many marine sediment profiles, thus $U_{37}^{K'}$ is preferred.

The tetraether index of 86 carbons, TEX_{86} shortly, is another proxy that are utilized for reconstructing SSTs in paleoceanography and paleoclimatology. It is an index based on the relative cyclization of glycerol dialkyl glycerol tetraethers (GDGTs) that contain 0 to 3 cyclopentane moieties. The degree of cyclization is affected by the growth temperature [Wuchter et al. (2004)]. There are two ways of defining this index [Schouten et al. (2002); Kim et al. (2010)]:

$$\begin{aligned} TEX_{86} &= \frac{[GDGT - 2] + [GDGT - 3] + [Cren']}{[GDGT - 1] + [GDGT - 2] + [GDGT - 3] + [Cren']} \\ TEX_{86}^L &= \frac{[GDGT - 2]}{[GDGT - 1] + [GDGT - 2] + [GDGT - 3]} \end{aligned} \quad (3.2)$$

, where $Cren'$ denotes the regioisomer of crenarchaeol that is a diagnostic biomarker for the Thaumarchaeota, the main producers of GDGTs [Sinninghe Damsté et al. (2002); Tierney & Tingley (2014)].

By definition, both $U_{37}^{K'}$ and TEX_{86} lie in the unit interval $[0, 1]$. Their underlying characteristics bring opportunities to utilize them for reconstructing past SSTs. A calibration model gives a relationship between the observations (known) and what we want to infer (unknown), either deterministic or probabilistic. Probabilistic calibration models aim at constructing distributions of the known given unknown, thus regarded as generative models (section 1.2.3): if additional information of the unknown independent of the known is available, then more reasonable inference on the unknown can be done in the framework of Bayesian statistics.

This chapter presents our ongoing work on the construction of probabilistic calibration mod-

els in various ways based on the Gaussian process regression (GPR) and their comparison. A novel multimodal GPR based on the Gibbs sampling (section A.4.1) and expectation-maximization (EM) algorithm [Dempster et al. (1977)] is also formulated in section 3.3.2.

3.1 RELATED WORK

For the past years, both alkenone and TEX_{86} proxies have been gathered and calibrated, and thus consequently validated as effective proxies for SSTs. This subsection briefly introduces the previous works in modelling.

3.1.1 ALKENONE CALIBRATION MODELS

In response to the apparent linearity of the $U_{37}^{K'}$ observations over SSTs, several works have been focusing on constructing simple or region-specific linear regression models from their available data. For examples, [Prahl et al. (1988)] used 5 observations in range of 8 to 25 °C and resulted in $U_{37}^{K'} = 0.034 \cdot \text{SST} + 0.039$; [Müller et al. (1998)] obtained $U_{37}^{K'} = 0.033 \cdot \text{SST} + 0.044$ from 370 observations from 0 to 29 °C; [Sonzogni et al. (1997)] suggested $U_{37}^{K'} = 0.023 \cdot \text{SST} + 0.316$ for 40 observations from 24 to 29 °C in Indian ocean.

More recently, [Conte et al. (2006)] suspected the nonlinearity relationship above 26 °C and constructed a set of region-specific nonlinear models: $U_{37}^{K'} = -1.004 \cdot 10^{-4} \cdot \text{SST}^3 + 5.744 \cdot 10^{-3} \cdot \text{SST}^2 - 6.207 \cdot 10^{-2} \cdot \text{SST} + 0.407$ for 451 observations from Atlantic, $U_{37}^{K'} = 0.0391 \cdot \text{SST} - 0.1364$ for 131 observations from Pacific. Among all data sources including Indian and Southern Ocean, the global model for 629 observations is given as $U_{37}^{K'} = -5.256 \cdot 10^{-5} \cdot \text{SST}^3 + 2.884 \cdot 10^{-3} \cdot \text{SST}^2 - 8.4933 \cdot 10^{-3} \cdot \text{SST} + 9.898$.

[Tierney & Tingley (2018)] developed a Bayesian B-spline regression model, BAYSPLINE, with the treatment of data seasonality, from 1344 globally-gathered observations. The spline transformation is simply a linear combination of B-splines, so the regression model becomes a multiple linear regression. B-spline coefficients and model variance have priors and inferred in the Bayesian framework. The resulted model is linear below 23.4 °C with slope 0.034 that is

indistinguishable from [Prahlet et al. (1988)], but the slope declines steadily to a minimum of 0.011 at the highest temperature 29.4 °C.

[Lee & Lawrence (2019)] constructed a GPR model as the calibration model from the same dataset of [Tierney & Tingley (2018)], after discarding those in the locations excluded from their analysis. The result of data with 1274 observations were used. After preprocessing on the remaining SST and $U_{37}^{K'}$ to be standardized, they applied the heteroscedastic GPR to them with a neural network kernel (section B.1.2) and constant mean prior function. Confirming the heteroscedasticity of data is one of their achievements.

3.1.2 TEX₈₆ CALIBRATION MODELS

Various simple models under certain conditions have been made so far. [Kim et al. (2008)] constructed a linear model $SST = 56.2 \cdot TEX_{86} - 10.78$ from 223 observations without data from Red Sea, those with $SST < 5$ °C and those with residuals bigger than 1 sigma; [Liu et al. (2009)] resulted in $SST = -16.332/TEX_{86} + 50.475$ from 287 observations that include all their available data; [Kim et al. (2010)] gives two models: one is $SST = 67.5 \cdot \log_{10} TEX_{86}^L + 49.9$ without Red Sea data and the other is $SST = 68.4 \cdot \log_{10} TEX_{86} + 38.6$ without data from Red Sea and those with $SST < 5$ °C.

As a more complicated model, [Tierney & Tingley (2014)] developed a Bayesian regression model, BAYSPAR, with the Gaussian process (GP) priors. Their region-specific linear models assume GP priors on the slope and intercept parameters over locations. They tested BAYSPAR under different conditions defined by either or some of Arctic data with $SST < 3$ °C, Red Sea data, or Antarctic data with $SST < 3$ °C were excluded. In either case, regional differences of slopes and intercepts were identified.

3.2 RESEARCH MOTIVE AND DATA

The goal is to construct robust calibration models from the training data that consist of pairs of present (known) SSTs and either $U_{37}^{K'}$ or TEX_{86} , and then apply the models for inferring past

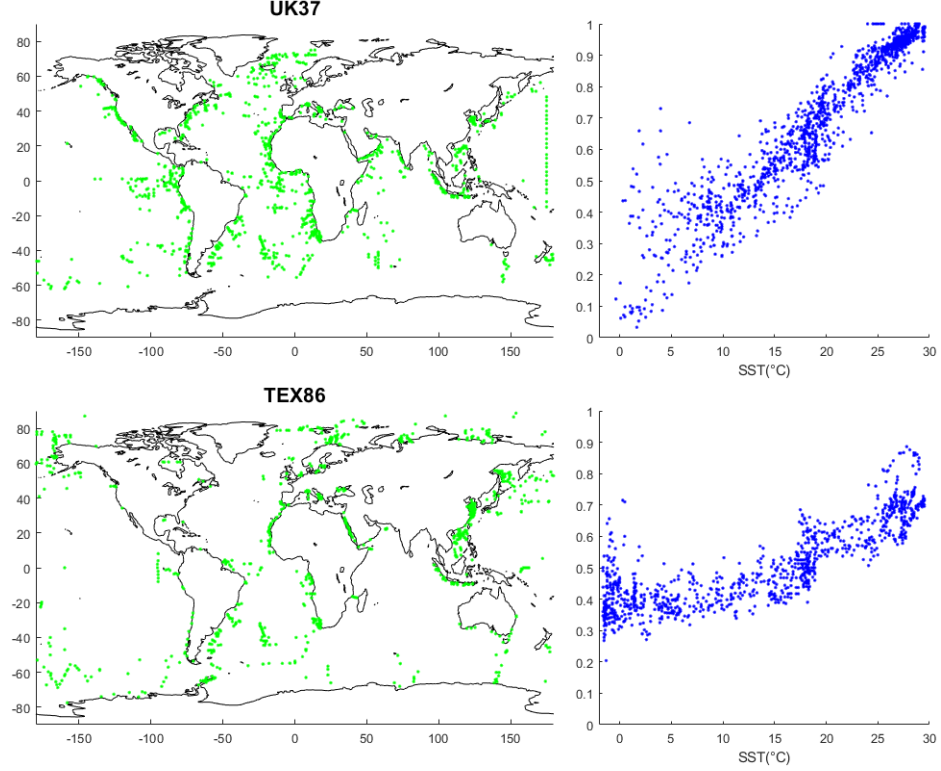


Figure 3.1: $U_{37}^{K'}$ and TEX_{86} data. Each left panel shows the data locations while right plots data values over SSTs.

SSTs given the ancient proxy observations. In more mathematical terminology, what we want to do is:

1. Constructing a calibration model $p(\text{Proxy}|\text{SST})$.
2. Applying it for $p(\text{SSTs}|\text{Proxies}) \propto p(\text{Proxies}|\text{SSTs})p(\text{SSTs}|\text{IK})$.

Here, IK stands for the independent knowledge of the SSTs from their locations or ages. Note that this simple framework could turn out to be a complicated probabilistic graphical model if $p(\text{SSTs}|\text{IK})$ forms a Markov chain (over ages), a spherical grid (over locations), or a more high-dimensional structure.

Figure 3.1 visualizes the dataset of $U_{37}^{K'}$ and TEX_{86} , respectively. We use the same dataset of [Tierney & Tingley (2018)] and [Tierney & Tingley (2014)]. The $U_{37}^{K'}$ dataset is comprised of [Arz et al. (2003); Bard et al. (2000); Barron et al. (2003); Barrows et al. (2007); Becker et al. (2015); Bendle & Rosell-Melé (2007); Cacho et al. (1999, 2001); Calvo et al. (2001, 2007); Caniupán et al. (2011); Chen et al. (2014); Conte et al. (2006); Benthien & Müller (2000); Chapman

et al. (1996); Doose et al. (1997); Dubois et al. (2009); Emeis et al. (2000); Fallet et al. (2012); Filippova et al. (2016); Fraser et al. (2014); Gutiérrez et al. (2011); Harada et al. (2006); Herbert et al. (1998, 2001); Ho et al. (2012); Horikawa et al. (2006); Huang et al. (1997); Ikehara et al. (1997); Jaeschke et al. (2007, 2017); Kaiser et al. (2005, 2008, 2014); Kennedy & Brassell (1992); Kienast & McKay (2001); Kienast et al. (2006, 2012); Kim et al. (2002a,b, 2003, 2004, 2007, 2015); Koutavas & Sachs (2008); Kudrass et al. (2001); Lamy et al. (2002); Leduc et al. (2007, 2010); Lee et al. (2001, 2008b,a); Leider et al. (2010); Liu & Herbert (2004); Lopes dos Santos et al. (2010); Lückge et al. (2009); Madureira et al. (1995); Marchal et al. (2002); Martrat et al. (2007); Max et al. (2012); McCaffrey et al. (1990); Mohtadi et al. (2011); Müller et al. (1998); Ohkouchi et al. (1999); Ostertag-Henning & Stax (2000); Pahnke & Sachs (2006); Pahnke et al. (2007); Pailler & Bard (2002); Pelejero & Grimalt (1997); Pelejero et al. (1999); Prah et al. (1989, 2006, 2010); Rein et al. (2005); Rodrigo-Gámiz et al. (2015); Rodrigues et al. (2009, 2010); Romero et al. (2008); Rosell-Melé et al. (1995); Rosell-Melé (1998); Rosell-Melé et al. (2000); Rühlemann et al. (1999); Sachs (2007); Sawada et al. (1996); Sawada & Handa (1998); Schefuss et al. (2005); Schulz et al. (2002); Seki et al. (2004); Sepúlveda et al. (2009); Sikes et al. (1991); Sikes & Keigwin (1994, 1996); Sikes et al. (1997, 2009); Sonzogni et al. (1997, 1998); Sperling et al. (2003); Tao et al. (2012); Ternois et al. (2000); Tierney & Tingley (2018); Weldeab et al. (2007); Zhao et al. (1995); Zhao et al. (2006)], whereas the TEX₈₆ dataset is consisting of [Castañeda et al. (2010); Chazen (2011); Chen et al. (2014); Fallet et al. (2012); Hernández-Sánchez et al. (2014); Ho et al. (2011, 2014); Hu et al. (2012); Jia et al. (2012); Kaiser et al. (2014); Kim et al. (2010); Leider et al. (2010); Lengger et al. (2014); Lü et al. (2014); Nieto-Moreno et al. (2013); Park et al. (2014); Richey et al. (2011); Seki et al. (2009, 2014); Shevenell et al. (2011); Shintani et al. (2011); Smith et al. (2013); Tierney & Tingley (2014); Trommer et al. (2009); Verleye (2011); Wei et al. (2011); Wu et al. (2012); Zell et al. (2014); Zhou et al. (2014)].

In either case, there are more than 1000 observations but less than 1500, thus the amounts of data are moderate. Though somewhat arbitrary linear regressions have been mostly applied so far, such relationships are not justified by theory, thus in fact there is no parametric model that not

only describes but also explains the relationship of SSTs and either $U_{37}^{K'}$ or TEX_{86} . As discussed in section 1.2.2, these characteristics fit well for the nonparametric approaches, such as GPR.

As shown in figure 3.1, both proxies show apparent heteroscedasticity over SSTs, thus it is reasonable that heteroscedastic GPR models tackle the problem. Though some of the previous research tend to either discard or adjust some observations based on their data-specific characteristics, here the full and raw dataset are used. To avoid the misspecification of GPR that assumes the monomodality of the data, we also apply a multimodal GPR model to the proxy datasets.

3.3 MODELS

As mentioned earlier, the GPR is a nonparametric regression model. However, it has a margin to reflect the parametric aspect in its structure: the mean prior function μ as discussed in Appendix B. For extending the previous works that are based on simple polynomial regressions (mostly linear), here we assume that μ is a cubic polynomial with unknown coefficient parameters to infer. Though the sizes of two datasets are between 1000 and 1500, a variational free energy (VFE) extension in section B.7 is applied in modelling. In this section, $\mathcal{X} = \{X_n\}_{n=1}^N$ and $\mathcal{Y} = \{Y_n\}_{n=1}^N$ are the sets of SSTs and proxy (either $U_{37}^{K'}$ or TEX_{86}) observations, respectively.

3.3.1 HOMOSCEDASTIC AND HETEROSCEDASTIC GAUSSIAN PROCESS REGRESSION

Here we extend the heteroscedastic GPR model in [Lee & Lawrence (2019)]. Detailed derivations can be found in section B.2. In either proxy, the observational model is given as follows:

$$p(\mathcal{Y}|\beta, \mathcal{X}) = \mathcal{N}(\mathcal{Y}|\beta, \Lambda_{\mathcal{X}}) = \prod_{n=1}^N \mathcal{N}(Y_n|\beta_n, \Lambda_{X_n}) \quad (3.3)$$

, where β is a latent state that represents the regression of $\mathcal{Y}$ on $\mathcal{X}$. Recall that the GPR is homoscedastic if the variance function Λ is given as a constant function and heteroscedastic if not (section B.2). The prior on β is given as follows:

$$p(\beta|\mathcal{X}) = \mathcal{GP}(\beta|\varphi(\mathcal{X})\alpha, \mathbb{K}_{\mathcal{X}\mathcal{X}}) \quad (3.4)$$

, where $\mathbb{K}$ is a neural network kernel (section B.1.2) with erf as the activation function, $\alpha = (\alpha_0, \alpha_1, \alpha_2, \alpha_3)^\top \in \mathbb{R}^4$ is a coefficient parameter and φ is a feature function defined as follows:

$$\varphi(\mathcal{X}) = \begin{pmatrix} 1 & X_1 & X_1^2 & X_1^3 \\ 1 & X_2 & X_2^2 & X_2^3 \\ \vdots & \vdots & \vdots & \vdots \\ 1 & X_N & X_N^2 & X_N^3 \end{pmatrix} \quad (3.5)$$

An exact GPR is specified by (3.3) and (3.4). However, we also consider the VFE extension (section B.7) here: recall that it returns a closed-form posterior predictive model in regression. Let $\mathcal{X}_0$ be the set of pseudo-inputs. Then, the regression model is given as follows:

$$\begin{aligned} p(y|x, \mathcal{X}, \mathcal{Y}, \mathcal{X}_0) &= \mathcal{N}(y|\varphi(x)\alpha + m(x), \Lambda_x + s(x)) \\ m(x) &= \mathbb{K}_{x\mathcal{X}_0} (\mathbb{K}_{\mathcal{X}_0\mathcal{X}_0} + \mathbb{K}_{\mathcal{X}_0\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}_0})^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} (\mathcal{Y} - \varphi(\mathcal{X})\alpha) \\ s(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}_0} \mathbb{K}_{\mathcal{X}_0\mathcal{X}_0}^{-1} \mathbb{K}_{\mathcal{X}_0x} + \mathbb{K}_{x\mathcal{X}_0} (\mathbb{K}_{\mathcal{X}_0\mathcal{X}_0} + \mathbb{K}_{\mathcal{X}_0\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}_0})^{-1} \mathbb{K}_{\mathcal{X}_0x} \end{aligned} \quad (3.6)$$

Thus, the resulting model $p(y|x, \mathcal{X}, \mathcal{Y}, \mathcal{X}_0)$ in (3.6) can be understood as an extended polynomial regression: $\varphi(x)\alpha$ is the parametric regression of y and $m(x)$ is a nonparametric bias function, and the regression function is given by their sum $\varphi(x)\alpha + m(x)$.

Here, not only the kernel hyperparameters but also the coefficient α are unknown, so to be tuned and learned. The objective function is the evidence lower bound (ELBO) $\mathcal{L}$ given by the following closed form:

$$\begin{aligned} \log p(\mathcal{Y}|\mathcal{X}) &\geq \mathcal{L} \triangleq \log \mathcal{N}(\mathcal{Y}|\varphi(\mathcal{X})\alpha, \Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}_0} \mathbb{K}_{\mathcal{X}_0\mathcal{X}_0}^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}}) \\ &\quad - \frac{1}{2} \text{trace}(\Lambda_{\mathcal{X}}^{-1} (\mathbb{K}_{\mathcal{X}\mathcal{X}} - \mathbb{K}_{\mathcal{X}\mathcal{X}_0} \mathbb{K}_{\mathcal{X}_0\mathcal{X}_0}^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}})) \end{aligned} \quad (3.7)$$

By the Woodbury matrix identity [Rasmussen & Williams (2006)], one can derive the follow-

ing:

$$(\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}_0} \mathbb{K}_{\mathcal{X}_0\mathcal{X}_0}^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}})^{-1} = \Lambda_{\mathcal{X}}^{-1} - \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}_0} (\mathbb{K}_{\mathcal{X}_0\mathcal{X}_0} + \mathbb{K}_{\mathcal{X}_0\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}_0})^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \quad (3.8)$$

Because $\Lambda_{\mathcal{X}}$ is a diagonal matrix, (3.8) is efficiently computed, so as $\mathcal{L}$. Moreover, the partial derivative of $\mathcal{L}$ with respect to α is given explicitly as follows:

$$\frac{\partial \mathcal{L}}{\partial \alpha} = \varphi(\mathcal{X})^{\top} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}_0} \mathbb{K}_{\mathcal{X}_0\mathcal{X}_0}^{-1} \mathbb{K}_{\mathcal{X}_0\mathcal{X}})^{-1} (\mathcal{Y} - \varphi(\mathcal{X}) \alpha) \quad (3.9)$$

Thus, one can learn α with a gradient descent using the partial derivative (3.9). Note that gradient descent on $\mathcal{L}$ for kernel hyperparameters is not a good idea: it is possible to express the partial derivatives in closed forms, but the time complexity increases so that the exact computation becomes intractable. The Markov-chain Monte Carlo (MCMC) algorithm on the Boltzmann distribution [Landau & Lifshitz (2013)] with $\mathcal{L}$ as the energy is a practical alternative.

The last issue is about the outlier classification: here the same kind of Bayesian modeling with that for benthic $\delta^{18}\text{O}$ in section 2.3.2 is applied.

Now we are prepared. Algorithm 2 summarizes the inference procedure.

Algorithm 2: GPR

- 1 **initialize** coefficient parameters and kernel hyperparameters
 - 2 **set** all data to be inliers
 - 3 **while** *converging* **do**
 - 4 **tune** kernel hyperparameters with inliers
 - 5 **learn** coefficient parameters with inliers
 - 6 **update** the variance function with inliers
 - 7 **classify** inliers and outliers
 - 8 **end**
 - 9 **compute** the posterior predictive as the regression model
 - 10 **return** parameters, outliers and the regression model
-

3.3.2 MULTIMODAL GAUSSIAN PROCESS REGRESSION

As discussed in section B.2.3, a GPR is misspecified if the data has multiple modes over inputs. To deal with such a drawback, several works have been presented. [Tresp (2001); Rasmussen & Ghahramani (2002)] consider multiple generative processes based on the input space separation and local expert models. [Lázaro-Gredilla et al. (2012); Kaiser et al. (2018)] depend on a set of GP priors on the multiple hidden states that stand for different generative processes and more flexible probabilistic association models. [Bodin et al. (2017)] extends the input domain with continuous latent inputs to be paired with inputs and eventually marginalized out.

Inspired by the idea of [Bodin et al. (2017)], here we present a new multimodal GPR that extends the input domain with discrete latent inputs to be sampled by the Gibbs sampling (section A.4.1) or estimated by the EM algorithm with an *explicit* M-step. Let $\mathcal{L} = \{L_n\}_{n=1}^N \subseteq \{0, 1, 2, \dots, C\}$ and $\mathcal{L}_0$ be the latent and *fixed* pseudo-latent inputs that are coupled with the real and pseudo-inputs $\mathcal{X} = \{X_n\}_{n=1}^N$ and $\mathcal{X}_0$, respectively. Then, the extended kernel $\tilde{\mathbb{K}}$ is defined as follows:

$$\tilde{\mathbb{K}}((x_1, l_1), (x_2, l_2)) \triangleq \mathbb{K}(x_1, x_2) \cdot \exp(-\delta^2 \cdot 1_{\{l_1 \neq l_2\}}) \quad (3.10)$$

, where δ is a hyperparameter that controls the cluster similarity. Note that $\tilde{\mathbb{K}}$ is a well-defined kernel because $d(l_1, l_2) = 1_{\{l_1 \neq l_2\}}$ is a metric over $\{0, 1, 2, \dots, C\}$.

Now, the full extended model is specified by the following observational model (3.11), GP prior (3.12) and cluster prior (3.13): note that assigning a far smaller value to ω_0 than the other ω_c 's implies that the cluster 0 is for outliers.

$$p(\mathcal{Y}|\beta, \mathcal{X}, \mathcal{L}) = \mathcal{N}(\mathcal{Y}|\beta, \Lambda_{\mathcal{X}}^{(\mathcal{L})}) \triangleq \prod_{n=1}^N \prod_{c=0}^C \mathcal{N}(Y_n|\beta_n, \Lambda_{X_n}^{(c)})^{1_{\{L_n=c\}}} \quad (3.11)$$

$$p(\beta|\mathcal{X}, \mathcal{L}) = \mathcal{GP}(\beta|\mu_{\mathcal{X}}^{(\mathcal{L})}, \mathbb{K}_{\mathcal{X}\mathcal{L}, \mathcal{X}\mathcal{L}}) \quad (3.12)$$

$$p(\mathcal{L}) = \prod_{n=1}^N \prod_{c=0}^C \omega_c^{1_{\{L_n=c\}}} \quad (3.13)$$

Once the label $\mathcal{L}$ is given, one can compute the corresponding posterior predictive distribution similarly with (3.6). If $\mathcal{L}$ is given as a set of samples, then the associated posterior predictive is obtained as the approximation that marginalizes $\mathcal{L}$ out by averaging them over samples.

What we are interested in is how to update $\mathcal{L}$. Suppose that all the unknowns are given. One would like to directly access to the posterior of $\mathcal{L}$ as follows:

$$\begin{aligned} p(\mathcal{L}|\mathcal{X}, \mathcal{Y}) &\propto p(\mathcal{Y}, \mathcal{L}|\mathcal{X}) = p(\mathcal{L}) \int p(\mathcal{Y}|\beta, \mathcal{X}, \mathcal{L}) p(\beta|\mathcal{X}, \mathcal{L}) d\beta \\ &= \left(\prod_{n=1}^N \prod_{c=0}^C \omega_c^{1_{\{L_n=c\}}} \right) \cdot \mathcal{N}\left(\mathcal{Y} \middle| \mu_{\mathcal{X}}^{(\mathcal{L})}, \mathbb{K}_{\mathcal{X}\mathcal{L}, \mathcal{X}\mathcal{L}} + \Lambda_{\mathcal{X}}^{(\mathcal{L})}\right) \end{aligned} \quad (3.14)$$

The problem is that (3.14) is hard to deal with for $\mathcal{L}$ because the latter term is indecomposable with respect to $\mathcal{L}$. Instead, let us think about the following: recall that γ is assumed to be a sufficient statistic in section B.7.

$$\begin{aligned} p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L}) &\approx \prod_{n=1}^N \int p(Y_n|b, X_n, L_n) p(b|X_n, L_n, \gamma, \mathcal{X}_0, \mathcal{L}_0) db \\ &= \prod_{n=1}^N \mathcal{N}(Y_n | \tilde{\mu}(X_n, L_n, \gamma), \tilde{\nu}(X_n, L_n)) \\ \tilde{\mu}(X_n, L_n, \gamma) &= \mu_{X_n}^{(L_n)} + \mathbb{K}_{X_n L_n, \mathcal{X}_0 \mathcal{L}_0} \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X}_0 \mathcal{L}_0}^{-1} \left(\gamma - \mu_{\mathcal{X}_0}^{(\mathcal{L}_0)} \right) \triangleq \tilde{\mu}_n \\ \tilde{\nu}(X_n, L_n) &= \Lambda_{X_n}^{(L_n)} + \mathbb{K}_{X_n L_n, X_n L_n} - \mathbb{K}_{X_n L_n, \mathcal{X}_0 \mathcal{L}_0} \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X}_0 \mathcal{L}_0}^{-1} \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, X_n L_n} \triangleq \tilde{\nu}_n \end{aligned} \quad (3.15)$$

Therefore, one can approximate $p(\mathcal{L}|\gamma, \mathcal{X}, \mathcal{Y})$ with following:

$$\begin{aligned} p(\mathcal{L}|\gamma, \mathcal{X}, \mathcal{Y}) &\propto p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L}) p(\mathcal{L}) \\ &\approx \prod_{n=1}^N \left(\mathcal{N}(Y_n | \tilde{\mu}(X_n, L_n, \gamma), \tilde{\nu}(X_n, L_n)) \prod_{c=0}^C \omega_c^{1_{\{L_n=c\}}} \right) \end{aligned} \quad (3.16)$$

Note that (3.16) is decomposable with respect to $\mathcal{L}$, thus sampling $\mathcal{L}$ can be done efficiently

given $\mathcal{X}$, $\mathcal{Y}$ and the latent sufficient statistic γ . Together with $p(\gamma|\mathcal{X}, \mathcal{Y}, \mathcal{L}) \approx \varphi(\gamma|\mathcal{L})$ in section B.7, Gibbs sampling (section A.4.1) efficiently samples γ given $\mathcal{L}$ and $\mathcal{L}$ given γ iteratively. Furthermore, one can estimate $\mathcal{L}$ by the following EM algorithm that marginalizes γ out:

E-step: Compute the following expectation:

$$\Psi(\mathcal{L}|\mathcal{L}^{(t)}) \triangleq \mathbb{E}_{\gamma|\mathcal{X}, \mathcal{Y}, \mathcal{L}^{(t)}} [\log p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L}) + \log p(\gamma|\mathcal{X}_0, \mathcal{L}_0) + \log p(\mathcal{L})] \quad (3.17)$$

M-step: Find $\mathcal{L}^{(t+1)} \triangleq \operatorname{argmax}_{\mathcal{L}} \Psi(\mathcal{L}|\mathcal{L}^{(t)})$.

With the approximation $p(\gamma|\mathcal{X}, \mathcal{Y}, \mathcal{L}) \approx \varphi(\gamma|\mathcal{L})$, we have:

$$\begin{aligned} \Psi(\mathcal{L}|\mathcal{L}^{(t)}) &= \int p(\gamma|\mathcal{X}, \mathcal{Y}, \mathcal{L}^{(t)}) [\log p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L}) + \log p(\gamma|\mathcal{X}_0, \mathcal{L}_0) + \log p(\mathcal{L})] d\gamma \\ &\approx \log p(\mathcal{L}) + \int \varphi(\gamma|\mathcal{L}^{(t)}) \log p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L}) d\gamma + \int \varphi(\gamma|\mathcal{L}^{(t)}) \log p(\gamma|\mathcal{X}_0, \mathcal{L}_0) d\gamma \end{aligned} \quad (3.18)$$

Because the last term in (3.18) is irrelevant to $\mathcal{L}$, we just need to find the maximizer of the sum of the first and second terms. Note that $p(\mathcal{Y}|\gamma, \mathcal{X}, \mathcal{L})$ is approximated with (3.15), thus:

$$\begin{aligned} &\Psi(\mathcal{L}|\mathcal{L}^{(t)}) - \int \varphi(\gamma|\mathcal{L}^{(t)}) \log p(\gamma|\mathcal{X}_0, \mathcal{L}_0) d\gamma \\ &\approx \log p(\mathcal{L}) \\ &\quad + \int \varphi(\gamma|\mathcal{L}^{(t)}) \log \prod_{n=1}^N \mathcal{N}\left(Y_n \middle| \mu_{X_n}^{(L_n)} + \mathbb{K}_{X_n L_n, X_0 L_0} \mathbb{K}_{X_0 L_0, X_0 L_0}^{-1} \left(\gamma - \mu_{X_0}^{(L_0)}\right), \tilde{v}_n\right) d\gamma \\ &= \log \prod_{n=1}^N \prod_{c=0}^C \omega_c^{1_{\{L_n=c\}}} \\ &\quad + \sum_{n=1}^N \int \varphi(\gamma|\mathcal{L}^{(t)}) \log \mathcal{N}\left(Y_n \middle| \mu_{X_n}^{(L_n)} + \mathbb{K}_{X_n L_n, X_0 L_0} \mathbb{K}_{X_0 L_0, X_0 L_0}^{-1} \left(\gamma - \mu_{X_0}^{(L_0)}\right), \tilde{v}_n\right) d\gamma \\ &= \sum_{n=1}^N \left[\sum_{c=0}^C 1_{\{L_n=c\}} \log \omega_c - \frac{1}{2} \left(\frac{A_n^2 + B_n}{\tilde{v}_n} + \log \tilde{v}_n \right) \right] \end{aligned} \quad (3.19)$$

, where:

$$\begin{aligned}
A_n &\triangleq Y_n - \mu_{X_n}^{(L_n)} - \mathbb{K}_{X_n L_n, \mathcal{X}_0 \mathcal{L}_0} \Sigma^{-1} \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X} \mathcal{L}^{(t)}} \left(\Lambda_{\mathcal{X}}^{(\mathcal{L}^{(t)})} \right)^{-1} \left(\mathcal{Y} - \mu_{\mathcal{X}}^{(\mathcal{L}^{(t)})} \right) \\
B_n &\triangleq \mathbb{K}_{X_n L_n^{(t)}, \mathcal{X}_0 \mathcal{L}_0} \Sigma^{-1} \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, X_n L_n^{(t)}} \\
\Sigma &\triangleq \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X}_0 \mathcal{L}_0} + \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X} \mathcal{L}^{(t)}} \left(\Lambda_{\mathcal{X}}^{(\mathcal{L}^{(t)})} \right)^{-1} \mathbb{K}_{\mathcal{X} \mathcal{L}^{(t)}, \mathcal{X}_0 \mathcal{L}_0}
\end{aligned} \tag{3.20}$$

Thus, $\Psi(\mathcal{L} | \mathcal{L}^{(t)})$ is decomposable with respect to $\mathcal{L}$. Thus, the M-step can be efficiently run by following:

$$L_n^{(t+1)} = \underset{L_n}{\operatorname{argmax}} \left[\log \omega_{L_n} - \frac{1}{2} \left(\frac{A_n^2 + B_n}{\tilde{v}_n} + \log \tilde{v}_n \right) \right] \tag{3.21}$$

What remains is to deal with the unknown label l for the query input x . In the learning procedure, all labels for the inputs are learned by the EM algorithm as described above. After learning, we have a pair of inputs and the learned labels $(\mathcal{X}, \hat{\mathcal{L}})$. Any classifier $p(l|x; \mathcal{X}, \hat{\mathcal{L}})$ including Gaussian process classifier (section B.3) that is learned from $(\mathcal{X}, \hat{\mathcal{L}})$ can be the association distribution of the label l given the query input x .

Finally, the posterior predictive distribution is given by the following Gaussian mixture:

$$p(y|x, \mathcal{X}, \mathcal{Y}, \mathcal{X}_0) = \sum_{l=0}^C p(l|x; \mathcal{X}, \hat{\mathcal{L}}) \mathcal{N}(y | \mu_x^{(l)} + m(x, l), \Lambda_x^{(l)} + s(x, l)) \tag{3.22}$$

, where: let $\mathbb{K}_{xl,0} = \mathbb{K}_{xl, \mathcal{X}_0 \mathcal{L}_0}$, $\mathbb{K}_{0,xl} = \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, xl}$, $\mathbb{K}_{00} = \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X}_0 \mathcal{L}_0}$, $\mathbb{K}_{01} = \mathbb{K}_{\mathcal{X}_0 \mathcal{L}_0, \mathcal{X} \hat{\mathcal{L}}}$ and $\mathbb{K}_{10} = \mathbb{K}_{\mathcal{X} \hat{\mathcal{L}}, \mathcal{X}_0 \mathcal{L}_0}$ for short.

$$\begin{aligned}
m(x, l) &\triangleq \mathbb{K}_{xl,0} \left(\mathbb{K}_{00} + \mathbb{K}_{01} \left(\Lambda_{\mathcal{X}}^{(\hat{\mathcal{L}})} \right)^{-1} \mathbb{K}_{10} \right)^{-1} \mathbb{K}_{01} \left(\Lambda_{\mathcal{X}}^{(\hat{\mathcal{L}})} \right)^{-1} \left(\mathcal{Y} - \mu_{\mathcal{X}}^{(\hat{\mathcal{L}})} \right) \\
s(x, l) &\triangleq \mathbb{K}_{xl,xl} - \mathbb{K}_{xl,0} \mathbb{K}_{00}^{-1} \mathbb{K}_{0,xl} + \mathbb{K}_{xl,0} \left(\mathbb{K}_{00} + \mathbb{K}_{01} \left(\Lambda_{\mathcal{X}}^{(\hat{\mathcal{L}})} \right)^{-1} \mathbb{K}_{10} \right)^{-1} \mathbb{K}_{0,xl}
\end{aligned} \tag{3.23}$$

Now we are again prepared. Algorithm 3 summarizes the inference procedure: note that this algorithm can be reduced to an unsupervised classifier by discarding last two steps before the last

line and not returning the association distribution and regression model (removable steps).

Algorithm 3: Multimodal GPR

```

1 initialize coefficient parameters and kernel hyperparameters
2 randomize the labels
3 while converging do
4   update labels by the EM algorithm
5   tune kernel hyperparameters
6   for each label do
7     learn the label-specific mean prior function with the classified data if needed
8     update the label-specific variance function with the classified data
9   end
10 end
11 learn association distribution given the learned labels (removable)
12 compute the posterior predictive as the regression model (removable)
13 return parameters, labels, association distribution (removable) and the regression
    model (removable)

```

Figure 3.2 visualizes the results obtained by running the reduced version of algorithm 2 on some examples. We could check that similar result was obtained with that of [Kaiser et al. (2018)] to classify outliers in the last panel of the figure. Though [Bodin et al. (2017)] deals with a more general setting that the extended inputs are continuous and to be marginalized out, our model so far is not only mathematically but also practically interesting because it efficiently updates the labels by the EM algorithm with an explicit M-step or samples them by the Gibbs sampling algorithm.

3.4 RESULTS AND DISCUSSION

For each proxy, we tried the following three models: homoscedastic GPR, heteroscedastic GPR and the reduced multimodal GPR for finding multiple clusters of the data if exist. Each set of results are visualized in the corresponding figure that consists of six panels. The first panel shows the resulting calibration model of the proxy over SST; the second panel presents the nonparametric bias function over SST; the third panel indicates the locations of inliers (green) and outliers (red) on the world map; the fourth panel gives the inferred standard deviation function over SST;

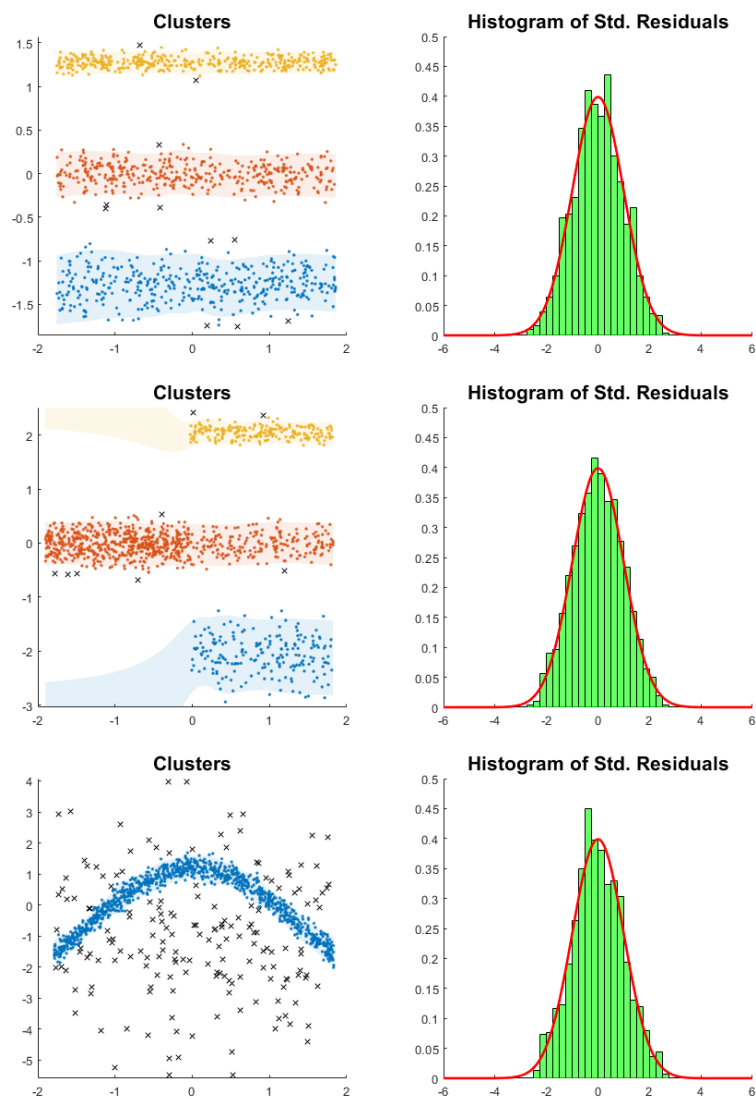


Figure 3.2: Some example results from our multimodal GPR.

the fifth and sixth panels are about the standardized residuals that are obtained by subtracting the proxy observations by their regressions and then dividing them by their inferred standard deviations.

3.4.1 ALKENONE CALIBRATION MODELS

Figure 3.3 visualizes the homoscedastic GPR results for the alkenone proxy data. As shown in the first and second panels, too many observations are classified as outliers, which implies that the model is not correctly specified. It is also supported by the histogram of the standardized residuals: inferred standard deviation function seems to be too narrow to encompass all the observations, but also too broad for the classified inliers. The classified outliers show somehow but not certain regional pattern in the world map. Even though the model was homoscedastic GPR, the inferred standard deviation function over SST is heteroscedastic: because observational variance is constant, this inconsistency stems from the model uncertainty. The unstable nonparametric bias function also supports the existence of extensive model uncertainty.

The heteroscedastic GPR results are visualized in figure 3.4. In contrast to the homoscedastic case, the nonparametric bias function is stable and not many outliers are classified. The inferred standard deviation function over SST clearly shows the heteroscedastic characteristic of the proxy data: $U_{37}^{K'}$ observations at higher temperatures are more certain than those at lower temperatures. The histogram that fits well for the standard normal distribution also supports that the classified inliers follow a single Gaussian model. However, there is one doubtful aspect that leads to the further investigation: outliers under the regression function seem to form another cluster, though they do not show a vivid regional pattern.

Some multimodal GRP clustering results are shown in figure 3.5: those results are preliminary because the label uncertainty was not considered. Here the cluster-specific mean prior functions are defined by different constant functions. The outlier cluster mentioned in the above paragraph gets absorbed into inliers as the number of clusters increases. Note that the analysis for 1 cluster (the top panel) suggests a nonparametric way of classifying outliers.

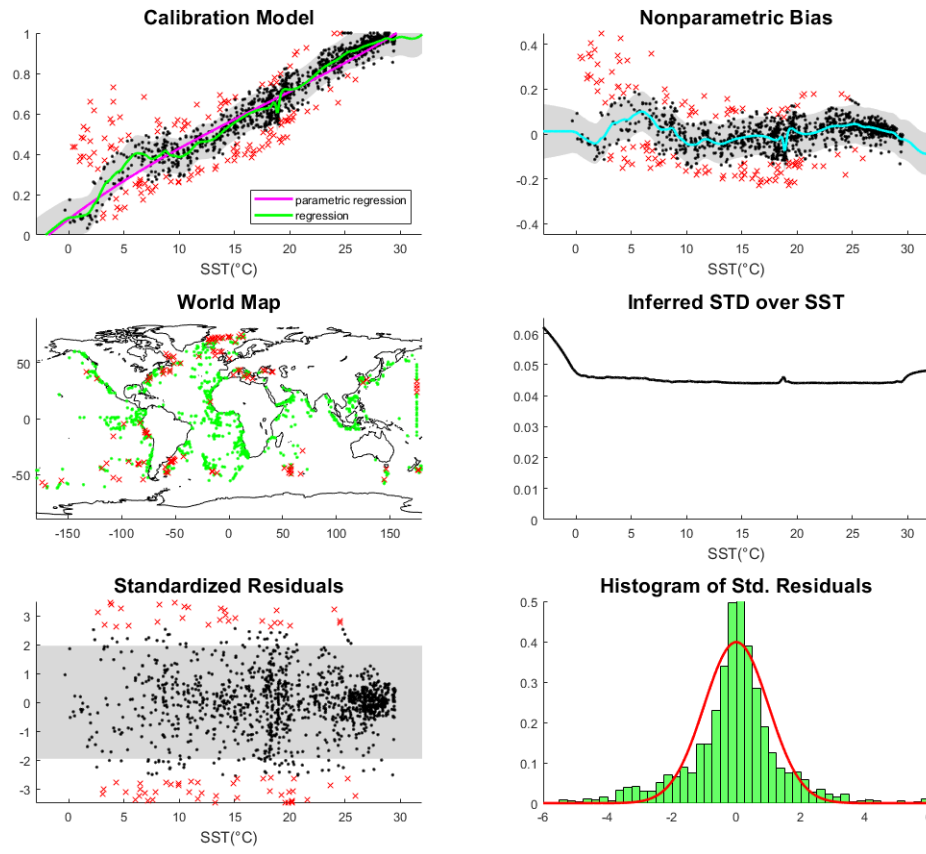


Figure 3.3: Homoscedastic Gaussian process regression of $U_{37}^{K'}$ on SST. In each panel, red crosses are the outliers and shaded region indicates the 95% confidence band.

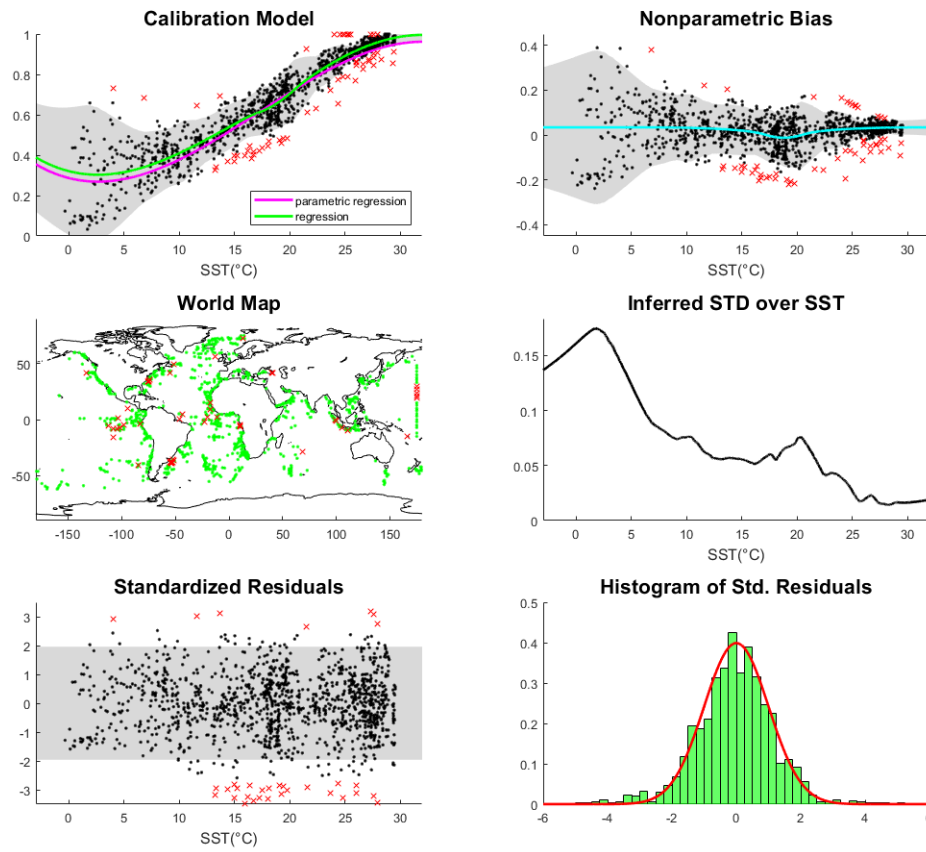


Figure 3.4: Heteroscedastic Gaussian process regression of $U_{37}^{K'}$ on SST.

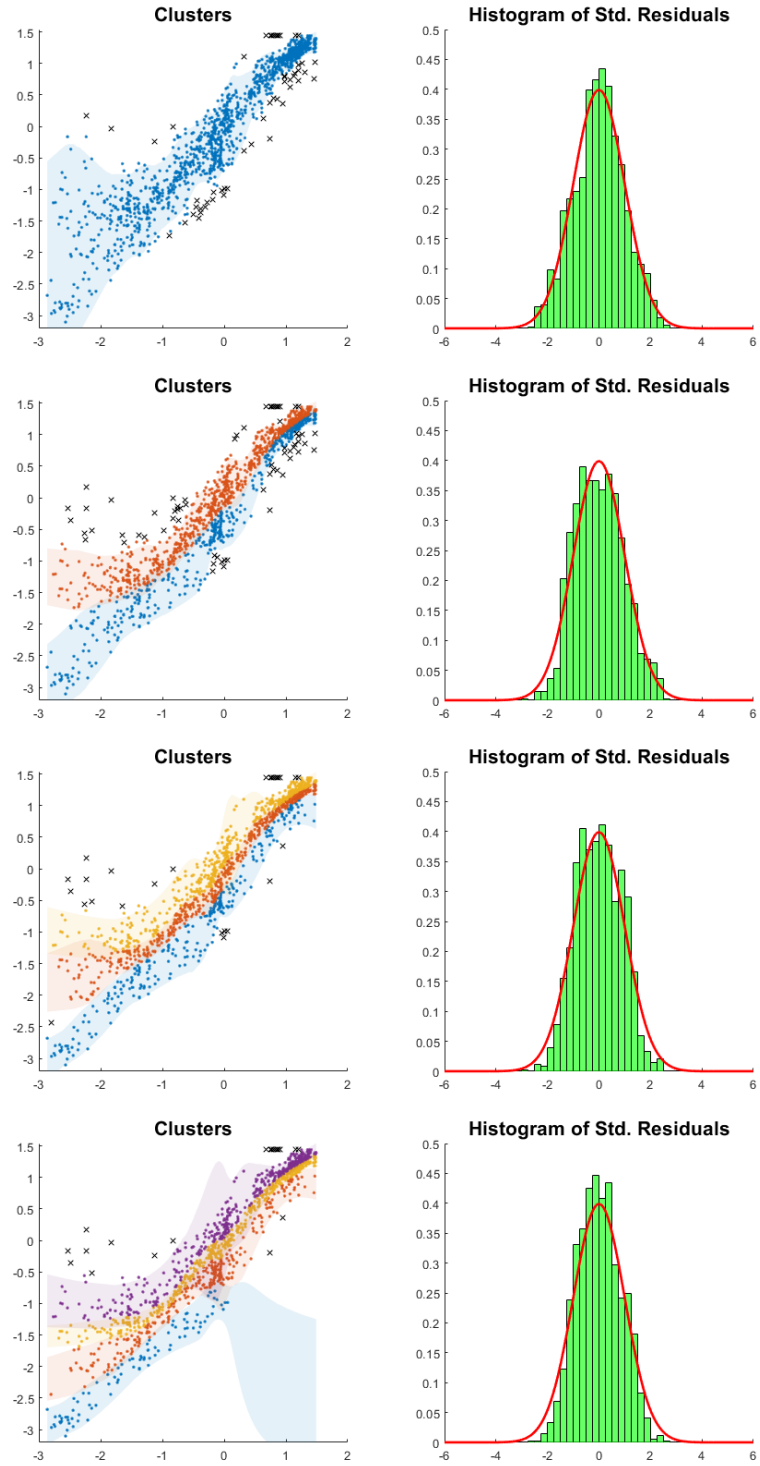


Figure 3.5: Multimodal Gaussian process regressions of $U_{37}^{K'}$ on SST.

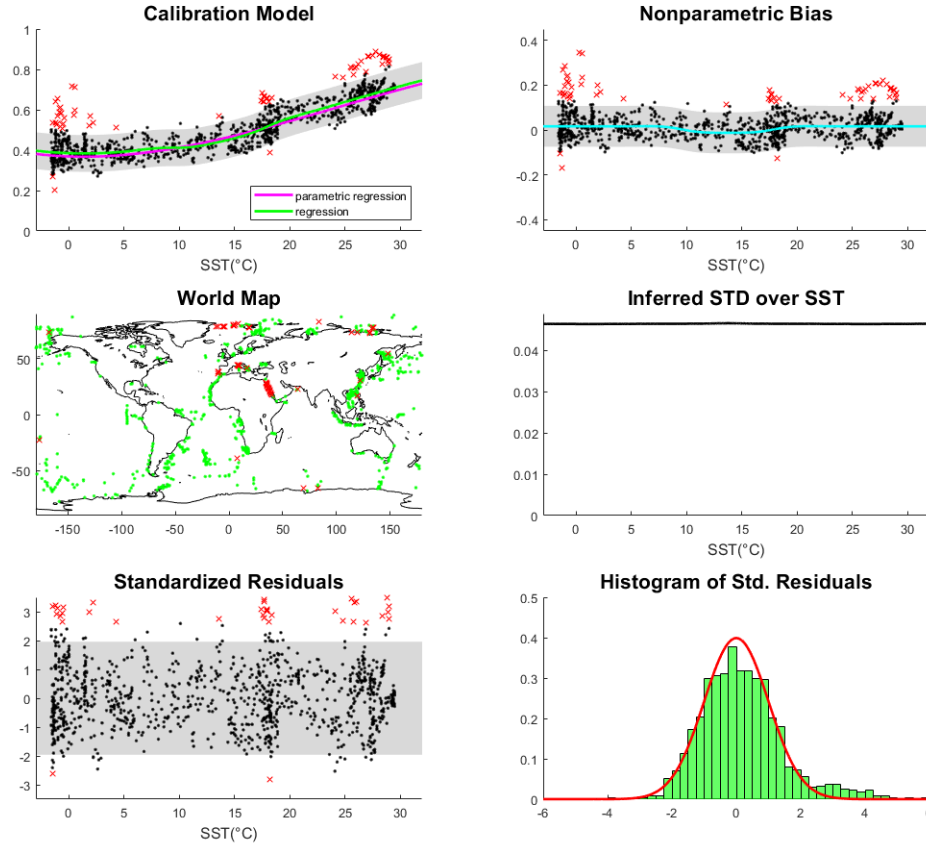


Figure 3.6: Homoscedastic Gaussian process regression of TEX_{86} on SST.

3.4.2 TEX_{86} CALIBRATION MODELS

Figure 3.6 visualizes the homoscedastic GPR results for the TEX_{86} proxy data. In the world map, the outliers show a vivid regional pattern: they are mainly from the Red Sea and the Arctic Ocean. It supports that the data are commonly follow the same calibration model. Unlike the homoscedastic GPR of the alkenone in section 3.4.1, the nonparametric bias function is stable over SST, so here the regression and the parametric regression of a cubic polynomial are almost identical. However, this calibration model is still not correctly specified: the shape of the histogram of standardized residuals implies that the inliers are more likely following a Gaussian mixture rather than a single Gaussian.

The heteroscedastic GPR results are visualized in figure 3.7. The calibration model gives a

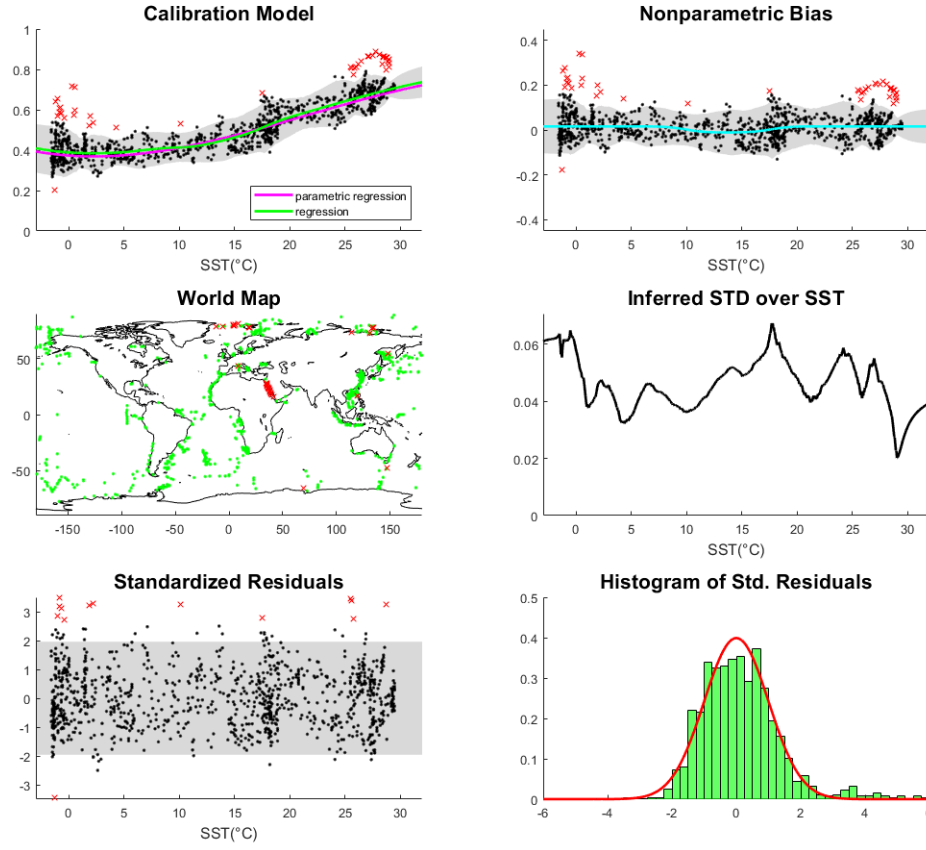


Figure 3.7: Heteroscedastic Gaussian process regression of TEX_{86} on SST.

more flexible variance function over SST to classify less observations as outliers. Thus, the regional pattern of the outliers become more vivid than the homoscedastic case. Unlike alkenone that show the diminishing pattern, the inferred standard deviation function of TEX_{86} over SST does not have a certain pattern here. Though outliers are addressed more carefully and sensitively, the histogram of standardized residuals still implies that the observations tend to follow a mixture of Gaussian distributions.

Some preliminary multimodal GRP clustering results are shown in figure 3.8. Histograms of the standardized residuals tend to fit well for the standard normal distribution as the number of clusters increases.

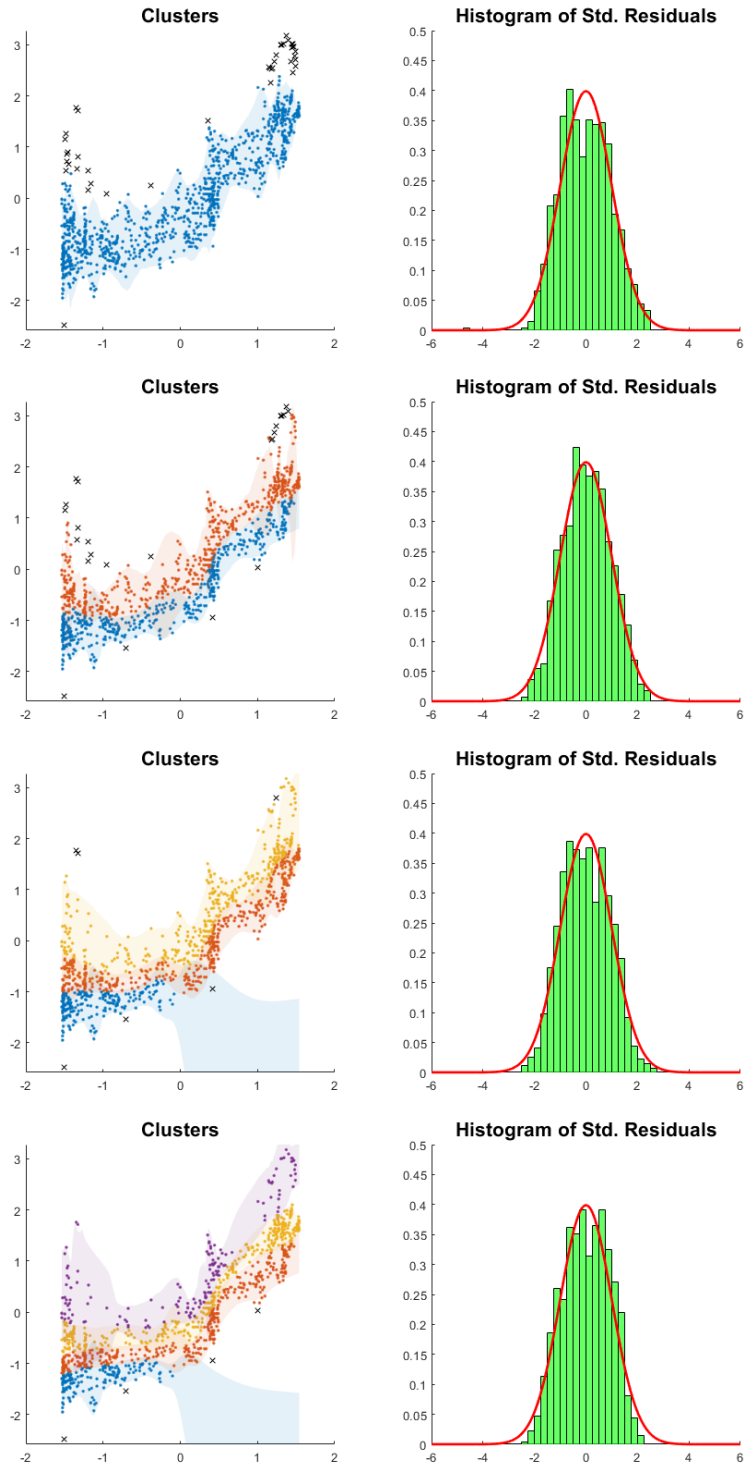


Figure 3.8: Multimodal Gaussian process regressions of TEX_{86} on SST.

3.5 FUTURE WORK

Many previous works of the calibration model construction have discarded or adjusted the proxy observations or SSTs based on the data-specific characteristics. This can be understood in the following framework: there are two categories of outliers. Outliers of the first category are just meaningless noises without any information. Thus, such outliers should be discarded in learning procedure, especially for the nonparametric models that are directly depending on the data. Outliers of the second category are, however, originally informative but structurally contaminated for some data-specific reasons. Classification into these two categories cannot be done without the aid of relevant experts in paleoceanography and/or paleoclimatology. Thus, to reflect such subtle differences of outliers, in future our work must be collaborated with the relevant experts to construct not only mathematically but also practically meaningful calibration models.

In relation to the previous criticism, an interdisciplinary work to deal with the lab-specific biases in alkenone is ongoing. One of the collaborators, Dr. Timothy Herbert, suspects that lab-specific biases exist and affect the pattern of data for the reason that the proxy dataset is a collection of the observations from multiple laboratories and each lab has its own procedure. This problem is similar to the multimodal GPR in section 3.3.2 but could be easier because it is supervised: lab associations are given a priori to be paired with the data.

As BAYSPAR of [Tierney & Tingley (2014)], including spatial information of the data to the inputs brings another interesting goal. GP models are easily extensible to take more inputs, such as spherical coordinates in this problem. In fact, we already have some preliminary results that are shown in figure 3.9: whereas BAYSPAR gives GP priors on the slope and intercept coefficients that are defined regionally with grids, our model is more nonparametric and returns the exact posterior predictive distributions instead of the slopes and intercepts. The site-specific regression model at each site tends to follow the observations near to the site over the SSTs that are in range of where those observations are made, but soon gets deviated to the general data as going outside of the range. One major problem of this approach is regularization by the guidance of experts

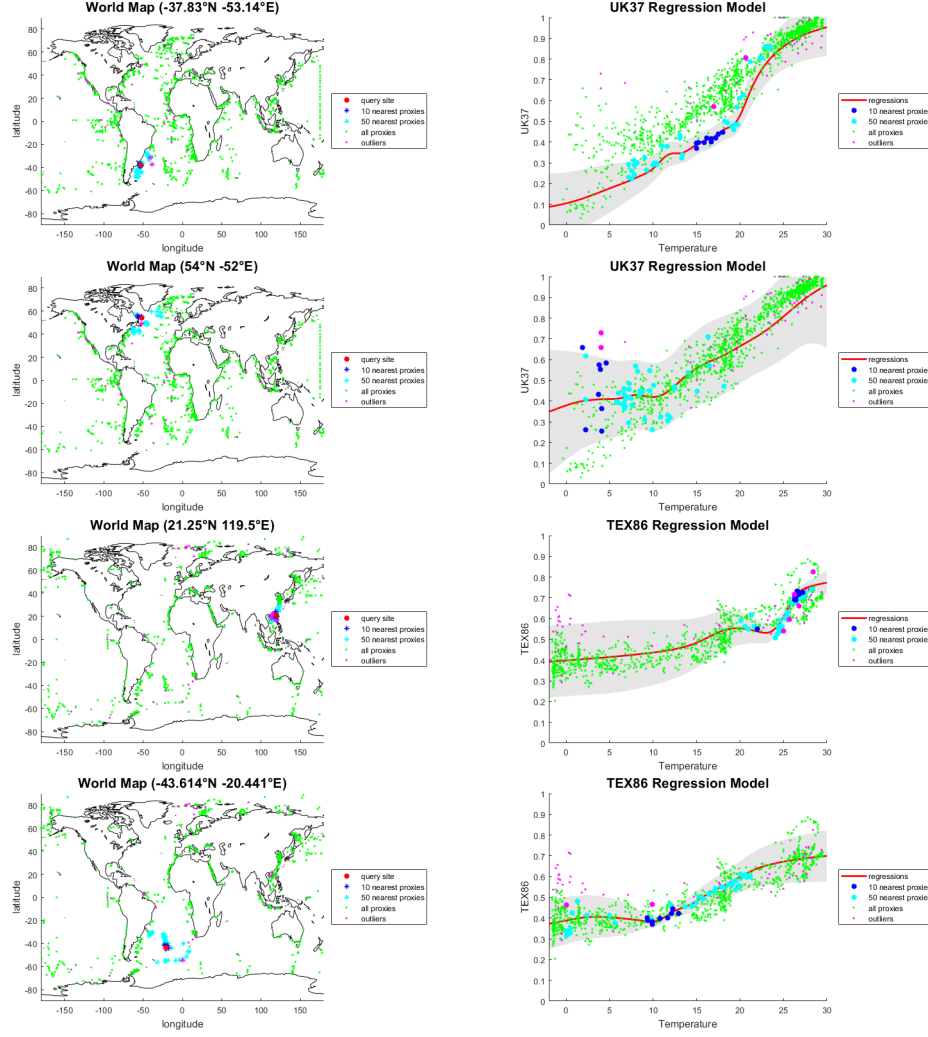


Figure 3.9: Some preliminary results from the site-specific GPR calibration model.

relevant to data: models that are too faithful to the data could lose something important such as its applicability or practicability by overfitting.

More abstract GP-based models, such as deep GP (section B.6), could be applied to the datasets as one of the future research directions. Like what deep neural network models do, deep GPs are simultaneously doing inference and extracting features. As discussed in section B.6, deep GPs contribute to overcoming the suboptimality of GP models, at the cost of interpretability: shallow GPR models are given explicitly and mean and variance terms of the posterior predictive distributions can be decomposed into either parametric regression and nonparametric bias or observational uncertainty and model uncertainty, as discussed earlier, but a deep GP regression

cannot be expressed in a closed form thus loses such a decomposable aspect.

This project ends at the calibration model construction: the obvious work to be followed is constructing a model for reconstructing SSTs given ancient proxy data. It could be done independently on the proxies, i.e., restoring an SST for each proxy one-by-one, but a better approach would be connecting SSTs over ages or locations, as discussed in section 3.2. If we define a GP state-space model over ages and locations, then the choice of kernels or their hyperparameters could be suggested by the relevant experts, based on their physical models from other related data.

4

Atmospheric CO₂ Reconstruction Using Boron Proxy

Carbon dioxide (CO₂) is a colorless gas in Earth's atmosphere, consists of a carbon atom covalently double bonded to two oxygen atoms. The current concentration is about 0.04% (412 ppm) by volume, having risen from pre-industrial levels of 280 ppm [[Eggleton \(2012\)](#)]. Though CO₂ gas occupies only a small portion of the atmosphere, its variations are key factors of climate change that has recently been a serious issue threatening all mankind.

The boron isotope, which is also a proxy for pH reconstruction, is one of the main marine proxies used for the atmospheric CO₂ reconstruction. Boron has two stable isotopes, ¹¹B and

^{10}B , in the forms of boric acid $\text{B}(\text{OH})_3$ and borate ion $\text{B}(\text{OH})_4^-$. Because boric acid has a higher $^{11}\text{B}/^{10}\text{B}$ ratio than borate ion by 1.0272 [Klochko et al. (2006)], the boron isotopic composition in marine carbonate is affected by the pH of the seawater in which the carbonate grows: the boron proxy is based on the observation that predominantly only borate ion is incorporated into marine calcium carbonate, such as the shells of foraminifera [Hemming & Hanson (1992)].

The composition of total dissolved carbon (DIC) to total alkalinity (ALK) is linked to the seawater pH that is closely related to CO_2 . Seawater CO_2 content is estimated by the nature of the CO_2 -system (with additional parameters) in seawater that considers the boron isotope composition [Lemarchand et al. (2002)] and minor vital effects relating to well-preserved foraminiferal physiology [Henehan et al. (2016)]. Atmospheric CO_2 is then reconstructed from the seawater CO_2 by taking the air-sea disequilibrium that should be accounted for as well*.

This chapter covers our preliminary work on the reconstruction of atmospheric CO_2 with boron proxy. In contrast to the previous work that accesses each proxy observation one-by-one, we apply various state-space models such as the particle smoother (section A.3.2) and Gaussian process state-space (GPST) model (section B.4).

4.1 RELATED WORK

While [Hönisch et al. (2009)] estimated a 2.1-million-year-long sea surface partial pressure of CO_2 ($p\text{CO}_2$) from boron isotopes in planktic foraminifer shells and could confirm that their estimates are consistent with the ice core records, [Foster & Rae (2016)] more carefully described how CO_2 is reconstructed from the boron isotopic composition ($\delta^{11}\text{B}$) of marine calcium carbonate. They reviewed the chemical principles underlying the proxy, summarized the available calibration data, and gave details of how boron isotopes can be utilized for estimating ocean pH and ultimately past atmospheric CO_2 .

While the above works have been focusing on the pointwise inference, [Garcia-Olivares & Herrero (2013)] constructed Local Stratification (LS) model, a dynamic model of global ice volume,

*I refer to <http://www.p-co2.org/boron> for the first three paragraphs in this chapter.

extent of the Antarctic ice sheet and atmospheric CO_2 over time. The LS model connects a set of different past events by the system of first-order ordinary differential equations (ODEs), just as PP04 models of [Paillard & Parrenin (2004); Garcia-Olivares & Herrero (2012)] that were fitted with the experimental time series of $\delta^{18}\text{O}$ and CO_2 for the last 800 kiloyears.

4.2 RESEARCH MOTIVE AND DATA

As discussed at the beginning of chapter A, the inference on the linked hidden states (past climate events over time) given the paired observations (proxy observations) can be effectively processed in the framework of state-space models. Thus, our goal is to design and run various state-space models for inferring the atmospheric CO_2 using boron proxy and verifying the results with the Antarctic ice core records [Monnin et al. (2001); Petit et al. (1999); Raynaud et al. (2005); Indermühle et al. (2000); Siegenthaler et al. (2005); Lüthi et al. (2008)] that have trapped air composed with CO_2 thus to be used as the “true” values. Because such “true” values are limited up to 800 kiloyears, our inference is also done up to at most that point. Let us call them “experimental time series”.

A transition model of the state-space model can be either parametric or nonparametric. In this section we consider both of the cases. Firstly, we construct a parametric state-space model with the transition model based on the dynamic model on multiple climate events including atmospheric CO_2 of the LS model [Garcia-Olivares & Herrero (2013)]: its parametric form is justified scientifically thus itself meaningful. This work was collaborated with my advisor Dr. Charles E. Lawrence and one of his students, Andrew Kaluzny and presented to the attendees of the 2nd Workshop on Cenozoic $p\text{CO}_2$ Reconstruction by my advisor on May 30th, 2018.

However, the state-space model with a parametric transition model would be problematic for the following reasons: 1) the inference highly depends on the estimated parameters of the transition model, but the underlying dynamics could be nonstationary due to the change of the other surrounding conditions over time, and 2) the transition of the hidden states is usually assumed to be made regularly (for example, for every kiloyear), but the boron proxy is not regularly spaced

over ages and too scarce to be properly rearranged regularly. The GPST is a nonparametric state-space model that deals with such irregularly-spaced inputs, thus could be perfect for this purpose. Therefore, we also try the GPST on the dataset.

However, the state-space model with a parametric transition model would be problematic for the following reasons: 1) the inference highly depends on the estimated parameters of the transition model, but the underlying dynamics could be nonstationary due to the change of the other surrounding conditions over time, and 2) the transition of the hidden states is usually assumed to be made regularly (for example, for every kiloyear), but the boron proxy is not regularly spaced over ages and too scarce to be properly rearranged regularly. The GPST is a nonparametric state-space model that deals with such irregularly-spaced inputs, thus could be perfect for this purpose. Therefore, we also try the GPST on the dataset.

We use the boron proxy data in ODP668 and ODP999 [Dyez et al. (2018)]. Our sea level data are referred to [Spratt & Lisiecki (2016)]. Benthic $\delta^{18}\text{O}$ proxy are from the ocean sediment cores MD95-2042 [Bard et al. (2004); Shackleton et al. (2000, 2004)], MD02-2575 [Nürnberg et al. (2008)] and MD02-2589 [Molyneux et al. (2007); Diz et al. (2007)].

Our work here is, however, strictly limited for the following reasons: 1) our emission (observation) models for proxies are just simple and fixed parametric models that are pre-trained based on the samples from other datasets, 2) some uncertain values such as ages of the proxies are regarded as the fixed ones according to the datasets, and 3) the number of boron proxy observations up to 800 kiloyears is only 86. Thus, the work is preliminary rather than complete, just for drawing the interdisciplinary experts' interests in future research opportunities.

4.3 MODELS

Two models are applied separately. The first model uses the particle smoother and the second one is GPST. In each model, its emission model is defined by a set of fixed observational models that are pre-trained by the data. However, the parametric transition model of the first model based is trained during the inference, and so as the kernel hyperparameters in the second model.

The first model is depending on the proxies of benthic $\delta^{18}\text{O}$ and the boron proxy, $\delta^{11}\text{B}$, which is a variation derived from the boron isotopes as follows [Foster & Rae (2016)]:

$$\delta^{11}\text{B} \triangleq \left(\left(\frac{{}^{11}\text{B}/{}^{10}\text{B}_{\text{sample}}}{{}^{11}\text{B}/{}^{10}\text{B}_{\text{standard}}} \right) - 1 \right) \times 1000 \quad (4.1)$$

, where ${}^{11}\text{B}/{}^{10}\text{B}_{\text{standard}} = 4.04367$ is the boron isotopic composition of National Institute of Standards and Technology Standard Reference Material 951 boric acid [Catanzaro et al. (1970)]. We assume that any proxy observation is paired with a given age. Unevenly gathered observations are adjusted to the grid of every 1 kiloyear. Here we limit the inference up to 260 kiloyears ago, where both benthic $\delta^{18}\text{O}$ and $\delta^{11}\text{B}$ observations in the cores exist.

Whereas the first model assumes that the ice volume is another hidden state to infer, it is given a priori as the values derived from the sea level data in [Spratt & Lisiecki (2016)] and the following relationship[†]:

$$\Delta V \cdot d_{\text{ice}} \equiv -\Delta b \cdot A \cdot d_{\text{seawater}} \quad (4.2)$$

, where V is the ice volume, $A = 3.618 \times 10^8 (\text{km}^2)$ is how much area the oceans cover, b is the height, and densities of ice and sea water are $d_{\text{ice}} = 0.9167 (\text{Gt}/\text{km}^3)$ and $d_{\text{seawater}} = 1.027 (\text{Gt}/\text{km}^3)$, respectively.

The same emission model of $\delta^{11}\text{B}$ given atmospheric CO_2 is used, whereas that of ice volume given atmospheric CO_2 is pre-trained by the paired values of the given ice volumes and given CO_2 of the Antarctic ice core records: this aspect makes the inference of this model preliminary. Here we try four kernels: squared-exponential (SE), Ornstein-Uhlenbeck (OU) and two other Matérn kernels with $d = 3/2$ and $5/2$ (section B.1.1) that correspond to the continuous autoregressive (AR) models $\text{AR}(\infty)$, $\text{AR}(1)$, $\text{AR}(2)$ and $\text{AR}(3)$, respectively. The inference is done up to 800 kiloyears ago.

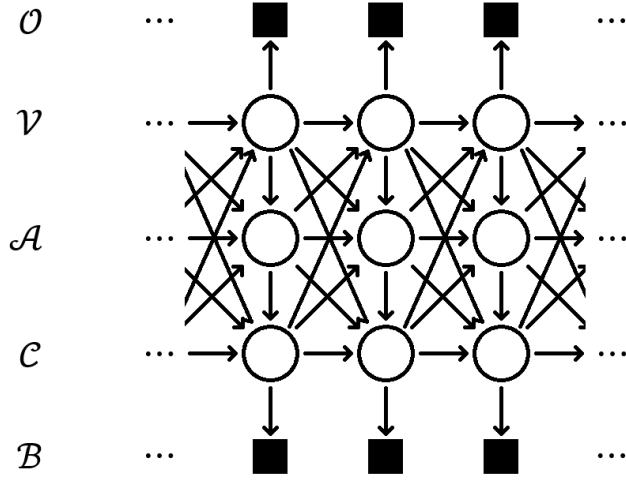


Figure 4.1: A graphical representation of our particle smoother for the CO_2 reconstruction.

4.3.1 PARTICLE SMOOTHER

First of all, we define the notations as follows: a graphical representation of our particle smoother is in figure 4.1.

- $\mathcal{V} = \{\mathbf{V}_t\}_{t=1}^T$: dimensionless indices of ice volumes.
- $\mathcal{A} = \{\mathbf{A}_t\}_{t=1}^T$: dimensionless indices of extents of the Antarctic ice sheet.
- $\mathcal{C} = \{\mathbf{C}_t\}_{t=1}^T$: dimensionless indices of atmospheric CO_2 .
- $\mathcal{S} = \{\mathbf{S}_t\}_{t=1}^T$: tuples of above dimensionless indices, where each $\mathbf{S}_t \triangleq (\mathbf{V}_t, \mathbf{A}_t, \mathbf{C}_t)^T$.
- $\mathcal{O} = \{\mathbf{O}^{(m)}\}_{m=1}^{M_{\mathcal{O}}}$: set of observed $\delta^{18}\text{O}$ proxies, where each $\mathbf{O}^{(m)} = \{\mathbf{O}_t^{(m)}\}_{t=1}^T$ from core m is the set of either interpolated values to be paired with $\mathcal{S}$ or $\emptyset$ (we do not extrapolate the data).
- $\mathcal{B} = \{\mathbf{B}^{(m)}\}_{m=1}^{M_{\mathcal{B}}}$: set of observed $\delta^{11}\text{B}$ proxies, where each $\mathbf{B}^{(m)} = \{\mathbf{B}_t^{(m)}\}_{t=1}^T$ from core m is the set of either interpolated values to be paired with $\mathcal{S}$ or $\emptyset$ (we do not extrapolate the data).

[†]<http://www.antarcticglaciers.org/glaciers-and-climate/estimating-glacier-contribution-to-sea-level-rise/>

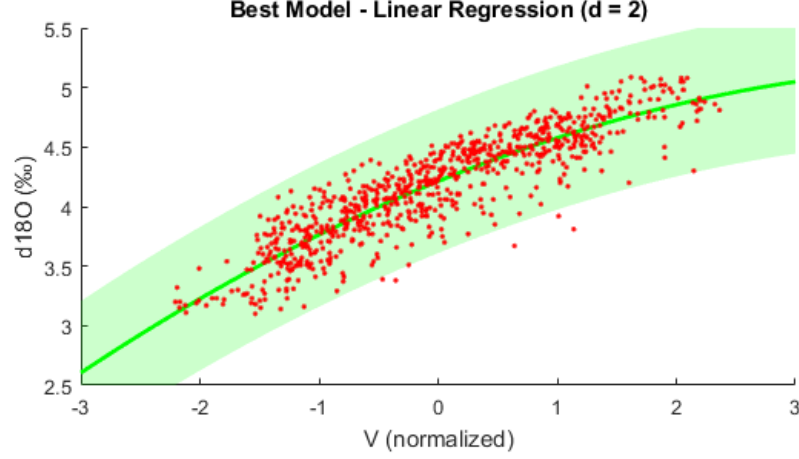


Figure 4.2: Ice volume and $\delta^{18}\text{O}$ pairs (red) and the chosen model with the 95% confidence band (green).

From the datasets, we first learn the following parametric observational model of $\delta^{18}\text{O}$ (O) given the dimensionless indices of ice volume (V):

$$p_{\varepsilon_V}(\text{O}|\text{V}) = \mathcal{N}(\text{O} | \beta_0 + \beta_1 V + \dots + \beta_d V^d, \exp(\gamma_0 + \gamma_1 V + \dots + \gamma_c V^c)) \quad (4.3)$$

Based on the Bayesian information criterion (BIC) [Schwarz (1978)], we choose the degree $d = 2$ and $c = 0$. The pre-trained emission model is given as follows: see figure 4.2 for its visualization with training data.

$$p_{\varepsilon_V}(\text{O}|\text{V}) = \mathcal{N}(\text{O} | 4.21 + 0.407V - 0.0429V^2, 0.197^2) \quad (4.4)$$

Also, we choose the observational model of $\delta^{11}\text{B}$ (B) given the dimensionless indices of atmospheric CO_2 (C) among the following two parametric models: these models with training samples were suggested and gathered by Andrew Kaluzny.

$$p_{\varepsilon_B}(\text{B}|\text{C}) = \mathcal{N}\left(\text{B} \middle| a - \sqrt{b + c \log(C + d)}, \exp(\delta_0 + \delta_1 C + \dots + \delta_q C^q)\right) \quad (4.5)$$

$$p_{\varepsilon_B}(\text{B}|\text{C}) = \mathcal{N}(\text{B} | a + bC + c \log(C + d), \exp(\delta_0 + \delta_1 C + \dots + \delta_q C^q))$$

By applying the BIC again, we choose the following model: see figure 4.3 for its visualization

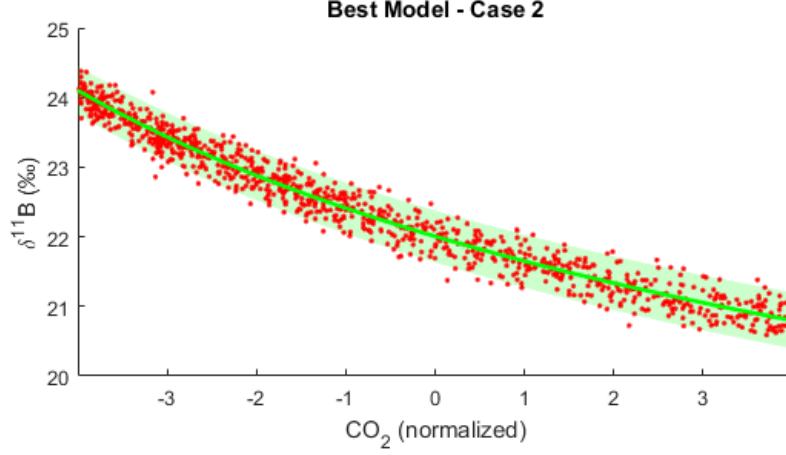


Figure 4.3: Atmospheric CO_2 and $\delta^{11}\text{B}$ pairs (red) and the chosen model with the 95% confidence band (green).

with training data.

$$p_{\mathcal{E}_B}(B|C) = \mathcal{N}(B|30.3 + 0.0552C - 3.80 \log(C + 8.76), \exp(-3.32 + 0.0392C)) \quad (4.6)$$

To apply the observational models (4.4) and (4.6) for the proxy observations, just as what we did in section 2.3.2, core-specific adjustments are necessary: we linearly adjust the models by core-specific coefficient parameters.

The emission model $p_e(O_t, B_t|S_t)$ is obtained as follows:

$$\begin{aligned} p_e(O_t^{(m)}, B_t^{(m)}|S_t) &\triangleq p_{\mathcal{E}_V}(O_t^{(m)}|V_t; \theta^{(m)})^{1_{\{O_t^{(m)} \neq \emptyset\}}} p_{\mathcal{E}_B}(B_t^{(m)}|C_t; \theta^{(m)})^{1_{\{B_t^{(m)} \neq \emptyset\}}} \\ p_e(O_t, B_t|S_t) &= \prod_{m=1}^M p_e(O_t^{(m)}, B_t^{(m)}|S_t) \end{aligned} \quad (4.7)$$

The transition model is originated by the LS model in [Garcia-Olivares & Herrero (2013)] that is defined as the system of the following first-order ODEs: $\tau_V^{(1)}$, $\tau_V^{(2)}$, τ_A , $\tau_C^{(1)}$ and $\tau_C^{(2)}$ are

parameters.

$$\begin{aligned}
\frac{dV}{dt} &= \frac{V_r - V}{\tau_V^{(1)} 1_{\{V < V_r\}} + \tau_V^{(2)} 1_{\{V \geq V_r\}}} \\
\frac{dA}{dt} &= \frac{-C - A}{\tau_A} \\
\frac{dC}{dt} &= \frac{C_r - C}{\tau_C^{(1)} 1_{\{C < C_r\}} + \tau_C^{(2)} 1_{\{C \geq C_r\}}}
\end{aligned} \tag{4.8}$$

, where, for the mean daily insolation I_{65} at 65°N on June 21st and parameters $x, y, z, \alpha, \beta, \gamma$ and δ ,

$$\begin{aligned}
V_r &\triangleq -xC - yI_{65} + z \\
C_r &\triangleq \alpha I_{65} - \beta V + \gamma 1_{\{bA + cC - d > 0\}} + \delta
\end{aligned} \tag{4.9}$$

By $dX/dt \approx (X_{n+1} - X_n)/(t_{n+1} - t_n)$, one can convert (4.8) to the probabilistic transition model as follows:

$$p_t(S_{t+1}|S_t, \zeta_t) \triangleq \mathcal{N}(S_{t+1} | W_{\zeta_t} S_t + B_{\zeta_t}, \Sigma) \tag{4.10}$$

, where Σ is a covariance matrix addressing errors and:

$$\begin{aligned}
W_{\zeta_t} &\triangleq \begin{pmatrix} 1 - \frac{\varepsilon}{\tau_V^{(\zeta_t)}} & 0 & -\frac{x\varepsilon}{\tau_V^{(\zeta_t)}} \\ 0 & 1 - \frac{\varepsilon}{\tau_A} & -\frac{\varepsilon}{\tau_A} \\ -\frac{\beta\varepsilon}{\tau_C^{(\zeta_t)}} & 0 & 1 - \frac{\varepsilon}{\tau_C^{(\zeta_t)}} \end{pmatrix}, \quad B_{\zeta_t} \triangleq \begin{pmatrix} -\frac{\gamma\varepsilon}{\tau_V^{(\zeta_t)}} I_t^{65} + \frac{z\varepsilon}{\tau_V^{(\zeta_t)}} \\ 0 \\ \frac{\alpha\varepsilon}{\tau_C^{(\zeta_t)}} I_t^{65} + \frac{\delta\varepsilon}{\tau_C^{(\zeta_t)}} + \frac{\gamma\varepsilon}{\tau_C^{(\zeta_t)}} \cdot 1_{\{bA_t + cC_t - d > 0\}} \end{pmatrix} \\
\zeta_t &\triangleq \begin{cases} 1, & V_t < -xC_t - yI_t^{65} + z & C_t < \alpha I_t^{65} - \beta V_t + \gamma 1_{\{bA_t + cC_t - d > 0\}} + \delta \\ 2, & V_t < -xC_t - yI_t^{65} + z & C_t \geq \alpha I_t^{65} - \beta V_t + \gamma 1_{\{bA_t + cC_t - d > 0\}} + \delta \\ 3, & V_t \geq -xC_t - yI_t^{65} + z & C_t < \alpha I_t^{65} - \beta V_t + \gamma 1_{\{bA_t + cC_t - d > 0\}} + \delta \\ 4, & V_t \geq -xC_t - yI_t^{65} + z & C_t \geq \alpha I_t^{65} - \beta V_t + \gamma 1_{\{bA_t + cC_t - d > 0\}} + \delta \end{cases} \\
\tau_V^{(\zeta_t)} &\triangleq \begin{cases} \tau_V^{(1)}, & \zeta_t = 1 \text{ or } 2 \\ \tau_V^{(2)}, & \zeta_t = 3 \text{ or } 4 \end{cases} \\
\tau_C^{(\zeta_t)} &\triangleq \begin{cases} \tau_C^{(1)}, & \zeta_t = 1 \text{ or } 3 \\ \tau_C^{(2)}, & \zeta_t = 2 \text{ or } 4 \end{cases}
\end{aligned} \tag{4.11}$$

Note that ζ_t is defined by the given parameters, I_t^{65} and S_t .

The particle smoother in section A.3.2 is completely specified by its transition (4.11) and emission (4.7) models. Suppose that we have drawn a bunch of sampled paths $\left\{\tilde{\mathcal{S}}^{(l)}\right\}_{l=1}^L = \left\{\tilde{\mathcal{V}}^{(l)}, \tilde{\mathcal{A}}^{(l)}, \tilde{\mathcal{C}}^{(l)}\right\}_{l=1}^L$. On the one hand, the covariance matrix Σ is easily updated by the following form:

$$\hat{\Sigma} = \frac{1}{L(T-1)} \sum_{l=1}^L \sum_{t=1}^{T-1} \sum_{k=1}^4 \left(\tilde{S}_{t+1}^{(n)} - \mathbf{W}_k \tilde{S}_t^{(n)} - \mathbf{B}_k \right) \left(\tilde{S}_{t+1}^{(n)} - \mathbf{W}_k \tilde{S}_t^{(n)} - \mathbf{B}_k \right)^{\mathbb{T}} \mathbf{1}_{\left\{\tilde{\zeta}_t^{(n)}=k\right\}} \quad (4.12)$$

For learning parameters, on the other hand, we have the following objective function:

$$f(\tilde{\mathcal{S}}; \Theta) \triangleq -\frac{1}{2} \sum_{n=1}^N \sum_{t=1}^{T-1} \sum_{k=1}^4 \left(\tilde{S}_{t+1}^{(n)} - \mathbf{W}_k \tilde{S}_t^{(n)} - \mathbf{B}_k \right)^{\mathbb{T}} \hat{\Sigma}^{-1} \left(\tilde{S}_{t+1}^{(n)} - \mathbf{W}_k \tilde{S}_t^{(n)} - \mathbf{B}_k \right) \mathbf{1}_{\left\{\tilde{\zeta}_t^{(n)}=k\right\}} \quad (4.13)$$

We do not have any closed form for updating any of the parameters with (4.13): instead, we can adopt gradient descent methods, after regularizing the only nondifferentiable term: for a large $\lambda > 0$ (we use $\lambda = 1000$),

$$\mathbf{1}_{\{bA_t + cC_t - d > 0\}} \approx \frac{1}{1 + \exp(-b\lambda A_t - c\lambda C_t + d\lambda)} \quad (4.14)$$

Note that each partial derivative of (4.13) is given as follows:

$$\frac{\partial f}{\partial \theta}(\tilde{\mathcal{S}}; \Theta) = \sum_{n=1}^N \sum_{t=1}^{T-1} \sum_{k=1}^4 \left(\tilde{S}_{t+1}^{(n)} - \mathbf{W}_k \tilde{S}_t^{(n)} - \mathbf{B}_k \right)^{\mathbb{T}} \hat{\Sigma}^{-1} \left(\frac{\partial \mathbf{W}_k}{\partial \theta} \tilde{S}_t^{(n)} + \frac{\partial \mathbf{B}_k}{\partial \theta} \right) \mathbf{1}_{\left\{\tilde{\zeta}_t^{(n)}=k\right\}} \quad (4.15)$$

By plugging the approximation (4.14) in the objective function (4.13), parameters are updated.

Now we are prepared. Algorithm 4 summarizes the inference procedure.

4.3.2 GAUSSIAN PROCESS STATE-SPACE MODEL

Although we already covered the GPST briefly in section B.4, here we give a more detailed inference procedure. Let us first define the notations as follows:

Algorithm 4: Particle Smoother Atmospheric CO₂ Reconstruction

```

1 initialize parameters and covariance matrix
2 while converging do
3   | sample dimensionless index paths by the particle smoother algorithm
4   | update the covariance matrix given parameters and sampled paths
5   | update the parameters given covariance matrix and sampled paths
6 end
7 restore atmospheric CO2 from the associated dimensionless indices
8 return a set of sampled paths of CO2

```

- $\mathcal{T} = T_{1:N}$: ages of observations.
- $\mathcal{X} = X_{1:N}$: atmospheric CO₂ at $\mathcal{T}$ to infer.
- $\mathcal{Y} = Y_{1:N}$: $\delta^{11}\text{B}$ at $\mathcal{T}$.
- $\mathcal{V} = V_{1:N}$: ice volumes at $\mathcal{T}$ somehow correlated with $\mathcal{X}$.

Then, we have the following models on them. First, what we model as a GP is the latent state $\mathcal{X}$ over $\mathcal{T}$: recall that Matérn kernels are corresponding to the continuous AR models (sections B.1.1 and B.4).

$$p(\mathcal{X}|\mathcal{T}) = \mathcal{N}\left(\mathcal{X} \middle| \underline{\mu}_{\mathcal{T}}, \mathbb{K}_{\mathcal{T}\mathcal{T}}\right) \quad (4.16)$$

Secondly, to deal with outliers, we consider a Student's t-distribution with 6 as the degree of freedom for the emission model of $\delta^{11}\text{B}$ given $\mathcal{X}$ and $\mathcal{T}$: the mean α and variance τ are the same as those of (4.6).

$$\begin{aligned}
p(\mathcal{Y}|\mathcal{X}, \mathcal{T}) &= \prod_{n=1}^N \mathcal{T}(Y_n | \alpha(X_n), \tau(X_n), 6) \\
\alpha(X_n) &= 30.3 + 0.0552X_n - 3.80 \log(X_n + 8.76) \\
\tau(X_n) &= \exp(-3.32 + 0.0392X_n)
\end{aligned} \quad (4.17)$$

Lastly, a Gaussian mixture model [Murphy (2012)] is considered for the emission model of ice

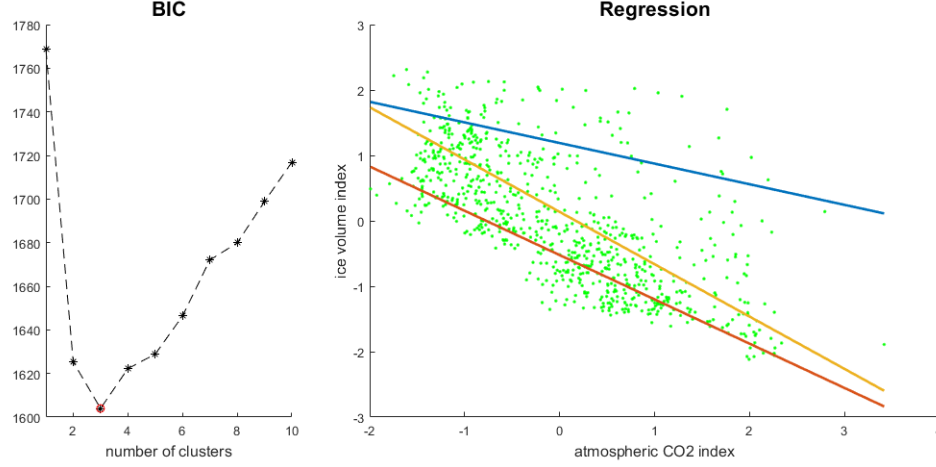


Figure 4.4: (Left) The BIC curve that selects 3 as the number of clusters and (Right) atmospheric CO₂ and ice volume index pairs (green) and the regressions of the clusters (in each color).

volumes given $\mathcal{X}$: figure 4.4 visualizes the training data and learned model.

$$p(\mathcal{V}|\mathcal{X}) = \prod_{n=1}^N \sum_{k=1}^K p_k \mathcal{N}(V_n | a_k X_n + b_k, \omega_k^2) \quad (4.18)$$

In the GP regression model, the goal is to obtain the marginalized distribution of the proxy given the associated input. Here, the model aims to infer $\mathcal{X}$ given $\mathcal{T}$, $\mathcal{V}$ and $\mathcal{Y}$. Note that the likelihoods (4.17) and (4.18) are not Gaussian, thus the posterior of $\mathcal{X}$ cannot be given as a known distribution.

A plausible approach is sampling $\mathcal{X}$ from the posterior $p(\mathcal{X}|\mathcal{T}, \mathcal{V}, \mathcal{Y})$ which is proportional to $p(\mathcal{Y}|\mathcal{X}, \mathcal{T})p(\mathcal{V}|\mathcal{X})p(\mathcal{X}|\mathcal{T})$ by a Markov chain Monte Carlo (MCMC) (section A.4). It is, however, limited because 1) $\mathcal{X}$ tends to be highly correlated one another so it would require a lot of iteration until reaching a burn-in period and 2) computing $p(\mathcal{X}|\mathcal{T})$ (or $p(X_n|\mathcal{X} \setminus X_n, \mathcal{T})$) becomes intractable quickly as the size of data N grows.

Instead, we consider the following variational method. Consider $q(\mathcal{X}|\Theta)$ as an approximation of $p(\mathcal{X}|\mathcal{T}, \mathcal{V}, \mathcal{Y})$ defined as follows: for $\mu = (\mu_1, \dots, \mu_N)^\top$ and $\Psi = \text{diag}(\psi_1^2, \dots, \psi_N^2)$,

$$q(\mathcal{X}|\Theta) = \mathcal{N}(\mathcal{X}|\mu, \Psi) = \prod_{n=1}^N \mathcal{N}(X_n | \mu_n, \psi_n^2) \quad (4.19)$$

A subsequent issue is about how to estimate $\Theta = \theta_{1:N} = \{\mu_{1:N}, \psi_{1:N}\}$. Let us consider the following Kullback-Leibler (KL) divergence:

$$\begin{aligned}
& \mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T}, \mathcal{V}, \mathcal{Y})) \\
&= \log p(\mathcal{V}, \mathcal{Y}|\mathcal{T}) + \mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T})) - \mathbb{E}_q[\log p(\mathcal{Y}|\mathcal{X}, \mathcal{T}) + \log p(\mathcal{V}|\mathcal{X})] \\
&= \log p(\mathcal{V}, \mathcal{Y}|\mathcal{T}) + \mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T})) - \sum_{n=1}^N \mathbb{E}_q[\log \mathcal{T}(Y_n|\alpha(X_n), \tau(X_n), \mathbf{G})] \\
&\quad - \sum_{n=1}^N \mathbb{E}_q \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n|a_k X_n + b_k, \omega_k^2) \right]
\end{aligned} \tag{4.20}$$

Since KL divergences are nonnegative, we have the following inequality giving the evidence lower bound (ELBO) to the log-marginal likelihood (log-ML) $\log p(\mathcal{V}, \mathcal{Y}|\mathcal{T})$:

$$\begin{aligned}
\log p(\mathcal{V}, \mathcal{Y}|\mathcal{T}) &\geq -\mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T})) + \sum_{n=1}^N \mathbb{E}_q[\log \mathcal{T}(Y_n|\alpha(X_n), \tau(X_n), \mathbf{G})] \\
&\quad + \sum_{n=1}^N \mathbb{E}_q \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n|a_k X_n + b_k, \omega_k^2) \right]
\end{aligned} \tag{4.21}$$

Because both $p(\mathcal{X}|\mathcal{T})$ and $q(\mathcal{X}|\Theta)$ are Gaussian, we can represent $\mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T}))$ in a closed form as follows:

$$\begin{aligned}
& \mathbb{D}_{\text{KL}}(q(\cdot|\Theta)||p(\cdot|\mathcal{T})) \\
&= \frac{1}{2} \left(\text{trace}(\mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} \Psi) + (\mu - \underline{\mu}_{\mathcal{T}})^{\top} \mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} (\mu - \underline{\mu}_{\mathcal{T}}) - 1 + \log \frac{|\mathbb{K}_{\mathcal{T}\mathcal{T}}|}{|\Psi|} \right)
\end{aligned} \tag{4.22}$$

The last term cannot be explicitly computed. One possible treatment is to use the following lower bound holds for each n by Jensen's inequality:

$$\begin{aligned}
& \mathbb{E}_q \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n|a_k X_n + b_k, \omega_k^2) \right] \geq \mathbb{E}_q \left[\sum_{k=1}^K p_k \log \mathcal{N}(V_n|a_k X_n + b_k, \omega_k^2) \right] \\
&= \sum_{k=1}^K p_k \left(\log \mathcal{N}(V_n|a_k \mu_n + b_k, \omega_k^2) - \frac{1}{2} a_k^2 \omega_k^{-2} \psi_n^2 \right)
\end{aligned} \tag{4.23}$$

However, it turns out that the gap in (4.23) is too big to guarantee a good variational approximation after running the real data. Therefore, a Monte Carlo integration should be considered instead. To reduce the variance, let us consider the reparameterization trick, as discussed in sections B.2.3 and B.4: for each n , let $\varepsilon_n \sim \mathcal{N}(0, 1^2)$. Then, the expectation can be represented in the following form:

$$\begin{aligned} & \mathbb{E}_q \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n | a_k X_n + b_k, \omega_k^2) \right] \\ &= \mathbb{E}_{\varepsilon_n} \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n | a_k (\mu_n + \psi_n \varepsilon_n) + b_k, \omega_k^2) \right] \end{aligned} \quad (4.24)$$

Therefore, its gradient with respect to θ_n is given as follows: in practice, we approximate it with the Monte Carlo integration.

$$\begin{aligned} & \nabla_{\theta_n} \mathbb{E}_q \left[\log \sum_{k=1}^K p_k \mathcal{N}(V_n | a_k X_n + b_k, \omega_k^2) \right] \\ &= \mathbb{E}_{\varepsilon_n} \left[\nabla_{\theta_n} \log \sum_{k=1}^K p_k \mathcal{N}(V_n | a_k (\mu_n + \psi_n \varepsilon_n) + b_k, \omega_k^2) \right] \end{aligned} \quad (4.25)$$

It is also not possible to express each $\mathbb{E}_q [\log \mathcal{T}(Y_n | \alpha(X_n), \tau(X_n), \mathcal{G})]$ explicitly. As similar with above, we have the following:

$$\begin{aligned} & \mathbb{E}_q [\log \mathcal{T}(Y_n | \alpha(X_n), \tau(X_n), \mathcal{G})] = \mathbb{E}_{\varepsilon_n} [\log \mathcal{T}(Y_n | \alpha(\mu_n + \psi_n \varepsilon_n), \tau(\mu_n + \psi_n \varepsilon_n), \mathcal{G})] \\ & \nabla_{\theta_n} \mathbb{E}_q [\log \mathcal{T}(Y_n | \alpha(X_n), \tau(X_n), \mathcal{G})] = \mathbb{E}_{\varepsilon_n} [\nabla_{\theta_n} \log \mathcal{T}(Y_n | \alpha(\mu_n + \psi_n \varepsilon_n), \tau(\mu_n + \psi_n \varepsilon_n), \mathcal{G})] \end{aligned} \quad (4.26)$$

Therefore, we can run a gradient ascend on the lower bound in (4.21) to variationally maximize $p(\mathcal{V}, \mathcal{Y} | \mathcal{T})$ under mild regularity conditions, including that α and τ are partially differentiable with respect to the first coordinates. Note that it also becomes intractable as N grows because (4.22) is involved in computing $\mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1}$. This can be dealt with by using small batches instead of the full data in feeding in the gradient ascent or considering another variational approximation on $p(\mathcal{X} | \mathcal{T}, \mathcal{V}, \mathcal{Y})$.

Suppose that Θ and $\mathbb{K}$ are estimated by so. Then, for an arbitrary age t , the associated atmospheric CO_2 x has the following posterior predictive:

$$\begin{aligned}
p(x|t, \mathcal{T}, \mathcal{Y}, \mathcal{V}) &= \int p(x|t, \mathcal{X}, \mathcal{T}) p(\mathcal{X}|\mathcal{T}, \mathcal{Y}, \mathcal{V}) d\mathcal{X} \\
&\approx \int p(x|t, \mathcal{X}, \mathcal{T}) q(\mathcal{X}|\Theta) d\mathcal{X} = \mathcal{N}(x|\rho(t), \delta(t)) \\
\rho(t) &= \mathbb{K}_{t\mathcal{T}} \mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} (\mu - \underline{\mu}_{\mathcal{T}}) + \underline{\mu}_t \\
\delta(t) &= \mathbb{K}_{tt} - \mathbb{K}_{t\mathcal{T}} \mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} \mathbb{K}_{\mathcal{T}t} + \mathbb{K}_{t\mathcal{T}} \mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} \Psi \mathbb{K}_{\mathcal{T}\mathcal{T}}^{-1} \mathbb{K}_{\mathcal{T}t}
\end{aligned} \tag{4.27}$$

However, we can do slightly better than this if ice volumes are given at any ages. Let v be that value at t . Then, the density of v given x is $\sum_{k=1}^K p_k \mathcal{N}(v|a_k x + b_k, \omega_k^2)$, which can be plugged in as a likelihood, so we have the following posterior of x , with (4.27) as the prior:

$$\begin{aligned}
p(x|t, v, \mathcal{T}, \mathcal{Y}, \mathcal{V}) &= \sum_{k=1}^K \bar{p}_k(t, v) \mathcal{N}(x|\bar{\mu}_k(t, v), \bar{v}_k(t, v)) \\
&\propto p(x|t, \mathcal{T}, \mathcal{Y}, \mathcal{V}) \sum_{k=1}^K p_k \mathcal{N}(v|a_k x + b_k, \omega_k^2) \\
\bar{p}_k(t, v) &= \frac{p_k \mathcal{N}(v|a_k \rho(t) + b_k, \omega_k^2 + a_k^2 \delta(t))}{\sum_{l=1}^K p_l \mathcal{N}(v|a_l \rho(t) + b_l, \omega_l^2 + a_l^2 \delta(t))} \\
\bar{\mu}_k(t, v) &= (\delta(t)^{-1} + a_k^2 \omega_k^{-2})^{-1} (a_k \omega_k^{-2} (v - b_k) + \delta(t)^{-1} \rho(t)) \\
\bar{v}_k(t, v) &= (\delta(t)^{-1} + a_k^2 \omega_k^{-2})^{-1}
\end{aligned} \tag{4.28}$$

To get the confidence bands with (4.28), note that, for any η , we have the following:

$$\begin{aligned}
\mathbb{P}(x \geq \eta|t, v, \mathcal{T}, \mathcal{Y}, \mathcal{V}) &= \sum_{k=1}^K \bar{p}_k(t, v) \left(1 - \Phi \left(\frac{\eta - \bar{\mu}_k(t, v)}{\sqrt{\bar{v}_k(t, v)}} \right) \right) \\
\mathbb{P}(x < \eta|t, v, \mathcal{T}, \mathcal{Y}, \mathcal{V}) &= \sum_{k=1}^K \bar{p}_k(t, v) \Phi \left(\frac{\eta - \bar{\mu}_k(t, v)}{\sqrt{\bar{v}_k(t, v)}} \right)
\end{aligned} \tag{4.29}$$

Now we just need to get numerical η 's to make the upper and lower term in (4.29) $\alpha/2$ for obtaining the 100 $(1 - \alpha)\%$ confidence level.

Now we are again prepared. Algorithm 5 summarizes the inference procedure.

Algorithm 5: Gaussian Process State-space Model for Atmospheric CO₂ Reconstruction

```

1 initialize kernel hyperparameters and variational parameters
2 while converging do
3   | tune kernel hyperparameters
4   | learn variational parameters
5 end
6 compute and return the posterior predictive mean and variance functions

```

4.4 RESULTS AND DISCUSSION

We tried two models for reconstructing atmospheric CO₂, one is the particle smoother and the other is the GPST. Each set of results are visualized in the corresponding figure that has three features. The black dashed curve and shaded region are that of reconstructed atmospheric CO₂ (the median of the sampled paths) and its 95% confidence band, respectively. The red curve indicates the “true” atmospheric CO₂ (experimental time series) from the Antarctic ice core records. The blue bars are the lower-upper intervals of the published CO₂ of [Dyez et al. (2018)].

4.4.1 PARTICLE SMOOTHER

We set the number of particles to be 200 and ran 100 iterations, where only the covariance matrix was supposed to be learned for the first 30 iterations. The learned covariance matrix and parameters are given in (4.30) and figure 4.5.

$$\hat{\Sigma} = \begin{pmatrix} 0.0464 & -0.0322 & -0.0298 \\ -0.0322 & 1.7145 & 0.0511 \\ -0.0298 & 0.0511 & 0.1850 \end{pmatrix} \quad (4.30)$$

Figure 4.6 shows three panels, one for each kind of hidden states. It is clear that the 95% confidence band of the reconstructed extent of the Antarctic ice sheet indices is far broader than

	Initial Value	Learned Parameter
α	0.237	-0.0618
β	0.793	1.0176
γ	1.955	0.0072
δ	0.146	1.1894
x	0.669	0.4840
y	0.527	0.8121
z	0.761	0.7416
b	0.936	0.8466
c	0.533	0.6288
d	0.069	-0.0779
$\tau_C^{(1)}$	2.793	11.5089
$\tau_C^{(2)}$	8.414	1.6338
τ_A	10.266	1.9672
$\tau_V^{(1)}$	11.425	10.1024
$\tau_V^{(2)}$	2.325	1.6374

Figure 4.5: Initial and learned values of the parameters in the particle smoother model. The initial values were selected from [Garcia-Olivares & Herrero (2013)].

the other indices because neither benthic $\delta^{18}\text{O}$ or $\delta^{11}\text{B}$ proxy is paired with it. The same logic is applied to the narrow band of the reconstructed ice volume indices because benthic $\delta^{18}\text{O}$ observations are plentiful compared to $\delta^{11}\text{B}$. Compared to the lower-upper intervals of the published ones, our 95% confidence band of the reconstructed atmospheric CO_2 is narrower thus more certain: it makes sense because we used not only $\delta^{11}\text{B}$ but also benthic $\delta^{18}\text{O}$ for the inference. Both reconstructions mostly cover the experimental time series as the “true” atmospheric CO_2 curve from the Antarctic ice core records, but the deviations appeared at 190-230 kiloyears: this could imply the model misspecification because the ice volume indices are also deviated from its experimental time series at the same range of ages.

Note that all the results become plausible under the assumption that the transition model is correctly specified and stationary over time. In fact, such assumptions are too strong, as criticized by a few attendees in the 2nd Workshop on Cenozoic $p\text{CO}_2$ Reconstruction. It means that the learned parameters cannot be reused for reconstructing the atmospheric CO_2 in different age ranges. Also, to make the data fit into the model, we adjusted their ages to the grid of every 1 kiloyear, so that the information loss is structurally inevitable.

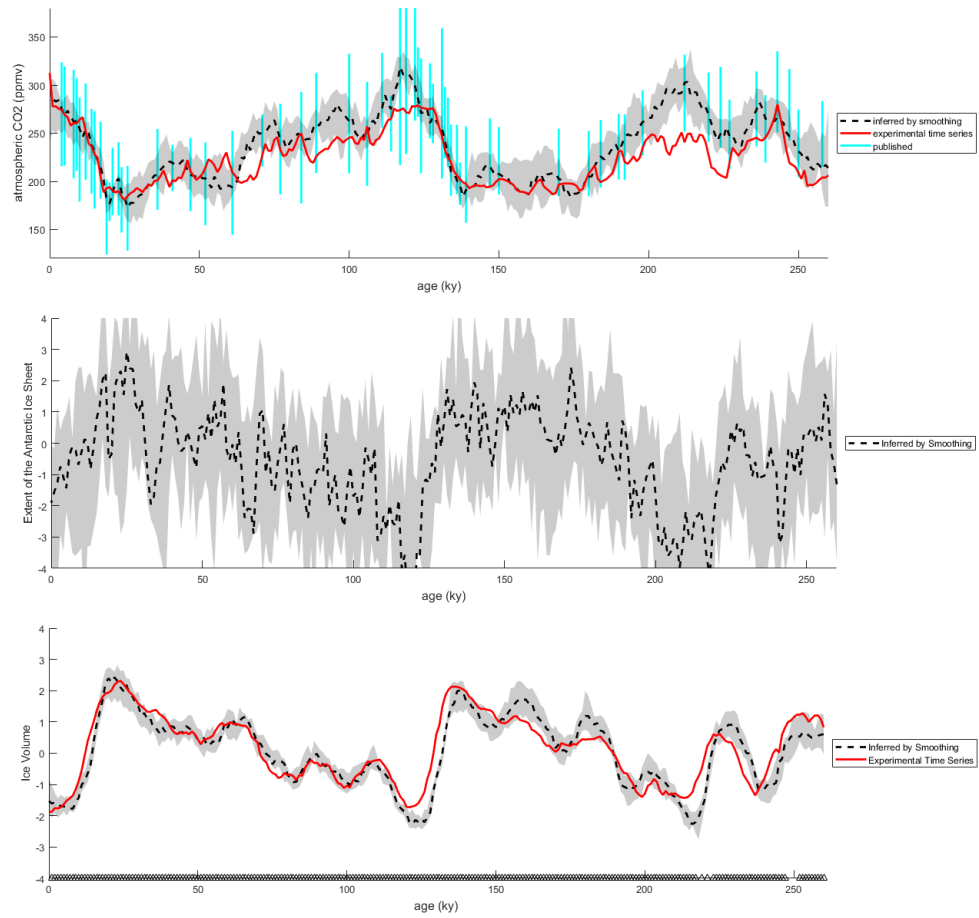


Figure 4.6: Results from the particle smoother. (Top) Reconstructed atmospheric CO₂ with the experimental time series and published lower-upper bands. (Middle) Reconstructed extent of the Antarctic ice sheet indices. (Bottom) Reconstructed ice volume indices with the experimental time series from [Spratt & Lisiecki (2016)].

4.4.2 GAUSSIAN PROCESS STATE-SPACE MODEL

Figures 4.7 to 4.10 visualize the results for the SE kernel, Matérn kernel with $d = 5/2$ (M25), Matérn kernel with $d = 3/2$ (M15) and OU kernel, respectively. Each figure consists of four panels: the first and second panel show the inference by the full model and the third and fourth panels present the results not using ice volumes as the proxy. Definitely the former gives narrower confidence bands to the latter because of more proxies. The second and fourth panels span up to 800 kiloyears while the first and third ones are up to 260 kiloyears to be coupled with the particle smoother results in section 4.4.1. The reconstructed atmospheric CO₂ series is now smoother and mostly covers the experimental time series, and the sensitiveness to the data is in order of SE, M25, M15 and OU, as expected by the kernel characteristics.

The nonparametric aspect of GPST means that the transition model is agnostic and dependent on the data directly. Also, it does not need regularly arranged data over ages. Thus, the model is now free from the strong assumptions on the transition model in the particle smoother (section 4.4.1). However, the inference is now not exact but variational. It is possible to overcome such aspect in particle smoother by increasing the number of particles, but GPST needs to adopt more general (complicated) parametric variational distributions for better approximation, that potentially leads to overfitting.

4.5 FUTURE WORK

As mentioned earlier, all the results are quite preliminary: each method adopts emission models that are constructed under very strong assumptions. Published ages and training data are given as true values and observational models do not reflect any particular data-specific characteristics but are just reduced to some simple parametric models. The dataset is also too scarce: there are less than 100 boron proxies up to 800 kiloyears. To overcome the hypothetical aspects stemming from such limits, the interdisciplinary research should be considered. With a help of the experts of the related fields, more reasonable transitional and observational models and more objective

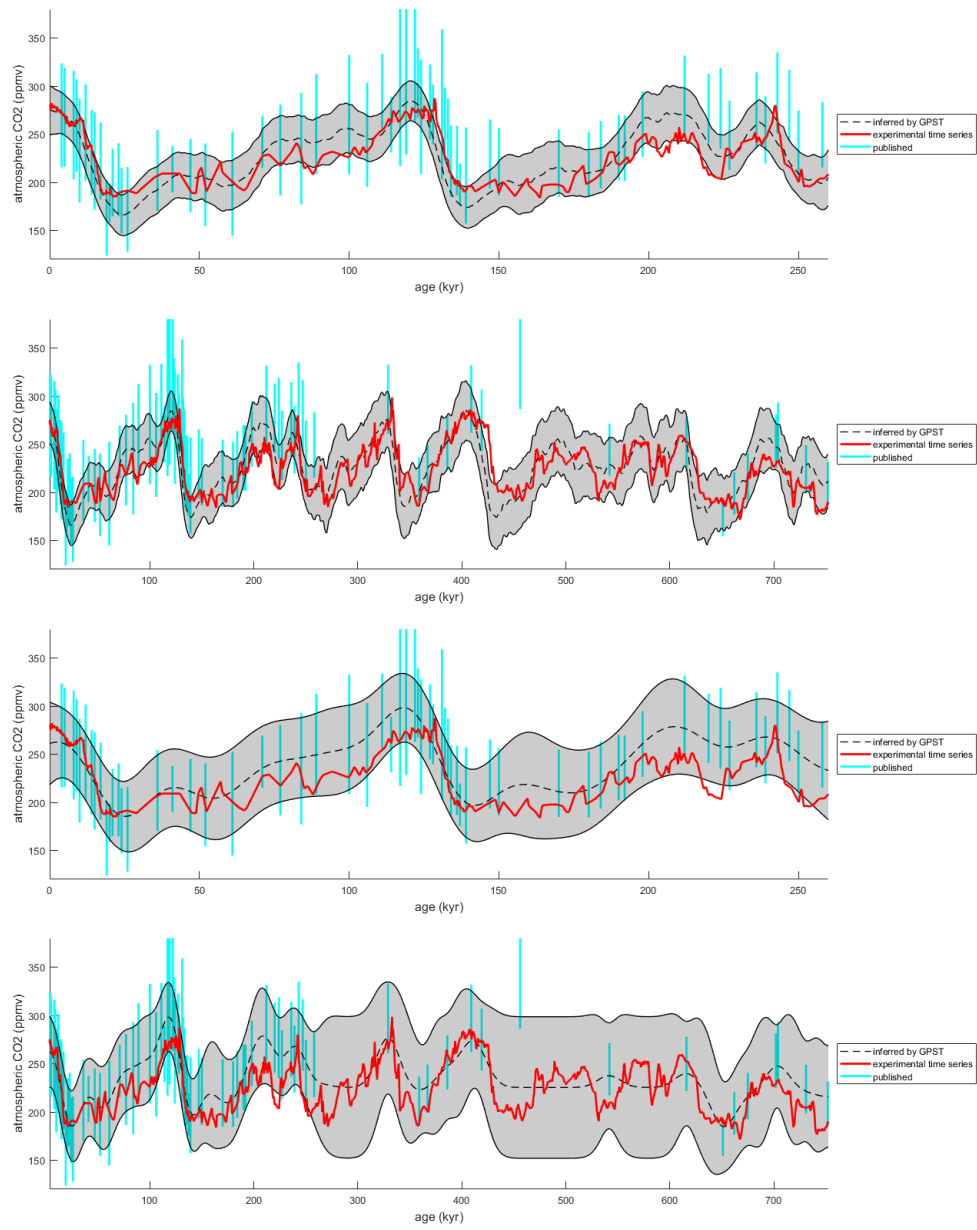


Figure 4.7: Results from the Gaussian process state-space model with the SE kernel.

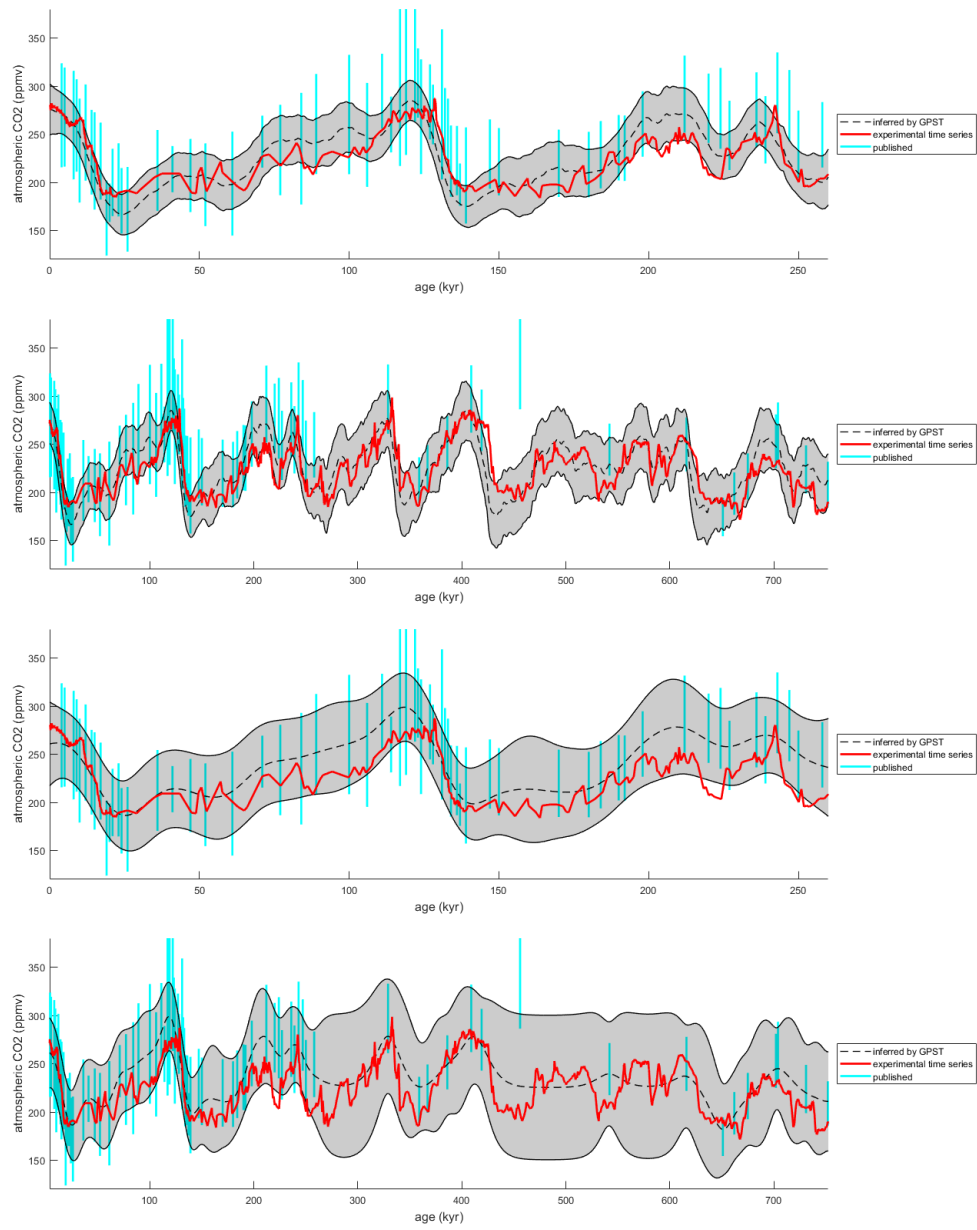


Figure 4.8: Results from the Gaussian process state-space model with the M25 kernel.

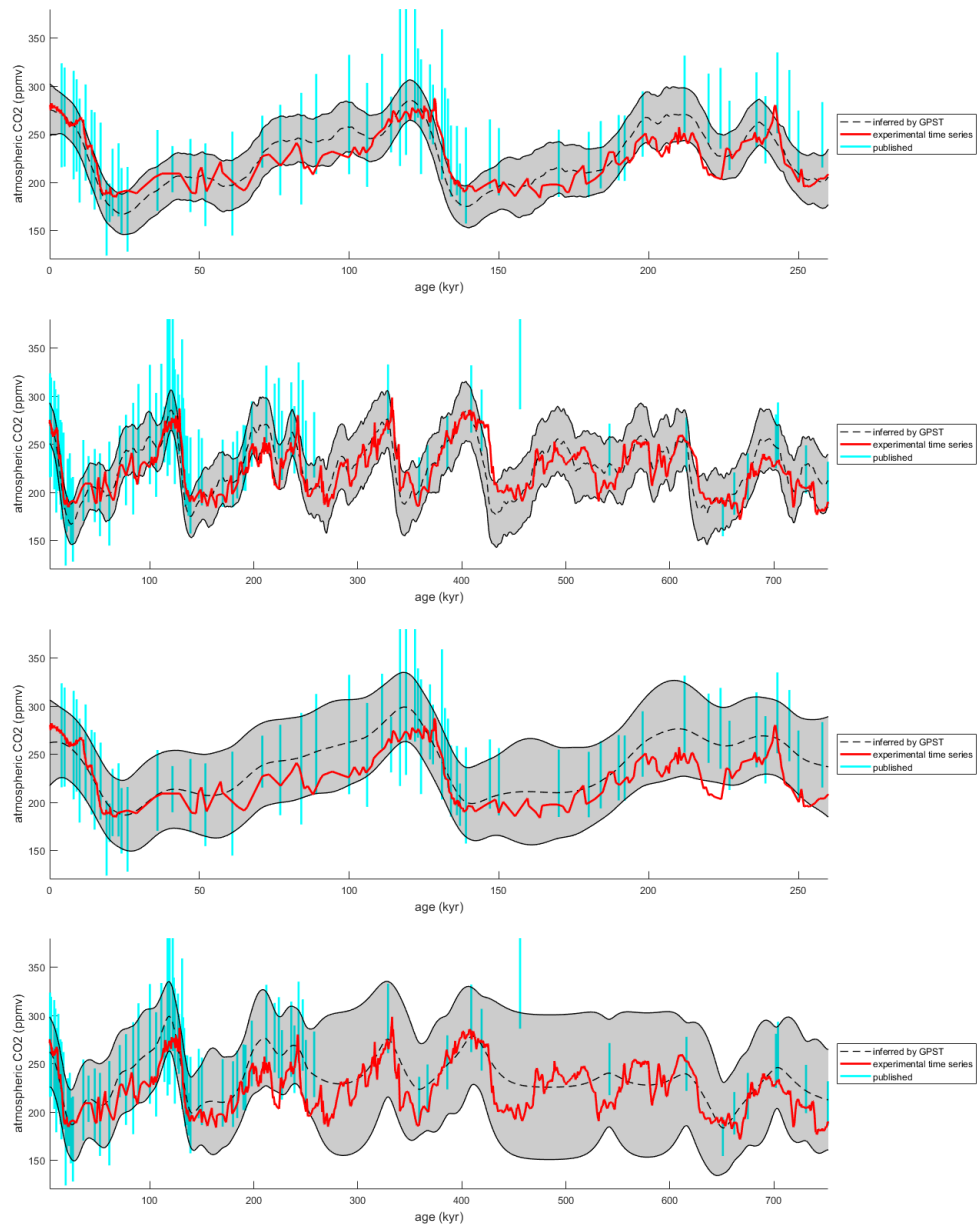


Figure 4.9: Results from the Gaussian process state-space model with the M15 kernel.

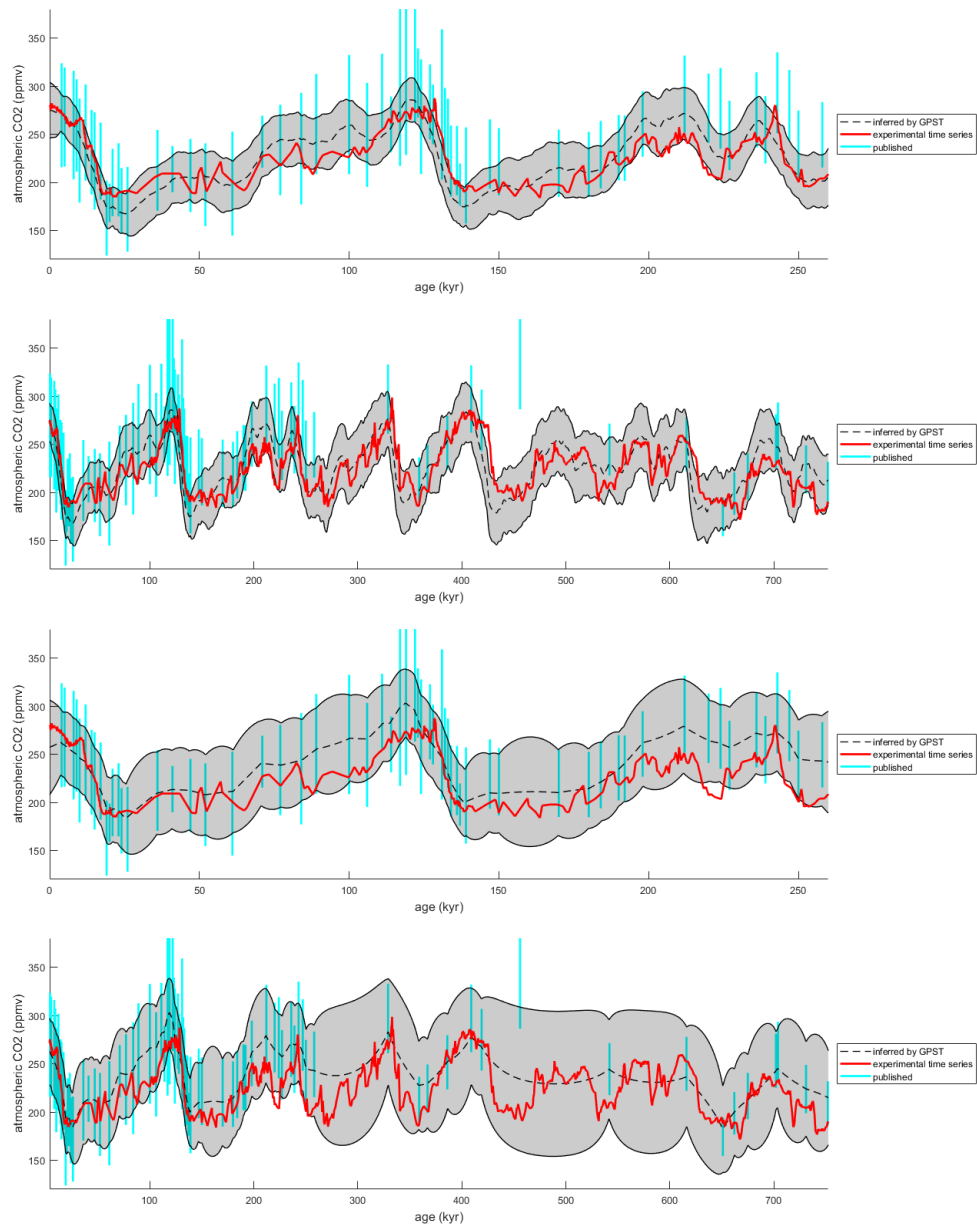


Figure 4.10: Results from the Gaussian process state-space model with the OU kernel.

training data will be available.

We tried four Matérn kernels in GPST modelling. As discussed in chapter B, the suboptimality can be overcome by adding multiple hidden layers (deep Gaussian process (GP) in section B.6) or adopting more complicated kernels. Though applying deep GP models to this problem is also intriguing, using weighted mixtures of kernels with periodic kernels (section B.1.1) would be an easy solution to deal with this problem, due to the apparently periodic pattern of the atmospheric CO₂ over time. More complicated transition models can be chosen in the particle smoother but the time complexity and curse of dimensionality on the particle smoother may preclude such attempts.

As discussed in chapter 2 about how the direct and indirect proxies are complementary, using various proxies that are correlated with the target ocean and climate events would be a good solution for overcoming the limited data availability. Figures 4.7 to 4.10 show the advantage of using multiple proxies vividly. Here we used benthic $\delta^{18}\text{O}$ and sea level data (that are transformed into the ice volumes) as such proxies, but did not use proxies for the other events, such as salinity and sea surface temperatures (SSTs), that also affect the CO₂ level. Constructing comprehensive models could be regarded as another future work, just as [Bowen et al. (2020)] did for reconstructing SSTs.

5

Conclusion

Here we presented some applications of the Bayesian statistical learning in connection with the fields of geoscience, paleoceanography and paleoclimatology. Various state-space models and Gaussian process models were chosen to utilize the proxies for reconstructing the relevant past ocean and climate events and for assigning ages to them. New alignment algorithm, improved calibration models, novel methodologies and promising preliminary works have been covered together with the results from the real data throughout the dissertation. The theoretical backgrounds of state-space models and Gaussian process models are extensively covered in Appendix.

My graduate research topics about paleoceanography and paleoclimatology were mainly covered through chapters 2 to 4. Chapter 2 is about the benthic $\delta^{18}\text{O}$ stack construction method,

a special case of SA-GPR algorithm for the multiple signal alignments, introduced in a submitted paper of mine, for inferring ages of cores. It is a motif-finding algorithm thus consisting of the core (signal) alignment algorithm and the stack (motif) construction algorithm. Each core is aligned to the target stack by a state-space model that is a hybrid of the particle smoother for initializing sampled age paths and Metropolis-Hastings algorithm for the refinements. The stack construction algorithm starts with utilizing the homoscedastic Gaussian process regression method that constructs a sample-specific model for each sampled age path. Then a core-specific stack is constructed from the affiliated sample-specific models by the model reduction method based on the moment-matching. Finally, a local stack is obtained by the variational method. This algorithm is for the cores in the same homogeneous set, which share the same local effect. In future, we expect to refine the model with differentiable transition models that allow the Hamiltonian Monte Carlo algorithm for efficient sampling, to extend it to the hierarchical stack construction method for handling multiple homogeneous sets simultaneously, and to add more age proxies for enhancing the age inference.

Chapter 3 covered the calibration model construction algorithms based on the Gaussian process regression, that extends the heteroscedastic Gaussian process regression model in a conference paper of mine. Both homoscedastic and heteroscedastic Gaussian process regression models with parametric mean prior functions were tried on the alkenone and TEX86 proxy data used for reconstructing sea surface temperatures individually. To deal with the large datasets, the variational extension based on the variational free energy was added to the exact models. For verifying the multimodality of the data, a novel multimodal Gaussian process regression algorithm that can be learned by the explicit expectation-maximization algorithm was designed and also applied for the datasets. In future, we will process the outliers more carefully by classifying them into two categories, one is the set of noninformative noises and the other is for the systematically contaminated data. For the latter case, we are working with some experts in the alkenone proxy for dealing with lab-specific biases. Constructing site-specific calibration models is another intriguing goal: the regression models are then depending not only on the sea surface temperatures but also on

the query location. Applying deep Gaussian process models and designing Bayesian inference models for reconstructing sea surface temperatures from the calibration models will follow.

Chapter 4 was consisting of two preliminary works for reconstructing atmospheric CO₂ by using the boron proxy. Two state-space models were considered: one was the particle smoother with a parametric transition model on multiple correlated climate events and the other was a nonparametric modelling, the Gaussian process state-space model. By plugging the pre-trained emission models in those models and using more proxies such as the benthic $\delta^{18}\text{O}$ and historical sea level, we obtained promising inference results that are mostly consistent with the “true” atmospheric CO₂ from the Antarctic ice core records. Four Matérn kernels that correspond to the continuous autoregressive models were tried. Due to the data scarcity, using boron proxy only did not result in a good inference with the Gaussian process state-space model. All inferences in this chapter were, however, limited due to the plain parametric observational models. More reliable and reasonable results will be available by collaborating with the experts in the $p\text{CO}_2$ community. Designing more detailed kernels for the Gaussian process state-space models is another way of increasing the opportunity. Clearly, extending the model to taking more proxies is helpful for the inference, no doubt.

Appendix A and B covered the base knowledge of probabilistic state-space models and Gaussian processes, respectively. The reasons why I organized such knowledge in separate chapters are not only that some tools were repeatedly appeared throughout my research topics but also that I wanted my dissertation to be regarded as a manual for the beginners who hope to know more about the statistical models that can be applied to their data which could be essentially similar to ours. Appendix A covered HMMs as the discrete state-space models, Kalman filter and smoother as the continuous models that can be inferred exactly, particle filter and smoother as the variational methods for more general settings, mixture Kalman filter as an example of the Rao-Blackwellized particle filter for reducing variances, and some of the popular Markov-chain Monte Carlo algorithms as the alternatives for the particle filter and smoother of which performance is dependent on the number of particles. Appendix B was designed for introducing Gaus-

sian process models. The nonparametric characteristic of Gaussian processes leads the affiliated continuous models more flexible than the parametric models, thus allows us to apply them for various purposes, such as regression, classification, state-space models and dimensionality reduction. Gaussian process regression models return the explicit posterior predictive distribution. Gaussian process classification methods reflect the density of data. Gaussian process state-space models can handle the inputs that are gathered irregularly without adjustments. The Gaussian process latent variable model is a nonparametric dimensionality reduction. Gaussian processes have two drawbacks: one is that the obtained models are usually suboptimal due to the misspecified kernels. The other is about the time complexity that precludes the exact inference quickly as the size of training data increases. These two weaknesses can be addressed by the deep Gaussian process models and variational inference such as variational free energy.

For the last decades, machine learning community has mainly focused on processing and handling large datasets with complicated parametric models such as deep neural networks relying on the exponential growth in computing power. These innovations are clearly helpful for the researches in certain areas where large data can be easily gathered or generated at cheap prices in the reality or laboratories. In contrast, some other fields, that are hard, or even impossible, to gather or generate data as much as being wanted, have been relatively alienated. Parametric models that are correctly specified by theory are often adopted for dealing with such kinds of datasets, but what if the dataset is out of any good model as well? State-space models and the Gaussian process are for dealing with such datasets. Both methods give certain structures to the hidden variables. While a state-space model depends on the parametric transition model, the Gaussian process assumes a nonparametric prior on the hidden variables. The inference is then done either exactly or variationally. These Bayesian methods lead to utilize the limited data more efficiently. One more advantage is that these tools also access to the large datasets as well. State-space models have been frequently adopted in geosciences, especially (nonlinear) geophysics, but the use of Gaussian process models is still limited. I hope that my dissertation will lead the geoscience community to pay more attention to the Gaussian process modelling.



State-space Models

Let us begin with a simple example of the Bayesian inference that consists of a set of data $\mathcal{Y} = \{Y_n\}_{n=1}^N$, parametric model p , a finite-dimensional parameter θ and their prior π . Under the assumption that the data Y_n 's are independent and identically distributed in (A.1), we have the following formulation:

$$p(\mathcal{Y}|\theta) = \prod_{n=1}^N p(Y_n|\theta) \quad (\text{A.1})$$

$$p(\theta|\mathcal{Y}) = \frac{\pi(\theta) p(\mathcal{Y}|\theta)}{p(\mathcal{Y})} \propto \pi(\theta) \prod_{n=1}^N p(Y_n|\theta) \quad (\text{A.2})$$

In other words, the posterior of θ given $\mathcal{Y}$ in (A.2) is proportional to the product of $p(Y_n|\theta)$'s

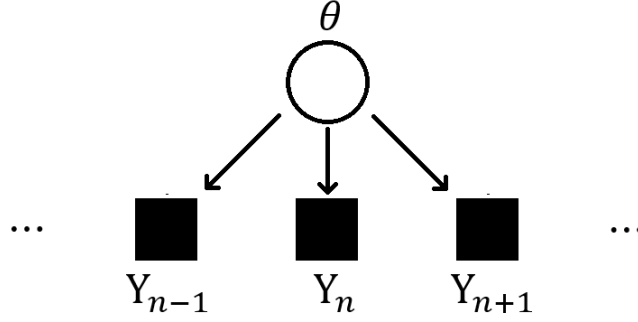


Figure A.1: A graphical representation of the typical Bayesian inference example. Here the observations Y_n 's are independent and identically distributed samples controlled by the parameter θ .

and the prior $\pi(\theta)$, thus intuitively it becomes more dependent on the data as the number of data N goes to infinity. Theoretically, there are consistency theorems such as the Doob's consistency theorem and Bernstein-von Mises theorem [[van der Vaart \(1998\)](#)] that give good asymptotic behaviors to (A.2) under regular conditions. Figure A.1 shows the graphical representation of this example.

Now let us further consider the following example under a different setting. We still have the same objects in the above example except a set of hidden variables $\mathcal{X} = \{X_n\}_{n=1}^N$ such that the parametric model is depending them instead of the parameter: (A.3) is often referred to as the conditional independence of $\mathcal{Y}$ given $\mathcal{X}$.

$$p(\mathcal{Y}|\mathcal{X}) = \prod_{n=1}^N p(Y_n|X_n) \quad (\text{A.3})$$

Suppose that the prior on $\mathcal{X}$ is of the naïve Bayes, i.e.,

$$p(\mathcal{X}) = \prod_{n=1}^N p(X_n) \quad (\text{A.4})$$

Then, the posterior distribution of $\mathcal{X}$ given $\mathcal{Y}$ is given as follows:

$$p(\mathcal{X}|\mathcal{Y}) = \frac{p(\mathcal{X})p(\mathcal{Y}|\mathcal{X})}{p(\mathcal{Y})} \propto p(\mathcal{X})p(\mathcal{Y}|\mathcal{X}) = \prod_{n=1}^N p(X_n)p(Y_n|X_n) \quad (\text{A.5})$$

(A.5) implies that each posterior of X_n is depending on only the associated observation Y_n , thus here we cannot utilize the consistency theorems and the inference is very sensitive to the choice of prior and potential outliers.

Some practical examples such as time series models give reasonable priors on $\mathcal{X}$ to utilize the dependency on X_n 's that makes each posterior distribution of X_n depend on the full $\mathcal{Y}$. If such a prior is also conjugate to the model, then the posterior is derived analytically. However, in general the posterior of $\mathcal{X}$ is not given in a closed form. The next best thing would be sampling the hidden variables from their posterior or doing a maximum a posteriori (MAP) estimation. Unfortunately, even such goals are usually not tractable straightforwardly because 1) $\mathcal{X}$ could be fully correlated (fully connected in the view of the graphical representation) one another over $p(\mathcal{X}|\mathcal{Y})$ thus Markov-chain Monte Carlo algorithm on it would not make sense in a tractable number of iterations and 2) Numerical methods for maximizing $p(\mathcal{X}|\mathcal{Y})$ could not work if either $\mathcal{X}$ is defined on a finite space or $p(\mathcal{X}|\mathcal{Y})$ is not differentiable.

The Markovian assumption on the prior $p(\mathcal{X})$ looks for the balance between the naïve Bayes and fully connected assumptions. It assumes that a subset of $\mathcal{X}$ is conditionally independent of those outside of their boundary given the boundary. This memoryless property could reduce the time complexity of the algorithm drastically: details can be found in [Barber (2012)]. If the data $\mathcal{Y}$ can be arranged over $\mathbb{R}$, time for instance, then $\mathcal{X}$ might be assumed as a Markov chain, i.e., it is a sequence of events where each of them depends only on the very previous event. The graphical representation of this case can be found in figure A.2.

Generally speaking, a state-space model (SSM) [Hangos et al. (2001)] is a mathematical model of a physical system consisting of a set of input, output and state variables that are connected by differential equations. State variables evolve through time and depend on the output variables.

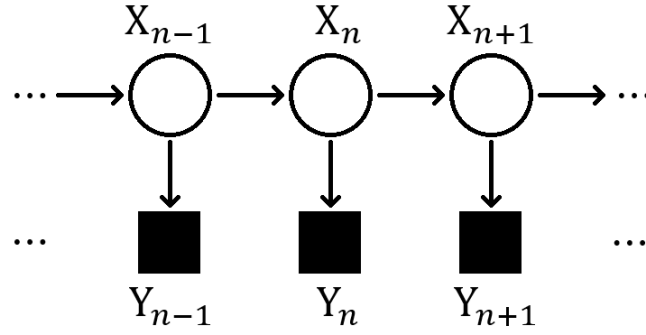


Figure A.2: A graphical representation of the typical example of the state-space model. Here the observations Y_n 's are conditionally independent given the hidden states X_n 's.

State-space models can be also understood as a class of probabilistic graphical model [Koller & Friedman (2009)] that describes the probabilistic dependence between the (latent) state variables and the output variables. These models are corresponding to the case where $\mathcal{X}$ is the set of latent state variables, $\mathcal{Y}$ is that of the output variables and the indices, say, $\mathcal{T} = \{T_n\}_{n=1}^N$, are input variables, in the above example.

This chapter briefly covers the probabilistic state-space models. All models in this chapter are involved with the Markovian assumption: section B.4 suggests a Gaussian process state-space model that does not take that assumption. Recurrent neural networks (RNNs) [Goodfellow et al. (2016)] are practically good choices in state-space modelling especially if data are plentiful, but this chapter does not deal with it.

A.1 HIDDEN MARKOV MODEL

The classification of state-space models is depending on the space in which each X_n is defined. Hidden Markov models (HMMs) [Durbin et al. (1998)] assume that each X_n is defined in a *finite* set while continuous state-space models that will be covered in section A.2 consider X_n to take continuous values. HMMs have been applied in the fields that are interested in finding latent states behind the observable sequences: gene prediction, protein folding, DNA motif discovery and chromatin state discovery are such examples in computational biology.

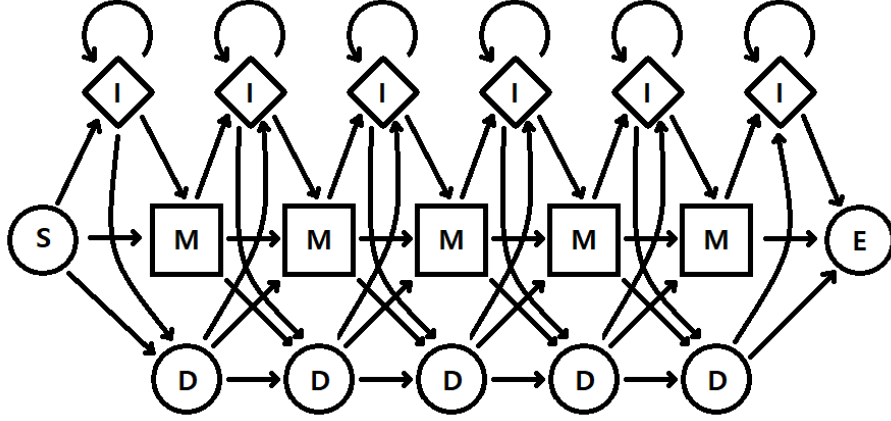


Figure A.3: Another type of the graphical representation of the hidden Markov model for the sequence alignment algorithm. Here, states S, E, I, D and M stand for start, end, insert, delete and match, respectively.

Before going further, let us first define the following notations:

- $\mathcal{Y} = \{Y_n\}_{n=1}^N$: observed sequence.
- $\mathcal{X} = \{X_n\}_{n=1}^N \subseteq \mathcal{S}$: hidden states defined in a finite set $\mathcal{S}$ that are associated with $\mathcal{Y}$.

An HMM is exactly specified by the transition and emission models. The transition model $p_t(X_{n+1}|X_n)$ determines how $\mathcal{X}$ is correlated a priori and the emission model $p_e(Y_n|X_n)$ gives the relationship between $\mathcal{Y}$ and $\mathcal{X}$. Note that those models might be defined nonstationary, i.e., depend on the index n , but here we abuse their notations as if they were stationary. Each X_n is to take values in a finite set, so the transition model can be represented as a transition matrix and the corresponding graphical representation focuses on the transition between candidates of the hidden states rather than hidden variables themselves, like figure A.3.

The transition and emission models give the full joint likelihood of the model as follows:

$$p(\mathcal{X}, \mathcal{Y}) = p(\mathcal{X})p(\mathcal{Y}|\mathcal{X}) = \prod_{n=1}^N p_t(X_n|X_{n-1}) \prod_{n=1}^N p_e(Y_n|X_n) \quad (\text{A.6})$$

The inference is involved with finding the MAP estimator of $\mathcal{X}$ and computing the posterior $p(\mathcal{X}|\mathcal{Y})$: both goals can be achieved exactly by the Viterbi algorithm (section A.1.1) and Forward-backward algorithm (section A.1.2), respectively. This HMM framework can be ex-

tended to a more general setting so that the transition model has an order more than 1 (e.g., $p_t(\mathbf{X}_{n+1}|\mathbf{X}_n, \mathbf{X}_{n-1})$) and the emission model is autoregressive (e.g., $p_e(\mathbf{Y}_n|\mathbf{X}_n, \mathbf{Y}_{n-1})$) for each n . Here, we just focus on the model with the vanilla setting for giving the intuition. Transition and emission models are either given a priori or to be learned together with $\mathcal{X}$. Section A.1.3 covers the latter case where such models are parametric.

A.1.1 VITERBI ALGORITHM

The Viterbi algorithm [Viterbi (1967); Durbin et al. (1998)] is a dynamic programming for finding the MAP sequence of hidden states in the HMM given observations. Because $p(\mathcal{X}|\mathcal{Y}) \propto p(\mathcal{X}, \mathcal{Y})$, thus maximizing $p(\mathcal{X}|\mathcal{Y})$ is equivalent to maximizing $p(\mathcal{X}, \mathcal{Y})$, so we have:

$$\begin{aligned} \operatorname{argmax}_{\mathcal{X}} p(\mathcal{X}|\mathcal{Y}) &= \operatorname{argmax}_{\mathcal{X}} p(\mathcal{X}, \mathcal{Y}) = \operatorname{argmax}_{\mathcal{X}} [\log p(\mathcal{X}) + \log p(\mathcal{Y}|\mathcal{X})] \\ &= \operatorname{argmax}_{\mathcal{X}} \left[\sum_{n=1}^N \log p_e(\mathbf{Y}_n|\mathbf{X}_n) + \sum_{n=1}^N \log p_t(\mathbf{X}_n|\mathbf{X}_{n-1}) \right] \end{aligned} \quad (\text{A.7})$$

Thus, finding the MAP sequence of $p(\mathcal{X}|\mathcal{Y})$ is again equivalent to obtaining the most probable path of $\mathcal{L}$ in (A.8):

$$\mathcal{L}(\mathcal{X}) \triangleq \sum_{n=1}^N \log p_e(\mathbf{Y}_n|\mathbf{X}_n) + \sum_{n=1}^N \log p_t(\mathbf{X}_n|\mathbf{X}_{n-1}) \quad (\text{A.8})$$

Let us define a recursive term $v_n(s)$ for each $s \in \mathcal{S}$ and $n = 2, 3, \dots, N$ as follows:

$$\begin{aligned} v_1(s) &\triangleq \log p_e(\mathbf{Y}_1|\mathbf{X}_1 = s) + \log p(\mathbf{X}_1 = s) \\ v_n(s) &\triangleq \log p_e(\mathbf{Y}_n|\mathbf{X}_n = s) + \max_{k \in \mathcal{S}} (v_{n-1}(k) + \log p_t(\mathbf{X}_n = s|\mathbf{X}_{n-1} = k)) \end{aligned} \quad (\text{A.9})$$

Then, one can observe that the following holds for each $m = 1, 2, \dots, N - 1$:

$$\begin{aligned} & \max_{X_{m:N}} \sum_{n=m+1}^N \log p_e(Y_n|X_n) + \sum_{n=m+1}^N \log p_t(X_n|X_{n-1}) + v_m(X_m) \\ &= \max_{X_{m+1:N}} \sum_{n=m+2}^N \log p_e(Y_n|X_n) + \sum_{n=m+2}^N \log p_t(X_n|X_{n-1}) + v_{m+1}(X_{m+1}) \end{aligned} \quad (\text{A.10})$$

Thus, $\max_{\mathcal{X}} \mathcal{L}(\mathcal{X}) = \max_{X_N} v_N(X_N)$. Note that computing each $v_n(s)$ requires the time complexity $\mathcal{O}(|\mathcal{S}|)$, so the time complexity of computing $\max_{\mathcal{X}} \mathcal{L}(\mathcal{X})$ is of $\mathcal{O}(|\mathcal{S}|^2 N)$ and the algorithm requires the storage of $\mathcal{O}(|\mathcal{S}| N)$.

The goal is to find the maximizer of $\mathcal{L}(\mathcal{X})$, i.e., $\hat{\mathcal{X}}$ such that $\mathcal{L}(\hat{\mathcal{X}}) = \max_{\mathcal{X}} \mathcal{L}(\mathcal{X})$. Let us define another recursive term $b_n(s)$ for each $s \in \mathcal{S}$ and $n = 1, 2, \dots, N - 1$ as follows:

$$b_n(s) \triangleq \operatorname{argmax}_{k \in \mathcal{S}} (v_n(k) + \log p_t(X_{n+1} = s | X_n = k)) \quad (\text{A.11})$$

Note that $\hat{X}_N = \operatorname{argmax}_{X_N} v_N(X_N)$ and $\hat{X}_{N-1} = b_{N-1}(\hat{X}_N)$ because:

$$\begin{aligned} \max_{\mathcal{X}} \mathcal{L}(\mathcal{X}) &= \max_{X_N} v_N(X_N) \\ &= \log p_e(Y_N | \hat{X}_N) + \max_{k \in \mathcal{S}} (v_{N-1}(k) + \log p_t(X_N = \hat{X}_N | X_{N-1} = k)) \end{aligned} \quad (\text{A.12})$$

Moreover, one can show that the following holds for each $m = 1, 2, \dots, N - 2$:

$$\begin{aligned} \mathcal{L}(\hat{\mathcal{X}}) &= \sum_{n=m+1}^N \log p_e(Y_n | \hat{X}_n) + \sum_{n=m+1}^{N-1} \log p_e(\hat{X}_{n+1} | \hat{X}_n) \\ &\quad + \max_{k \in \mathcal{S}} (v_m(k) + \log p_t(X_{m+1} = \hat{X}_{m+1} | X_m = k)) \end{aligned} \quad (\text{A.13})$$

Thus, $\hat{X}_n = b_n(\hat{X}_{n+1})$ holds for all $n = 1, 2, \dots, N - 1$. Since each $b_n(s)$ can be stored in the step for updating v_n 's with time complexity $\mathcal{O}(|\mathcal{S}|)$, the time complexity of computing $\operatorname{argmax}_{\mathcal{X}} \mathcal{L}(\mathcal{X}) = \operatorname{argmax}_{\mathcal{X}} p(\mathcal{X}|\mathcal{Y})$ is still $\mathcal{O}(|\mathcal{S}|^2 N)$.

So far, we have observed how the MAP estimator of the hidden states given observations is

obtained by the Viterbi algorithm with time complexity $\mathcal{O}(|\mathcal{S}|^2 N)$, i.e., linear to the length of a sequence and quadratic to the size of a hidden state space. Though this algorithm is still popular, it has the following drawbacks:

1. This deterministic algorithm only returns the most probable path, thus is not proper if multiple probable paths are expected and to be considered one-by-one.
2. It assumes that both transition and emission models are given a priori, which is not common in practice.

The above downside will be addressed in the following subsections.

A.I.2 FORWARD-BACKWARD ALGORITHM

The forward-backward algorithm [Durbin et al. (1998)] is a probabilistic approach to HMMs. It focuses on computing the posterior distribution of hidden states and drawing sample paths from the posterior. As shown in its name, the forward-backward algorithm consists of the forward and backward algorithm. The forward algorithm recursively computes the forward sum $p(Y_{1:n}, X_n = s)$ while the backward algorithm updates $p(Y_{n+1:N}|X_n = s)$ for each n and $s \in \mathcal{S}$.

Note that the following holds for any $1 \leq l \leq m \leq N$: let $p(Y_{N+1:N}|X_m) \equiv 1$ in notation for saving the space.

$$\begin{aligned} p(X_{l:m}|Y_{1:N}) &\propto p(X_{l:m}, Y_{1:N}) \\ &= p(Y_{m+1:N}|X_m) p(Y_{1:l}, X_l) \prod_{n=l+1}^m p(X_n|X_{n-1}) p(Y_n|X_n) \end{aligned} \quad (\text{A.14})$$

, and, as a special case, $p(X_n|Y_{1:N}) \propto p(Y_{n+1:N}|X_n) p(Y_{1:n}, X_n)$.

For the forward algorithm, the following holds for each $n = 2, 3, \dots, N$ to run the algorithm

recursively:

$$\begin{aligned} p(Y_1, X_1 = s) &= p(Y_1|X_1 = s)p(X_1 = s) \\ p(Y_{1:n}, X_n = s) &= p(Y_n|X_n = s) \sum_{k \in \mathcal{S}} p(X_n = s|X_{n-1} = k)p(Y_{1:n-1}, X_{n-1} = k) \end{aligned} \quad (\text{A.15})$$

One can compute the marginal likelihood $p(Y_{1:N}) = \sum_{s \in \mathcal{S}} p(Y_{1:N}, X_N = s)$ from the last forward term.

Similarly, the backward algorithm is iterated based on the following equality for each $n = N - 1, N - 2, \dots, 1$:

$$p(Y_{n+1:N}|X_n = s) = \sum_{k \in \mathcal{S}} p(X_{n+1} = k|X_n = s)p(Y_{n+1}|X_{n+1} = k)p(Y_{n+2:N}|X_{n+1} = k) \quad (\text{A.16})$$

Note that computing each forward/backward term has time complexity of $\mathcal{O}(|\mathcal{S}|)$, thus that of the algorithm is $\mathcal{O}(|\mathcal{S}|^2 N)$ in time and $\mathcal{O}(|\mathcal{S}| N)$ in storage. Thus, the forward-backward algorithm gives a way to compute the posterior of each hidden state efficiently.

The backward sampling algorithm is a variant of the forward-backward algorithm. Instead of computing the backward terms $p(Y_{n+1:N}|X_n = s)$'s, the algorithm samples $\mathcal{X}$ from the posterior $p(\mathcal{X}|\mathcal{Y})$ backwards by utilizing the forward terms with the following relationships:

$$\begin{aligned} p(X_N = s|Y_{1:N}) &\propto p(Y_{1:N}, X_N = s) \\ p(X_n = s|X_{n+1} = \tilde{X}_{n+1}, Y_{1:N}) &\propto p(X_{n+1} = \tilde{X}_{n+1}|X_n = s)p(Y_{1:n}, X_n = s) \end{aligned} \quad (\text{A.17})$$

Note that the backward sampling in (A.17) is of time complexity $\mathcal{O}(|\mathcal{S}| N)$, but the full algorithm is still $\mathcal{O}(|\mathcal{S}|^2 N)$ because it uses the forward terms that are obtained in the forward algorithm.

In practice, the algorithm is usually challenged by the underflow, i.e., the forward and backward terms are too small not to be regarded as 0 in computer if the length of sequence N is large. One solution is to store the logarithmic values of the forward and backward terms instead of the

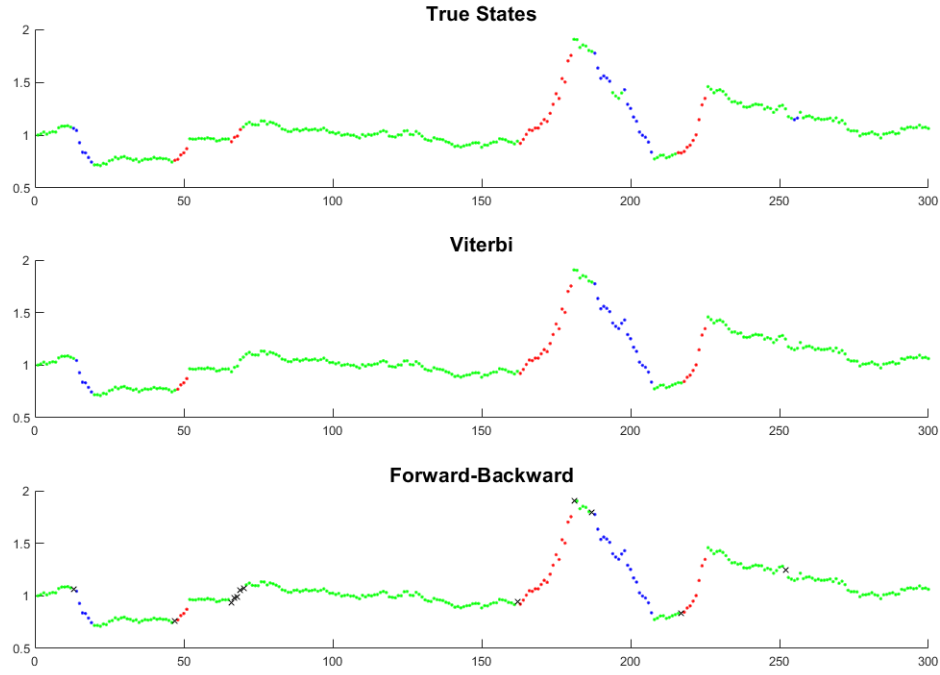


Figure A.4: An example of HMM and the inference by the Viterbi and forward-backward algorithm.

original values. It utilizes the following log-sum-exp trick:

$$\log \sum_{n=1}^N \exp(x_n) = \max_m x_m + \log \sum_{n=1}^N \exp\left(x_n - \max_m x_m\right) \quad (\text{A.18})$$

Figure A.4 shows an example of the HMM for the time series analysis. The upper panel shows the true hidden states in colors: blue dots are the observations suffering from the crash, red dots are those going on soaring, and green dots move sideways. The middle and bottom panels show the inferred hidden states by the Viterbi and forward-backward algorithms, respectively. One advantage of the forward-backward algorithm is that the inference is probabilistic so does more careful analysis: black crosses indicate that the associated MAP probabilities are less than 0.7, i.e., less certain than the others.

A.1.3 BAUM-WELCH ALGORITHM

So far, we have assumed that the transition and emission models are given. However, in practice both transition and emission models are not given a priori. Here those models are assumed to be parametric.

The Baum-Welch algorithm [Durbin et al. (1998)] is an expectation-maximization (EM) algorithm [Dempster et al. (1977)] that infers unknown parameters of HMMs. Let us formulate the full joint likelihood of $\mathcal{X}$ and $\mathcal{Y}$, where transition and emission models are parametrized by θ and ψ , respectively.

$$p(\mathcal{X}, \mathcal{Y}, \theta, \psi) = \pi(\theta, \psi) p(\mathcal{X}, \mathcal{Y} | \theta, \psi) = \pi(\theta, \psi) \prod_{n=1}^N p_t(X_n | X_{n-1}, \theta) \prod_{n=1}^N p_e(Y_n | X_n, \psi) \quad (\text{A.19})$$

, where $\pi(\theta, \psi)$ is the prior on the parameters.

The EM algorithm is consisting of the expectation step (E-step) and maximization step (M-step). Here $\mathcal{X}$ is the hidden state, so the E-step computes the following expectation for the previously updated parameters $\theta^{(t)}, \psi^{(t)}$:

$$\begin{aligned} \mathcal{Q}(\theta, \psi | \theta^{(t)}, \psi^{(t)}) &\triangleq \mathbb{E}_{\mathcal{X} | \mathcal{Y}, \theta^{(t)}, \psi^{(t)}} [\log p(\mathcal{X}, \mathcal{Y}, \theta, \psi)] \\ &= \mathbb{E}_{\mathcal{X} | \mathcal{Y}, \theta^{(t)}, \psi^{(t)}} \left[\log \pi(\theta, \psi) + \sum_{n=1}^N \log p_t(X_n | X_{n-1}, \theta) + \sum_{n=1}^N \log p_e(Y_n | X_n, \psi) \right] \\ &= \log \pi(\theta, \psi) + \sum_{n=1}^N \sum_{X_n \in \mathcal{S}} \sum_{X_{n-1} \in \mathcal{S}} \mathbb{P}(X_{n-1}, X_n) \log p_t(X_n | X_{n-1}, \theta) \\ &\quad + \sum_{n=1}^N \sum_{X_n \in \mathcal{S}} \mathbb{P}(X_n) \log p_e(Y_n | X_n, \psi) \end{aligned} \quad (\text{A.20})$$

, where $\mathbb{P}(X_n) \triangleq p(X_n | Y_{1:N}, \theta^{(t)}, \psi^{(t)})$ and $\mathbb{P}(X_{n-1}, X_n) \triangleq p(X_{n-1:n} | Y_{1:N}, \theta^{(t)}, \psi^{(t)})$ for each n . Note that $\mathbb{P}$ is given by the forward-backward algorithm in section A.1.2.

The M-step finds the parameters that maximize $\mathcal{Q}(\theta, \psi | \theta^{(t)}, \psi^{(t)})$, i.e.,

$$(\theta^{(t+1)}, \psi^{(t+1)}) = \underset{\theta, \psi}{\operatorname{argmax}} \mathcal{Q}(\theta, \psi | \theta^{(t)}, \psi^{(t)}) \quad (\text{A.21})$$

Both models and hidden states can be inferred by iterating the E-step that computes $\mathbb{P}$ from the forward-backward algorithm and the M-step that optimizes (A.21) until convergence. Note that, however, the EM algorithm only guarantees the local optimum, thus limited in such a sense.

The optimization of (A.21) can be done by either numerically (e.g., gradient descent) or analytically, which depends on how transition and emission models are parametrized. For example, suppose that $\pi \equiv 1$, the observations $\mathcal{Y}$ are also defined in a finite set $\mathcal{V}$, both transition and emission models are stationary and θ and ψ indicate transition and emission tables, respectively. Then, one can derive the following updating formulations analytically:

$$\begin{aligned} \theta^{(t+1)}(s, k) &= \frac{\sum_{n=1}^N \mathbb{P}(\mathbf{X}_{n-1} = s, \mathbf{X}_n = k)}{\sum_{n=1}^N \mathbb{P}(\mathbf{X}_{n-1} = s)} \\ \psi^{(t+1)}(s, v) &= \frac{\sum_{n=1}^N \mathbf{1}_{Y_n=v} \cdot \mathbb{P}(\mathbf{X}_n = s)}{\sum_{n=1}^N \mathbb{P}(\mathbf{X}_n = s)} \end{aligned} \quad (\text{A.22})$$

, where $\theta(s, k) \triangleq p_t(\mathbf{X}_n = k | \mathbf{X}_{n-1} = s)$ and $\psi(s, v) \triangleq p_e(Y_n = v | \mathbf{X}_n = s)$. The Baum-Welch algorithm infers both models and hidden states by iterating the E-step that computes $\mathbb{P}$ from the forward-backward algorithm and the M-step that updates the transition and emission models by (A.22) until convergence.

A.I.4 VARIATIONAL APPROXIMATION

As discussed earlier, the exact inference of a fully connected hidden state model is usually intractable. One detour is to approximate it with a tractable Markovian model. This subsection briefly introduces the core idea regarding that methodology.

Let us assume that the real joint distribution of $(\mathcal{X}, \mathcal{Y})$ is $p(\mathcal{X}, \mathcal{Y}) = p(\mathcal{Y} | \mathcal{X})p(\mathcal{X})$ and let $q(\mathcal{X}) = q(\mathcal{X} | \theta)$ be a tractable Markov chain parametrized by θ for approximating $p(\mathcal{X} | \mathcal{Y})$.

Then, because the Kullback-Leibler (KL) divergence is always nonnegative, we have the following:

$$\mathbb{D}_{\text{KL}}(q(\cdot|\theta)||p(\cdot|\mathcal{Y})) = \log p(\mathcal{Y}) + \mathbb{D}_{\text{KL}}(q(\cdot|\theta)||p) - \mathbb{E}_{q(\mathcal{X}|\theta)}[\log p(\mathcal{Y}|\mathcal{X})] \geq 0 \quad (\text{A.23})$$

Therefore, the log-marginal likelihood $\log p(\mathcal{Y})$ has the evidence lower bound (ELBO) $\mathcal{L}$ defined as follows:

$$\log p(\mathcal{Y}) \geq \mathbb{E}_{q(\mathcal{X}|\theta)}[\log p(\mathcal{Y}|\mathcal{X})] - \mathbb{D}_{\text{KL}}(q(\cdot|\theta)||p) = \mathcal{L} \quad (\text{A.24})$$

Now the goal is to maximize $\mathcal{L}$ instead of $\log p(\mathcal{Y})$ directly: this is the core idea of variational approximation [Bishop(2006)]. Note that $\log p(\mathcal{Y}|\mathcal{X})$ is easily decomposed into the summation if $\mathcal{Y}$ is assumed to be conditionally independent given $\mathcal{X}$ and $\mathbb{D}_{\text{KL}}(q(\cdot|\theta)||p)$ can be expressed analytically as a summation in discrete modelling.

Once q is learned, one can utilize it for the sampling purpose. Note that the sample paths drawn from the variational approximation q is not the exact ones but can be used for the initialization in the Markov-chain Monte Carlo (MCMC) algorithm that will be covered in section A.4. If obtaining posterior distributions of each $\mathbf{X}_n$ is interesting only, defining q as a naïve Bayes could be an alternative: it directly approximates the posteriors then.

A.2 KALMAN FILTER AND SMOOTHER

In section A.1, we discussed the case that hidden states are defined in a finite set. Briefly speaking, Viterbi and forward-backward algorithms for HMMs give ways of considering all possible assignments to the hidden states efficiently. However, this cannot be extended to the case that hidden states are defined continuously: there are infinitely many possible assignments. Nonetheless, computing the analogues of forward and backward terms in forms of continuous distributions is still feasible under some conditions. This section covers Kalman filter (section A.2.1) and smoother (section A.2.2) for the cases that both transition and emission models are Gaussian.

In practice, Gaussian assumption could be too strong to model the data: mixture Kalman filter (section A.2.3) would be an alternative to alleviate such strong conditions.

A.2.1 KALMAN FILTER

The Kalman filter [Kalman (1960); Murphy (2012)] is an algorithm that uses a series of noisy observations to estimate the underlying hidden states that are associated to them. Similar to HMMs, it consists of the transition and emission models: sometimes the model also assumes the control-input model that intentionally affect the hidden states in each step, but it is not considered in this dissertation. The following notations are assumed in the rest of section A.2.

- $\mathcal{X} = \{\mathbf{X}_n\}_{n=0}^N$: a series of hidden states where each $\mathbf{X}_n \in \mathbb{R}^D$ and $\mathbf{X}_0$ is an (hypothetical) initial state.
- $\mathcal{Y} = \{\mathbf{Y}_n\}_{n=1}^N$: a series of observations where each $\mathbf{Y}_n \in \mathbb{R}^L$.
- $\mathcal{F} = \{\mathbf{F}_n\}_{n=1}^N$: a set of transition models where a $D \times D$ matrix $\mathbf{F}_n$ affects the transition from $\mathbf{X}_{n-1}$ to $\mathbf{X}_n$.
- $\mathcal{H} = \{\mathbf{H}_n\}_{n=1}^N$: a set of emission models where a $L \times D$ matrix $\mathbf{H}_n$ affects the emission from $\mathbf{X}_n$ to $\mathbf{Y}_n$.
- $\mathcal{Q} = \{\mathbf{Q}_n\}_{n=1}^N$: a set of covariance matrices in transition models where a $D \times D$ matrix $\mathbf{Q}_n$ is involved with the transition from $\mathbf{X}_{n-1}$ to $\mathbf{X}_n$.
- $\mathcal{B} = \{\mathbf{b}_n\}_{n=1}^N$: a set of bias vectors in transition models where each $\mathbf{b}_n \in \mathbb{R}^D$ affects the transition from $\mathbf{X}_{n-1}$ to $\mathbf{X}_n$.
- $\mathcal{R} = \{\mathbf{R}_n\}_{n=1}^N$: a set of covariance matrices in emission models where a $L \times L$ matrix $\mathbf{R}_n$ is involved with the emission from $\mathbf{X}_n$ to $\mathbf{Y}_n$.
- $\mathcal{C} = \{\mathbf{c}_n\}_{n=1}^N$: a set of bias vectors in emission models where each $\mathbf{c}_n \in \mathbb{R}^L$ affects the emission from $\mathbf{X}_n$ to $\mathbf{Y}_n$.

The transition $p_t = \{p_t^{(n)}\}_{n=1}^N$ and emission $p_e = \{p_e^{(n)}\}_{n=1}^N$ models are given as follows:

$$\begin{aligned} p_t^{(n)}(X_n|X_{n-1}) &= \mathcal{N}(X_n|F_n X_{n-1} + b_n, Q_n) \\ p_e^{(n)}(Y_n|X_n) &= \mathcal{N}(Y_n|H_n X_n + c_n, R_n) \end{aligned} \quad (\text{A.25})$$

The Kalman filter aims at computing the prediction term $p(X_n|Y_{1:n-1})$, forward posterior $p(X_n|Y_{1:n})$ and posterior predictive $p(Y_n|Y_{1:n-1})$ for each n : this approach is more realistic if the data are being received on real time and one is supposed to make predictions promptly. Suppose that $p(X_{n-1}|Y_{1:n-1}) = \mathcal{N}(X_{n-1}|\mu_{n-1}, \Sigma_{n-1})$. Then, one can derive the following:

$$\begin{aligned} p(X_n|Y_{1:n-1}) &= \mathcal{N}(X_n|\mu_{n|n-1}, \Sigma_{n|n-1}) \\ &\triangleq \mathcal{N}(X_n|F_n \mu_{n-1} + b_n, Q_n + F_n \Sigma_{n-1} F_n^\top) \\ &= \int p_t^{(n)}(X_n|X_{n-1}) p(X_{n-1}|Y_{1:n-1}) dX_{n-1} \end{aligned} \quad (\text{A.26})$$

$$\begin{aligned} p(Y_n|Y_{1:n-1}) &= \mathcal{N}(X_n|\hat{\mu}_n, \hat{\Sigma}_n) \\ &\triangleq \mathcal{N}(Y_n|H_n \mu_{n|n-1} + c_n, R_n + H_n \Sigma_{n|n-1} H_n^\top) \\ &= \int p_e^{(n)}(Y_n|X_n) p(X_n|Y_{1:n-1}) dX_n \end{aligned} \quad (\text{A.27})$$

$$\begin{aligned} p(X_n|Y_{1:n}) &= \mathcal{N}(X_n|\mu_n, \Sigma_n) \\ &\triangleq \mathcal{N}(X_n|\mu_{n|n-1} + \Sigma_{n|n-1} H_n^\top \hat{\Sigma}_n^{-1} (Y_n - \hat{\mu}_n), \Sigma_{n|n-1} - \Sigma_{n|n-1} H_n^\top \hat{\Sigma}_n^{-1} H_n \Sigma_{n|n-1}) \\ &\propto p_e^{(n)}(Y_n|X_n) p(X_n|Y_{1:n-1}) \end{aligned} \quad (\text{A.28})$$

The terms to compute can be recursively updated by iterating (A.26), (A.27) and (A.28). The initial term $p(X_1|Y_1) = \mathcal{N}(X_1|\mu_1, \Sigma_1) \propto p_e^{(1)}(Y_1|X_1) p_t^{(1)}(X_1|X_0)$ is given as follows:

$$\begin{aligned} \mu_1 &= F_1 X_0 + b_1 + Q_1 H_1^\top (R_1 + H_1 Q_1 H_1^\top)^{-1} (Y_1 - H_1 (F_1 X_0 + b_1) - c_1) \\ \Sigma_1 &= Q_1 - Q_1 H_1^\top (R_1 + H_1 Q_1 H_1^\top)^{-1} H_1 Q_1 \end{aligned} \quad (\text{A.29})$$

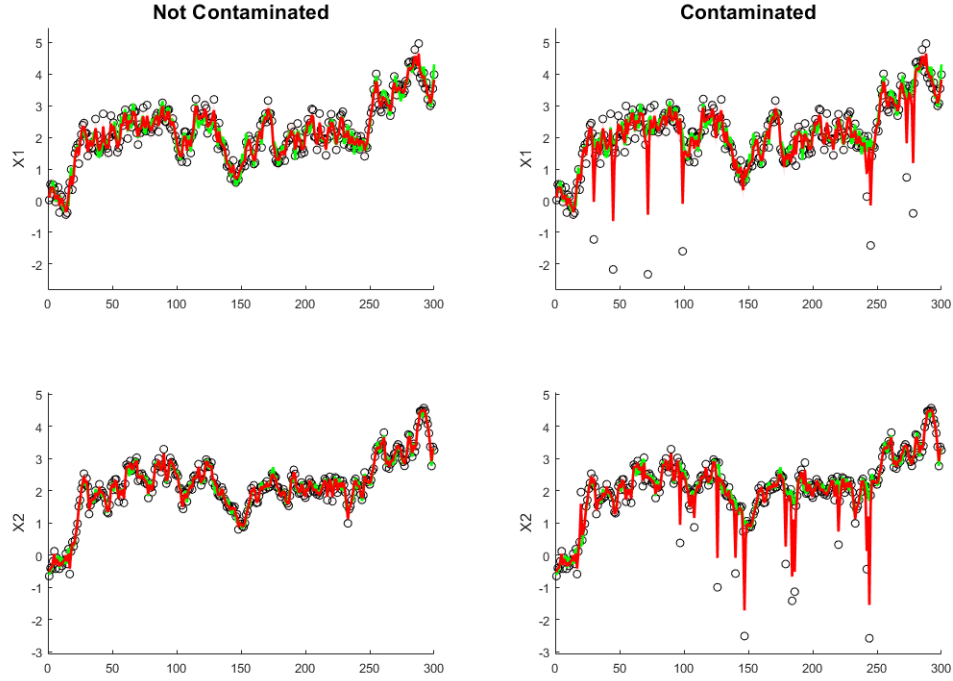


Figure A.5: An example of Kalman smoother where the 2-dimensional noisy observations are filtered. For each panel, black circles are the observations, the green curve is the true denoised observations (hidden states) and red curve is the inferred hidden states from the noisy observations. The left and right panels use the same observations except for 5% of the contamination in the right panels.

Note that the terms in (A.26), (A.27), (A.28) and (A.29) are expressed in closed forms because their integrations and proportionality are involved with Gaussian terms, thus not extendable to the cases with general transition and observation models.

A.2.2 KALMAN SMOOTHER

The Kalman smoother [Rauch et al. (1965)] is an analogue of the forward-backward algorithm for HMMs, i.e., it computes the posterior distribution $p(X_n|Y_{1:N})$ for each n *after* observing all $Y_{1:N}$. There are two derivations: one updates $p(X_n|Y_{1:N})$ from $p(X_{n+1}|Y_{1:N})$ one-by-one backwards [Murphy (2012)] and the other is two-filter smoothing [Kitagawa (1994); Briers et al. (2010)] that recursively computes $p(X_n|Y_{n+1:N})$ backwards and restores each $p(X_n|Y_{1:N})$ from $p(X_n|Y_{n+1:N})$ and $p(X_n|Y_{n:N})$, just as the forward-backward algorithm. This section follows the former one.

From (A.25) to (A.28), we have the following:

$$\begin{aligned}
p(\mathbf{X}_n, \mathbf{X}_{n+1} | \mathbf{Y}_{1:n}) &= p_t^{(n)}(\mathbf{X}_{n+1} | \mathbf{X}_n) p(\mathbf{X}_n | \mathbf{Y}_{1:n}) \\
&= \mathcal{N}(\mathbf{X}_{n+1} | \mathbf{F}_{n+1} \mathbf{X}_n + b_{n+1}, \mathbf{Q}_{n+1}) \mathcal{N}(\mathbf{X}_n | \mu_n, \Sigma_n) \\
&= \mathcal{N}\left(\begin{pmatrix} \mathbf{X}_n \\ \mathbf{X}_{n+1} \end{pmatrix} \middle| \begin{pmatrix} \mu_n \\ \mu_{n+1|n} \end{pmatrix}, \begin{pmatrix} \Sigma_n & \Sigma_n \mathbf{F}_{n+1}^\top \\ \mathbf{F}_{n+1} \Sigma_n & \Sigma_{n+1|n} \end{pmatrix}\right)
\end{aligned} \tag{A.30}$$

, thus:

$$\begin{aligned}
p(\mathbf{X}_n | \mathbf{X}_{n+1}, \mathbf{Y}_{1:n}) &= \mathcal{N}\left(\mathbf{X}_n \middle| \mu_n + \Sigma_n \mathbf{F}_{n+1}^\top \Sigma_{n+1|n}^{-1} (\mathbf{X}_{n+1} - \mu_{n+1|n}), \Sigma_n - \Sigma_n \mathbf{F}_{n+1}^\top \Sigma_{n+1|n}^{-1} \mathbf{F}_{n+1} \Sigma_n\right)
\end{aligned} \tag{A.31}$$

Suppose that $p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:N}) = \mathcal{N}(\mathbf{X}_{n+1} | \mu_{n+1|N}, \Sigma_{n+1|N})$ is given. Then, one can derive the following with (A.31):

$$\begin{aligned}
p(\mathbf{X}_n | \mathbf{Y}_{1:N}) &= \mathcal{N}(\mathbf{X}_n | \mu_{n|N}, \Sigma_{n|N}) \\
&\triangleq \int p(\mathbf{X}_n | \mathbf{X}_{n+1}, \mathbf{Y}_{1:n}) p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:N}) d\mathbf{X}_{n+1}
\end{aligned} \tag{A.32}$$

, where:

$$\begin{aligned}
\mu_{n|N} &= \mu_n + \Sigma_n \mathbf{F}_{n+1}^\top \Sigma_{n+1|n}^{-1} (\mu_{n+1|N} - \mu_{n+1|n}) \\
\Sigma_{n|N} &= \Sigma_n + \Sigma_n \mathbf{F}_{n+1}^\top \Sigma_{n+1|n}^{-1} (\Sigma_{n+1|N} - \Sigma_{n+1|n}) \Sigma_{n+1|n}^{-1} \mathbf{F}_{n+1} \Sigma_n
\end{aligned} \tag{A.33}$$

In addition, the initial term $p(\mathbf{X}_N | \mathbf{Y}_{1:N}) = \mathcal{N}(\mathbf{X}_N | \mu_{N|N}, \Sigma_{N|N})$ is given directly from (A.28). To sum up, the posterior distribution of each $\mathbf{X}_n$ given the full observation $\mathbf{Y}_{1:N}$ is again a Gaussian distribution with explicit mean vector and covariance matrix that can be computed backwards recursively.

Since $p(\mathbf{X}_n | \mathbf{X}_{n+1}, \mathbf{Y}_{1:N}) = p(\mathbf{X}_n | \mathbf{X}_{n+1}, \mathbf{Y}_{1:n})$, backward sampling is easily done by (A.31): $\tilde{\mathbf{X}}_N$ is drawn from $p(\mathbf{X}_N | \mathbf{Y}_{1:N}) = \mathcal{N}(\mathbf{X}_N | \mu_{N|N}, \Sigma_{N|N})$ and $\tilde{\mathbf{X}}_n$ from $p(\mathbf{X}_n | \tilde{\mathbf{X}}_{n+1}, \mathbf{Y}_{1:n})$ backwards one-by-one recursively.

A.2.3 MIXTURE KALMAN FILTER

So far, we have derived the Kalman filter and smoother that assume Gaussian transition and emission models. In reality, it is natural that either or both of transition and emission models are non-Gaussian and figure A.5 shows that even a small portion of contamination can ruin the inference by Kalman smoother. There have been a lot of approximation methods for such cases, and here the mixture Kalman filter [Chen & Liu (2000)], an example of the Rao-Blackwellized particle filters that reduce the sample variance much, is discussed as one of them.

First of all, let us redefine the notations:

- $\mathcal{X} = \{\mathbf{X}_n\}_{n=0}^N$: a series of hidden states where each $\mathbf{X}_n \in \mathbb{R}^D$ and $\mathbf{X}_0$ is an (hypothetical) initial state.
- $\mathcal{Y} = \{\mathbf{Y}_n\}_{n=1}^N$: a series of observations where each $\mathbf{Y}_n \in \mathbb{R}^L$.
- $\mathcal{Z} = \{\mathbf{Z}_n\}_{n=1}^N$: a series of latent states where $\mathbf{Z}_n \in \Theta$ could be either discrete or continuous.
- $\mathcal{F} = \{\mathbf{F}_\theta\}_{\theta \in \Theta}$: the parametrized transition model where $\mathbf{F}_\theta$ is a $D \times D$ matrix.
- $\mathcal{H} = \{\mathbf{H}_\theta\}_{\theta \in \Theta}$: the parametrized emission model where $\mathbf{H}_\theta$ is a $L \times D$ matrix.
- $\mathcal{Q} = \{\mathbf{Q}_\theta\}_{\theta \in \Theta}$: the parametrized covariance matrix in transition models where each $\mathbf{Q}_\theta$ is a $D \times D$ matrix.
- $\mathcal{B} = \{\mathbf{b}_\theta\}_{\theta \in \Theta}$: the parametrized bias vector in transition models where each $\mathbf{b}_\theta \in \mathbb{R}^D$.
- $\mathcal{R} = \{\mathbf{R}_\theta\}_{\theta \in \Theta}$: the parametrized covariance matrix in emission models where each $\mathbf{R}_\theta$ is a $L \times L$ matrix.
- $\mathcal{C} = \{\mathbf{c}_\theta\}_{\theta \in \Theta}$: the parametrized bias vector in emission models where each $\mathbf{c}_\theta \in \mathbb{R}^L$.

The major difference from the vanilla version in section A.2.1 and A.2.2 is that the models are now parametrized by latent states, as follows:

$$\begin{aligned} p_t(X_n|X_{n-1}, Z_n) &= \mathcal{N}(X_n|F_{Z_n}X_{n-1} + b_{Z_n}, Q_{Z_n}) \\ p_c(Y_n|X_n, Z_n) &= \mathcal{N}(Y_n|H_{Z_n}X_n + c_{Z_n}, R_{Z_n}) \end{aligned} \quad (\text{A.34})$$

Note that each Z_n is a latent state thus not given a priori. The intuition comes from the following marginal transition and emission models that are expressed as mixtures of Gaussian:

$$\begin{aligned} p_t(X_n|X_{n-1}) &= \int \mathcal{N}(X_n|F_\theta X_{n-1} + b_\theta, Q_\theta) \pi(\theta) d\theta \\ p_c(Y_n|X_n) &= \int \mathcal{N}(Y_n|H_\theta X_n + c_\theta, R_\theta) \pi(\theta) d\theta \end{aligned} \quad (\text{A.35})$$

, where π is a prior on θ . Thus, this model can be applied as an approximation method to non-Gaussian cases. The goal is to obtain the forward posterior $p(X_n|Y_{1:n})$ for each n . Suppose that $p(X_{n-1}|Y_{1:n-1})$ is given as follows:

$$\begin{aligned} p(X_{n-1}|Y_{1:n-1}, \tilde{Z}_{1:n-1}^{(m)}) &= \mathcal{N}(X_{n-1}|\mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)}) \\ p(X_{n-1}|Y_{1:n-1}) &= \sum_{m=1}^M \omega_{n-1}^{(m)} \mathcal{N}(X_{n-1}|\mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)}) \end{aligned} \quad (\text{A.36})$$

, where $\tilde{Z}_{1:n-1}^{(m)}$ is the m th sample path of $Z_{1:n-1}$.

Then, one can derive the following approximation to $p(X_n|Y_{1:n})$ by the importance sampling

with a proposal q :

$$\begin{aligned}
p(X_n | Y_{1:n}) &= \iint \frac{p(X_n, X_{n-1}, Z_n, Y_n | Y_{1:n-1})}{p(Y_n | Y_{1:n-1})} dZ_n dX_{n-1} \\
&= \sum_{m=1}^M \omega_n^{(m)} \mathcal{N}(X_n | \mu_n^{(m)}, \Sigma_n^{(m)}) \\
&\propto \sum_{m=1}^M \omega_{n-1}^{(m)} \int \mathcal{N}(X_{n-1} | \mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)}) \int p_e(Y_n | X_n, Z_n) p_t(X_n | X_{n-1}, Z_n) \pi(Z_n) dZ_n dX_{n-1} \\
&\approx \sum_{m=1}^M \omega_{n-1}^{(m)} \int \mathcal{N}(X_{n-1} | \mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)}) \frac{p_e(Y_n | X_n, \tilde{Z}_n^{(m)}) p_t(X_n | X_{n-1}, \tilde{Z}_n^{(m)}) \pi(\tilde{Z}_n^{(m)})}{q(\tilde{Z}_n^{(m)} | \tilde{Z}_{1:n-1}^{(m)}, \mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)})} dX_{n-1}
\end{aligned} \tag{A.37}$$

, where $\tilde{Z}_n^{(m)}$ is a sample drawn from the proposal $q(Z_n | \tilde{Z}_{1:n-1}^{(m)}, \mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)})$ and:

$$\begin{aligned}
\mu_{n|n-1}^{(m)} &\triangleq F_{\tilde{Z}_n^{(m)}} \mu_{n-1}^{(m)} + b_{\tilde{Z}_n^{(m)}} \\
\Sigma_{n|n-1}^{(m)} &\triangleq Q_{\tilde{Z}_n^{(m)}} + F_{\tilde{Z}_n^{(m)}} \Sigma_{n-1}^{(m)} F_{\tilde{Z}_n^{(m)}}^\top \\
\mu_n^{(m)} &= \mu_{n|n-1}^{(m)} + \Sigma_{n|n-1}^{(m)} H_{\tilde{Z}_n^{(m)}}^\top \left(R_{\tilde{Z}_n^{(m)}} + H_{\tilde{Z}_n^{(m)}} \Sigma_{n|n-1}^{(m)} H_{\tilde{Z}_n^{(m)}}^\top \right)^{-1} \left(Y_n - H_{\tilde{Z}_n^{(m)}} \mu_{n|n-1}^{(m)} - c_{\tilde{Z}_n^{(m)}} \right) \\
\Sigma_n^{(m)} &= \Sigma_{n|n-1}^{(m)} - \Sigma_{n|n-1}^{(m)} H_{\tilde{Z}_n^{(m)}}^\top \left(R_{\tilde{Z}_n^{(m)}} + H_{\tilde{Z}_n^{(m)}} \Sigma_{n|n-1}^{(m)} H_{\tilde{Z}_n^{(m)}}^\top \right)^{-1} H_{\tilde{Z}_n^{(m)}} \Sigma_{n|n-1}^{(m)} \\
\omega_n^{(m)} &\propto \frac{\omega_{n-1}^{(m)} \pi(\tilde{Z}_n^{(m)}) / q(\tilde{Z}_n^{(m)} | \tilde{Z}_{1:n-1}^{(m)}, \mu_{n-1}^{(m)}, \Sigma_{n-1}^{(m)})}{\mathcal{N}(Y_n | H_{\tilde{Z}_n^{(m)}} \mu_{n|n-1}^{(m)} + c_{\tilde{Z}_n^{(m)}}, R_{\tilde{Z}_n^{(m)}} + H_{\tilde{Z}_n^{(m)}} \Sigma_{n|n-1}^{(m)} H_{\tilde{Z}_n^{(m)}}^\top)}
\end{aligned} \tag{A.38}$$

Thus, one can approximate each $p(X_n | Y_{1:n})$ with a Gaussian mixture recursively by (A.37). Resampling on $\{\tilde{Z}_{1:n}^{(m)}\}_{m=1}^M$ with respect to their weights $\{\omega_n^{(m)}\}_{m=1}^M$ is recommended for stability. [Chen & Liu (2000)] also derive the extended version of the model which is not covered in this subsection.

There are other variants that deal with nonlinearity or non-Gaussian model aspects, as suggested in [Murphy (2012)]: extended Kalman filter (ENK) [McElhoe (1966)], unscented Kalman filter (UKF) [Julier & Uhlmann (1997)], assumed density filtering (ADF) [Maybeck (1979)] and ensemble Kalman filter [Evensen (1994); Houtekamer & Mitchell (1998)]. Note that those models do not guarantee the optimal inference as well.

A.3 PARTICLE FILTER AND SMOOTHER

Though section A.2.3 introduced an approximation variant of the Kalman filter, it is not enough to cover all cases: for example, $\mathcal{Y}$ is defined on a finite set while $\mathcal{X}$ still take continuous values.

Let us go back to the forward posterior updating formulation that is analogous to the forward algorithm of HMMs:

$$p(\mathbf{X}_{n+1}|\mathbf{Y}_{1:n+1}) = \frac{p(\mathbf{Y}_{n+1}|\mathbf{X}_{n+1})}{p(\mathbf{Y}_{n+1}|\mathbf{Y}_{1:n})} \int p(\mathbf{X}_{n+1}|\mathbf{X}_n) p(\mathbf{X}_n|\mathbf{Y}_{1:n}) d\mathbf{X}_n \quad (\text{A.39})$$

The main obstacle that prevents the exact formulation is that the integration in (A.39) cannot be done analytically in general. The Monte Carlo methods approximate general integrations and the importance sampling [[Gelman et al. \(2013\)](#)] reduces the variance that could be very high in the naïve version. The core idea is to bring a proposal distribution q that is believed to get high density near to the mode of the posterior distribution $p(\mathbf{X}_n|\mathbf{Y}_{1:n}, \mathbf{X}_{n+1})$ and approximate the integration as follows:

$$\int p(\mathbf{X}_{n+1}|\mathbf{X}_n) p(\mathbf{X}_n|\mathbf{Y}_{1:n}) d\mathbf{X}_n \approx \frac{1}{M} \sum_{m=1}^M \frac{p(\mathbf{X}_{n+1}|\tilde{\mathbf{X}}_n^{(m)}) p(\tilde{\mathbf{X}}_n^{(m)}|\mathbf{Y}_{1:n})}{q(\tilde{\mathbf{X}}_n^{(m)})} \quad (\text{A.40})$$

, where $\tilde{\mathbf{X}}_n^{(m)}$'s are independent and identically distributed samples drawn from q .

The particle filter (section A.3.1) and smoother (section A.3.2) [[Doucet et al. \(2001\)](#); [Klaas et al. \(2006\)](#)] are variational algorithms based on the sequential importance sampling. While the particle filter aims at updating the forward posterior $p(\mathbf{X}_n|\mathbf{Y}_{1:n})$ at each n , the particle smoother is for computing the posterior $p(\mathbf{X}_n|\mathbf{Y}_{1:N})$ after observing all $\mathbf{Y}_{1:N}$.

A.3.1 PARTICLE FILTER

As section A.2.1, let us first define the following notations that will be used in section A.3.

- $\mathcal{X} = \{\mathbf{X}_n\}_{n=0}^N$: a series of hidden states where each $\mathbf{X}_n$ and $\mathbf{X}_0$ is an (hypothetical) initial

state.

- $\mathcal{Y} = \{Y_n\}_{n=1}^N$: a series of observations where each Y_n .

One major difference of the above definitions from those in section A.2.1 is that now X_n and Y_n are not restricted to be finite-dimensional vectors. Also, the transition and emission models $p_t^{(n)}(X_n|X_{n-1})$ and $p_e^{(n)}(Y_n|X_n)$ need not to be Gaussian here. In section A.2.1, each $p(X_n|Y_{1:n})$ is expressed as a Gaussian distribution that can be updated one-by-one recursively; section A.2.3 approximates it with a Gaussian mixture. Here, suppose that $p(X_n|Y_{1:n})$ is approximated with a weighted empirical density [Wasserman (2006)] as follows:

$$p(X_n|Y_{1:n}) = \sum_{m=1}^M \omega_n^{(m)} 1_{\{X_n=x_n^{(m)}\}}(X_n) \quad (\text{A.41})$$

The updating formulation depends on the proposal distribution q to sample X_{n+1} . We have the following:

$$\begin{aligned} p(X_{n+1}|Y_{1:n+1}) &= \int \frac{p_e^{(n+1)}(Y_{n+1}|X_{n+1}) p_t^{(n+1)}(X_{n+1}|X_n) p(X_n|Y_{1:n})}{p(Y_{n+1}|Y_{1:n})} dX_n \\ &= \frac{p_e^{(n+1)}(Y_{n+1}|X_{n+1})}{p(Y_{n+1}|Y_{1:n})} \int p_t^{(n+1)}(X_{n+1}|X_n) p(X_n|Y_{1:n}) dX_n \\ &= \frac{p_e^{(n+1)}(Y_{n+1}|X_{n+1})}{p(Y_{n+1}|Y_{1:n})} \sum_{k=1}^M \int p_t^{(n+1)}(X_{n+1}|X_n) \omega_n^{(k)} 1_{\{X_n=x_n^{(k)}\}}(X_n) dX_n \\ &= \frac{1}{p(Y_{n+1}|Y_{1:n})} \sum_{k=1}^M \omega_n^{(k)} p_e^{(n+1)}(Y_{n+1}|X_{n+1}) p_t^{(n+1)}(X_{n+1}|X_n = x_n^{(k)}) \\ &= \sum_{k=1}^M \frac{\omega_n^{(k)} p_e^{(n+1)}(Y_{n+1}|X_{n+1}) p_t^{(n+1)}(X_{n+1}|X_n = x_n^{(k)})}{q(X_{n+1}) p(Y_{n+1}|Y_{1:n})} q(X_{n+1}) \end{aligned} \quad (\text{A.42})$$

The core idea is to approximate the proposal distribution q by the following empirical density estimation:

$$q(X_{n+1}) \approx \frac{1}{M} \sum_{m=1}^M 1_{\{X_{n+1}=x_{n+1}^{(m)}\}}(X_{n+1}) \quad (\text{A.43})$$

, where $x_{n+1}^{(m)}$'s are the independent and identically distributed samples drawn from q . Here we

particularly call those samples “particles” that the algorithm is named after.

With (A.43), the below formulation follows:

$$\begin{aligned}
p(X_{n+1}|Y_{1:n+1}) &= \sum_{m=1}^M \omega_{n+1}^{(m)} 1_{\{X_{n+1}=x_{n+1}^{(m)}\}}(X_{n+1}) \\
&\approx \sum_{k=1}^M \frac{\omega_n^{(k)} p_e^{(n+1)}(Y_{n+1}|X_{n+1}) p_t^{(n+1)}(X_{n+1}|X_n = x_n^{(k)})}{q(X_{n+1}) p(Y_{n+1}|Y_{1:n})} \frac{1}{M} \sum_{m=1}^M 1_{\{X_{n+1}=x_{n+1}^{(m)}\}}(X_{n+1}) \\
&\propto \sum_{m=1}^M \sum_{k=1}^M \frac{\omega_n^{(k)} p_e^{(n+1)}(Y_{n+1}|X_{n+1} = x_{n+1}^{(m)}) p_t^{(n+1)}(X_{n+1} = x_{n+1}^{(m)}|X_n = x_n^{(k)})}{q(X_{n+1} = x_{n+1}^{(m)}) p(Y_{n+1}|Y_{1:n})} 1_{\{X_{n+1}=x_{n+1}^{(m)}\}}(X_{n+1})
\end{aligned} \tag{A.44}$$

, where the weights $\omega_{n+1}^{(m)}$ ’s satisfy the following:

$$\omega_{n+1}^{(m)} \propto \frac{p_e^{(n+1)}(Y_{n+1}|X_{n+1} = x_{n+1}^{(m)})}{q(X_{n+1} = x_{n+1}^{(m)})} \left(\sum_{k=1}^M \omega_n^{(k)} p_t^{(n+1)}(X_{n+1} = x_{n+1}^{(m)}|X_n = x_n^{(k)}) \right) \tag{A.45}$$

So far, we have obtained an approximation (A.44) of $p(X_{n+1}|Y_{1:n+1})$ that has the same form of $p(X_n|Y_{1:n})$ in (A.41). For each n , the particle filter draws particles and computes their weights to recursively approximate the forward posteriors as above.

A.3.2 PARTICLE SMOOTHER

The particle smoother is analogous to the Kalman smoother in section A.2.2. The main goal is to approximate the posterior $p(X_n|Y_{1:N})$ with a weighted empirical density, as usual. Note that the particle filter returns $p(X_N|Y_{1:N})$. Let us assume that $p(X_{n+1}|Y_{1:N})$ has the following approximation:

$$p(X_{n+1}|Y_{1:N}) \approx \sum_{m=1}^M \hat{\omega}_{n+1}^{(m)} 1_{\{X_{n+1}=x_{n+1}^{(m)}\}}(X_n) \tag{A.46}$$

Then, one can derive the following relationship:

$$\begin{aligned}
p(\mathbf{X}_n | \mathbf{Y}_{1:N}) &= \int p(\mathbf{X}_n, \mathbf{X}_{n+1} | \mathbf{Y}_{1:N}) d\mathbf{X}_{n+1} \\
&= \int p(\mathbf{X}_n | \mathbf{X}_{n+1}, \mathbf{Y}_{1:n}) p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:N}) d\mathbf{X}_{n+1} \\
&= \int \frac{p_t^{(n)}(\mathbf{X}_{n+1} | \mathbf{X}_n) p(\mathbf{X}_n | \mathbf{Y}_{1:n})}{p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:n})} p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:N}) d\mathbf{X}_{n+1}
\end{aligned} \tag{A.47}$$

$p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:N})$ is from (A.46) and the particle filter returns an approximation of $p(\mathbf{X}_n | \mathbf{Y}_{1:n})$ for each n . For $p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:n})$,

$$\begin{aligned}
p(\mathbf{X}_{n+1} | \mathbf{Y}_{1:n}) &= \int p_t^{(n)}(\mathbf{X}_{n+1} | \mathbf{X}_n) p(\mathbf{X}_n | \mathbf{Y}_{1:n}) d\mathbf{X}_n \\
&\approx \int p_t^{(n)}(\mathbf{X}_{n+1} | \mathbf{X}_n) \sum_{m=1}^M \omega_n^{(m)} 1_{\{\mathbf{X}_n = \mathbf{x}_n^{(m)}\}}(\mathbf{X}_n) d\mathbf{X}_n \\
&= \sum_{m=1}^M \omega_n^{(m)} p_t^{(n)}(\mathbf{X}_{n+1} | \mathbf{X}_n = \mathbf{x}_n^{(m)})
\end{aligned} \tag{A.48}$$

Thus, we have the following approximation again for n :

$$p(\mathbf{X}_n | \mathbf{Y}_{1:N}) \approx \sum_{m=1}^M \hat{\omega}_n^{(m)} 1_{\{\mathbf{X}_n = \mathbf{x}_n^{(m)}\}}(\mathbf{X}_n) \tag{A.49}$$

, where the weights $\hat{\omega}_n^{(m)}$'s are defined as follows:

$$\hat{\omega}_n^{(m)} \propto \omega_n^{(m)} \left(\frac{\sum_{k=1}^M \hat{\omega}_{n+1}^{(k)} p_t^{(n)}(\mathbf{X}_{n+1} = \mathbf{x}_{n+1}^{(k)} | \mathbf{X}_n = \mathbf{x}_n^{(m)})}{\sum_{l=1}^M \omega_n^{(l)} p_t^{(n)}(\mathbf{X}_{n+1} = \mathbf{x}_{n+1}^{(k)} | \mathbf{X}_n = \mathbf{x}_n^{(l)})} \right) \tag{A.50}$$

Note that particles that were sampled in the particle filter are again used in the approximations. Thus, weights $\hat{\omega}_n^{(m)}$'s could be of high variance in practice.

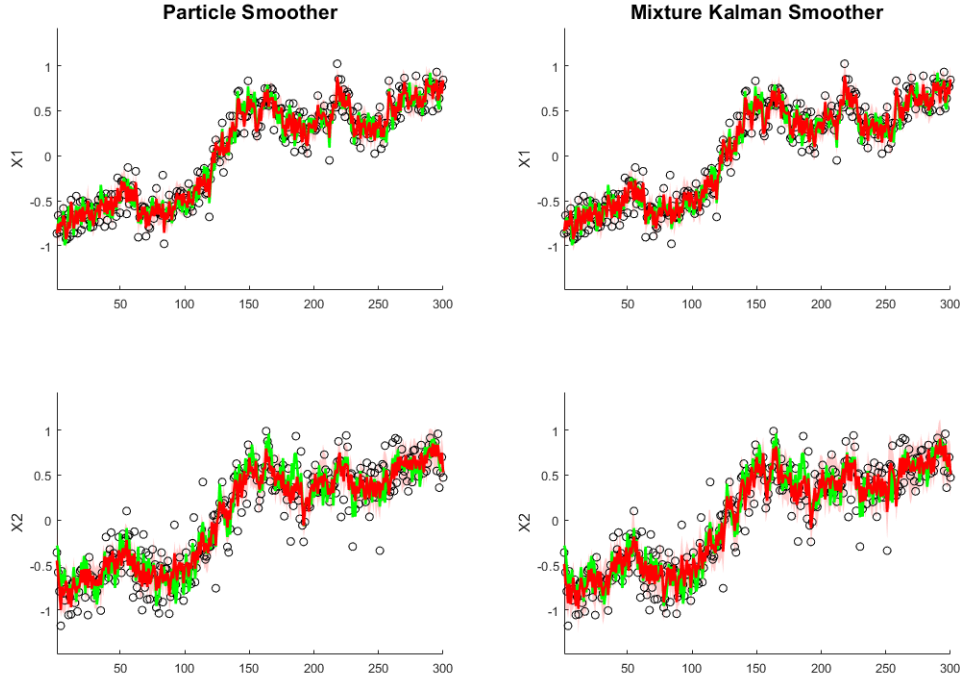


Figure A.6: Examples of the particle smoother and mixture Kalman smoother on the same 2-dimensional dataset. Black circles are the noisy observations and green and red curves are the true and inferred denoised hidden states, respectively.

Backward sampling directly comes from the following:

$$\begin{aligned}
 p(X_n | X_{n+1} = \tilde{x}_{n+1}, Y_{1:N}) &= p(X_n | X_{n+1} = \tilde{x}_{n+1}, Y_{1:n}) \\
 &\propto p_t^{(n)}(X_{n+1} = \tilde{x}_{n+1} | X_n) p(X_n | Y_{1:n}) \\
 &\approx \sum_{m=1}^M \omega_n^{(m)} p_t^{(n)}(X_{n+1} = \tilde{x}_{n+1} | X_n = x_n^{(m)}) 1_{\{X_n = x_n^{(m)}\}}(X_n)
 \end{aligned} \tag{A.51}$$

Thus, given $X_{n+1} = \tilde{x}_{n+1}$, X_n is drawn from the particles $\{x_n^{(m)}\}_{m=1}^M$ with probabilities proportional to $\left\{ \omega_n^{(m)} p_t^{(n)}(X_{n+1} = \tilde{x}_{n+1} | X_n = x_n^{(m)}) \right\}_{m=1}^M$.

Figure A.6 shows examples of the particle smoother and mixture Kalman smoother applied to the same dataset. The performance of particle filtering (and smoother) is highly dependent on the choice of the proposal distribution q . If most of the particles bring 0 to the weights, the approximations on the forward posteriors are eventually controlled by only a small number of particles. Thus, the proposal distribution is often designed as a conditional distribution given

the particles at the previous indices and the associated observations. The resampling step [Murphy (2012)] supplements the algorithmic efficiency. However, all these attempts become to no purpose as the dimensionality of hidden states goes higher. Note that increasing the number of particles soon becomes intractable because the time complexity is quadratic to it, thus is not an ultimate solution. The Equivalent-weights particle filter [Ades & van Leeuwen (2013)] and local particle filter [Penny & Miyoshi (2016); Poterjoy (2016)] challenge such drawbacks.

A.4 MARKOV-CHAIN MONTE CARLO ALGORITHM

Section A.3 discussed about variational approximation methods for filtering and smoothing that can be applied in any transition and emission models. However, the sampling efficiency of particle smoother is still an unresolved challenge and, moreover, they do not return exact posteriors and sample paths but their approximation.

One might consider Markov-chain Monte Carlo (MCMC) [Liu (2001); Peters (2008)] algorithms that generate a series of sample paths from the given distributions (that have been either normalized or not). Under regular conditions, those sample paths converge to the exact ones, though the convergence is slow in practice especially if the initial sample path is in a region of low density. Because hidden states in state-space models are usually highly correlated one another, applying MCMC alone to the state-space models is not enough in general, but instead it can be utilized to supplement the particle filter and smoother for the refinement purpose: initializing sample paths by the particle smoother, and then running a MCMC algorithm from them. To avoid the dependency on the sample paths drawn from the MCMC algorithm, preparing multiple initial sample paths by the particle smoother is recommended.

There are various MCMC algorithms: this section briefly covers three of them, Gibbs sampling (section A.4.1), Metropolis-Hastings algorithm (section A.4.2) and Hamiltonian Monte Carlo algorithm (section A.4.3). The whole section A.4 take the notations in section A.3.

A.4.1 GIBBS SAMPLING

The basic idea of Gibbs sampling algorithm [Geman & Geman (1984)] is to sample each hidden state given all the others in turn. This algorithm is good if sampling each one given the others can be done efficiently while drawing together from the full posterior is not tractable. In mathematical terminology, X_n is drawn from the element-wise posterior $p(X_n | X_{\sim n}, Y_{1:N})$ for each n . Note that one does not need to compute the full posterior one-by-one by considering the structure of state-space models as follows:

$$\begin{aligned} p(X_n | X_{\sim n}, Y_{1:N}) &\propto p(X_n, Y_n, X_{n+1} | X_{1:n-1}, X_{n+2:N}, Y_{\sim n}) \\ &= p_e^{(n)}(Y_n | X_n) p_t^{(n)}(X_{n+1} | X_n) p_t^{(n)}(X_n | X_{n-1}) \end{aligned} \quad (\text{A.52})$$

The Gibbs sampling algorithm starts with an initial path $\tilde{\mathcal{X}} = \tilde{\mathcal{X}}^{(0)} = \tilde{X}_{1:N}^{(0)}$. Then it iterates the following for multiple times expected to cover all $\{1, 2, \dots, N\}$: firstly, sample n from $\{1, 2, \dots, N\}$ randomly, then sample X_n from $p(X_n | \tilde{\mathcal{X}}_{\sim n}, Y_{1:N})$ given $\tilde{\mathcal{X}} = \tilde{\mathcal{X}}^{(t)}$ and replace $\tilde{\mathcal{X}}_n$ with it. After that, define the next $\tilde{\mathcal{X}}^{(t+1)} = \tilde{\mathcal{X}}$ and go back to the first step for obtaining a series $\{\tilde{\mathcal{X}}^{(0)}, \tilde{\mathcal{X}}^{(1)}, \dots, \tilde{\mathcal{X}}^{(T)}\}$ of T until the series has “burned in”. Finally, return $\tilde{\mathcal{X}}^{(T)}$ as a random sample from $p(\mathcal{X} | \mathcal{Y})$ and go back to the very previous step with $\tilde{\mathcal{X}} = \tilde{\mathcal{X}}^{(T)}$ if another sample path is needed.

Note that $p(X_n | X_{\sim n}, Y_{1:N})$ in (A.52) is depending only on the adjacent neighborhoods X_{n+1} , X_{n-1} and Y_n . Thus, one can run the blocked Gibbs sampling that draws $p(X_{\mathcal{E}} | X_{\sim \mathcal{E}}, Y_{1:N})$ simultaneously where $\mathcal{E}$ is a set of either odd or even indices. Though it does not reduce the (theoretical) time complexity, some programming languages, such as MATLAB, are designed to run faster on this blocked version.

However, Gibbs sampling is usually inefficient because (A.52) is not represented explicitly for general transition and emission models. Besides rejection sampling, a practical variant is to fix a finite set of candidate assignments to the hidden state (by discretization) a priori and then to compute (A.52) to those candidates one-by-one for assigning weights used for the discrete sampling:

it is then efficient because storing a table for transition and emission models is now available.

A.4.2 METROPOLIS-HASTINGS ALGORITHM

As mentioned in the previous subsection, Gibbs sampling could be inefficient if exact sampling from the element-wise posteriors is not expressed explicitly. The Metropolis-Hastings algorithm [Metropolis et al. (1953); Hastings (1970); Martino et al. (2015)] would be more efficient in such a case. It additionally requires a proposal distribution q for drawing a candidate (proposal state) that could replace the previous one: defining it to be conditioned on the previously sampled hidden state makes the algorithm more efficient in practice.

The Metropolis-Hastings algorithm starts with an initial path $\tilde{\mathcal{X}} = \tilde{\mathcal{X}}^{(0)} = \tilde{\mathbf{X}}_{1:N}^{(0)}$. The iteration consists of the following steps: it first samples a candidate $\mathcal{X}^*$ from $q(\mathcal{X}|\tilde{\mathcal{X}})$ given $\tilde{\mathcal{X}} = \tilde{\mathcal{X}}^{(t)}$. Then compute the following acceptance ratio α :

$$\alpha = \min \left\{ 1, \frac{p(\mathcal{X}^*|\mathcal{Y})}{p(\tilde{\mathcal{X}}|\mathcal{Y})} \cdot \frac{q(\tilde{\mathcal{X}}|\mathcal{X}^*)}{q(\mathcal{X}^*|\tilde{\mathcal{X}})} \right\} \quad (\text{A.53})$$

Lastly, replace $\tilde{\mathcal{X}}$ with $\mathcal{X}^*$ in probability α (“accept” $\mathcal{X}^*$ in probability α) and define the next $\tilde{\mathcal{X}}^{(t+1)} = \tilde{\mathcal{X}}$.

After repeating the above steps for T times to get a series (chain) $\{\tilde{\mathcal{X}}^{(0)}, \tilde{\mathcal{X}}^{(1)}, \dots, \tilde{\mathcal{X}}^{(T)}\}$ until the series has “burned in”, return $\tilde{\mathcal{X}}^{(T)}$ as the sample drawn from $p(\mathcal{X}|\mathcal{Y})$. Multiple samples are drawn by repeating the above steps.

Under regular conditions on the proposal distribution q , the Markov chain constructed by the Metropolis-Hastings algorithm is irreducible and aperiodic thus has (and converges to) the unique stationary distribution by the Fundamental Theorem of Markov chains [Norris (1997)]. Note that this theorem does not offer the mixing time that determines the number of iterations until the series (chain) has burned in. A practical way is to define q to change (at most) only a small number of hidden states per iteration.

Note that it is not required to compute the ratio of $p(\mathcal{X}^*|\mathcal{Y})$ and $p(\tilde{\mathcal{X}}|\mathcal{Y})$ by computing

them exactly: if q only replaces at most one state per iteration just like Gibbs sampling, then the ratio is given as follows: let X_n is the target to sample.

$$\frac{p(X^*|\mathcal{Y})}{p(\tilde{X}|\mathcal{Y})} = \frac{p_c^{(n)}(Y_n|X_n^*) p_t^{(n)}(X_{n+1}|X_n^*) p_t^{(n)}(X_n^*|X_{n-1})}{p_c^{(n)}(Y_n|\tilde{X}_n) p_t^{(n)}(X_{n+1}|\tilde{X}_n) p_t^{(n)}(\tilde{X}_n|X_{n-1})} \quad (\text{A.54})$$

A similar simplification could also be obtained in computing the ratio of $q(\tilde{X}|X^*)$ and $q(X^*|\tilde{X})$. If q is moreover designed to be symmetric, i.e., $q(\tilde{X}|X^*) \equiv q(X^*|\tilde{X})$, then it is only required to compute (A.54) for obtaining the acceptance ratio α in (A.53). Note that Gibbs sampling is a special case of the Metropolis-Hastings algorithm where $q(\tilde{X}|X^*) = p(\tilde{X}|\mathcal{Y})$, which results in $\alpha \equiv 1$, i.e., always accept the candidate to replace the previous one.

A.4.3 HAMILTONIAN MONTE CARLO ALGORITHM

One drawback of the Metropolis-Hastings algorithm is that the proposal distribution q affects the mixing time. Thus, it is natural to consider about the improvements involved with efficient sampling, or more comprehensively, better chains. The Hamiltonian (Hybrid) Monte Carlo (HMC) algorithm [Duane et al. (1987); Neal (1993); MacKay (2003); Neal (2010); Girolami & Calderhead (2011)] is such an outcome.

An HMC regards the parameters and hidden states as a “particle” and assumes auxiliary variables that represent the momentum of that particle. It updates not only the hidden states (as the models in sections A.4.1 and A.4.2) but also the momentum. Thus, the HMC requires the gradient of the log-posterior density of hidden states, better in the computational efficiency if it is analytically computable. The following steps are referring to [Gelman et al. (2013)].

Step 0. Initialization:

Step 0.1. Fix a mass covariance matrix Σ and learning rate $\varepsilon > 0$.

Step 0.2. Draw the momentum vector $\mathcal{Z}^{(0)}$ from $\mathcal{N}(\vec{0}, \Sigma)$ and the hidden state $\mathcal{X}^{(0)}$.

Step 1. Repeat the following for T times, until the burn-in status is reached:

Step 1.1. Initialize $(\mathcal{Z}^{(t,0)}, \mathcal{X}^{(t,0)}) = (\mathcal{Z}^{(t)}, \mathcal{X}^{(t)})$ and repeat the following ‘leapfrog’ steps for

L times:

Step 1.1.1. Update the momentum vector by following:

$$\mathcal{Z}^{(t,l+0.5)} = \mathcal{Z}^{(t,l)} + \frac{1}{2}\varepsilon \cdot \nabla_{\mathcal{X}} \log p(\mathcal{X}^{(t,l)}|\mathcal{Y}) \quad (\text{A.55})$$

Step 1.1.2. Update the hidden states by following:

$$\mathcal{X}^{(t,l+1)} = \mathcal{X}^{(t,l)} + \varepsilon \cdot \Sigma^{-1} \mathcal{Z}^{(t,l+0.5)} \quad (\text{A.56})$$

Step 1.1.3. Update the momentum vector again by following:

$$\mathcal{Z}^{(t,l+1)} = \mathcal{Z}^{(t,l+0.5)} + \frac{1}{2}\varepsilon \cdot \nabla_{\mathcal{X}} \log p(\mathcal{X}^{(t,l+1)}|\mathcal{Y}) \quad (\text{A.57})$$

Step 1.2. Compute the following acceptance rate r :

$$r = \min \left\{ 1, \frac{p(\mathcal{X}^{(t,L)}|\mathcal{Y}) \mathcal{N}(\mathcal{Z}^{(t,L)}|\vec{0}, \Sigma)}{p(\mathcal{X}^{(t)}|\mathcal{Y}) \mathcal{N}(\mathcal{Z}^{(t)}|\vec{0}, \Sigma)} \right\} \quad (\text{A.58})$$

Step 1.3. Define $(\mathcal{X}^{(t+1)}, \mathcal{Z}^{(t+1)}) = (\mathcal{X}^{(t,L)}, \mathcal{Z}^{(t,L)})$ with probability r .

Otherwise, $(\mathcal{X}^{(t+1)}, \mathcal{Z}^{(t+1)}) = (\mathcal{X}^{(t)}, \mathcal{Z}^{(t)})$.

Step 2. Return $\mathcal{X}^{(T)}$ as the sample path drawn from $p(\mathcal{X}|\mathcal{Y})$.

Step 3. Go back to step 1 until enough sample paths are drawn.

Because $\log p(\mathcal{X}|\mathcal{Y}) = \log p(\mathcal{Y}|\mathcal{X}) + \log p(\mathcal{X}) - \log p(\mathcal{Y})$ and $\nabla_{\mathcal{X}} \log p(\mathcal{Y}) \equiv 0$, one can compute $\nabla_{\mathcal{X}} \log p(\mathcal{X}|\mathcal{Y})$ as follows:

$$\begin{aligned} \nabla_{\mathcal{X}} \log p(\mathcal{X}|\mathcal{Y}) &= \nabla_{\mathcal{X}} \log p(\mathcal{Y}|\mathcal{X}) + \nabla_{\mathcal{X}} \log p(\mathcal{X}) \\ &= \sum_{n=1}^N \nabla_{\mathcal{X}} \log p_e^{(n)}(\mathbf{Y}_n|\mathbf{X}_n) + \sum_{n=1}^N \nabla_{\mathcal{X}} \log p_t^{(n)}(\mathbf{X}_n|\mathbf{X}_{n-1}) \end{aligned} \quad (\text{A.59})$$

Thus, it is not required to compute $\nabla_{\mathcal{X}} \log p(\mathcal{X}|\mathcal{Y})$ from $p(\mathcal{X}|\mathcal{Y})$ directly but just enough

to compute the two sums in the right-hand of (A.59).

If both transition and emission models are analytically differentiable, computing (A.59) can be done efficiently. Also, giving a diagonal mass covariance matrix allows to compute (A.56) fast. [Hoffman & Gelman (2014)] suggests a way of choosing the number of steps L and learning rate ε which the performance is sensitive to automatically.

However, the HMC is a bit limited in practice. It cannot be applied to the state-space model when $\mathcal{X}$ is defined in a discrete space. Strictly positive posteriors are required because it depends on the gradient of log-posteriors. Clearly, it is prohibited if either the transition or emission model is nondifferentiable. For example, the core alignment algorithm in section 2.2 has a non-differentiable transition model thus the HMC cannot be applied to that problem. Despite of such limitation, the HMC is a good variant of the Metropolis-Hastings for dealing with high-dimensional and highly correlated hidden states.

B

Gaussian Process Models

The Gaussian Process (GP) [[Rasmussen & Williams \(2006\)](#)] is a stochastic process where any finite collection of random variables follows a (multivariate) Gaussian distribution. More formally, let $\mathcal{X} = \{\mathbf{X}_t\}_{t \in \mathcal{T}}$ be a stochastic process defined over a set of indices $\mathcal{T}$. Then, $\mathcal{X}$ is a GP if and only if $\mathbf{X}_T \triangleq (\mathbf{X}_{t_n})_{n=1}^N$ is a multivariate Gaussian random variable for any finite subset $T = \{t_n\}_{n=1}^N \subseteq \mathcal{T}$. More constructively, a GP is entirely specified by its mean function μ and covariance function $\mathbb{K}$ that are defined on $\mathcal{T}$ and $\mathcal{T} \times \mathcal{T}$, respectively: $p(\mathbf{X}_T) = \mathcal{N}(\mathbf{X}_T | \mu_T, \mathbb{K}_{TT})$: it is sometimes denoted by $\mathcal{GP}(\mathbf{X}_T | \mu_T, \mathbb{K}_{TT})$, to emphasize that $\mathbf{X}_T$ follows a GP. Rest of the notations are defined in section 1.4.1. The covariance function $\mathbb{K}$ should be not only positive semi-definite and symmetric but also consistently extendable as T varies, thus a set of functions

that satisfy some certain conditions are considered: those are called ‘kernel’ covariance functions. Details can be found in section B.1.

In statistical learning, GPs are popular for nonparametric priors on the continuous random functions. Suppose that we have a set of inputs $\mathcal{X}$ and the associated outputs $\mathcal{Y}$ and want to give a relationship between $\mathcal{X}$ and $\mathcal{Y}$. Instead of constructing a direct function that maps $\mathcal{X}$ to $\mathcal{Y}$, $p(\mathcal{Y}|\mathcal{X})$, here we assume a latent state β that might explain $\mathcal{Y}$ more directly and then consider the relationship between $\mathcal{X}$ and β , $p(\beta|\mathcal{X})$, and $(\mathcal{X}, \beta)$ and $\mathcal{Y}$, $p(\mathcal{Y}|\mathcal{X}, \beta)$. Specifically, β follows a GP with $\mathcal{X}$ as an input: $p(\beta|\mathcal{X}) = \mathcal{GP}(\beta|\mu_{\mathcal{X}}, \mathbb{K}_{\mathcal{X}\mathcal{X}})$. The goal is constructing a function which maps an arbitrary query input x to a distribution of the associated output y given $\mathcal{X}$ and $\mathcal{Y}$ by marginalizing the latent variables out, i.e.,

$$p(y|x, \mathcal{X}, \mathcal{Y}) = \int p(y|x, b) \int p(b|\beta, x, \mathcal{X}) p(\beta|\mathcal{X}, \mathcal{Y}) d\beta db \quad (\text{B.1})$$

By assumption, (β, b) also follows a GP with $(\mathcal{X}, x)$ as an input. Thus, one can easily derive the following:

$$p(b|\beta, x, \mathcal{X}) = \mathcal{N}(b|\mu_x + \mathbb{K}_{x\mathcal{X}}\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1}(\beta - \mu_{\mathcal{X}}), \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}}\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1}\mathbb{K}_{\mathcal{X}x}) \quad (\text{B.2})$$

$p(y|x, b)$ is defined by the assumed relationship $p(\mathcal{Y}|\mathcal{X}, \beta)$, and $p(\beta|\mathcal{X}, \mathcal{Y})$ is a posterior with $p(\beta|\mathcal{X})$ as the prior and $p(\mathcal{Y}|\mathcal{X}, \beta)$ as the likelihood. The characteristics of a GP model are determined by the choice of $p(\mathcal{Y}|\mathcal{X}, \beta)$, i.e., how outputs are explained by not only the associated inputs but also the latent state, and the choice of mean and kernel covariance function that define $p(\beta|\mathcal{X})$, which makes the model very flexible. Sections B.2 to B.5 cover the specific applications of the GP models.

Because the latent state is marginalized out, the final model is depending only on the training data and kernel covariance function, thus overfitting is not caused by the excessive number of parameters. The nonparametric aspect of GPs gives a structure-free inference. Tuning hyperparameters is often a convoluted issue in nonparametric modelling because the objective function

is not given certainly. However, the marginal likelihood (ML) of $\mathcal{Y}$ given $\mathcal{X}$, $p(\mathcal{Y}|\mathcal{X})$, is often chosen as the objective function to maximize in GP models for tuning kernel covariance functions.

There are two principal drawbacks of the GP models. One is that the model is often sub-optimal due to the inflexible kernel covariance function [Remes et al. (2017); Blomqvist et al. (2019)]. The other is about the computational issue: the inference is involved with at least one matrix inversion $\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1}$ of size $|\mathcal{X}| \times |\mathcal{X}|$ thus quickly becomes intractable ($\mathcal{O}(|\mathcal{X}|^{2.373})$) [Le Gall (2014)] as the size of training data $|\mathcal{X}|$ increases. Section B.6 introduces a deep Gaussian process to deal with the first issue and section B.7 shows a variational approximation method to address the second one.

B.1 KERNELS

Kernel covariance functions (kernels) play a key role in the inference based on GPs: roughly speaking, kernels give some prior assumptions on the characteristics such as smoothness, non-stationary, periodicity, and multiscale or hierarchical structures [Gelman et al. (2013)]. Not every function can be a kernel. A formal definition is given as follows: a symmetric function $\mathbb{K} : \mathcal{T} \times \mathcal{T} \rightarrow \mathbb{R}$ is called a (positive semi-definite) kernel on a nonempty set $\mathcal{T}$ if the following holds for any $N \in \mathbb{N}$, $\{t_n\}_{n=1}^N \subseteq \mathcal{T}$ and $\{c_n\}_{n=1}^N \subseteq \mathbb{R}$:

$$\sum_{m=1}^N \sum_{n=1}^N c_m c_n \mathbb{K}(t_m, t_n) \geq 0 \quad (\text{B.3})$$

More intuitively, we can say that the kernel function $\mathbb{K}(u, v)$ measures how much the associated outputs X_u and X_v are correlated. Each kernel is parametrized by a small set of kernel hyperparameters thus ‘tuning’ kernels means to tune involved kernel hyperparameters. The misspecification problem stems from the facts that there are infinitely many kernels and no perfect and objective rule exists for the real data that potentially include outliers, though there have been some publications to deal with it: see [Duvenaud et al. (2013); Ben Abdesslem et al. (2019)] for

examples. Note that any sums and products of kernels are again kernels: sums are straightforward whereas products are by the Mercer's theorem.

Kernels can be defined on any set $\mathcal{T}$ (e.g. string kernels [Lodhi et al. (2002)] and fisher kernels [Jaakkola et al. (2000)]), but here we assume that the kernel functions are computed at vector realizations, i.e., $\mathcal{T} = \mathbb{R}^D$ for some $D \in \mathbb{N}$. The following subsections cover some popular kernels [Rasmussen & Williams (2006); Barber (2012); Duvenaud (2014)].

B.1.1 STATIONARY KERNELS

A stationary kernel is a kernel function $\mathbb{K} : \mathbb{R}^D \times \mathbb{R}^D \rightarrow \mathbb{R}$ that satisfies $\mathbb{K}(u, v) \equiv f(u - v)$ for some function $f : \mathbb{R}^D \rightarrow \mathbb{R}$. Thus, the functional value depends only on the relative difference between the inputs: it implies that the prior assumptions on the characteristics are universal in the domain.

It is natural to start with the most trivial one. A constant kernel $\mathbb{K}_{\text{const}}$ with a kernel hyperparameter σ_0^2 is defined by following:

$$\mathbb{K}_{\text{const}}(u, v) = \sigma_0^2 \cdot \mathbb{1}_{\{u=v\}} \quad (\text{B.4})$$

, which implies that $\mathbb{K}_{\mathcal{X}\mathcal{X}} \equiv \sigma_0^2 \cdot \mathbb{I}$ and $\mathbb{K}_{x\mathcal{X}} \equiv 0$, so the $p(b|\beta, x, \mathcal{X})$ in (B.2) is reduced to the following:

$$p(b|\beta, x, \mathcal{X}) = \mathcal{N}(b|\mu_x, \sigma_0^2) \quad (\text{B.5})$$

Here (B.5) is independent of β and $\mathcal{X}$, so $p(y|x, \mathcal{X}, \mathcal{Y}) = \int p(y|x, b) \mathcal{N}(b|\mu_x, \sigma_0^2) db$: a naïve Bayes modelling. Because any sum of kernels is again a kernel, constant kernels are usually added for the numerical stability and overlapping inputs.

Now let us discuss about less trivial cases. A squared-exponential (SE) kernel $\mathbb{K}_{\text{SE}}$ with kernel hyperparameters σ_0^2 and $l > 0$ has the following form:

$$\mathbb{K}_{\text{SE}}(u, v) = \sigma_0^2 \cdot \exp\left(-\frac{|u - v|^2}{2l^2}\right) \quad (\text{B.6})$$

Here, σ_0^2 determines the average distance of the GP from its mean and the characteristic length-scale l scales the distance $|u - v|$. If $D > 1$, adopting (B.6) is usually suboptimal: the natural extension is the following:

$$\mathbb{K}_{\text{SE}}(u, v) = \sigma_0^2 \cdot \exp\left(-\frac{1}{2} |L(u - v)|^2\right) \quad (\text{B.7})$$

, where L is a $D \times D$ matrix.

The SE kernel makes the associated GP infinitely mean square (MS) differentiable thus the associated GP becomes very smooth, which makes it the most popular kernel in the field, though [Stein (1999)] argues that SE kernels are too smooth to model real physical processes. Note that the SE kernel is also called as the radial basis function (RBF) kernel.

In fact, the above SE kernel is a special case of the Matérn class of kernel functions [Matérn (1986); Stein (1999); Genton (2002)]. A Matérn kernel with kernel hyperparameters $\sigma_0^2, d > 0$ and $l > 0$ has the following form:

$$\mathbb{K}_{\text{Matern}}^{(d)}(u, v) = \sigma_0^2 \cdot \frac{2^{1-d}}{\Gamma(d)} \left(\frac{\sqrt{2d}}{l} |u - v|\right)^d K_d\left(\frac{\sqrt{2d}}{l} |u - v|\right) \quad (\text{B.8})$$

, where K_d is a modified Bessel function [Abramowitz & Stegun (1964)]. Each kernel is specified by d , and it can be shown that $\mathbb{K}_{\text{Matern}}^{(d)}(u, v)$ converges to $\mathbb{K}_{\text{SE}}(u, v)$ as $d \rightarrow \infty$ [Rasmussen & Williams (2006)].

The form in (B.8) can be simplified if $d = n - 1/2$ for $n \in \mathbb{N}$. The examples are following:

1. $d = 1/2$: (Ornstein-Uhlenbeck (OU))

$$\mathbb{K}_{\text{OU}}(u, v) \triangleq \mathbb{K}_{\text{Matern}}^{(1/2)}(u, v) = \sigma_0^2 \cdot \exp\left(-\frac{|u - v|}{l}\right) \quad (\text{B.9})$$

2. $d = 3/2$:

$$\mathbb{K}_{\text{Matern}}^{(3/2)}(u, v) = \sigma_0^2 \cdot \left(1 + \frac{\sqrt{3} |u - v|}{l}\right) \exp\left(-\frac{\sqrt{3} |u - v|}{l}\right) \quad (\text{B.10})$$

3. $d = 5/2$:

$$\mathbb{K}_{\text{Matern}}^{(5/2)}(u, v) = \sigma_0^2 \cdot \left(1 + \frac{\sqrt{5}|u-v|}{l} + \frac{5|u-v|^2}{3l^2} \right) \exp\left(-\frac{\sqrt{5}|u-v|}{l}\right) \quad (\text{B.11})$$

In practice, $\mathbb{K}_{\text{Matern}}^{(5/2)}(u, v)$ is quite similar with the SE kernel, thus $\mathbb{K}_{\text{Matern}}^{(d)}(u, v)$ for $d > 5/2$ is often not considered. A GP with the Matérn kernel $\mathbb{K}_{\text{Matern}}^{(d)}$ is k -times MS differentiable if and only if $k < d$.

As discussed in the SE kernel, a natural extension of the Matérn kernels for $D > 1$ is given as follows for a $D \times D$ matrix L :

$$\mathbb{K}_{\text{Matern}}^{(d)}(u, v) = \sigma_0^2 \cdot \frac{2^{1-d}}{\Gamma(d)} \left(\sqrt{2d} \cdot |L(u-v)| \right)^d K_d \left(\sqrt{2d} \cdot |L(u-v)| \right) \quad (\text{B.12})$$

One noteworthy characteristic of the Matérn kernels is that each of them is in connection with a continuous version of the autoregressive (AR) model [Stein (1999); Papoulis & Pillai (2001); Rasmussen & Williams (2006)] where $\mathbb{K}_{\text{Matern}}^{(d)}$ corresponds to AR $(d + 1/2)$. We will revisit this characteristic in section B.4.

The SE kernel is also in another family of kernels, the γ -exponential (GE) kernels [Schoenberg (1938); Stein (1999)], that with kernel hyperparameters σ_0^2 , $0 < \gamma \leq 2$ and $l > 0$ have the following form:

$$\mathbb{K}_{\text{GE}}^{(\gamma)}(u, v) = \sigma_0^2 \cdot \exp\left(-\left(\frac{|u-v|}{\sqrt{2}l}\right)^\gamma\right) \quad (\text{B.13})$$

Note that $\mathbb{K}_{\text{GE}}^{(2)}(u, v) \equiv \mathbb{K}_{\text{SE}}(u, v)$. One limitation of these kernels is that the corresponding GPs are not MS differentiable unless $\gamma = 2$. A natural extension for $D > 1$ is given as follows for a $D \times D$ matrix L :

$$\mathbb{K}_{\text{GE}}^{(\gamma)}(u, v) = \sigma_0^2 \cdot \exp\left(-\left(\frac{|L(u-v)|}{\sqrt{2}}\right)^\gamma\right) \quad (\text{B.14})$$

The rational quadratic (RQ) kernel [Rasmussen & Williams (2006)] is derived by a scale mixture of SE kernels with a certain gamma prior on the inverse of the square of characteristic length-

scales: now σ_0^2 , $\alpha > 0$ and $l > 0$ are the kernel hyperparameters.

$$\begin{aligned}
\mathbb{K}_{\text{RQ}}^{(\alpha)}(u, v) &= \sigma_0^2 \int \Gamma(\tau | \alpha, \alpha l^2) \exp\left(-\frac{|u - v|^2}{2} \tau\right) d\tau \\
&= \sigma_0^2 \int \frac{(\alpha l^2)^\alpha}{\Gamma(\alpha)} \tau^{\alpha-1} \exp\left(-\alpha l^2 \tau - \frac{|u - v|^2}{2} \tau\right) d\tau \\
&= \sigma_0^2 \left(1 + \frac{|u - v|^2}{2\alpha l^2}\right)^{-\alpha}
\end{aligned} \tag{B.15}$$

Note that $\mathbb{K}_{\text{RQ}}^{(\alpha)}(u, v)$ converges to $\mathbb{K}_{\text{SE}}(u, v)$ as $\alpha \rightarrow \infty$. One advantage of the RQ kernel over the Matérn and GE kernels shown above is that the corresponding GP is now infinitely MS differentiable for every α . One can show that the following variant for a $D \times D$ matrix L is also a valid kernel:

$$\mathbb{K}_{\text{RQ}}^{(\alpha)}(u, v) = \sigma_0^2 \left(1 + \frac{|L(u - v)|^2}{2\alpha}\right)^{-\alpha} \tag{B.16}$$

The kernels that have been covered so far are just stationary. [MacKay (1998)] derives the following periodic kernel which allows to model functions repeating themselves exactly:

$$\mathbb{K}_{\text{Period}}^{(p)}(u, v) = \sigma_0^2 \cdot \exp\left(-\frac{2 \sin^2(\pi |u - v|/p)}{l^2}\right) \tag{B.17}$$

Like the SE kernel, still σ_0^2 determines the average distance of the GP from its mean and the characteristic length-scale l scales the distance $|u - v|$. The period p decides the frequency of function. Note that the above periodic kernel is defined on $\mathbb{R} \times \mathbb{R}$ only. The extension to the multidimensional case is straightforwardly given in [Haji Ghassemi (2013)]:

$$\mathbb{K}_{\text{Period}}^{(p)}(u, v) = \sigma_0^2 \cdot \exp\left(-\sum_{d=1}^D \frac{2 \sin^2(\pi |u_d - v_d|/p)}{l_d^2}\right) \tag{B.18}$$

Figure B.1 shows three functions that are drawn from the GP with each kernel described in this subsection. CONST, SE, OU, M15, M25, GE, RQ and PERIOD stand for the constant kernel, squared-exponential kernel, Ornstein-Uhlenbeck kernel, Matérn kernel with $d = 3/2$,

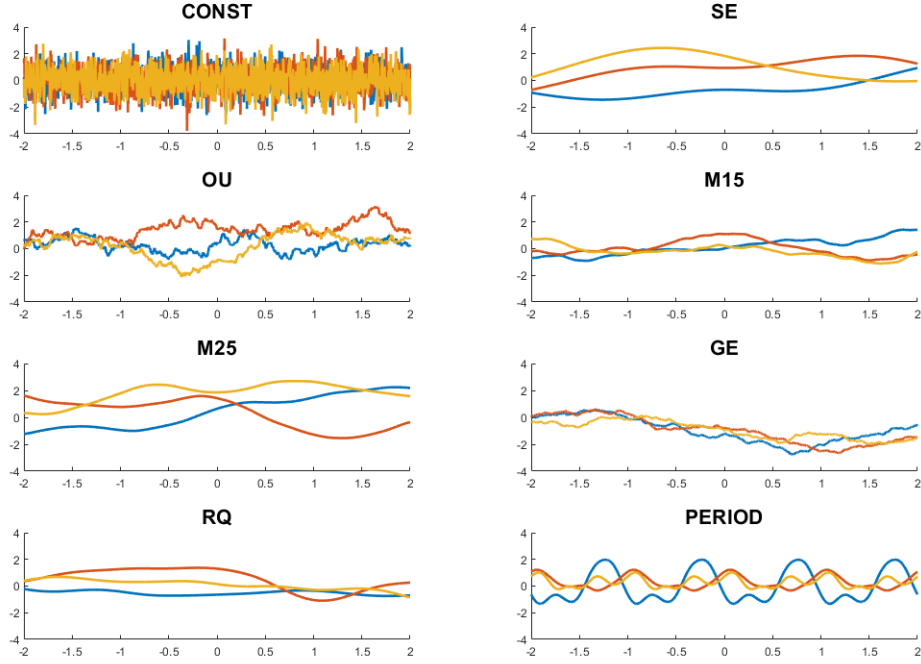


Figure B.1: Functions sampled from the Gaussian process with different stationary kernels.

Matérn kernel with $d = 5/2$, γ -exponential kernel with $\gamma = 1.5$, rational quadratic kernel with $\alpha = 0.5$ and periodic kernel with $p = 1$, respectively: we set $\sigma_0^2 = 1$ and $l = 1$ for each case.

B.1.2 NONSTATIONARY KERNELS

For a given kernel $\mathbb{K}(u, v)$ and a feature function $\varphi : \mathbb{R}^D \rightarrow \mathbb{R}^E$, $\tilde{\mathbb{K}}(u, v) \triangleq \mathbb{K}(\varphi(u), \varphi(v))$ is again a kernel function, by definition. This feature function gives the GP models more flexibility. Nonstationary kernels are often adorned with feature functions because they are not needed to be dependent only on the input difference as the stationary cases. In this section, we assume that φ is an arbitrary feature function.

As we did in section B.1.1, let us start again with a trivial case. A linear kernel with Σ is defined by following:

$$\mathbb{K}_{\text{linear}}(u, v) = \varphi(u)^\top \Sigma \varphi(v) \quad (\text{B.19})$$

, where Σ is a covariance function. Note that the GP regression with $\mathbb{K}_{\text{linear}}(u, v)$ results in a

linear regression with the generalized least squares (GLS) method.

Scaling on the stationary kernels on high-dimensional spaces is done on their input distances. A similar thing can be done in the nonstationary cases. A Gibbs kernel [Gibbs (1997)] is defined by following:

$$\mathbb{K}_{\text{Gibbs}}(u, v) = \prod_{d=1}^D \left(\frac{2l_d(u)l_d(v)}{l_d^2(u) + l_d^2(v)} \right)^{1/2} \cdot \exp \left(- \sum_{d=1}^D \frac{(u_d - v_d)^2}{l_d^2(u) + l_d^2(v)} \right) \quad (\text{B.20})$$

, where each l_d is a positive function on $\mathbb{R}^D$ that scales inputs.

Lastly, neural network kernels are derived as limiting cases of (shallow) neural networks with infinite hidden units with particular activation functions: σ_0^2 and Σ are the kernel hyperparameters.

1. The erf activation function [Williams (1998)]:

$$\mathbb{K}_{\text{NN}}^{\text{erf}}(u, v) = \sigma_0^2 \cdot \sin^{-1} \left(\frac{2\varphi(u)^{\top} \Sigma \varphi(v)}{\sqrt{(1 + 2\varphi(u)^{\top} \Sigma \varphi(u)) (1 + 2\varphi(v)^{\top} \Sigma \varphi(v))}} \right) \quad (\text{B.21})$$

2. The step activation function [Cho & Saul (2009)]:

$$\mathbb{K}_{\text{NN}}^{\text{step}}(u, v) = \sigma_0^2 (\pi - \xi(u, v)) \quad (\text{B.22})$$

3. The rectified linear unit (ReLU) activation function [Cho & Saul (2009)]:

$$\mathbb{K}_{\text{NN}}^{\text{ReLU}}(u, v) = \sigma_0^2 (\sin \xi(u, v) + (\pi - \xi) \cos \xi(u, v)) \quad (\text{B.23})$$

, where:

$$\xi(u, v) \triangleq \cos^{-1} \left(\frac{\varphi(u)^{\top} \Sigma \varphi(v)}{\sqrt{\varphi(u)^{\top} \Sigma \varphi(u) \cdot \varphi(v)^{\top} \Sigma \varphi(v)}} \right) \quad (\text{B.24})$$

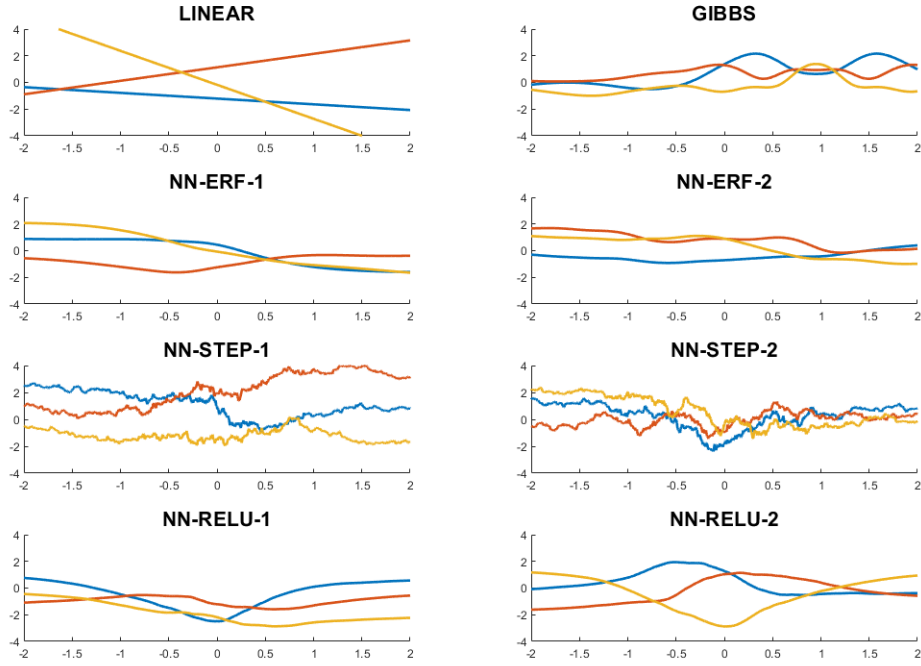


Figure B.2: Functions sampled from the Gaussian process with different nonstationary kernels.

In other words, some shallow neural networks that are involved with a lot of parameters can be converted to the GP models that are controlled by only a few hyperparameters.

Figure B.2 shows three functions that are drawn from the GP with each kernel described in this subsection. LINEAR, GIBBS, NN-ERF, NN-STEP, NN-RELU stand for the linear, Gibbs, neural network kernel with the erf activation function, neural network kernel with the step activation function and neural network kernel with ReLU activation function, respectively: we set $\sigma_0^2 = 1$ for each case. Two panels are assigned to each neural network kernel: the first one assumes $\varphi(u) = (1, u)$ and the second one takes $\varphi(u) = (1, u, u^2)$.

In applications, kernels are often designed by the combinations of the sums and products of the above stationary and nonstationary ones: constructing the “best” kernel for the training data is still an open problem. Choosing a proper feature function is another knotty issue: see [Calandra et al. (2016)].

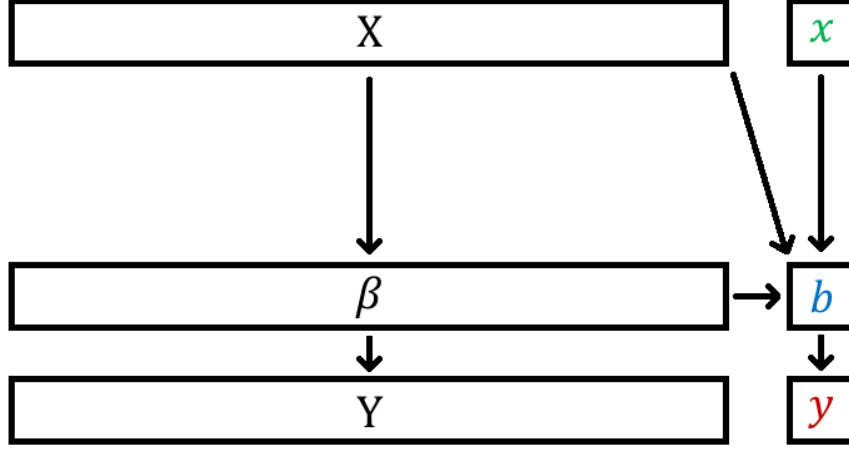


Figure B.3: A graphical representation of GPR and GPC. Capital X and Y are the training inputs and outputs, respectively, and β indicates the latent state, that is either regression (GPR) or log-odd (GPC). Green x and red y are the query input and the associated output, respectively, and blue b stands for the associated latent variable.

B.2 GAUSSIAN PROCESS REGRESSION

Regression aims at finding a simpler model that describes the given data for the explanatory and/or predictive purposes. In particular, statistical regression constructs probabilistic models as the regression functions. Gaussian process regression (GPR) [Rasmussen & Williams (2006)] is a nonparametric regression method that reduces the training data to a probabilistic model: its graphical representation is given in figure B.3.

Let us recall the general setting at the beginning of this chapter: the GP model is determined by $p(\mathcal{Y}|\mathcal{X}, \beta)$. A GPR assumes that the latent state β implies the regression of $\mathcal{Y}$, thus here $p(\mathcal{Y}|\mathcal{X}, \beta)$ is called an observational model. For now, the observational model is assumed to be Gaussian and each $Y_n \in \mathbb{R}$:

$$p(\mathcal{Y}|\mathcal{X}, \beta) = \mathcal{N}(\mathcal{Y}|\beta, \Lambda_{\mathcal{X}}) = \prod_{n=1}^N \mathcal{N}(Y_n|\beta_n, \Lambda_{X_n}) \quad (\text{B.25})$$

$$p(y|x, b) = \mathcal{N}(y|b, \Lambda_x)$$

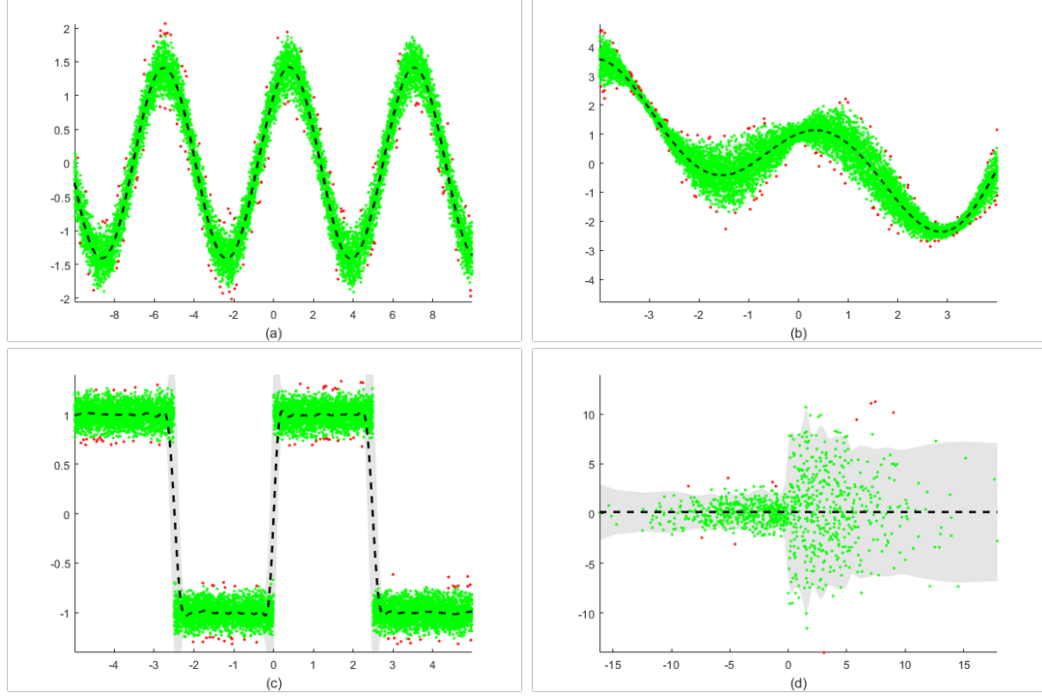


Figure B.4: Examples of the Gaussian process regression. Panel (a) and (c) are the homoscedastic cases whereas (b) and (d) are for the heteroscedastic data.

Then, one can show that the posterior $p(\beta|\mathcal{X}, \mathcal{Y})$ is analytically given as follows:

$$\begin{aligned}
 p(\beta|\mathcal{X}, \mathcal{Y}) &= \mathcal{N}(\beta | \bar{\mu}(\mathcal{X}, \mathcal{Y}), \bar{\Sigma}(\mathcal{X}, \mathcal{Y})) \\
 \bar{\mu}(\mathcal{X}, \mathcal{Y}) &= \mu_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
 \bar{\Sigma}(\mathcal{X}, \mathcal{Y}) &= \mathbb{K}_{\mathcal{X}\mathcal{X}} - \mathbb{K}_{\mathcal{X}\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}}
 \end{aligned} \tag{B.26}$$

By plugging (B.26) together with (B.25) and (B.2) into (B.1), we have the following posterior predictive:

$$\begin{aligned}
 p(y|x, \mathcal{X}, \mathcal{Y}) &= \mathcal{N}(y | \hat{\mu}(x), \Lambda_x + \hat{v}(x)) \\
 \hat{\mu}(x) &= \mu_x + \mathbb{K}_{x\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
 \hat{v}(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}x}
 \end{aligned} \tag{B.27}$$

In other words, a GPR returns the posterior predictive distribution as the regression model in a closed form: we will revisit this form later in section B.7 that discusses the actual tractability of (B.27) in practice.

Let us look deeper into (B.27). The variance term $\Lambda_x + \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}x}$ is decomposed into Λ_x and $\mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}x}$. Λ_x comes from the observational uncertainty in (B.25) while the rest of them stem from the model uncertainty obtained by marginalizing the latent states out: because the latter varies over x unless the kernel is constant (B.4), all GPR is heteroscedastic over inputs in principle. However, the resulting model is often suboptimal if the heteroscedasticity is let depend on $\mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}x}$ only, thus one might consider the observational heteroscedasticity as well. Homoscedastic GPR (section B.2.1) assumes that Λ is given as a constant function; heteroscedastic GPR (section B.2.2), however, regards Λ as a function that could change over the inputs. Figure B.4 gives some examples of the GPR.

There is one golden rule in modeling based on the Gaussian distributions including GP: they are very sensitive to the outliers! In the real data, it is often the case that some of them are just noninformative noises. Because GP is a nonparametric method that depends directly on the data, classifying (and discarding) outliers is indispensable: see [Kaiser et al. (2018); Lee & Lawrence (2019)] for more details. Another way of dealing with outliers is to take non-Gaussian observational models that will be discussed in section B.2.3.

B.2.1 HOMOSCEDASTIC GAUSSIAN PROCESS REGRESSION

A homoscedastic GPR assumes that $\Lambda \equiv \lambda_0^2 \mathbb{I}$ for some constant $\lambda_0 > 0$. Thus, the previous formulations are rewritten as follows:

1. Observational model:

$$\begin{aligned}
 p(\mathcal{Y}|\mathcal{X}, \beta) &= \mathcal{N}(\mathcal{Y}|\beta, \lambda_0^2 \mathbb{I}) = \prod_{n=1}^N \mathcal{N}(Y_n|\beta_n, \lambda_0^2) \\
 p(y|x, b) &= \mathcal{N}(y|b, \lambda_0^2)
 \end{aligned} \tag{B.28}$$

2. Posterior distribution of the latent state:

$$\begin{aligned}
p(\beta|\mathcal{X}, \mathcal{Y}) &= \mathcal{N}(\beta|\bar{\mu}(\mathcal{X}, \mathcal{Y}), \bar{\Sigma}(\mathcal{X}, \mathcal{Y})) \\
\bar{\mu}(\mathcal{X}, \mathcal{Y}) &= \mu_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}} (\lambda_0^2 \mathbb{I} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
\bar{\Sigma}(\mathcal{X}, \mathcal{Y}) &= \mathbb{K}_{\mathcal{X}\mathcal{X}} - \mathbb{K}_{\mathcal{X}\mathcal{X}} (\lambda_0^2 \mathbb{I} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}}
\end{aligned} \tag{B.29}$$

3. Posterior predictive as a regression model:

$$\begin{aligned}
p(y|x, \mathcal{X}, \mathcal{Y}) &= \mathcal{N}(y|\hat{\mu}(x), \lambda_0^2 + \hat{v}(x)) \\
\hat{\mu}(x) &= \mu_x + \mathbb{K}_{x\mathcal{X}} (\lambda_0^2 \mathbb{I} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
\hat{v}(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\lambda_0^2 \mathbb{I} + \mathbb{K}_{\mathcal{X}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}x}
\end{aligned} \tag{B.30}$$

Here, the objective function to be maximized for tuning the kernel hyperparameters is given as the following ML:

$$p(\mathcal{Y}|\mathcal{X}) = \int p(\mathcal{Y}|\mathcal{X}, \beta) p(\beta|\mathcal{X}) d\beta = \mathcal{N}(\mathcal{Y}|\mu_{\mathcal{X}}, \lambda_0^2 \mathbb{I} + \mathbb{K}_{\mathcal{X}\mathcal{X}}) \tag{B.31}$$

Under the assumption that the mean function μ has already been identified, the hyperparameters to tune are λ_0 and those in the kernel $\mathbb{K}$. Gradient descent algorithms on the log-ML are considered for the optimization, or, if only a few and low-dimensional hyperparameters are given, picking the most probable tuple among the finite candidates could be an alternative.

The mean prior function μ reflects a prior belief on the hidden state (the regression here). It could be given as a constant vector (including $\vec{0}$) or parametric function that had been either specified (with known parameters) or not (with parameters to learn). Even if μ is a parametric function to be learned, it can be done by maximizing (B.31).

The leave-one-out cross-validation (LOO-CV) is an alternative to define the tuning procedure. Instead of maximizing the ML in (B.31), it attempts to optimize the LOO log predictive likeli-

hood $\mathcal{L}_{\text{LOO}}$ as follows: details can be found in [Rasmussen & Williams (2006)].

$$\mathcal{L}_{\text{LOO}} = \sum_{n=1}^N \log p(Y_n | \mathcal{X}, \mathcal{Y} \setminus Y_n) \quad (\text{B.32})$$

The ML uses the likelihood of the observations under the assumption that the GP model is correctly specified. Thus, there is nothing strange about the argument that LOO-CV is more robust against model misspecification [Wahba (1990)].

B.2.2 HETEROSCEDASTIC GAUSSIAN PROCESS REGRESSION

Contrary to the homoscedastic GPR, a heteroscedastic GPR assumes the general variance function Λ . If Λ is given a priori, then what remains is to tune kernel hyperparameters with the following objective function to maximize:

$$p(\mathcal{Y} | \mathcal{X}) = \int p(\mathcal{Y} | \mathcal{X}, \beta) p(\beta | \mathcal{X}) d\beta = \mathcal{N}(\mathcal{Y} | \mu_{\mathcal{X}}, \Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{X}}) \quad (\text{B.33})$$

If Λ is also supposed to be learned, however, the problem becomes more complicated: first of all, there is no target variance, unlike the observations to be regressed are given. Defining Λ to be parametric and learning those parameters by maximizing (B.33) would be an answer, but it could be inflexible or misspecified.

Though there is no perfect answer that addresses this problem, a lot of relevant literatures have been published. [Goldberg et al. (1998); Le et al. (2005)] consider the noise itself to be a random variable and estimate it along with the mean of the distribution. [Kersting et al. (2007); Wang & Chen (2016)] independently infer Λ from the sampled variances given the regression model, step-by-step. [Lázaro-Gredilla & Titsias (2011)] parametrizes Λ as the exponential function where the exponent follows a GP prior. [Binois et al. (2018)] assumes a latent variable for constructing an apparatus to jointly infer a spatial field of simulation variances along with the mean-field GP.

[Lee & Lawrence (2019)] focuses on the structure of the observational model: $p(y|x, b) = \mathcal{N}(y|b, \Lambda_x)$ implies that the variance function Λ_x is supposed to be the variance of y given b as

the mean, which is equivalent to the expectation of $(y - b)^2$. Though y and b are not given at an arbitrary query input x , $(y - b)^2$ can be regressed given $(\mathcal{X}, \beta, \mathcal{Y})$ by the following Nadaraya-Watson kernel regression [Nadaraya (1964); Watson (1964); Bierens (1994); Langrené & Warin (2019)]:

$$(y - b)^2 \approx \frac{\sum_{n=1}^N (Y_n - \beta_n)^2 \mathcal{K}_b(x - X_n)}{\sum_{n=1}^N \mathcal{K}_b(x - X_n)} \quad (\text{B.34})$$

, where $\mathcal{K}$ and b are the density kernel and bandwidth parameter, respectively.

One advantage of the GP regression model is that the posterior $p(\beta|\mathcal{X}, \mathcal{Y})$ is explicitly given. Thus, one can compute the following expectation analytically:

$$\begin{aligned} \Lambda_x &\triangleq \mathbb{E}_{\beta|\mathcal{X}, \mathcal{Y}} \left[\frac{\sum_{n=1}^N (Y_n - \beta_n)^2 \mathcal{K}_b(x - X_n)}{\sum_{n=1}^N \mathcal{K}_b(x - X_n)} \right] \\ &= \frac{\sum_{n=1}^N ((Y_n - \hat{\mu}(X_n))^2 + \hat{v}(X_n)) \mathcal{K}_b(x - X_n)}{\sum_{n=1}^N \mathcal{K}_b(x - X_n)} \end{aligned} \quad (\text{B.35})$$

The subsequent problem is that computing each Λ_x requires the values of $\Lambda_{\mathcal{X}}$ at $\mathcal{X} = \{X_n\}_{n=1}^N$ a priori in (B.35). This can be addressed by iterating (B.35) to update $\Lambda_{\mathcal{X}}$ until convergence. The performance of the heteroscedastic GPR is substantially dependent on the bandwidth b : one suggestion is the K-nearest neighbor bandwidth [Loftsgaarden & Quesenberry (1965); Terrell & Scott (1992)] with the LOO-CV for choosing K.

As mentioned at the beginning of the section, tuning hyperparameters is straightforward once $\Lambda_{\mathcal{X}}$ is specified. Thus, the heteroscedastic GP regression can be done by iterating the following steps:

1. Choose b and K given $\mathbb{K}$ and $\Lambda_{\mathcal{X}}$.
2. Update $\Lambda_{\mathcal{X}}$ given b , K and $\mathbb{K}$.
3. Tune $\mathbb{K}$ given $\Lambda_{\mathcal{X}}$.

B.2.3 NON-GAUSSIAN OBSERVATIONAL MODELS

So far, observational models have been assumed to be Gaussian. As discussed earlier in this section, it could be too strong to assume that the standardized residuals $\mathcal{R} \triangleq \{r_n\}_{n=1}^N$ would follow the standard normal distribution in practice, where:

$$r_n \triangleq \frac{Y_n - \hat{\mu}(X_n)}{\sqrt{\Lambda_{X_n} + \hat{v}(X_n)}} \quad (\text{B.36})$$

The model misspecification could be from the multimodality of the data [Lázaro-Gredilla et al. (2012); Bodin et al. (2017); Kaiser et al. (2018)], but here we just focus on the misspecification of the observational models.

Let us go back to the posterior predictive in (B.1): $p(y|x, \mathcal{X}, \mathcal{Y})$ is expressed as a double integration over the hidden states β and b . If the observational model is assumed to be Gaussian, then not only $p(y|x, b)$ but also $p(\beta|\mathcal{X}, \mathcal{Y})$ are also Gaussian, thus computing (B.1) is done explicitly. If not, it is even not guaranteed to get $p(\beta|\mathcal{X}, \mathcal{Y})$ as a known distribution. One of the next best things is to sample a number of β from $p(\beta|\mathcal{X}, \mathcal{Y}) \propto p(\beta|\mathcal{X})p(\mathcal{Y}|\beta, \mathcal{X})$ and approximate (B.1) as follows: if $p(\mathcal{Y}|\beta, \mathcal{X})$ is differentiable with respect to β , then the Hamiltonian Monte Carlo (HMC) algorithm in section A.4.3 can be applied for the sampling. Let $\{\tilde{\beta}^{(m)}\}_{m=1}^M$ be such samples, then:

$$p(y|x, \mathcal{X}, \mathcal{Y}) \approx \frac{1}{M} \sum_{m=1}^M \int p(y|x, b) \mathcal{N}\left(b \middle| \mu_x + \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} (\tilde{\beta}^{(m)} - \mu_{\mathcal{X}}), \hat{v}(x)\right) db \quad (\text{B.37})$$

The other approach corresponds to the approximation of $p(\beta|\mathcal{X}, \mathcal{Y})$ with a known distribution, $q(\beta|\theta)$, where θ is a variational parameter. Similar to the discussion in A.1.4, we might consider the following Kullback-Leibler (KL) divergence:

$$\begin{aligned} \mathbb{D}_{\text{KL}}(q(\cdot|\theta) || p(\cdot|\mathcal{X}, \mathcal{Y})) \\ = \log p(\mathcal{Y}|\mathcal{X}) + \mathbb{D}_{\text{KL}}(q(\cdot|\theta) || p(\cdot|\mathcal{X})) - \mathbb{E}_{q(\beta|\theta)} [\log p(\mathcal{Y}|\beta, \mathcal{X})] \geq 0 \end{aligned} \quad (\text{B.38})$$

Therefore, the log-ML $\log p(\mathcal{Y}|\mathcal{X})$ has the ELBO $\mathcal{L}$ as follows:

$$\log p(\mathcal{Y}|\mathcal{X}) \geq \mathbb{E}_{q(\beta|\theta)} [\log p(\mathcal{Y}|\beta, \mathcal{X})] - \mathbb{D}_{\text{KL}}(q(\cdot|\theta) || p(\cdot|\mathcal{X})) = \mathcal{L} \quad (\text{B.39})$$

Let us assume further that $q(\beta|\theta)$ is given a fully independent Gaussian distribution:

$$q(\beta|\theta) \triangleq \mathcal{N}(\beta | \ddot{\mu}, \text{diag}(\ddot{\sigma}^2)) = \prod_{n=1}^N \mathcal{N}(\beta_n | \ddot{\mu}_n, \ddot{\sigma}_n^2) \quad (\text{B.40})$$

Then, the second term $\mathbb{D}_{\text{KL}}(q(\cdot|\theta) || p(\cdot|\mathcal{X}))$ in (B.39) is explicitly given as follows:

$$\begin{aligned} & \mathbb{D}_{\text{KL}}(q(\cdot|\theta) || p(\cdot|\mathcal{X})) \\ &= \frac{1}{2} \left(\sum_{n=1}^N \frac{\ddot{\sigma}_n^2}{\mathbb{K}_{\mathbf{X}_n \mathbf{X}_n}} + (\mu_{\mathcal{X}} - \ddot{\mu})^T \mathbb{K}_{\mathcal{X} \mathcal{X}}^{-1} (\mu_{\mathcal{X}} - \ddot{\mu}) - 1 + \log |\mathbb{K}_{\mathcal{X} \mathcal{X}}| - \sum_{n=1}^N \log \ddot{\sigma}_n^2 \right) \end{aligned} \quad (\text{B.41})$$

The first term is decomposed as follows:

$$\mathbb{E}_{q(\beta|\theta)} [\log p(\mathcal{Y}|\beta, \mathcal{X})] = \sum_{n=1}^N \mathbb{E}_{\mathcal{N}(\beta_n | \ddot{\mu}_n, \ddot{\sigma}_n^2)} [\log p(Y_n | \beta_n, \mathbf{X}_n)] \quad (\text{B.42})$$

The reparameterization trick [Kingma & Welling (2013)] gives a way of reducing variances in computing the numerical gradient of $\mathbb{E}_{q(\beta|\theta)} [\log p(\mathcal{Y}|\beta, \mathcal{X})]$ with respect to the variational parameters:

$$\nabla_{\theta} \mathbb{E}_{q(\beta|\theta)} [\log p(\mathcal{Y}|\beta, \mathcal{X})] = \sum_{n=1}^N \mathbb{E}_{\mathcal{N}(\varepsilon | 0, 1)} [\nabla_{\theta} \log p(Y_n | \ddot{\mu}_n + \ddot{\sigma}_n \varepsilon, \mathbf{X}_n)] \quad (\text{B.43})$$

Once learning/tuning the variational parameters/kernel hyperparameters with $\mathcal{L}$ as the objective function are done, we can approximate $p(y|x, \mathcal{X}, \mathcal{Y})$ by replacing $p(\beta|\mathcal{X}, \mathcal{Y})$ with $q(\beta|\theta)$ in (B.1):

$$\begin{aligned} p(y|x, \mathcal{X}, \mathcal{Y}) &\approx \int p(y|x, b) \int p(b|\beta, x, \mathcal{X}) q(\beta|\theta) d\beta db \\ &= \int p(y|x, b) \mathcal{N}(b | \tilde{\mu}(x), \tilde{v}(x)) db \end{aligned} \quad (\text{B.44})$$

, where:

$$\begin{aligned}\tilde{\mu}(x) &= \mu_x + \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} (\ddot{\mu} - \mu_{\mathcal{X}}) \\ \tilde{\nu}(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}x} + \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \text{diag}(\ddot{\sigma}^2) \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}x}\end{aligned}\tag{B.45}$$

Note that the above analytic computation is available if q is a (multivariate) Gaussian. (B.40) is just the simplest choice for the variational approximation: we will discuss further about the choice of q in section B.4.

B.3 GAUSSIAN PROCESS CLASSIFICATION

Classification is a function that maps each query input to one (or more) of the discrete categories (labels). Specifically, statistical classification gives a probabilistic model as the classification function. Gaussian process classification (GPC) [Rasmussen & Williams (2006); Barber (2012)] is a nonparametric classification method that assigns a discrete probability distribution to each query input, based on the data. Here we consider the binary classification only (in other words, each $Y_n \in \{-1, +1\}$), though multi-class extension is straightforward [Rasmussen & Williams (2006); Hernández-Lobato et al. (2011); Murphy (2012); Villacampa-Calvo et al. (2020)]. A graphical representation is given in figure B.3.

Contrary to the GPR models in section B.2, a GPC regards the latent state β as a predictor (log-odd) function (over $\mathcal{X}$) that is squashed through the logistic function, so $p(\mathcal{Y}|\mathcal{X}, \beta) = p(\mathcal{Y}|\beta)$ in (B.1) is replaced with the following:

$$\begin{aligned}p(\mathcal{Y}|\beta) &= \prod_{n=1}^N \sigma(Y_n \beta_n) \\ p(y|b) &= \sigma(yb)\end{aligned}\tag{B.46}$$

, whereas we assume the same prior on β , $p(\beta|\mathcal{X}) = \mathcal{GP}(\beta|\vec{0}, \mathbb{K}_{\mathcal{X}\mathcal{X}})$ as usual with $\mu \equiv \vec{0}$ (note that β now represents a log-odd), and $\sigma: \mathbb{R} \rightarrow (0, 1)$ is a logistic function that is also defined as follows:

$$\sigma(t) = \frac{1}{1 + \exp(-t)}\tag{B.47}$$

Note that (B.46) that is non-Gaussian prevents from computing both posterior predictive $p(y|x, \mathcal{X}, \mathcal{Y})$ and posterior $p(\beta|\mathcal{X}, \mathcal{Y})$ analytically. Thus, it is inevitable to consider either Monte Carlo sampling or approximation for accessing $p(\beta|\mathcal{X}, \mathcal{Y})$, as discussed in section B.2.3. In addition, characteristics of the classification bring other approximation methods such as Laplace approximation (section B.3.1) and expectation propagation (section B.3.2). Moreover, we will cover one of the most recent variational approach that uses the variable augmentation technique in section B.3.3.

B.3.1 LAPLACE APPROXIMATION

The Laplace approximation [Williams & Barber (1998)] aims at approximating the posterior $p(\beta|\mathcal{X}, \mathcal{Y})$ that is non-Gaussian with a Gaussian distribution $q(\beta)$ that is obtained by doing a second-order Taylor expansion of $\log p(\beta|\mathcal{X}, \mathcal{Y})$ around its maximizer $\hat{\beta}$:

$$q(\beta) = \mathcal{N}\left(\beta \middle| \hat{\beta}, \left(-\nabla\nabla \log p(\beta|\mathcal{X}, \mathcal{Y})\right)|_{\beta=\hat{\beta}}\right)^{-1} \quad (\text{B.48})$$

Because $\log p(\beta|\mathcal{X}, \mathcal{Y}) = \log p(\mathcal{Y}|\beta) + \log p(\beta|\mathcal{X}) - \log p(\mathcal{Y}|\mathcal{X}) \triangleq \mathcal{L}(\beta) - \log p(\mathcal{Y}|\mathcal{X})$,

$$\begin{aligned} \hat{\beta} &= \underset{\beta}{\operatorname{argmax}} \log p(\beta|\mathcal{X}, \mathcal{Y}) = \underset{\beta}{\operatorname{argmax}} \mathcal{L}(\beta) \\ \nabla\nabla \log p(\beta|\mathcal{X}, \mathcal{Y}) &= \nabla\nabla \mathcal{L}(\beta) \end{aligned} \quad (\text{B.49})$$

Since $Y_n^2 \equiv 1$, we have:

$$\begin{aligned} \mathcal{L}(\beta) &= -\sum_{n=1}^N \log(1 + \exp(-Y_n \beta_n)) - \frac{1}{2} \beta^T \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \beta - \frac{1}{2} \log |\mathbb{K}_{\mathcal{X}\mathcal{X}}| \\ \nabla \mathcal{L}(\beta) &= \nabla \log p(\mathcal{Y}|\beta) - \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \beta \\ \nabla\nabla \mathcal{L}(\beta) &= \nabla\nabla \log p(\mathcal{Y}|\beta) - \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \end{aligned} \quad (\text{B.50})$$

, where:

$$\begin{aligned}\nabla \log p(\mathcal{Y}|\beta) &= \begin{pmatrix} Y_1 (1 + \exp(Y_1 \beta_1))^{-1} \\ Y_2 (1 + \exp(Y_2 \beta_2))^{-1} \\ \vdots \\ Y_N (1 + \exp(Y_N \beta_N))^{-1} \end{pmatrix} \\ \nabla \nabla \log p(\mathcal{Y}|\beta) &= \text{diag} \left(\begin{pmatrix} -\sigma(Y_1 \beta_1) (1 - \sigma(Y_1 \beta_1)) \\ -\sigma(Y_2 \beta_2) (1 - \sigma(Y_2 \beta_2)) \\ \vdots \\ -\sigma(Y_N \beta_N) (1 - \sigma(Y_N \beta_N)) \end{pmatrix} \right)\end{aligned}\tag{B.51}$$

Finding $\hat{\beta}$ can be done numerically by the Newton's method, i.e., iterating the following until convergence:

$$\begin{aligned}\beta^{(t+1)} &\triangleq \beta^{(t)} - \left(\nabla \nabla \mathcal{L}(\beta^{(t)}) \right)^{-1} \left(\nabla \mathcal{L}(\beta^{(t)}) \right) \\ &= \beta^{(t)} + \left(\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} - \nabla \nabla \log p(\mathcal{Y}|\beta)|_{\beta=\beta^{(t)}} \right)^{-1} \left(\nabla \log p(\mathcal{Y}|\beta)|_{\beta=\beta^{(t)}} - \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \beta^{(t)} \right)\end{aligned}\tag{B.52}$$

Once $\hat{\beta}$ is obtained, $q(\beta) = \mathcal{N}(\beta | \hat{\beta}, (\nabla \nabla \mathcal{L}(\hat{\beta}))^{-1})$ is derived from (B.48) and (B.50), and the posterior predictive distribution as the classifier is approximated as follows:

$$\begin{aligned}p(y|x, \mathcal{X}, \mathcal{Y}) &\approx \int p(y|b) \int p(b|\beta, x, \mathcal{X}) q(\beta) d\beta db \\ &= \int \sigma(yb) \mathcal{N}(b | \tilde{\mu}(x), \tilde{v}(x)) db\end{aligned}\tag{B.53}$$

, where:

$$\begin{aligned}\tilde{\mu}(x) &= \mu_x + \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} (\hat{\beta} - \mu_{\mathcal{X}}) \\ \tilde{v}(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}x} + \mathbb{K}_{x\mathcal{X}} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \left(\nabla \nabla \mathcal{L}(\hat{\beta}) \right)^{-1} \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}x}\end{aligned}\tag{B.54}$$

Tuning kernel hyperparameters can be done by maximizing the approximated marginal likelihood with $q(\beta)$: details are in [Barber (2012)].

B.3.2 EXPECTATION PROPAGATION

Just as the Laplace approximation in the previous subsection, the expectation propagation (EP) [Minka (2001)] method also seeks the approximation of $p(\beta|\mathcal{X}, \mathcal{Y})$ with a Gaussian $q(\beta)$. The difference comes from the form of $q(\beta)$: it is designed to be proportional to the product of the GP prior $p(\beta|\mathcal{X})$ and pointwise inferences on β .

The core idea starts from approximating each $p(Y_n|\beta_n)$ with $\mathcal{N}(\beta_n|\tilde{\mu}_n, \tilde{\nu}_n)$ up to constant product of $\tilde{Z}_n$ as follows:

$$p(Y_n|\beta_n) \approx \tilde{Z}_n \mathcal{N}(\beta_n|\tilde{\mu}_n, \tilde{\nu}_n) \quad (\text{B.55})$$

Then define $q(\beta)$ from (B.55) and $p(\beta|\mathcal{X})$ as follows:

$$\begin{aligned} q(\beta) &= \mathcal{N}\left(\beta \middle| \left(\mathbb{I} + \text{diag}(\tilde{\nu}) \mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1}\right)^{-1} \tilde{\mu}, \left(\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1} + \text{diag}(\tilde{\nu})^{-1}\right)^{-1}\right) \\ &\propto p(\beta|\mathcal{X}) \prod_{n=1}^N \mathcal{N}(\beta_n|\tilde{\mu}_n, \tilde{\nu}_n) \end{aligned} \quad (\text{B.56})$$

The inference on $\tilde{\mu}$ and $\tilde{\nu}$ is done by sequential optimization: details can be found in [Rasmussen & Williams (2006)]. Note that defining $p(\mathcal{Y}|\beta)$ with the probit likelihood $\Phi(t) = \int_{-\infty}^t \mathcal{N}(x|0, 1) dx$ gives not only the explicit updating form for $\tilde{\mu}$ and $\tilde{\nu}$ but also the explicit posterior predictive approximation (that is advantageous over the Laplace approximation) as follows:

$$p(y = +1|x, \mathcal{X}, \mathcal{Y}) \approx \Phi\left(\frac{\mathbb{K}_{x\mathcal{X}} (\mathbb{K}_{\mathcal{X}\mathcal{X}} + \text{diag}(\tilde{\nu}))^{-1} \tilde{\mu}}{\sqrt{1 + \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{X}} (\mathbb{K}_{\mathcal{X}\mathcal{X}} + \text{diag}(\tilde{\nu}))^{-1} \mathbb{K}_{\mathcal{X}x}}}\right) \quad (\text{B.57})$$

B.3.3 POLYA-GAMMA DATA AUGMENTATION

The Polya-Gamma data augmentation [Wenzel et al. (2018)] begins with exploiting the following property of the logistic function:

$$\sigma(z) = \frac{1}{1 + \exp(-z)} = \frac{1}{2} \int \exp\left(\frac{z}{2} - \frac{z^2}{2}w\right) \text{PG}(w|1, 0) dw \quad (\text{B.58})$$

, where $\text{PG}(w|1, c)$ is the Polya-Gamma distribution defined by the following moment generating function and proportionality:

$$\begin{aligned}\mathbb{E}_{\text{PG}(w|1,0)} [\exp(-wt)] &= \frac{1}{\cosh(\sqrt{t/2})} \\ \text{PG}(w|1, c) &\propto \exp\left(-\frac{c^2}{2}w\right) \text{PG}(w|1, 0)\end{aligned}\tag{B.59}$$

From (B.58), we have the following proportionality with an augmented state $\mathcal{W} = \mathbf{W}_{1:N}$:

$$p(\mathcal{Y}, \mathcal{W}|\beta, \mathcal{X}) \propto \exp\left(-\frac{1}{2}\mathcal{Y}^\top\beta - \frac{1}{2}\beta^\top \text{diag}(\mathcal{W})\beta\right) \prod_{n=1}^N \text{PG}(\mathbf{W}_n|1, 0)\tag{B.60}$$

The key idea is to approximate $p(\beta, \mathcal{W}|\mathcal{X}, \mathcal{Y}) \propto p(\mathcal{Y}, \mathcal{W}|\beta, \mathcal{X}) p(\beta|\mathcal{X})$ with the following $q(\beta, \mathcal{W}|\mu, \Sigma, \mathcal{C})$ with variational parameters μ, Σ and $\mathcal{C} = c_{1:N}$:

$$q(\beta, \mathcal{W}|\mu, \Sigma, \mathcal{C}) \triangleq \mathcal{N}(\beta|\mu, \Sigma) \prod_{n=1}^N \text{PG}(\mathbf{W}_n|1, c_n)\tag{B.61}$$

The KL divergence $\mathbb{D}_{\text{KL}}(q(\cdot|\mu, \Sigma, \mathcal{C})||p(\cdot|\mathcal{X}, \mathcal{Y}))$ results in the following ELBO $\mathcal{L}$:

$$\begin{aligned}\log p(\mathcal{Y}|\mathcal{X}) &\geq \mathcal{L} \\ &\triangleq \mathbb{E}_{q(\beta, \mathcal{W}|\mu, \Sigma, \mathcal{C})} [\log p(\mathcal{Y}|\beta, \mathcal{W}, \mathcal{X})] - \mathbb{D}_{\text{KL}}\left(q(\cdot|\mu, \Sigma, \mathcal{C}) \left\| p(\cdot|\mathcal{X}) \prod_{n=1}^N \text{PG}(\mathbf{W}_n|1, 0)\right.\right)\end{aligned}\tag{B.62}$$

The variational parameters are learned by maximizing $\mathcal{L}$. One advantage of this approach is that $\mathcal{L}$ is given in a closed form and $\mathcal{C} = c_{1:N}$ can be updated analytically: see [Wenzel et al. (2018)] for more details. Once variational parameters μ, Σ and $\mathcal{C}$ are learned, (B.61) gives the marginal $q(\beta|\mu, \Sigma, \mathcal{C}) = \mathcal{N}(\beta|\mu, \Sigma)$ directly. What remains is to approximate the posterior predictive by replacing $p(\beta|\mathcal{X}, \mathcal{Y})$ with q , as done in section B.3.1.

Figure B.5 presents two examples of the GP classification based on the Polya-Gamma data augmentation. The left panel uses 10,000 observations whereas the right one has 1,000. Note

that the right example results in less certain posteriors than the left one: this suggests an advantage of the nonparametric classification algorithm that the inference directly depends on the quality, quantity and sparsity of the data.

B.4 GAUSSIAN PROCESS STATE-SPACE MODELS

State-space models in chapter A assume the Markov property and well-established transition models. Those are too strong assumptions to cover all real problems: such a ‘memoryless’ property has potential of losing too many information from the previously observed states and observations would not be made regularly. Recurrent neural networks (RNNs) [Goodfellow et al. (2016)] straightforwardly address the first case, however, usually require too many parameters thus would not be appropriate if only a small number of observations are given.

Gaussian process state-space models (GPSTs) [Frigola et al. (2014); Eleftheriadis et al. (2017)] are good alternatives to those models that depend on the Markov assumption. It gives a GP prior on the series of each predictor: a typical graphical representation is described in figure B.6. In this section, we follow the same notations in chapter A. Suppose that the hidden states $\mathcal{X} = \{\mathbf{X}_n\}_{n=1}^N$ are associated with the pair of indices and outputs $(\mathcal{T}, \mathcal{Y})$, and each $\mathbf{X}_n \in \mathbb{R}^D$ for $D \in \mathbb{N}$. Then we could consider $\mathcal{X}$ as a $N \times D$ matrix where each row indicates $\mathbf{X}_n$ and the following is assumed for each $d = 1, 2, \dots, D$:

$$p(\mathbf{X}_{:,d}|\mathcal{T}) = \mathcal{GP}\left(\mathbf{X}_{:,d} \middle| \boldsymbol{\mu}_{\mathcal{T}}^{(d)}, \mathbb{K}_{\mathcal{T}\mathcal{T}}^{(d)}\right) \quad (\text{B.63})$$

Just as the other state-space models in chapter A, $\mathcal{Y}$ is also defined on a high-dimensional space $\mathbb{R}^L$ and the emission model is given as $p_e(Y_n|T_n, \mathbf{X}_n)$ for each n : now the possible nonstationary characteristic is controlled by the associated index (time in common) T_n .

Thus, the prior (B.64) and likelihood (B.65) are given as follows:

$$p(\mathcal{X}|\mathcal{T}) = \prod_{d=1}^D p(\mathbf{X}_{:,d}|\mathcal{T}) = \prod_{d=1}^D \mathcal{GP}\left(\mathbf{X}_{:,d} \middle| \boldsymbol{\mu}_{\mathcal{T}}^{(d)}, \mathbb{K}_{\mathcal{T}\mathcal{T}}^{(d)}\right) \quad (\text{B.64})$$

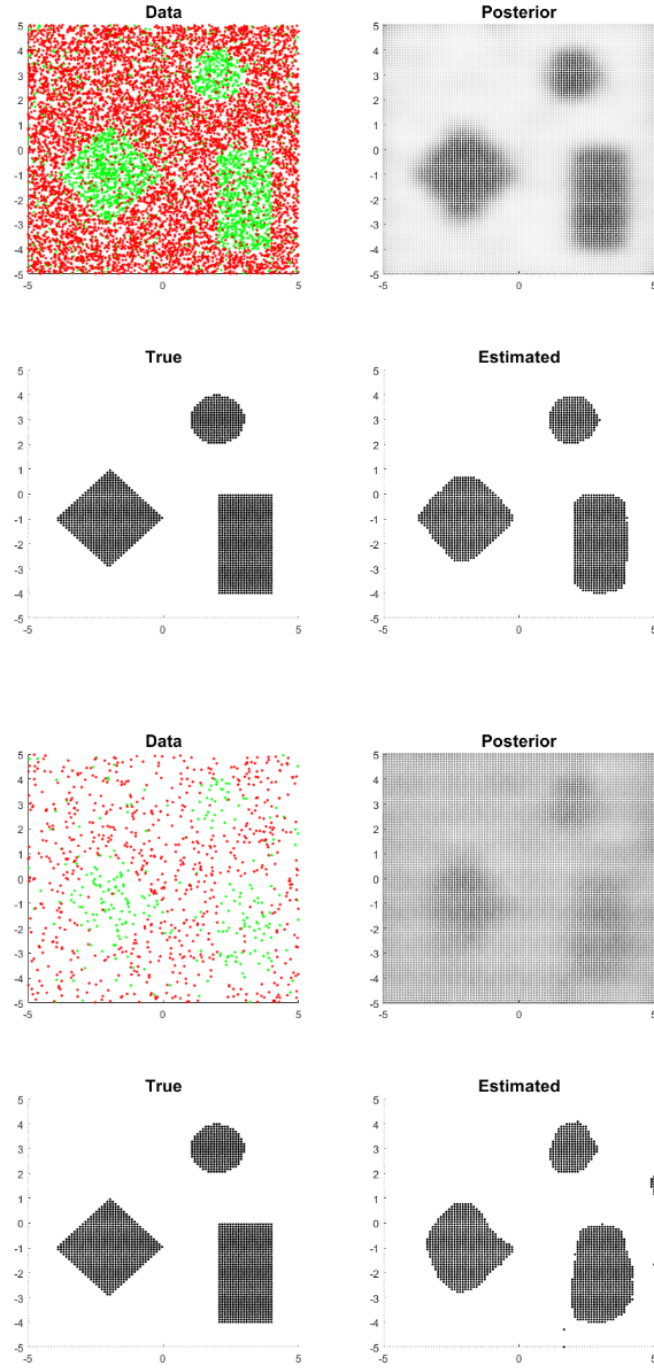


Figure B.5: Examples of the GP classification. Each panel consists of four sub-panels. The first sub-panel show the artificial data that are drawn from the distribution depicted in the third sub-panel with 10% contamination. The second and fourth sub-panels show the posterior distributions of the labels with representing their levels with the brightness and the corresponding estimation, respectively.

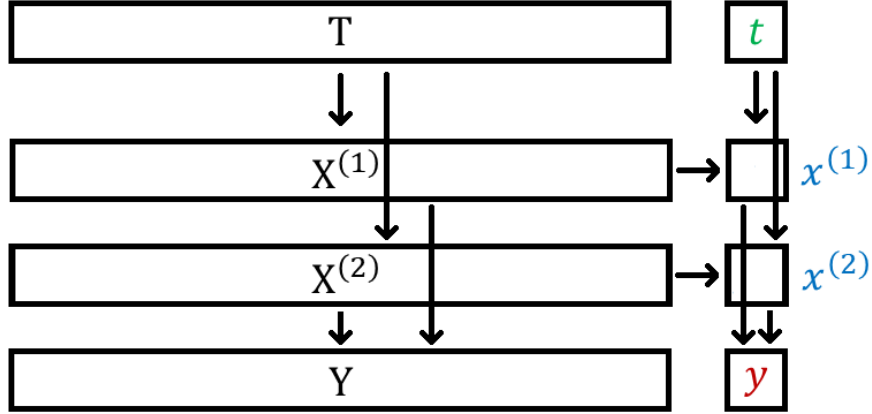


Figure B.6: A graphical representation of GPST. Capital T and Y are the training indices (time) and outputs, respectively, and $X^{(1)}$ and $X^{(2)}$ indicate the latent states, that are hidden states to infer as those of the state-space models in chapter A. Green t and red y are the query index (time) and the associated output, respectively, and blue $x^{(1)}$ and $x^{(2)}$ stand for the associated latent variables.

$$p(\mathcal{Y}|\mathcal{T}, \mathcal{X}) = \prod_{n=1}^N p_e(Y_n|T_n, X_n) \quad (\text{B.65})$$

From above, the posterior of $\mathcal{X}$ is given as follows:

$$p(\mathcal{X}|\mathcal{T}, \mathcal{Y}) \propto p(\mathcal{X}|\mathcal{T})p(\mathcal{Y}|\mathcal{T}, \mathcal{X}) \quad (\text{B.66})$$

In state-space models, the major concern is to compute the prediction of the hidden state x , $p(x|t, \mathcal{T}, \mathcal{Y})$, at an arbitrary query index t .

$$p(x|t, \mathcal{T}, \mathcal{X}, \mathcal{Y}) = \int p(x|t, \mathcal{T}, \mathcal{X})p(\mathcal{X}|\mathcal{T}, \mathcal{Y}) d\mathcal{X} \quad (\text{B.67})$$

, where $p(x|t, \mathcal{T}, \mathcal{X})$ is derived from the assumption (B.64):

$$\begin{aligned} & p(x|t, \mathcal{T}, \mathcal{X}) \\ &= \prod_{d=1}^D \mathcal{N}\left(x_d \middle| \mu_t^{(d)} + \mathbb{K}_{t\mathcal{T}}^{(d)} \left(\mathbb{K}_{\mathcal{T}\mathcal{T}}^{(d)}\right)^{-1} \left(\mathbf{X}_{:,d} - \mu_{\mathcal{T}}^{(d)}\right), \mathbb{K}_{tt}^{(d)} - \mathbb{K}_{t\mathcal{T}}^{(d)} \left(\mathbb{K}_{\mathcal{T}\mathcal{T}}^{(d)}\right)^{-1} \mathbb{K}_{\mathcal{T}t}^{(d)}\right) \end{aligned} \quad (\text{B.68})$$

The choice of kernel function $\mathbb{K}$ characterizes the underlying dynamics of the hidden states. As discussed in section B.1.1, Matérn kernels correspond to the continuous version of AR models. This property is quite useful. Suppose that one wants to design a state-space model for a given time series under the assumption that the series of hidden states follows AR (2) and in addition the indices in the time series are not regular: it is quite common in paleoceanography and paleoclimatology, as discussed in section 1.1.3. Rearrangement of the time series by interpolation would be required to fit the data into the parametric framework, but the GPST with a kernel $\mathbb{K}_{\text{Matern}}^{(3/2)}$ is working properly regardless of the index regularity. Moreover, it does not need the Markovian assumption whereas parametric state-space models in chapter A are intractable without it. Lastly, GPST is nonparametric thus can be called data-driven AR model, i.e., we can easily solve the problem that would be picky in the parametric setting by moving it to the nonparametric world.

In most of the cases, however, $p(\mathcal{X}|\mathcal{T}, \mathcal{Y})$ is not given in a closed form. The variational inference discussed in section A.1.4 is still valid in this case: infer the tractable approximation $q(\mathcal{X}|\theta)$ together with kernel hyperparameters that maximize the ELBO $\mathcal{L}$:

$$\log p(\mathcal{Y}|\mathcal{T}) \geq \mathbb{E}_{q(\mathcal{X}|\theta)} [\log p(\mathcal{Y}|\mathcal{T}, \mathcal{X})] - \mathbb{D}_{\text{KL}}(q(\cdot|\theta) \| p(\cdot|\mathcal{T})) \quad (\text{B.69})$$

Once q is learned, $p(x|t, \mathcal{T}, \mathcal{X}, \mathcal{Y})$ in (B.67) is approximated by following:

$$p(x|t, \mathcal{T}, \mathcal{X}, \mathcal{Y}) \approx \int p(x|t, \mathcal{T}, \mathcal{X}) q(\mathcal{X}|\theta) d\mathcal{X} \quad (\text{B.70})$$

This section ends with the discussion about the form of $q(\mathcal{X}|\theta)$. Because $\mathcal{X}$ is not a vector here, roughly we have the following ways of defining it:

1. Markov chain: it is good for sampling paths.

$$q(\mathcal{X}|\theta) = q(\mathbf{X}_1|\theta) \prod_{n=2}^N \pi(\mathbf{X}_n|\mathbf{X}_{n-1}, \theta_n) \quad (\text{B.71})$$

2. Predictor-wise model: we will see this structure again in section B.6.

$$q(\mathcal{X}|\theta) = \prod_{d=1}^D q\left(\mathbf{X}_{:,d} \middle| \theta^{(d)}\right) \quad (\text{B.72})$$

3. The full Gaussian model: not practical if $p(\mathcal{X}|\mathcal{T}, \mathcal{Y})$ is multimodal.

$$q(\mathcal{X}|\theta) = \mathcal{N}\left(\begin{pmatrix} \mathbf{X}_{:,1} \\ \vdots \\ \mathbf{X}_{:,D} \end{pmatrix} \middle| \mu, \Sigma\right) \quad (\text{B.73})$$

Recently, [Eleftheriadis et al. (2017)] considers the RNN models for modelling q . With the reparameterization trick, the gradient of $\mathcal{L}$ is efficiently approximated thus variational parameters can be learned numerically regardless of the parametrization. If q is Gaussian, then the second term $\mathbb{D}_{\text{KL}}(q(\cdot|\theta)||p(\cdot|\mathcal{T}))$ is expressed in a closed form.

B.5 GAUSSIAN PROCESS LATENT VARIABLE MODELS

Dimensionality reduction is essential for dealing with the curse of dimensionality in the preprocessing step before applying nonparametric models. The Gaussian process latent variable model (GPLVM) [Lawrence (2005); Damianou et al. (2016)] designs a nonparametric dimensionality reduction model based on GP: [Titsias & Lawrence (2010)] revisit the GPLVM to show that a variational inference framework can marginalize the input variables. More extensions can be found in [Li & Chen (2016)]. However, in this subsection we just cover the original model which aims at optimizing the latent inputs. The below formulations are referred to [Lawrence (2005); Murphy (2012)].

Unlike the settings so far, only outputs $\mathcal{Y} = \{\mathbf{Y}_n\}_{n=1}^N \subseteq \mathbb{R}^L$ are given here. The goal is to find $\mathcal{X} = \{\mathbf{X}_n\}_{n=1}^N \subseteq \mathbb{R}^D$ as the low dimension reduction of $\mathcal{Y}$ where $D \ll L$. Thus, it is an unsupervised learning. To be specific, we have the following relationship with a $L \times D$ matrix $\mathbf{W}$

and a positive constant σ_0 :

$$p(Y_n|X_n, W, \sigma_0^2) = \mathcal{N}(Y_n|WX_n, \sigma_0^2\mathbb{I}) \quad (\text{B.74})$$

The probabilistic principal component analysis (PPCA) [Tipping & Bishop (1999)] assumes a prior $p(X_n) = \mathcal{N}(X_n|\vec{0}, \mathbb{I})$ on X_n to get the marginal likelihood $p(Y_n|W, \sigma_0^2)$ and infer W that maximizes the following likelihood:

$$\begin{aligned} p(\mathcal{Y}|W, \sigma_0^2) &= \prod_{n=1}^N p(Y_n|W, \sigma_0^2) \\ &= \prod_{n=1}^N \int p(Y_n|X_n, W, \sigma_0^2) p(X_n) dX_n \\ &= \prod_{n=1}^N \mathcal{N}(Y_n|\vec{0}, \sigma_0^2\mathbb{I} + W^T W) \end{aligned} \quad (\text{B.75})$$

The GPLVM starts from the dual problem of PPCA, that is, to assume a prior $p(W_{l,:}^T) = \mathcal{N}(W_{l,:}^T|\vec{0}, \alpha^{-1}\mathbb{I})$ with a positive constant $\alpha > 0$ for each $l = 1, 2, \dots, L$ and obtain the marginal likelihood $p(\mathcal{Y}|\mathcal{X}, \sigma_0^2, \alpha)$ that has the following form: let us regard Y as a $N \times L$ matrix where each row indicates Y_n and X as a $N \times D$ matrix where each row indicates X_n .

$$\begin{aligned} p(\mathcal{Y}|\mathcal{X}, \sigma_0^2, \alpha) &= \int \prod_{l=1}^L p(W_{l,:}^T) \prod_{n=1}^N p(Y_n|X_n, W, \sigma_0^2) dW \\ &= \int \prod_{l=1}^L \mathcal{N}(W_{l,:}^T|\vec{0}, \alpha^{-1}\mathbb{I}) \prod_{n=1}^N \mathcal{N}(Y_n|WX_n, \sigma_0^2\mathbb{I}) dW \\ &= \int \prod_{l=1}^L \mathcal{N}(W_{l,:}^T|\vec{0}, \alpha^{-1}\mathbb{I}) \prod_{l=1}^L \mathcal{N}(Y_{:,l}|XW_{l,:}^T, \sigma_0^2\mathbb{I}) dW \\ &= \prod_{l=1}^L \mathcal{N}(Y_{:,l}|\vec{0}, \sigma_0^2\mathbb{I} + \alpha XX^T) \end{aligned} \quad (\text{B.76})$$

Note that $\sigma_0^2\mathbb{I} + \alpha XX^T = \mathbb{K}_{\mathcal{X}\mathcal{X}}$ for the sum of constant and linear kernels $\mathbb{K} = \mathbb{K}_{\text{const}} + \mathbb{K}_{\text{linear}}$ in (B.4) and (B.19) with $\Sigma = \alpha\mathbb{I}$ and $\varphi(u) \equiv u$: this approach is called the kernel trick [Theodor-

idis & Koutroumbas (2008)]. By perceiving (B.76) as a special case where the linear kernel is chosen, the GPLVM has the general objective function for an arbitrary kernel $\mathbb{K}$ as follows:

$$p(\mathcal{Y}|\mathcal{X}) = \prod_{l=1}^L \mathcal{N}(\mathbf{Y}_{:,l}|\vec{0}, \mathbb{K}_{\mathcal{X}\mathcal{X}}) \propto |\mathbb{K}_{\mathcal{X}\mathcal{X}}|^{-L/2} \exp\left(-\frac{1}{2}\text{trace}(\mathbb{K}_{\mathcal{X}\mathcal{X}}^{-1}\mathcal{Y}\mathcal{Y}^T)\right) \quad (\text{B.77})$$

Now the goal is to find $\mathcal{X}$ and kernel hyperparameters of $\mathbb{K}$ that maximize (B.77) as the objective function: the maximizer $\mathcal{X}$ is then regarded as the low dimension reduction of $\mathcal{Y}$. It can be done by a gradient descent on the logarithm of (B.77) whereas (B.76) is maximized explicitly via eigenvalue methods.

B.6 DEEP GAUSSIAN PROCESS

As discussed at the top of this section, the GP models are often suboptimal from the kernel misspecification. Because there is no gold rule for finding the optimal kernel to the given data, one cannot avoid the following dilemma: making kernels complicated has a potential to make them more misspecified whereas choosing simple kernels is itself misspecified by their simple structures.

Regression (section B.2) and classification (section B.3) models adopt a column vector (let us call it as a one-dimensional hidden layer) as β while state-space models (section B.4) allow it to have a matrix form (a multi-dimensional hidden layer). Regardless of the associations, however, each model covered so far assumes only one hidden state β in the model (call them ‘shallow’ DP models) If the model with one hidden layer is working, why is the model not to have more than one?

Deep Gaussian process (DGP) [Damianou & Lawrence (2013)] challenges such a drawback by stacking multiple multi-dimensional hidden layers: figure B.7 shows its example graphical representation. Each hidden layer is stacked with the very previous one as the input. For simplicity, suppose that in the DPG model there are two hidden layers, $N \times H_1$ and $N \times H_2$ matrices β and γ , with $\mathcal{X} = \{\mathbf{X}_n\}_{n=1}^N \subseteq \mathbb{R}^D$ and $\mathcal{Y} = \{\mathbf{Y}_n\}_{n=1}^N \subseteq \mathbb{R}^L$ as the input and output of the training

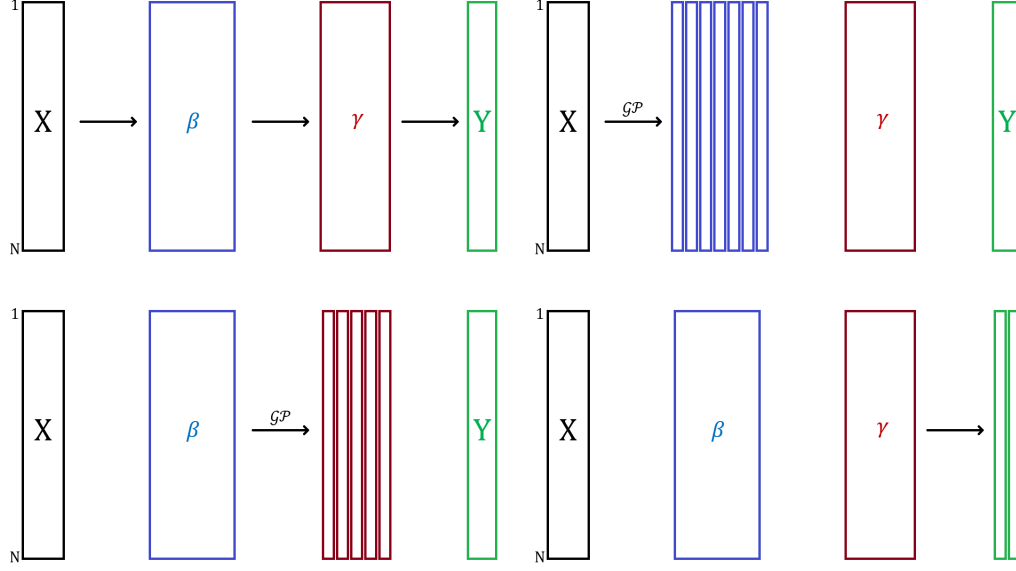


Figure B.7: A graphical representation of DGP. X and Y are the training inputs and outputs, respectively, and β and γ are the latent states (upper-left). Each column vector of β is assumed to follow a GP with X as the input (upper-right). Each column vector of γ is then assumed to follow another GP with β as the input (lower-left). Finally, Y is modelled given γ (lower-right).

data, as usual. Then, the priors on β and γ are given as follows:

$$\begin{aligned}
 p(\beta|\mathcal{X}) &= \prod_{b=1}^{H_1} \mathcal{GP}(\beta_{:,b} | \vec{0}, \mathbb{K}_{\mathcal{X}\mathcal{X}}^{(1)}) \\
 p(\gamma|\beta) &= \prod_{b=1}^{H_2} \mathcal{GP}(\gamma_{:,b} | \vec{0}, \mathbb{K}_{\beta\beta}^{(2)})
 \end{aligned} \tag{B.78}$$

With the observational model $p(\mathcal{Y}|\gamma)$, the posterior of hidden layers and posterior predictive (B.79) are given as follows:

$$\begin{aligned}
 p(\beta, \gamma | \mathcal{X}, \mathcal{Y}) &\propto p(\mathcal{Y}|\gamma) p(\gamma|\beta) p(\beta|\mathcal{X}) \\
 p(y|x, \mathcal{X}, \mathcal{Y}) &= \int p(y|c) \iiint p(c|\gamma, b, \beta) p(b|\beta, x, \mathcal{X}) p(\beta, \gamma | \mathcal{X}, \mathcal{Y}) d\beta d\gamma db dc
 \end{aligned} \tag{B.79}$$

Due to the nonlinearity of β in (B.78), the forms (B.79) cannot be expressed in closed forms. We instead consider the variational inference. Let $q(\beta, \gamma) \triangleq q_1(\beta) q_2(\gamma)$ the variation distribution to approximate $p(\beta, \gamma | \mathcal{X}, \mathcal{Y})$. Then, the following ELBO $\mathcal{L}$ is derived from the nonnegative

KL divergence $\mathbb{D}_{\text{KL}}(q||p(\cdot|\mathcal{X}, \mathcal{Y}))$:

$$\log p(\mathcal{Y}|\mathcal{X}) \geq \mathbb{E}_{q_2(\gamma)} [\log p(\mathcal{Y}|\gamma)] - \mathbb{E}_{q_1(\beta)} [\mathbb{D}_{\text{KL}}(q_2||p(\cdot|\beta))] - \mathbb{D}_{\text{KL}}(q_1||p(\cdot|\mathcal{X})) = \mathcal{L} \quad (\text{B.80})$$

Note that the KL divergences in (B.80) can be expressed in closed forms if q_1 and q_2 are given to be Gaussian. Once KL divergence are given explicitly, the reparameterization trick assists in computing the numerical gradient of $\mathcal{L}$. The variational distribution q can replace $p(\beta, \gamma|\mathcal{X}, \mathcal{Y})$ in (B.79) for the approximation after learning.

Details of the straightforward extension of the above formulations to the cases where the number of hidden layers is bigger than 2 with more complicated variational distributions can be found in [Damianou & Lawrence (2013); Salimbeni & Deisenroth (2017)]: they also deal with the cases where the input $\mathcal{X}$ is also assumed to be latent.

B.7 VARIATIONAL INFERENCE

So far, what we have covered are the exact modelling, i.e., all arithmetic operations are assumed to be feasible. In practice, however, matrix inversion quickly becomes intractable as the size of the matrix increases, as discussed earlier. Thus, the GP models have been limited and moreover often misunderstood as an anachronism in the era of big data.

Various attempts have been made to process big data with GP-based models: [Liu et al. (2020)] reviews them briefly. This subsection summarizes one of them, the variational free energy (VFE) proposed by [Titsias (2009)], as a variational approach. Note that all GP models that have been introduced in sections B.2 to B.4 can be supplemented by this VFE technique, thus modifiable to the tractable models, regardless of the data size: most of the cited publications assign parts to demonstrate such extensions. The reason why such extensions have been intentionally ignored in the previous subsections is to avoid the distraction that prevents the readers from grasping the essence of the GP models.

The core idea of VFE is assuming additional pair of pseudo-inputs $\mathcal{Z}$ and the associated latent

state γ that follows a GP with an input $\mathcal{Z}$ where γ is the sufficient statistic for β . With that pair, we have the following extended formulation of the posterior predictive:

$$\begin{aligned}
p(y|x, \mathcal{X}, \mathcal{Y}) &= \int p(y|x, b) \iint p(b|\beta, \gamma, x, \mathcal{X}, \mathcal{Z}) p(\beta|\gamma, \mathcal{X}, \mathcal{Z}) p(\gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y}) d\gamma d\beta db \\
&= \int p(y|x, b) \iint p(b|\gamma, x, \mathcal{Z}) p(\beta|\gamma, \mathcal{X}, \mathcal{Z}) p(\gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y}) d\gamma d\beta db \\
&= \int p(y|x, b) \int p(b|\gamma, x, \mathcal{Z}) p(\gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y}) d\gamma db
\end{aligned} \tag{B.81}$$

Now the problem is converted to computing $p(\gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y})$ that is also an uneasy work. One resolution is to utilize the variational distribution $q(\beta, \gamma|\mathcal{X}, \mathcal{Z}) = p(\beta|\gamma, \mathcal{X}, \mathcal{Z}) \varphi(\gamma)$ that approximates $p(\beta, \gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y})$. The following ELBO is then derived from the nonnegative KL divergence $\mathbb{D}_{\text{KL}}(q(\cdot|\mathcal{X}, \mathcal{Z})||p(\cdot|\mathcal{X}, \mathcal{Z}, \mathcal{Y}))$:

$$\begin{aligned}
\log p(\mathcal{Y}|\mathcal{X}) &\geq \int \varphi(\gamma) \left(\int p(\beta|\gamma, \mathcal{X}, \mathcal{Z}) \log p(\mathcal{Y}|\mathcal{X}, \beta) d\beta - \log \frac{\varphi(\gamma)}{p(\gamma|\mathcal{Z})} \right) d\gamma \\
&= \int \varphi(\gamma) \int p(\beta|\gamma, \mathcal{X}, \mathcal{Z}) \log p(\mathcal{Y}|\mathcal{X}, \beta) d\beta d\gamma - \mathbb{D}_{\text{KL}}(\varphi||p(\cdot|\mathcal{Z}))
\end{aligned} \tag{B.82}$$

In the regression setting that assumes $p(\mathcal{Y}|\mathcal{X}, \beta)$ to follow (B.25), [Titsias (2009)] shows that the optimal φ that maximizes (B.82) can be derived explicitly and is given as follows:

$$\begin{aligned}
\varphi(\gamma) &= \mathcal{N}(\gamma|\ddot{\mu}, \ddot{\Sigma}) \\
\ddot{\mu} &= \mu_{\mathcal{Z}} + \mathbb{K}_{\mathcal{Z}\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{Z}} \mathbb{K}_{\mathcal{Z}\mathcal{Z}}^{-1} \mathbb{K}_{\mathcal{Z}\mathcal{X}})^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
\ddot{\Sigma} &= \mathbb{K}_{\mathcal{Z}\mathcal{Z}} - \mathbb{K}_{\mathcal{Z}\mathcal{X}} (\Lambda_{\mathcal{X}} + \mathbb{K}_{\mathcal{X}\mathcal{Z}} \mathbb{K}_{\mathcal{Z}\mathcal{Z}}^{-1} \mathbb{K}_{\mathcal{Z}\mathcal{X}})^{-1} \mathbb{K}_{\mathcal{X}\mathcal{Z}}
\end{aligned} \tag{B.83}$$

Because $q(\gamma|\mathcal{X}, \mathcal{Z}) = \int q(\beta, \gamma|\mathcal{X}, \mathcal{Z}) d\beta = \varphi(\gamma)$ approximates $p(\gamma|\mathcal{X}, \mathcal{Z}, \mathcal{Y})$, the posterior predictive (B.81) is given by the following closed form:

$$\begin{aligned}
p(y|x, \mathcal{X}, \mathcal{Y}) &= \mathcal{N}(y|m(x), \Lambda_x + s(x)) \\
m(x) &= \mu_x + \mathbb{K}_{x\mathcal{Z}} (\mathbb{K}_{\mathcal{Z}\mathcal{Z}} + \mathbb{K}_{\mathcal{Z}\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{Z}})^{-1} \mathbb{K}_{\mathcal{Z}\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} (\mathcal{Y} - \mu_{\mathcal{X}}) \\
s(x) &= \mathbb{K}_{xx} - \mathbb{K}_{x\mathcal{Z}} \mathbb{K}_{\mathcal{Z}\mathcal{Z}}^{-1} \mathbb{K}_{\mathcal{Z}x} + \mathbb{K}_{x\mathcal{Z}} (\mathbb{K}_{\mathcal{Z}\mathcal{Z}} + \mathbb{K}_{\mathcal{Z}\mathcal{X}} \Lambda_{\mathcal{X}}^{-1} \mathbb{K}_{\mathcal{X}\mathcal{Z}})^{-1} \mathbb{K}_{\mathcal{Z}x}
\end{aligned} \tag{B.84}$$

Note that $\Lambda_{\mathcal{X}}^{-1}$ is efficiently computed regardless of $|\mathcal{X}|$ since here $\Lambda_{\mathcal{X}}$ is a diagonal matrix.

If $p(\mathcal{Y}|\mathcal{X}, \beta)$ is not Gaussian, then φ is not optimized in a closed form. Thus, one needs to consider the reparameterization trick to efficiently approximate the gradient of ELBO in (B.82), after parameterizing φ .

Though [Bauer et al. (2016)] process the comparison and recommend the VFE rather than the fully independent training conditional (FITC) [Csató & Opper (2002); Quiñonero-Candela & Rasmussen (2005)] for dealing with huge datasets, VFE has a certain drawback: the performance is depending on the choice of the pseudo-inputs $\mathcal{Z}$. With an ill-posed $\mathcal{Z}$, ELBO in (B.82) could be deviated far from the exact log-ML $\log p(\mathcal{Y}|\mathcal{X})$. In principle, optimizing $\mathcal{Z}$ is not straightforward. [Bui et al. (2017)] attempt both of the pseudo-input optimization and the unification of VFE and FITC by power expectation propagation.

References

- Abramowitz, M. & Stegun, I. A. (1964). *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*. New York: Dover.
- Ades, M. & van Leeuwen, P. J. (2013). An exploration of the equivalent weights particle filter. *Quarterly Journal of the Royal Meteorological Society*, 139(672), 820–840.
- Adkins, J. F., McIntyre, K., & Schrag, D. P. (2002). The salinity, temperature, and $\delta^{18}\text{O}$ of the glacial deep ocean. *Science*, 298(5599), 1769–1773.
- Ahn, S. (2016). *Bayesian Inference in Statistical Analysis of Paleoclimate Records*. PhD thesis, Brown University.
- Arz, H. W., Pätzold, J., Müller, P. J., & Moammar, M. O. (2003). Influence of northern hemisphere climate and global sea level rise on the restricted red sea marine environment during termination i. *Paleoceanography*, 18(2).
- Barber, D. (2012). *Bayesian Reasoning and Machine Learning*. Cambridge University Press.
- Bard, E. (2001). Comparison of alkenone estimates with other paleotemperature proxies. *Geochemistry, Geophysics, Geosystems*, 2(1).
- Bard, E., Rostek, F., & Ménot-Combes, G. (2004). Radiocarbon calibration beyond 20,000 ^{14}C yr b.p. by means of planktonic foraminifera of the iberian margin. *Quaternary Research*, 61(2), 204–214.
- Bard, E., Rostek, F., Turon, J.-L., & Gendreau, S. (2000). Hydrological impact of heinrich events in the subtropical northeast atlantic. *Science*, 289(5483), 1321–1324.
- Barron, J. A., Heusser, L., Herbert, T., & Lyle, M. (2003). High-resolution climatic evolution of coastal northern california during the past 16,000 years. *Paleoceanography*, 18(1).
- Barrows, T. T., Juggins, S., De Deckker, P., Calvo, E., & Pelejero, C. (2007). Long-term sea surface temperature and climate change in the australian–new zealand region. *Paleoceanography*, 22(2).
- Bauer, M., van der Wilk, M., & Rasmussen, C. E. (2016). Understanding probabilistic sparse gaussian process approximations. In *Proceedings of the 30th International Conference on Neural Information Processing Systems*, NIPS’16 (pp. 1533–1541). USA: Curran Associates Inc.

- Becker, K. W., Lipp, J. S., Versteegh, G. J., Wörmer, L., & Hinrichs, K.-U. (2015). Rapid and simultaneous analysis of three molecular sea surface temperature proxies and application to sediments from the sea of marmara. *Organic Geochemistry*, 85, 42 – 53.
- Ben Abdesslem, A., Dervilis, N., Wagg, D., & Worden, K. (2019). Model selection and parameter estimation of dynamical systems using a novel variant of approximate bayesian computation. *Mechanical Systems and Signal Processing*, 122, 364 – 386.
- Bendle, J. & Rosell-Mele, A. (2004). Distributions of uk37 and uk37' in the surface waters and sediments of the nordic seas: Implications for paleoceanography. *Geochemistry Geophysics Geosystems*.
- Bendle, J. A. P. & Rosell-Melé, A. (2007). High-resolution alkenone sea surface temperature variability on the north icelandic shelf: implications for nordic seas palaeoclimatic development during the holocene. *The Holocene*, 17(1), 9–24.
- Benthien, A. & Müller, P. J. (2000). Anomalously low alkenone temperatures caused by lateral particle and sediment transport in the malvinas current region, western argentine basin. *Deep Sea Research Part I: Oceanographic Research Papers*, 47(12), 2369 – 2393.
- Bierens, H. J. (1994). *Topics in Advanced Econometrics: Estimation, Testing, and Specification of Cross-Section and Time Series Models*. Cambridge University Press.
- Binois, M., Gramacy, R. B., & Ludkovski, M. (2018). Practical heteroscedastic gaussian process modeling for large simulation experiments. *Journal of Computational and Graphical Statistics*, 27(4), 808–821.
- Bishop, C. M. (2006). *Pattern Recognition and Machine Learning (Information Science and Statistics)*. Springer.
- Blaauw, M. & Christen, J. A. (2005). Radiocarbon peat chronologies and environmental change. *Journal of the Royal Statistical Society. Series C (Applied Statistics)*, 54(4), 805–816.
- Blaauw, M. & Christen, J. A. (2011). Flexible paleoclimate age-depth models using an autoregressive gamma process. *Bayesian Anal.*, 6(3), 457–474.
- Blomqvist, K., Kaski, S., & Heinonen, M. (2019). Deep convolutional gaussian process. In *Proceedings of the European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases*.
- Bodin, E., Campbell, N. D. F., & Ek, C. H. (2017). Latent gaussian process regression. arXiv:1707.05534v2 [stat.ML].
- Bowen, G. J., Fischer-Femal, B., Reichart, G.-J., Sluijs, A., & Lear, C. H. (2020). Joint inversion of proxy system models to reconstruct paleoenvironmental time series from heterogeneous data. *Climate of the Past*, 16(1), 65–78.
- Briers, M., Doucet, A., & Maskell, S. (2010). Smoothing algorithms for state–space models. *Annals of the Institute of Statistical Mathematics*, 62(1), 61–89.

- Bui, T. D., Yan, J., & Turner, R. E. (2017). A unifying framework for gaussian process pseudo-point approximations using power expectation propagation. *Journal of Machine Learning Research*, 18(104), 1–72.
- Cacho, I., Grimalt, J. O., Canals, M., Sbaffi, L., Shackleton, N. J., Schönfeld, J., & Zahn, R. (2001). Variability of the western mediterranean sea surface temperature during the last 25,000 years and its connection with the northern hemisphere climatic changes. *Paleoceanography*, 16(1), 40–52.
- Cacho, I., Grimalt, J. O., Pelejero, C., Canals, M., Sierro, F. J., Flores, J. A., & Shackleton, N. (1999). Dansgaard-oeschger and heinrich event imprints in alboran sea paleotemperatures. *Paleoceanography*, 14(6), 698–705.
- Calandra, R., Peters, J., Rasmussen, C. E., & Deisenroth, M. P. (2016). Manifold gaussian processes for regression. In *International Joint Conference on Neural Networks (IJCNN)* (pp. 3338–3345).: IEEE.
- Calvo, E., Pelejero, C., De Deckker, P., & Logan, G. A. (2007). Antarctic deglacial pattern in a 30 kyr record of sea surface temperature offshore south australia. *Geophysical Research Letters*, 34(13).
- Calvo, E., Pelejero, C., Herguera, J. C., Palanques, A., & Grimalt, J. O. (2001). Insolation dependence of the southeastern subtropical pacific sea surface temperature over the last 400 kyrs. *Geophysical Research Letters*, 28(12), 2481–2484.
- Caniupán, M., Lamy, F., Lange, C. B., Kaiser, J., Arz, H., Kilian, R., Baeza Urrea, O., Aracena, C., Hebbeln, D., Kissel, C., Laj, C., Mollenhauer, G., & Tiedemann, R. (2011). Millennial-scale sea surface temperature and patagonian ice sheet changes off southernmost chile (53°s) over the past ~60 kyr. *Paleoceanography*, 26(3).
- Castañeda, I. S., Schefuß, E., Pätzold, J., Sinninghe Damsté, J. S., Weldeab, S., & Schouten, S. (2010). Millennial-scale sea surface temperature changes in the eastern mediterranean (nile river delta region) over the last 27,000 years. *Paleoceanography*, 25(1).
- Catanzaro, E. J., Champion, C. E., Garner, E. L., Marinenko, G., Sappenfield, K. M., & Shields, W. R. (1970). *Standard Reference Materials: Boric acid; isotopic, and assay standard reference materials*. National Bureau of Standards (U.S.).
- Chapman, M. R., Shackleton, N. J., Zhao, M., & Eglinton, G. (1996). Faunal and alkenone reconstructions of subtropical north atlantic surface hydrography and paleotemperature over the last 28 kyr. *Paleoceanography*, 11(3), 343–357.
- Chazen, C. (2011). *Holocene climate evolution of the eastern tropical Pacific told from high resolution climate records from the Peru margin and equatorial upwelling regions*. PhD thesis, Brown University.
- Chen, R. & Liu, J. S. (2000). Mixture kalman filters. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 62(3), 493–508.

- Chen, W., Mohtadi, M., Schefuß, E., & Mollenhauer, G. (2014). Organic-geochemical proxies of sea surface temperature in surface sediments of the tropical eastern indian ocean. *Deep Sea Research Part I: Oceanographic Research Papers*, 88, 17 – 29.
- Cho, Y. & Saul, L. K. (2009). Kernel methods for deep learning. In *Proceedings of the 22nd International Conference on Neural Information Processing Systems*, NIPS'09 (pp. 342–350). Red Hook, NY, USA: Curran Associates Inc.
- Christen, J. & Sergio, E. (2009). A new robust statistical model for radiocarbon data. *Radiocarbon*, 51(3), 1047–1059.
- Clifford, P. (1990). Markov random fields in statistics. In G. Grimmett & D. Welsh (Eds.), *Disorder in Physical Systems: A Volume in Honour of John M. Hammersley* (pp. 19–32). Oxford: Oxford University Press.
- Collins, J., Schefuß, E., Heslop, D., Mulitza, S., Prange, M., Zabel, M., Tjallingii, R., Dokken, T., Huang, E., Mackensen, A., Schulz, M., Tian, J., Zarriess, M., & Wefer, G. (2011). Interhemispheric symmetry of the tropical african rainbelt over the past 23,000 years. *Nature Geoscience*, 4, 42–45.
- Conte, M. H., Sicre, M.-A., Rühlemann, C., Weber, J. C., Schulte, S., Schulz-Bull, D., & Blanz, T. (2006). Global temperature calibration of the alkenone unsaturation index (uk'37) in surface waters and comparison with surface sediments. *Geochemistry, Geophysics, Geosystems*, 7(2).
- Csató, L. & Oppor, M. (2002). Sparse on-line gaussian processes. *Neural Computation*, 14(3), 641–668.
- Damianou, A. & Lawrence, N. (2013). Deep gaussian processes. In C. M. Carvalho & P. Ravikumar (Eds.), *Proceedings of the Sixteenth International Conference on Artificial Intelligence and Statistics*, volume 31 of *Proceedings of Machine Learning Research* (pp. 207–215). Scottsdale, Arizona, USA: PMLR.
- Damianou, A. C., Titsias, M. K., & Lawrence, N. D. (2016). Variational inference for latent variables and uncertain inputs in gaussian processes. *Journal of Machine Learning Research*, 17(42), 1–62.
- Demenocal, P., Ortiz, J., Guilderson, T., Adkins, J., Sarnthein, M., Baker, L., & Yarusinsky, M. (2000). Abrupt onset and termination of the african humid period. *Quaternary Science Reviews - QUATERNARY SCIENCE REV*, 19, 347–361.
- Dempster, A. P., Laird, N. M., & Rubin, D. B. (1977). Maximum likelihood from incomplete data via the em algorithm. *Journal of the Royal Statistical Society, Series B*, 39(1), 1–38.
- Dickson, A. (2010). The carbon dioxide system in seawater: Equilibrium chemistry and measurements. *Guide to Best Practices for Ocean Acidification Research and Data Reporting*, (pp. 17–40).

- Diz, P., Hall, I. R., Zahn, R., & Molyneux, E. G. (2007). Paleooceanography of the southern agulhas plateau during the last 150 ka: Inferences from benthic foraminiferal assemblages and multispecies epifaunal carbon isotopes. *Paleoceanography*, 22(4).
- Doose, H., Prahl, F. G., & Lyle, M. W. (1997). Biomarker temperature estimates for modern and last glacial surface waters of the california current system between 33° and 42°n. *Paleoceanography*, 12(4), 615–622.
- Doucet, A., de Freitas, N., & Gordon, N. E. (2001). *Sequential Monte Carlo methods in practice*. Springer.
- Duane, S., Kennedy, A., Pendleton, B., & Roweth, D. (1987). Hybrid Monte Carlo. *Physics Letters B*, 195, 216–222.
- Dubois, N., Kienast, M., Normandeau, C., & Herbert, T. D. (2009). Eastern equatorial pacific cold tongue during the last glacial maximum as seen from alkenone paleothermometry. *Paleoceanography*, 24(4).
- Durbin, R., Eddy, S. R., Krogh, A., & Mitchison, G. (1998). *Biological Sequence Analysis: Probabilistic Models of Proteins and Nucleic Acids*. Cambridge University Press.
- Duvenaud, D., Lloyd, J. R., Grosse, R., Tenenbaum, J. B., & Ghahramani, Z. (2013). Structure discovery in nonparametric regression through compositional kernel search. In *Proceedings of the 30th International Conference on International Conference on Machine Learning - Volume 28*, ICML'13 (pp. III–1166–III–1174).: JMLR.org.
- Duvenaud, D. K. (2014). *Automatic model construction with Gaussian processes*. PhD thesis, University of Cambridge.
- Dyez, K. A., Hönisch, B., & Schmidt, G. A. (2018). Early pleistocene obliquity-scale pco2 variability at 1.5 million years ago. *Paleoceanography and Paleoclimatology*, 33(11), 1270–1291.
- Eggleton, T. (2012). *A Short Introduction to Climate Change*. Cambridge University Press.
- Eleftheriadis, S., Nicholson, T. F., Deisenroth, M. P., & Hensman, J. (2017). Identification of gaussian process state space models. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, NIPS'17 (pp. 5315–5325). USA: Curran Associates Inc.
- Emeis, K.-C., Struck, U., Schulz, H.-M., Rosenberg, R., Bernasconi, S., Erlenkeuser, H., Sakamoto, T., & Martinez-Ruiz, F. (2000). Temperature and salinity variations of mediterranean sea surface waters over the last 16,000 years from records of planktonic stable oxygen isotopes and alkenone unsaturation ratios. *Palaeogeography, Palaeoclimatology, Palaeoecology*, 158(3), 259 – 280.
- Epstein, S., Buchsbaum, R., Lowenstam, H. A., & Urey, H. C. (1953). Revised carbonate-water isotopic temperature scale. *GSA Bulletin*, 64(11), 1315–1326.

- Evensen, G. (1994). Sequential data assimilation with a nonlinear quasi-geostrophic model using monte carlo methods to forecast error statistics. *Journal of Geophysical Research: Oceans*, 99(C5), 10143–10162.
- Fallet, U., Castañeda, I. S., Henry-Edwards, A., Richter, T. O., Boer, W., Schouten, S., & Brummer, G.-J. (2012). Sedimentation and burial of organic and inorganic temperature proxies in the mozambique channel, sw indian ocean. *Deep Sea Research Part I: Oceanographic Research Papers*, 59, 37 – 53.
- Filippova, A., Kienast, M., Frank, M., & Schneider, R. R. (2016). Alkenone paleothermometry in the north atlantic: A review and synthesis of surface sediment data and calibrations. *Geochemistry, Geophysics, Geosystems*, 17(4), 1370–1382.
- Foster, G. L. & Rae, J. W. (2016). Reconstructing ocean pH with boron isotopes in foraminifera. *Annual Review of Earth and Planetary Sciences*, 44(1), 207–237.
- Fraser, N., Kuhnt, W., Holbourn, A., Bolliet, T., Andersen, N., Blanz, T., & Beaufort, L. (2014). Precipitation variability within the west pacific warm pool over the past 120ka: Evidence from the davao gulf, southern philippines. *Paleoceanography*, 29(11), 1094–1110.
- Frigola, R., Chen, Y., & Rasmussen, C. E. (2014). Variational gaussian process state-space models. In *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 2*, NIPS’14 (pp. 3680–3688). Cambridge, MA, USA: MIT Press.
- Garcia-Olivares, A. & Herrero, C. (2012). Fitting the last pleistocene $\delta^{18}O$ and CO_2 time series with simple box models. *Scientia Marina*, 76S1, 209–218.
- Garcia-Olivares, A. & Herrero, C. (2013). Simulation of glacial-interglacial cycles by simple relaxation models: Consistency with observational results. *Climate Dynamics*, 41, 1307–1331.
- Gebbie, G. & Huybers, P. (2012). The mean age of ocean waters inferred from radiocarbon observations: Sensitivity to surface sources and accounting for mixing histories. *Journal of Physical Oceanography*, 42(2), 291–305.
- Geenens, G. (2011). Curse of dimensionality and related issues in nonparametric functional regression. *Statistics Surveys*, 5, 30–43.
- Gelman, A. (2008). Objections to bayesian statistics. *Bayesian Analysis*, 3(3), 445–449.
- Gelman, A., Carlin, J., Stern, H., Dunson, D., Vehtari, A., & Rubin, D. (2013). *Bayesian Data Analysis*. Chapman & Hall/CRC Texts in Statistical Science. CRC Press.
- Geman, S. & Geman, D. (1984). Stochastic relaxation, gibbs distributions, and the bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-6(6), 721–741.
- Genton, M. G. (2002). Classes of kernels for machine learning: A statistics perspective. *Journal of Machine Learning Research*, 2, 299–312.

- Gibbs, M. N. (1997). *Bayesian Gaussian Processes for Regression and Classification*. PhD thesis, University of Cambridge.
- Girolami, M. & Calderhead, B. (2011). Riemann manifold langevin and hamiltonian monte carlo methods. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 73(2), 123–214.
- Goldberg, P. W., Williams, C. K. I., & Bishop, C. M. (1998). Regression with input-dependent noise: A gaussian process treatment. In *Proceedings of the 1997 Conference on Advances in Neural Information Processing Systems 10*, NIPS '97 (pp. 493–499). Cambridge, MA, USA: MIT Press.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. The MIT Press.
- Grousset, F. (1993). Patterns of ice-rafted detritus in the glacial north atlantic (40- 55° n). *Paleoceanography*, 8, 175–192.
- Gutiérrez, D., Bouloubassi, I., Sifeddine, A., Purca, S., Goubanova, K., Graco, M., Field, D., Méjanelle, L., Velazco, F., Lorre, A., Salvatelli, R., Quispe, D., Vargas, G., Dewitte, B., & Ortlieb, L. (2011). Coastal cooling and increased productivity in the main upwelling zone off peru since the mid-twentieth century. *Geophysical Research Letters*, 38(7).
- Haji Ghassemi, N. (2013). Analytic long-term forecasting with periodic gaussian processes. Master's thesis, Blekinge Institute of Technology.
- Hangos, K., Szederkényi, G., Lakner, R., & Gerzson, M. (2001). *Intelligent Control Systems: An Introduction with Examples*. Applied Optimization. Springer.
- Harada, N., Ahagon, N., Sakamoto, T., Uchida, M., Ikehara, M., & Shibata, Y. (2006). Rapid fluctuation of alkenone temperature in the southwestern okhotsk sea during the past 120 ky. *Global and Planetary Change*, 53(1), 29 – 46. Late Quaternary paleoceanography of the north-western Pacific: Results from IMAGES program.
- Haslett, J. & Parnell, A. (2008). A simple monotone process with application to radiocarbon-dated depth chronologies. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 57(4), 399–418.
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. Springer series in statistics. Springer.
- Hastings, W. K. (1970). Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, 57(1), 97–109.
- Hemming, N. & Hanson, G. (1992). Boron isotopic composition and concentration in modern marine carbonates. *Geochimica et Cosmochimica Acta*, 56(1), 537 – 543.
- Henehan, M. J., Foster, G. L., Bostock, H. C., Greenop, R., Marshall, B. J., & Wilson, P. A. (2016). A new boron isotope-ph calibration for orbitolites, with implications for understanding and accounting for 'vital effects'. *Earth and Planetary Science Letters*, 454, 282 – 292.

- Herbert, T. D., Schuffert, J. D., Andreasen, D., Heusser, L., Lyle, M., Mix, A., Ravelo, A. C., Stott, L. D., & Herguera, J. C. (2001). Collapse of the california current during glacial maxima linked to climate change on land. *Science*, 293(5527), 71–76.
- Herbert, T. D., Schuffert, J. D., Thomas, D., Lange, C., Weinheimer, A., Peleo-Alampay, A., & Herguera, J.-C. (1998). Depth and seasonality of alkenone production along the california margin inferred from a core top transect. *Paleoceanography*, 13(3), 263–271.
- Hernández-Lobato, D., Hernández-Lobato, J. M., & Dupont, P. (2011). Robust multi-class gaussian process classification. In *Proceedings of the 24th International Conference on Neural Information Processing Systems*, NIPS'11 (pp. 280–288). Red Hook, NY, USA: Curran Associates Inc.
- Hernández-Sánchez, M., Woodward, E., Taylor, K., Henderson, G., & Pancost, R. (2014). Variations in gdgt distributions through the water column in the south east atlantic ocean. *Geochimica et Cosmochimica Acta*, 132, 337 – 348.
- Ho, S. L., Mollenhauer, G., Fietz, S., Martínez-Garcia, A., Lamy, F., Rueda, G., Schipper, K., Méheust, M., Rosell-Melé, A., Stein, R., & Tiedemann, R. (2014). Appraisal of tex86 and tex86l thermometries in subpolar and polar regions. *Geochimica et Cosmochimica Acta*, 131, 213 – 226.
- Ho, S. L., Mollenhauer, G., Lamy, F., Martínez-Garcia, A., Mohtadi, M., Gersonde, R., Hebbeln, D., Nunez-Ricardo, S., Rosell-Melé, A., & Tiedemann, R. (2012). Sea surface temperature variability in the pacific sector of the southern ocean over the past 700 kyr. *Paleoceanography*, 27(4).
- Ho, S. L., Yamamoto, M., Mollenhauer, G., & Minagawa, M. (2011). Core top tex86 values in the south and equatorial pacific. *Organic Geochemistry*, 42(1), 94 – 99.
- Hoffman, M. D. & Gelman, A. (2014). The no-u-turn sampler: Adaptively setting path lengths in hamiltonian monte carlo. *Journal of Machine Learning Research*, 15(47), 1593–1623.
- Hogg, A. G., Hua, Q., Blackwell, P. G., Niu, M., Buck, C. E., Guilderson, T. P., Heaton, T. J., Palmer, J. G., Reimer, P. J., & Reimer, R. W. (2013). Shcal13 southern hemisphere calibration, 0–50,000 years cal bp. *Radiocarbon*, 55(4), 1889–1903.
- Hönisch, B., Hemming, N. G., Archer, D., Siddall, M., & McManus, J. F. (2009). Atmospheric carbon dioxide concentration across the mid-pleistocene transition. *Science*, 324(5934), 1551–1554.
- Horikawa, K., Minagawa, M., Murayama, M., Kato, Y., & Asahi, H. (2006). Spatial and temporal sea-surface temperatures in the eastern equatorial pacific over the past 150 kyr. *Geophysical Research Letters*, 33(13).
- Houtekamer, P. L. & Mitchell, H. L. (1998). Data assimilation using an ensemble kalman filter technique. *Monthly Weather Review*, 126(3), 796–811.

- Hu, J., Meyers, P. A., Chen, G., Peng, P., & Yang, Q. (2012). Archaeal and bacterial glycerol dialkyl glycerol tetraethers in sediments from the eastern liao spreading center, south pacific ocean. *Organic Geochemistry*, 43, 162 – 167.
- Huang, C.-Y., Wu, S.-F., Zhao, M., Chen, M.-T., Wang, C.-H., Tu, X., & Yuan, P. B. (1997). Surface ocean and monsoon climate variability in the south china sea since the last glaciation. *Marine Micropaleontology*, 32(1), 71 – 94. *Marine Micropaleontology in China*.
- Ikehara, M., Kawamura, K., Ohkouchi, N., Kimoto, K., Murayama, M., Nakamura, T., Oba, T., & Taira, A. (1997). Alkenone sea surface temperature in the southern ocean for the last two deglaciations. *Geophysical Research Letters*, 24(6), 679–682.
- Indermühle, A., Monnin, E., Stauffer, B., Stocker, T. F., & Wahlen, M. (2000). Atmospheric co₂ concentration from 60 to 20 kyr bp from the taylor dome ice core, antarctica. *Geophysical Research Letters*, 27(5), 735–738.
- Jaakkola, T., Diekhans, M., & Haussler, D. (2000). A discriminative framework for detecting remote protein homologies. *Journal of Computational Biology*, 7(1-2), 95–114.
- Jaeschke, A., Rühlemann, C., Arz, H., Heil, G., & Lohmann, G. (2007). Coupling of millennial-scale changes in sea surface temperature and precipitation off northeastern brazil with high-latitude climate shifts during the last glacial period. *Paleoceanography*, 22(4).
- Jaeschke, A., Wengler, M., Hefter, J., Ronge, T. A., Geibert, W., Mollenhauer, G., Gersonde, R., & Lamy, F. (2017). A biomarker perspective on dust, productivity, and sea surface temperature in the pacific sector of the southern ocean. *Geochimica et Cosmochimica Acta*, 204, 120 – 139.
- Jia, G., Zhang, J., Chen, J., Peng, P., & Zhang, C. L. (2012). Archaeal tetraether lipids record subsurface water temperature in the south china sea. *Organic Geochemistry*, 50, 68 – 77.
- Julier, S. J. & Uhlmann, J. K. (1997). New extension of the Kalman filter to nonlinear systems. In I. Kadar (Ed.), *Signal Processing, Sensor Fusion, and Target Recognition VI*, volume 3068 (pp. 182 – 193): International Society for Optics and Photonics SPIE.
- Kaiser, J., Lamy, F., & Hebbeln, D. (2005). A 70-kyr sea surface temperature record off southern chile (ocean drilling program site 1233). *Paleoceanography*, 20(4).
- Kaiser, J., Ruggieri, N., Hefter, J., Siegel, H., Mollenhauer, G., Arz, H. W., & Lamy, F. (2014). Lipid biomarkers in surface sediments from the gulf of genoa, ligurian sea (nw mediterranean sea) and their potential for the reconstruction of palaeo-environments. *Deep Sea Research Part I: Oceanographic Research Papers*, 89, 68 – 83.
- Kaiser, J., Schefuß, E., Lamy, F., Mohtadi, M., & Hebbeln, D. (2008). Glacial to holocene changes in sea surface temperature and coastal vegetation in north central chile: high versus low latitude forcing. *Quaternary Science Reviews*, 27(21), 2064 – 2075.
- Kaiser, M., Otte, C., Runkler, T., & Ek, C. H. (2018). Data association with gaussian processes. arXiv:1810.07158v3 [stat.ML].

- Kalman, R. E. (1960). A new approach to linear filtering and prediction problems. *Journal of Fluids Engineering*, 82(1), 35–45.
- Kennedy, J. A. & Brassell, S. C. (1992). Molecular stratigraphy of the santa barbara basin: comparison with historical records of annual climate change. *Organic Geochemistry*, 19(1), 235 – 244. Proceedings of the 15th International Meeting on Organic Geochemistry.
- Kersting, K., Plagemann, C., Pfaff, P., & Burgard, W. (2007). Most likely heteroscedastic gaussian process regression. In *Proceedings of the 24th International Conference on Machine Learning*, ICML '07 (pp. 393–400). New York, NY, USA: ACM.
- Kienast, M., Kienast, S., Calvert, S., Eglinton, T., Mollenhauer, G., François, R., & Mix, A. (2006). Eastern pacific cooling and atlantic overturning circulation during the last deglaciation. *Nature*, 443, 846–9.
- Kienast, M., MacIntyre, G., Dubois, N., Higginson, S., Normandeau, C., Chazen, C., & Herbert, T. D. (2012). Alkenone unsaturation in surface sediments from the eastern equatorial pacific: Implications for sst reconstructions. *Paleoceanography*, 27(1).
- Kienast, S. S. & McKay, J. L. (2001). Sea surface temperatures in the subarctic northeast pacific reflect millennial-scale climate oscillations during the last 16 kyrs. *Geophysical Research Letters*, 28(8), 1563–1566.
- Kim, J.-H., Meggers, H., Rimbu, N., Lohmann, G., Freudenthal, T., Müller, P. J., & Schneider, R. R. (2007). Impacts of the North Atlantic gyre circulation on Holocene climate off northwest Africa. *Geology*, 35(5), 387–390.
- Kim, J.-H., Rimbu, N., Lorenz, S. J., Lohmann, G., Nam, S.-I., Schouten, S., Rühlemann, C., & Schneider, R. R. (2004). North pacific and north atlantic sea-surface temperature variability during the holocene. *Quaternary Science Reviews*, 23(20), 2141 – 2154. Holocene climate variability - a marine perspective.
- Kim, J.-H., Schneider, R. R., Hebbeln, D., Müller, P. J., & Wefer, G. (2002a). Last deglacial sea-surface temperature evolution in the southeast pacific compared to climate changes on the south american continent. *Quaternary Science Reviews*, 21(18), 2085 – 2097.
- Kim, J.-H., Schneider, R. R., Mulitza, S., & Müller, P. J. (2003). Reconstruction of se trade-wind intensity based on sea-surface temperature gradients in the southeast atlantic over the last 25 kyr. *Geophysical Research Letters*, 30(22).
- Kim, J.-H., Schneider, R. R., Müller, P. J., & Wefer, G. (2002b). Interhemispheric comparison of deglacial sea-surface temperature patterns in atlantic eastern boundary currents. *Earth and Planetary Science Letters*, 194(3), 383 – 393.
- Kim, J.-H., Schouten, S., Hopmans, E. C., Donner, B., & Damsté, J. S. S. (2008). Global sediment core-top calibration of the tex86 paleothermometer in the ocean. *Geochimica et Cosmochimica Acta*, 72(4), 1154 – 1173.

- Kim, J.-H., van der Meer, J., Schouten, S., Helmke, P., Willmott Puig, V., Sangiorgi, F., Koc, N., Hopmans, E., & Sinninghe-Damste, J. (2010). New indices and calibrations derived from the distribution of crenarchaeal isoprenoid tetraether lipids: Implications for past sea surface temperature reconstructions. *Geochimica et Cosmochimica Acta*, 74, 4639–4654.
- Kim, R. A., Lee, K. E., & Bae, S. W. (2015). Sea surface temperature proxies (alkenones, foraminiferal mg/ca, and planktonic foraminiferal assemblage) and their implications in the okinawa trough. *Prog. in Earth and Planet. Sci.*, 2(43).
- Kingma, D. P. & Welling, M. (2013). Auto-encoding variational bayes. arXiv:1312.6114v10 [stat.ML].
- Kitagawa, G. (1994). The two-filter formula for smoothing and an implementation of the gaussian-sum smoother. *Annals of the Institute of Statistical Mathematics*, 46(4), 605–623.
- Klaas, M., Briers, M., de Freitas, N., Doucet, A., Maskell, S., & Lang, D. (2006). Fast particle smoothing: if i had a million particles. In *ICML*.
- Klochko, K., Kaufman, A. J., Yao, W., Byrne, R. H., & Tossell, J. A. (2006). Experimental measurement of boron isotope fractionation in seawater. *Earth and Planetary Science Letters*, 248(1), 276 – 285.
- Koller, D. & Friedman, N. (2009). *Probabilistic Graphical Models: Principles and Techniques*. Adaptive computation and machine learning. MIT Press.
- Koutavas, A. & Sachs, J. P. (2008). Northern timing of deglaciation in the eastern equatorial pacific from alkenone paleothermometry. *Paleoceanography*, 23(4).
- Kudrass, H., Hofmann, A., Dose, H., Emeis, K., & Erlenkeuser, H. (2001). Modulation and amplification of climatic changes in the Northern Hemisphere by the Indian summer monsoon during the past 80 k.y. *Geology*, 29(1), 63–66.
- Lamy, F., Rühlemann, C., Hebbeln, D., & Wefer, G. (2002). High- and low-latitude climate control on the position of the southern peru-chile current during the holocene. *Paleoceanography*, 17(2), 16–1–16–10.
- Landau, L. D. & Lifshitz, E. M. (2013). *Statistical Physics: Volume 5*. Elsevier Science.
- Langrené, N. & Warin, X. (2019). Fast and stable multivariate kernel density estimation by fast sum updating. *Journal of Computational and Graphical Statistics*, 28(3), 596–608.
- Lawrence, N. (2005). Probabilistic non-linear principal component analysis with gaussian process latent variable models. *Journal of Machine Learning Research*, 6, 1783–1816.
- Lázaro-Gredilla, M. & Titsias, M. K. (2011). Variational heteroscedastic gaussian process regression. In *Proceedings of the 28th International Conference on International Conference on Machine Learning* (pp. 841–848).
- Lázaro-Gredilla, M., Van Vaerenbergh, S., & Lawrence, N. D. (2012). Overlapping mixtures of gaussian processes for the data association problem. *Pattern Recognition*, 45(4), 1386–1395.

- Le, Q. V., Smola, A. J., & Canu, S. (2005). Heteroscedastic gaussian process regression. In *Proceedings of the 22nd International Conference on Machine Learning, ICML '05* (pp. 489–496).: ACM.
- Le Gall, F. (2014). Powers of tensors and fast matrix multiplication. In *Proceedings of the 39th International Symposium on Symbolic and Algebraic Computation, ISSAC '14* (pp. 296–303). New York, NY, USA: ACM.
- Leduc, G., Schneider, R., Kim, J.-H., & Lohmann, G. (2010). Holocene and eemian sea surface temperature trends as revealed by alkenone and mg/ca paleothermometry. *Quaternary Science Reviews*, 29(7), 989 – 1004.
- Leduc, G., Vidal, L., Tachikawa, K., Rostek, F., Sonzogni, C., Beaufort, L., & Bard, E. (2007). Moisture transport across central america as a positive feedback on abrupt climatic changes. *Nature*, 445, 908–11.
- Lee, K. E., Bahk, J. J., & Choi, J. (2008a). Alkenone temperature estimates for the east sea during the last 190,000 years. *Organic Geochemistry*, 39(6), 741 – 753.
- Lee, K. E., Kim, J.-H., Wilke, I., Helmke, P., & Schouten, S. (2008b). A study of the alkenone, tex86, and planktonic foraminifera in the benguela upwelling system: Implications for past sea surface temperature estimates. *Geochemistry, Geophysics, Geosystems*, 9(10).
- Lee, K. E., Slowey, N. C., & Herbert, T. D. (2001). Glacial sea surface temperatures in the subtropical north pacific: A comparison of u37 k', $\delta^{18}O$, and foraminiferal assemblage temperature estimates. *Paleoceanography*, 16(3), 268–279.
- Lee, T. & Lawrence, C. E. (2019). Heteroscedastic gaussian process regression on the alkenone over sea surface temperatures. In *Proceedings of the 9th International Workshop on Climate Informatics* (pp. 269–274).
- Lee, T., Lisiecki, L. E., Rand, D., Gebbie, G., & Lawrence, C. E. (2019). A multiple continuous signal alignment algorithm with gaussian process profiles and an application to paleoceanography. arXiv:1907.08738v3 [stat.AP].
- Leider, A., Hinrichs, K.-U., Mollenhauer, G., & Versteegh, G. (2010). Core-top calibration of the lipid-based u37k' and tex86 temperature proxies on the southern italian shelf (sw adriatic sea, gulf of taranto). *Earth and Planetary Science Letters*, 300, 112–124.
- Lemarchand, D., Gaillardet, J., Lewin, E., & Allègre, C. J. (2002). Boron isotope systematics in large rivers: implications for the marine boron budget and paleo-ph reconstruction over the cenozoic. *Chemical Geology*, 190(1), 123 – 140. Geochemistry of Crustal Fluids-Fluids in the Crust and Chemical Fluxes at the Earth's Surface.
- Lengger, S. K., Hopmans, E. C., Sinninghe Damsté, J. S., & Schouten, S. (2014). Impact of sedimentary degradation and deep water column production on gdgt abundance and distribution in surface sediments in the arabian sea: Implications for the tex86 paleothermometer. *Geochimica et Cosmochimica Acta*, 142, 386 – 399.

- Li, P. & Chen, S. (2016). A review on gaussian process latent variable models. *CAAI Transactions on Intelligence Technology*, 1(4), 366–376.
- Lin, L., Khider, D., Lisiecki, L. E., & Lawrence, C. E. (2014). Probabilistic sequence alignment of stratigraphic records. *Paleoceanography*, 29(10).
- Lisiecki, L. E. & Lisiecki, P. A. (2002). Application of dynamic programming to the correlation of paleoclimate records. *Paleoceanography*, 17(4), 1–1–1–12.
- Lisiecki, L. E. & Raymo, M. E. (2005). A pliocene-pleistocene stack of 57 globally distributed benthic $\delta^{18}\text{O}$ records. *Paleoceanography*, 20.
- Lisiecki, L. E. & Stern, J. V. (2016). Regional and global benthic $\delta^{18}\text{O}$ stacks for the last glacial cycle. *Paleoceanography*, 31(10), 1368–1394.
- Liu, H., Ong, Y.-S., Shen, X., & Cai, J. (2020). When gaussian process meets big data: A review of scalable gps. *IEEE transactions on neural networks and learning systems*.
- Liu, J. S. (2001). *Monte Carlo Strategies in Scientific Computing*. Springer Series in Statistics. Springer.
- Liu, Z. & Herbert, T. (2004). High-latitude influence on the eastern equatorial pacific climate in the early pleistocene epoch. *Nature*, 427, 720–723.
- Liu, Z., Pagani, M., Zinniker, D., DeConto, R., Huber, M., Brinkhuis, H., Shah, S. R., Leckie, R. M., & Pearson, A. (2009). Global cooling during the eocene-oligocene climate transition. *Science*, 323(5918), 1187–1190.
- Lodhi, H., Saunders, C., Shawe-Taylor, J., Cristianini, N., & Watkins, C. (2002). Text classification using string kernels. *Journal of Machine Learning Research*, 2, 419–444.
- Loftsgaarden, D. O. & Quesenberry, C. P. (1965). A nonparametric estimate of a multivariate density function. *Ann. Math. Statist.*, 36(3), 1049–1051.
- Lopes dos Santos, R. A., Prange, M., Castañeda, I. S., Schefuß, E., Mulitza, S., Schulz, M., Niedermeyer, E. M., Sinninghe Damsté, J. S., & Schouten, S. (2010). Glacial–interglacial variability in atlantic meridional overturning circulation and thermocline adjustments in the tropical north atlantic. *Earth and Planetary Science Letters*, 300(3), 407 – 414.
- Lü, X., Yang, H., Song, J., Versteegh, G. J., Li, X., Yuan, H., Li, N., Yang, C., Yang, Y., Ding, W., & Xie, S. (2014). Sources and distribution of isoprenoid glycerol dialkyl glycerol tetraethers (gdgts) in sediments from the east coastal sea of china: Application of gdgt-based paleothermometry to a shallow marginal sea. *Organic Geochemistry*, 75, 24 – 35.
- Lückge, A., Mohtadi, M., Rühlemann, C., Scheeder, G., Vink, A., Reinhardt, L., & Wiedicke, M. (2009). Monsoon versus ocean circulation controls on paleoenvironmental conditions off southern sumatra during the past 300,000 years. *Paleoceanography*, 24(1).

- Lüthi, D., Floch, M., Bereiter, B., Blunier, T., Barnola, J.-M., Siegenthaler, U., Raynaud, D., Jouzel, J., Fischer, H., Kawamura, K., & Stocker, T. (2008). High-resolution carbon dioxide concentration record 650,000-800,000 years before present. *Nature*, 453, 379–82.
- MacKay, D. J. (1998). Introduction to gaussian processes. *NATO ASI Series F Computer and Systems Sciences*, 168, 133–166.
- MacKay, D. J. C. (2003). *Information Theory, Inference and Learning Algorithms*. Cambridge University Press.
- Madureira, L. A. S., Conte, M. H., & Eglinton, G. (1995). Early diagenesis of lipid biomarker compounds in north atlantic sediments. *Paleoceanography*, 10(3), 627–642.
- Marchal, O., Cacho, I., Stocker, T. F., Grimalt, J. O., Calvo, E., Martrat, B., Shackleton, N., Vautravers, M., Cortijo, E., van Kreveld, S., Andersson, C., Koç, N., Chapman, M., Sbaffi, L., Duplessy, J.-C., Sarnthein, M., Turon, J.-L., Duprat, J., & Jansen, E. (2002). Apparent long-term cooling of the sea surface in the northeast atlantic and mediterranean during the holocene. *Quaternary Science Reviews*, 21(4), 455 – 483.
- Martino, L., Read, J., & Luengo, D. (2015). Independent doubly adaptive rejection metropolis sampling within gibbs sampling. *IEEE Transactions on Signal Processing*, 63(12), 3123–3138.
- Martrat, B., Grimalt, J. O., Shackleton, N. J., de Abreu, L., Hutterli, M. A., & Stocker, T. F. (2007). Four climate cycles of recurring deep and surface water destabilizations on the iberian margin. *Science*, 317(5837), 502–507.
- Matérn, B. (1986). *Spatial Variation*. Springer-Verlag.
- Max, L., Riethdorf, J.-R., Tiedemann, R., Smirnova, M., Lembke-Jene, L., Fahl, K., Nürnberg, D., Matul, A., & Mollenhauer, G. (2012). Sea surface temperature variability and sea-ice extent in the subarctic northwest pacific during the past 15,000 years. *Paleoceanography*, 27(3).
- Maybeck, P. (1979). *Stochastic Models, Estimation and Control*. Number pt. 1 in Mathematics in science and engineering. Academic Press.
- McCaffrey, M. A., Farrington, J. W., & Repeta, D. J. (1990). The organic geochemistry of peru margin surface sediments: I. a comparison of the c37 alkenone and historical el niño records. *Geochimica et Cosmochimica Acta*, 54(6), 1671 – 1682.
- McElhoe, B. A. (1966). An assessment of the navigation and course corrections for a manned flyby of mars or venus. *IEEE Transactions on Aerospace and Electronic Systems*, AES-2(4), 613–623.
- McManus, J. F., Francois, R., Gherardi, J.-M., Keigwin, L. D., & Brown-Leger, S. (2004). Collapse and rapid resumption of atlantic meridional circulation linked to deglacial climate changes. *Nature*, 428, 834–837.
- Metropolis, N., Rosenbluth, A. W., Rosenbluth, M. N., Teller, A. H., & Teller, E. (1953). Equation of State Calculations by Fast Computing Machines. *Journal of Chemical Physics*, 21, 1087–1092.

- Minka, T. P. (2001). *A family of algorithms for approximate Bayesian inference*. PhD thesis, MIT.
- Mohtadi, M., Oppo, D. W., Lückge, A., DePol-Holz, R., Steinke, S., Groeneveld, J., Hemme, N., & Hebbeln, D. (2011). Reconstructing the thermal structure of the upper ocean: Insights from planktic foraminifera shell chemistry and alkenones in modern sediments of the tropical eastern indian ocean. *Paleoceanography*, 26(3).
- Molyneux, E. G., Hall, I. R., Zahn, R., & Diz, P. (2007). Deep water variability on the southern agulhas plateau: Interhemispheric links over the past 170 ka. *Paleoceanography*, 22(4).
- Monnin, E., Indermühle, A., Dällenbach, A., Flückiger, J., Stauffer, B., Stocker, T. F., Raynaud, D., & Barnola, J.-M. (2001). Atmospheric co₂ concentrations over the last glacial termination. *Science*, 291(5501), 112–114.
- Mulitza, S., Prange, M., Stuut, J.-B., Zabel, M., von Dobeneck, T., Itambi, A. C., Nizou, J., Schulz, M., & Wefer, G. (2008). Sahel megadroughts triggered by glacial slowdowns of atlantic meridional overturning. *Paleoceanography*, 23(4).
- Munich, M. E. & Perona, P. (1999). Continuous dynamic time warping for translation-invariant curve alignment with applications to signature verification. In *Proceedings of the Seventh IEEE International Conference on Computer Vision*, volume 1 (pp. 108–115).
- Murphy, K. P. (2012). *Machine Learning: A Probabilistic Perspective*. The MIT Press.
- Müller, P. J., Kirst, G., Ruhland, G., [von Storch], I., & Rosell-Melé, A. (1998). Calibration of the alkenone paleotemperature index u_{37k'} based on core-tops from the eastern south atlantic and the global ocean (60°n-60°s). *Geochimica et Cosmochimica Acta*, 62(10), 1757 – 1772.
- Nadaraya, E. A. (1964). On estimating regression. *Theory of Probability and its Applications*, 9(1), 141–142.
- Neal, R. M. (1993). *Probabilistic inference using Markov chain Monte Carlo methods*. Technical report, University of Toronto.
- Neal, R. M. (2010). MCMC using Hamiltonian dynamics. *Handbook of Markov Chain Monte Carlo*, 54, 113–162.
- Nieto-Moreno, V., Martínez-Ruiz, F., Willmott, V., García-Orellana, J., Masqué, P., & Damsté, J. S. (2013). Climate conditions in the westernmost mediterranean over the last two millennia: An integrated biomarker approach. *Organic Geochemistry*, 55, 1 – 10.
- Norris, J. R. (1997). *Markov Chains*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press.
- Nürnberg, D., Ziegler, M., Karas, C., Tiedemann, R., & Schmidt, M. W. (2008). Interacting loop current variability and mississippi river discharge over the past 400 kyr. *Earth and Planetary Science Letters*, 272(1), 278 – 289.

- Ohkouchi, N., Kawamura, K., Kawahata, H., & Okada, H. (1999). Depth ranges of alkenone production in the central pacific ocean. *Global Biogeochemical Cycles - GLOBAL BIOGEOCHEM CYCLE*, 13, 695–704.
- Ostertag-Henning, C. & Stax, R. (2000). Data report: Carbonate records from sites 1012, 1013, 1017, and 1019 and alkenone-based sea-surface temperatures from site 1017. *Proceedings of the Ocean Drilling Program: Scientific Results*, 167, 297–302.
- Pahnke, K. & Sachs, J. P. (2006). Sea surface temperatures of southern midlatitudes 0–160 kyr b.p. *Paleoceanography*, 21(2).
- Pahnke, K., Sachs, J. P., Keigwin, L., Timmermann, A., & Xie, S.-P. (2007). Eastern tropical pacific hydrologic changes during the past 27,000 years from d/h ratios in alkenones. *Paleoceanography*, 22(4).
- Paillard, D. & Parrenin, F. (2004). The antarctic ice sheet and the triggering of deglaciations. *Earth and Planetary Science Letters*, 227(3), 263 – 271.
- Pailler, D. & Bard, E. (2002). High frequency palaeoceanographic changes during the past 140000 yr recorded by the organic matter in sediments of the iberian margin. *Palaeogeography, Palaeoclimatology, Palaeoecology*, 181(4), 431 – 452.
- Papoulis, A. & Pillai, U. (2001). *Probability, Random Variables and Stochastic Processes*. McGraw-Hill, 4th edition.
- Park, Y.-H., Yamamoto, M., Nam, S.-I., Irino, T., Polyak, L., Harada, N., Nagashima, K., Khim, B.-K., Chikita, K., & Saitoh, S.-I. (2014). Distribution, source and transportation of glycerol dialkyl glycerol tetraethers in surface sediments from the western arctic ocean and the northern bering sea. *Marine Chemistry*, 165, 10 – 24.
- Paterne, M., Kallel, N., Labeyrie, L., Vautravers, M., Duplessy, J.-C., Rossignol-Strick, M., Cortijo, E., Arnold, M., & Fontugne, M. (1999). Hydrological relationship between the north atlantic ocean and the mediterranean sea during the past 15-75 kyr. *Paleoceanography*, 14(5), 626–638.
- Pelejero, C. & Grimalt, J. O. (1997). The correlation between the 37k index and sea surface temperatures in the warm boundary: The south china sea. *Geochimica et Cosmochimica Acta*, 61(22), 4789 – 4797.
- Pelejero, C., Grimalt, J. O., Heilig, S., Kienast, M., & Wang, L. (1999). High-resolution uk 37 temperature reconstructions in the south china sea over the past 220 kyr. *Paleoceanography*, 14(2), 224–231.
- Penny, S. G. & Miyoshi, T. (2016). A local particle filter for high-dimensional geophysical systems. *Nonlinear Processes in Geophysics*, 23(6), 391–405.
- Peters, G. (2008). Markov chain monte carlo: stochastic simulation for bayesian inference. *Statistics in Medicine*, 27(16), 3213–3214.

- Petit, J. R., Jouzel, J., Raynaud, D., Barkov, N. I., Barnola, J.-M., BASILE-DOELSCH, I., Bender, M., Chappellaz, J., Davis, M., Delaygue, G., Delmotte, M., Kotlyakov, V. M., Legrand, M., Lipenkov, V. Y., Lorius, C., Pepin, L., Ritz, C., Saltzman, E., & Stievenard, M. (1999). Climate and atmospheric history of the past 420,000 years from the Vostok ice core, Antarctica. *Nature*, 399(6735), 429–436.
- Petitjean, F., Forestier, G., Webb, G. I., Nicholson, A. E., Chen, Y., & Keogh, E. J. (2014). Dynamic time warping averaging of time series allows faster and more accurate classification. *2014 IEEE International Conference on Data Mining*, (pp. 470–479).
- Petitjean, F., Ketterlin, A., & Gançarski, P. (2011). A global averaging method for dynamic time warping, with applications to clustering. *Pattern Recognition*, 44(3), 678–693.
- Poterjoy, J. (2016). A localized particle filter for high-dimensional nonlinear systems. *Monthly Weather Review*, 144(1), 59–76.
- Prahl, F., Rontani, J.-F., Zabeti, N., Walinsky, S., & Sparrow, M. (2010). Systematic pattern in $\delta^{18}O$ - temperature residuals for surface sediments from high latitude and other oceanographic settings. *Geochimica et Cosmochimica Acta*, 74, 131–143.
- Prahl, F. G., Mix, A. C., & Sparrow, M. A. (2006). Alkenone paleothermometry: Biological lessons from marine sediment records off western south america. *Geochimica et Cosmochimica Acta*, 70(1), 101 – 117.
- Prahl, F. G., Muehlhausen, L. A., & Lyle, M. (1989). An organic geochemical assessment of oceanographic conditions at manop site c over the past 26,000 years. *Paleoceanography*, 4(5), 495–510.
- Prahl, F. G., Muehlhausen, L. A., & Zahnle, D. L. (1988). Further evaluation of long-chain alkenones as indicators of paleoceanographic conditions. *Geochimica et Cosmochimica Acta*, 52(9), 2303–2310.
- Quiñonero-Candela, J. & Rasmussen, C. E. (2005). A unifying view of sparse approximate gaussian process regression. *Journal of Machine Learning Research*, 6, 1939–1959.
- Rasmussen, C. E. & Ghahramani, Z. (2002). Infinite mixtures of gaussian process experts. In T. G. Dietterich, S. Becker, & Z. Ghahramani (Eds.), *Advances in Neural Information Processing Systems 14* (pp. 881–888). MIT Press.
- Rasmussen, C. E. & Williams, C. K. I. (2006). *Gaussian Processes for Machine Learning (Adaptive Computation and Machine Learning)*. The MIT Press.
- Rauch, H. E., Tung, F., & Striebel, C. T. (1965). Maximum likelihood estimates of linear dynamic systems. *AIJA Journal*, 3(8), 1445–1450.
- Raynaud, D., Barnola, J.-M., Souchez, R., Lorrain, R., Petit, J.-R., Duval, P., & Lipenkov, V. Y. (2005). The record for marine isotopic stage 11. *Nature*, 436, 39–40.

- Reimer, P. J., Bard, E., Bayliss, A., Beck, J. W., Blackwell, P. G., Ramsey, C. B., Buck, C. E., Cheng, H., Edwards, R. L., & Friedrich, M. (2013). Intcal13 and marine13 radiocarbon age calibration curves 0–50,000 years cal bp. *Radiocarbon*, 55(4), 1869–1887.
- Rein, B., Lückge, A., Reinhardt, L., Sirocko, F., Wolf, A., & Dullo, W.-C. (2005). El niño variability off peru during the last 20,000 years. *Paleoceanography*, 20(4).
- Remes, S., Heinonen, M., & Kaski, S. (2017). Non-stationary spectral kernels. In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, & R. Garnett (Eds.), *Advances in Neural Information Processing Systems 30* (pp. 4642–4651). Curran Associates, Inc.
- Richey, J. N., Hollander, D. J., Flower, B. P., & Eglinton, T. I. (2011). Merging late holocene molecular organic and foraminiferal-based geochemical records of sea surface temperature in the gulf of mexico. *Paleoceanography*, 26(1).
- Rodrigo-Gámiz, M., Rampen, S. W., de Haas, H., Baas, M., Schouten, S., & Sinninghe Damsté, J. S. (2015). Constraints on the applicability of the organic temperature proxies $u_{37}^{K'}$, tex_{86} and l_{di} in the subpolar region around iceland. *Biogeosciences*, 12(22), 6573–6590.
- Rodrigues, T., Grimalt, J. O., Abrantes, F., Naughton, F., & Flores, J.-A. (2010). The last glacial–interglacial transition (lgit) in the western mid-latitudes of the north atlantic: Abrupt sea surface temperature change and sea level implications. *Quaternary Science Reviews*, 29(15), 1853 – 1862. Special Theme: Arctic Palaeoclimate Synthesis (PP. 1674-1790).
- Rodrigues, T., Grimalt, J. O., Abrantes, F. G., Flores, J. A., & Lebreiro, S. M. (2009). Holocene interdependences of changes in sea surface temperature, productivity, and fluvial inputs in the iberian continental shelf (tagus mud patch). *Geochemistry, Geophysics, Geosystems*, 10(7).
- Romero, O. E., Kim, J.-H., & Donner, B. (2008). Submillennial-to-millennial variability of diatom production off mauritania, nw africa, during the last glacial cycle. *Paleoceanography*, 23(3).
- Rosell-Melé, A. (1998). Interhemispheric appraisal of the value of alkenone indices as temperature and salinity proxies in high-latitude locations. *Paleoceanography*, 13, 694–703.
- Rosell-Melé, A., Comes, P., Muller, P., & Ziveri, P. (2000). Alkenone fluxes and anomalous $u_{37}^{K'}$ values during 1989–1990 in the northeast atlantic (48°n 21°w). *Marine Chemistry*, 71, 251–264.
- Rosell-Melé, A., Eglinton, G., Pflaumann, U., & Sarnthein, M. (1995). Atlantic core-top calibration of the $u_{37}^{K'}$ index as a sea-surface palaeotemperature indicator. *Geochimica et Cosmochimica Acta*, 59(15), 3099 – 3107.
- Rühlemann, C., Mulitza, S., Müller, P. J., Wefer, G., & Zahn, R. (1999). Warming of the tropical atlantic ocean and slowdown of thermohaline circulation during the last deglaciation. *Nature*, 402(6761), 511–514.
- Sachs, J. P. (2007). Cooling of northwest atlantic slope waters during the holocene. *Geophysical Research Letters*, 34(3).

- Salimbeni, H. & Deisenroth, M. P. (2017). Doubly stochastic variational inference for deep gaussian processes. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, NIPS'17 (pp. 4591–4602). Red Hook, NY, USA: Curran Associates Inc.
- Sarnthein, M., Winn, K., Jung, S. J. A., Duplessy, J.-C., Labeyrie, L., Erlenkeuser, H., & Ganssen, G. (1994). Changes in east atlantic deepwater circulation over the last 30,000 years: Eight time slice reconstructions. *Paleoceanography*, 9(2), 209–267.
- Sawada, K. & Handa, N. (1998). Variability of the path of the kuroshio ocean current over the past 25,000 years. *Nature*, 392, 592–595.
- Sawada, K., Handa, N., Shiraiwa, Y., Danbara, A., & Montani, S. (1996). Long-chain alkenones and alkyl alkenoates in the coastal and pelagic sediments of the northwest north pacific, with special reference to the reconstruction of emiliania huxleyi and gephyrocapsa oceanica ratios. *Organic Geochemistry*, 24(8), 751 – 764.
- Schefuss, E., Schouten, S., & Schneider, R. (2005). Climatic controls of central african hydrology during the last 20,000 years. *Nature*, 437, 1003–6.
- Schoenberg, I. J. (1938). Metric spaces and positive definite functions. *Transactions of the American Mathematical Society*, 44, 522–536.
- Schouten, S., Hopmans, E. C., Schefuß, E., & Damsté, J. S. S. (2002). Distributional variations in marine crenarchaeotal membrane lipids: a new tool for reconstructing ancient sea water temperatures? *Earth and Planetary Science Letters*, 204(1), 265 – 274.
- Schulz, H., Emeis, K.-C., Erlenkeuser, H., [von Rad], U., & Rolf, C. (2002). The toba volcanic event and interstadial/stadial climates at the marine isotopic stage 5 to 4 transition in the northern indian ocean. *Quaternary Research*, 57(1), 22 – 31.
- Schwarz, G. (1978). Estimating the dimension of a model. *Ann. Statist.*, 6(2), 461–464.
- Scott, G. H. (1963). Uniformitarianism, the uniformity of nature, and paleoecology. *New Zealand Journal of Geology and Geophysics*, 6(4), 510–527.
- Seki, O., Bendle, J. A., Harada, N., Kobayashi, M., Sawada, K., Moossen, H., Inglis, G. N., Nagao, S., & Sakamoto, T. (2014). Assessment and calibration of tex86 paleothermometry in the sea of okhotsk and sub-polar north pacific region: Implications for paleoceanography. *Progress in Oceanography*, 126, 254 – 266. Biogeochemical and physical processes in the Sea of Okhotsk and the linkages to the Pacific Ocean.
- Seki, O., Kawamura, K., Ikehara, M., Nakatsuka, T., & Oba, T. (2004). Variation of alkenone sea surface temperature in the sea of okhotsk over the last 85 kyrs. *Organic Geochemistry*, 35(3), 347 – 354. Selected papers from the Eleventh International Humic Substances Society Conference.
- Seki, O., Sakamoto, T., Sakai, S., Schouten, S., Hopmans, E. C., Sinninghe Damste, J. S., & Pancost, R. D. (2009). Large changes in seasonal sea ice distribution and productivity in the sea of okhotsk during the deglaciations. *Geochemistry, Geophysics, Geosystems*, 10(10).

- Sepúlveda, J., Pantoja, S., Hughen, K. A., Bertrand, S., Figueroa, D., León, T., Drenzek, N. J., & Lange, C. (2009). Late holocene sea-surface temperature and precipitation variability in northern patagonia, chile (jaca fjord, 44°s). *Quaternary Research*, 72(3), 400 – 409.
- Shackleton, N. J., Fairbanks, R. G., Chiu, T.-C., & Parrenin, F. (2004). Absolute calibration of the greenland time scale: implications for antarctic time scales and for $\Delta^{14}\text{C}$. *Quaternary Science Reviews*, 23(14), 1513–1522.
- Shackleton, N. J., Hall, M. A., & Edith, V. (2000). Phase relationships between millennial-scale events 64,000-24,000 years ago. *Paleoceanography*, 15.
- Shevenell, A., Ingalls, A., Domack, E., & Kelly, C. (2011). Holocene southern ocean surface temperature variability west of the antarctic peninsula. *Nature*, 470, 250–4.
- Shintani, T., Yamamoto, M., & Chen, M.-T. (2011). Paleoenvironmental changes in the northern south china sea over the past 28,000years: A study of tex86-derived sea surface temperatures and terrestrial biomarkers. *Journal of Asian Earth Sciences*, 40(6), 1221 – 1229. Quaternary Paleoclimate of the Western Pacific and East Asia: State of the Art and New Discovery.
- Siegenthaler, U., Stocker, T. F., Monnin, E., Lüthi, D., Schwander, J., Stauffer, B., Raynaud, D., Barnola, J.-M., Fischer, H., Masson-Delmotte, V., & Jouzel, J. (2005). Stable carbon cycle - climate relationship during the late pleistocene. *Science*, 310(5752), 1313–1317.
- Sikes, E., Farrington, J., & Keigwin, L. (1991). Use of the alkenone unsaturation ratio u_{37k} to determine past sea surface temperatures: core-top sst calibrations and methodology considerations. *Earth and Planetary Science Letters*, 104(1), 36 – 47.
- Sikes, E. L., Howard, W. R., Samson, C. R., Mahan, T. S., Robertson, L. G., & Volkman, J. K. (2009). Southern ocean seasonal temperature and subtropical front movement on the south tasman rise in the late quaternary. *Paleoceanography*, 24(2).
- Sikes, E. L. & Keigwin, L. D. (1994). Equatorial atlantic sea surface temperature for the last 30 kyr: A comparison of u_{37k} , $\delta^{18}\text{O}$ and foraminiferal assemblage temperature estimates. *Paleoceanography*, 9(1), 31–45.
- Sikes, E. L. & Keigwin, L. D. (1996). A reexamination of northeast atlantic sea surface temperature and salinity over the last 16 kyr. *Paleoceanography*, 11(3), 327–342.
- Sikes, E. L., Volkman, J. K., Robertson, L. G., & Pichon, J.-J. (1997). Alkenones and alkenes in surface waters and sediments of the southern ocean: Implications for paleotemperature estimation in polar regions. *Geochimica et Cosmochimica Acta*, 61(7), 1495 – 1505.
- Sinninghe Damsté, J., Schouten, S., Hopmans, E., Van Duin, A., & Geenevasen, J. (2002). Crenarchaeol: The characteristic core glycerol dibiphytanyl glycerol tetraether membrane lipid of cosmopolitan pelagic crenarchaeota. *Journal of Lipid Research*, 43(10), 1641–1651.
- Skinner, L. C., Elderfield, H., & Hall, M. (2013). *Phasing of Millennial Climate Events and Northeast Atlantic Deep-Water Temperature Change Since 50 Ka Bp*, (pp. 197–208). American Geophysical Union (AGU).

- Skinner, L. C. & Shackleton, N. J. (2004). Rapid transient changes in northeast atlantic deep water ventilation age across termination i. *Paleoceanography*, 19(2).
- Skinner, L. C. & Shackleton, N. J. (2005). An atlantic lead over pacific deep-water change across termination i: implications for the application of the marine isotope stage stratigraphy. *Quaternary Science Reviews*, 24(5), 571–580.
- Skinner, L. C., Shackleton, N. J., & Elderfield, H. (2003). Millennial-scale variability of deep-water temperature and $\delta^{18}\text{O}_{dw}$ indicating deep-water source variations in the northeast atlantic, 0–34 cal. ka bp. *Geochemistry Geophysics Geosystems - GEOCHEM GEOPHYS GEOSYST*, 4.
- Smith, M., De Deckker, P., Rogers, J., Brocks, J., Hope, J., Schmidt, S., Lopes dos Santos, R., & Schouten, S. (2013). Comparison of u37k', tex86h and ldi temperature proxies for reconstruction of south-east australian ocean temperatures. *Organic Geochemistry*, 64, 94 – 104.
- Sonzogni, C., Bard, E., & Rostek, F. (1998). Tropical sea-surface temperatures during the last glacial period: A view based on alkenones in indian ocean sediments. *Quaternary Science Reviews*, 17(12), 1185 – 1201.
- Sonzogni, C., Bard, E., Rostek, F., Dollfus, D., Rosell-Melé, A., & Eglinton, G. (1997). Temperature and salinity effects on alkenone ratios measured in surface sediments from the indian ocean. *Quaternary Research*, 47(3), 344–355.
- Sperling, M., Schmiedl, G., Hemleben, C., Emeis, K., Erlenkeuser, H., & Grootes, P. (2003). Black sea impact on the formation of eastern mediterranean sapropel s1? evidence from the marmara sea. *Palaeogeography, Palaeoclimatology, Palaeoecology*, 190, 9 – 21.
- Spratt, R. M. & Lisiecki, L. E. (2016). A late pleistocene sea level stack. *Climate of the Past*, 12(4), 1079–1092.
- Stein, M. L. (1999). *Interpolation of Spatial Data: Some Theory for Kriging*. Springer-Verlag New York.
- Stern, J. V. & Lisiecki, L. E. (2014). Termination 1 timing in radiocarbon-dated regional benthic $\delta^{18}\text{O}$ stacks. *Paleoceanography*, 29(12), 1127–1142.
- Stuart, A. & Ord, K. (1994). *Kendall's Advanced Theory of Statistics: Volume 1: Distribution Theory*. Wiley.
- Tao, S., Xing, L., Luo, X., Wei, H., & Liu, Y. (2012). Alkenone distribution in surface sediments of the southern yellow sea and implications for the u 37k' thermometer. *Geo-marine Letters - GEO-MAR LETT*, 32, 61–71.
- Ternois, Y., Kawamura, K., Ohkouchi, N., & Keigwin, L. (2000). Alkenone sea surface temperature in the okhotsk sea for the last 15 kyr. *GEOCHEMICAL JOURNAL*, 34(4), 283–293.
- Terrell, G. R. & Scott, D. W. (1992). Variable kernel density estimation. *Ann. Statist.*, 20(3), 1236–1265.

- Theodoridis, S. & Koutroumbas, K. (2008). *Pattern Recognition, Fourth Edition*. USA: Academic Press, Inc., 4th edition.
- Tierney, J. E. & Tingley, M. P. (2014). A bayesian, spatially-varying calibration model for the tex86 proxy. *Geochimica et Cosmochimica Acta*, 127, 83–106.
- Tierney, J. E. & Tingley, M. P. (2018). Bayspline: A new calibration for the alkenone paleothermometer. *Paleoceanography and Paleoclimatology*, 33(3), 281–301.
- Tipping, M. E. & Bishop, C. M. (1999). Probabilistic principal component analysis. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 61(3), 611–622.
- Titsias, M. (2009). Variational learning of inducing variables in sparse gaussian processes. In D. van Dyk & M. Welling (Eds.), *Proceedings of the 12th International Conference on Artificial Intelligence and Statistics*, volume 5 of *Proceedings of Machine Learning Research* (pp. 567–574). Hilton Clearwater Beach Resort, Clearwater Beach, Florida USA: PMLR.
- Titsias, M. & Lawrence, N. D. (2010). Bayesian gaussian process latent variable model. In Y. W. Teh & M. Titterton (Eds.), *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics*, volume 9 of *Proceedings of Machine Learning Research* (pp. 844–851). Chia Laguna Resort, Sardinia, Italy: PMLR.
- Tjallingii, R., Claussen, M., Stuut, J.-B. W., Fohlmeister, J., Jahn, A., Bickert, T., Lamy, F., & Röhl, U. (2008). Coherent high- and low-latitude control of the northwest african hydrological balance. *Nature Geoscience*, 1, 670–676.
- Tresp, V. (2001). Mixtures of gaussian processes. In *Advances in Neural Information Processing Systems 13* (pp. 654–660).: MIT Press.
- Trommer, G., Siccha, M., [van der Meer], M. T., Schouten, S., Damsté, J. S. S., Schulz, H., Hemleben, C., & Kucera, M. (2009). Distribution of crenarchaeota tetraether membrane lipids in surface sediments from the red sea. *Organic Geochemistry*, 40(6), 724 – 731.
- van der Vaart, A. W. (1998). *Asymptotic Statistics*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press.
- Verleye, T. (2011). *The late Quaternary palaeoenvironmental changes along the western South-American continental slope: a reconstruction based on dinoflagellate cysts and TEX86*. PhD thesis, Ghent University.
- Villacampa-Calvo, C., Zaldivar, B., Garrido-Merchán, E. C., & Hernández-Lobato, D. (2020). Multi-class gaussian process classification with noisy inputs. arXiv:2001.10523v2 [stat.ML].
- Viterbi, A. (1967). Error bounds for convolutional codes and an asymptotically optimum decoding algorithm. *IEEE Transactions on Information Theory*, 13(2), 260–269.
- Waelbroeck, C., Skinner, L. C., Labeyrie, L., Duplessy, J.-C., Michel, E., Vazquez Riveiros, N., Gherardi, J.-M., & Dewilde, F. (2011). The timing of deglacial circulation changes in the atlantic. *Paleoceanography*, 26(3).

- Wahba, G. (1990). *Spline Models for Observational Data*. CBMS-NSF Regional Conference Series in Applied Mathematics. Society for Industrial and Applied Mathematics.
- Wang, W. & Chen, X. (2016). The effects of estimation of heteroscedasticity on stochastic kriging. In *2016 Winter Simulation Conference (WSC)* (pp. 326–337).
- Wasserman, L. (2004). *All of Statistics: A Concise Course in Statistical Inference*. Springer Texts in Statistics. Springer.
- Wasserman, L. (2006). *All of Nonparametric Statistics*. Springer Texts in Statistics. Springer.
- Watson, G. S. (1964). Smooth regression analysis. *Sankhya: The Indian Journal of Statistics, Series A*, 26, 359–372.
- Wei, Y., Wang, J., Liu, J., Dong, L., Li, L., Wang, H., Wang, P., Zhao, M., & Zhang, C. L. (2011). Spatial variations in archaeal lipids of surface water and core-top sediments in the south china sea and their implications for paleoclimate studies. *Applied and Environmental Microbiology*, 77(21), 7479–7489.
- Weldeab, S., Schneider, R. R., & Müller, P. (2007). Comparison of mg/ca- and alkenone-based sea surface temperature estimates in the fresh water-influenced gulf of guinea, eastern equatorial atlantic. *Geochemistry, Geophysics, Geosystems*, 8(5).
- Wenzel, F., Galy-Fajou, T., Donner, C., Kloft, M., & Oppel, M. (2018). Efficient gaussian process classification using polygamma data augmentation. *Proceedings of the AAAI Conference on Artificial Intelligence*, 33.
- Williams, C. K. I. (1998). Computation with infinite neural networks. *Neural Computation*, 10(5), 1203–1216.
- Williams, C. K. I. & Barber, D. (1998). Bayesian classification with gaussian processes. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 20(12), 1342–1351.
- Wu, W., Tan, W., Zhou, L., Yang, H., & Xu, Y. (2012). Sea surface temperature variability in southern okinawa trough during last 2700 years. *Geophysical Research Letters*, 39(14).
- Wuchter, C., Schouten, S., Coolen, M. J. L., & Sinninghe Damsté, J. S. (2004). Temperature-dependent variation in the distribution of tetraether membrane lipids of marine crenarchaeota: Implications for tex86 paleothermometry. *Paleoceanography*, 19(4).
- Zarriess, M., Johnstone, H., Prange, M., Steph, S., Groeneveld, J., Mulitza, S., & Mackensen, A. (2011). Bipolar seesaw in the northeastern tropical atlantic during heinrich stadials. *Geophysical Research Letters*, 38(4).
- Zarriess, M. & Mackensen, A. (2011). Testing the impact of seasonal phytodetritus deposition on $\delta^{13}\text{C}$ of epibenthic foraminifer cibicidoides wuellerstorfi: A 31,000 year high-resolution record from the northwest african continental slope. *Paleoceanography*, 26(2).

- Zell, C., Kim, J.-H., Hollander, D., Lorenzoni, L., Baker, P., Silva, C. G., Nittrouer, C., & Damsté, J. S. S. (2014). Sources and distributions of branched and isoprenoid tetraether lipids on the amazon shelf and fan: Implications for the use of gdgt-based proxies in marine sediments. *Geochimica et Cosmochimica Acta*, 139, 293 – 312.
- Zhao, M., Beveridge, N. A. S., Shackleton, N. J., Sarnthein, M., & Eglinton, G. (1995). Molecular stratigraphy of cores off northwest africa: Sea surface temperature history over the last 80 ka. *Paleoceanography*, 10(3), 661–675.
- Zhao, M., Huang, C.-Y., Wang, C.-C., & Wei, G. (2006). A millennial-scale U37K' sea-surface temperature record from the South China Sea (8°N) over the last 150 kyr: Monsoon and sea-level influence. *Palaeogeography Palaeoclimatology Palaeoecology*, 236(1-2), 39–55.
- Zhou, H., Hu, J., Spiro, B., Peng, P., & Tang, J. (2014). Glycerol dialkyl glycerol tetraethers in surficial coastal and open marine sediments around china: Indicators of sea surface temperature and effects of their sources. *Palaeogeography, Palaeoclimatology, Palaeoecology*, 395, 114 – 121.



THIS THESIS WAS TYPESET using $\text{\LaTeX}$, originally developed by Leslie Lamport and based on Donald Knuth's $\text{\TeX}$. The body text is set in 11 point Egenolff-Berner Garamond, a revival of Claude Garamont's humanist typeface. The above illustration, "Science Experiment 02", was created by Ben Schlitter and released under CC BY-NC-ND 3.0. A template that can be used to format a PhD thesis with this look and feel has been released under the permissive MIT (x11) license, and can be found online at github.com/suchow/Dissertate or from its author, Jordan Suchow, at suchow@post.harvard.edu.