
Pseudo-Marginal Inference for Gaussian Process Classification with Large Datasets

Daniel Williams `daniel.williams@bristol.ac.uk`

Dom Owens `dom.owens@bristol.ac.uk`

Jake Spiteri `jake.spiteri@bristol.ac.uk`

COMPASS CDT, Institute of Statistical Science, University of Bristol

August 18, 2026

Gaussian Process (GP) models are a powerful tool for non-parametric supervised learning tasks, popular in both machine learning and statistics. Following the methods of Filippone and Girolami, 2014, we outline and implement a GP for binary classification with a fully-Bayesian specification, using a pseudo-marginal Metropolis-Hastings algorithm to sample the model hyperparameters from their posteriors. Due to the computational expense of this approach, we extend this model to handle large datasets using the *Informative Vector Machine* of Herbrich, Lawrence, and Seeger, 2003. We detail each step of the modelling process, and compare the performance of the classifier against a Bayesian logistic regression model on the e-mail spam dataset (Dua and Graff, 2017). Our implementation is available as an R package at github.com/dannyjameswilliams/gpc.

1 Introduction

Consider the probabilistic binary classification problem: each input datum x may belong to one of two classes, $\mathcal{C}_-$ or $\mathcal{C}_+$. Membership is non-deterministic, and we denote class membership for the datum with $y = \pm 1$. After observing pairs $(x_i, y_i), i = 1, \dots, n$, we wish to assign a new observation x_* to a class. We first predict the probability of the data belonging to each class, then assign a class by following some decision rule, chosen to take into account the possible loss incurred by each decision.

Our focus here is on Gaussian processes. These are a flexible non-parametric method, which can closely approximate a rich class of non-linear data generating processes, and the probabilistic nature of these models permits us to quantify how uncertain we are in the predictions we make.

To understand Gaussian process classification (GPC), we should begin by understanding the related problem of Gaussian process regression (GPR), where the aim is to approximate a continuous-response process; a function $f(\cdot)$ defines the mean, and a kernel function $k(\cdot, \cdot)$ defines the covariance. A Gaussian likelihood and a Gaussian prior distribution give rise to a Gaussian posterior, and thus all the necessary computations are analytically tractable. In the classification problem, however, a Gaussian likelihood would be nonsensical, which would mean we lack an analytical form for the posterior distribution. To overcome this, we can approximate the posterior over the latent variables f with a Gaussian distribution (via the Laplace approximation). This will provide us with conjugacy, and allows us to integrate out the latent variables.

In much the same manner that GPR can be cast as a generalisation of linear regression (Williams and Rasmussen, 2006), GPC is a generalisation of linear logistic regression. Data are mapped to class probabilities by applying a sigmoid function to a predicted continuous response; in the general case we allow this response function to be a smooth non-linear function of the data, defined by the chosen hyperparameters.

The standard approach to inferring hyperparameters would be to integrate out the latent variables using some approximation method, then optimise the resulting approximate marginal likelihood; this is called *type II Maximum Likelihood* (ML). Many deterministic approximations for integrating out $\mathbf{f}$ exist, including the Laplace Approximation (Section 2.2) used in this report.

It would of course be desirable to extend this to a fully Bayesian approach. We would like to (i) infer all parameters and latent variables from the data, and (ii) make predictions by integrating out latent variables and hyperparameters with respect to their posterior distributions. To do this, we require the joint posterior distribution of the latent variables and the model hyperparameters, which is intractable. Thus we must approximate this quantity using a Markov chain Monte Carlo (MCMC) method, which is itself a non-trivial task given the strong coupling of the latent variables and hyperparameters. To overcome this problem we use a *pseudo-marginal* (PM) approach as proposed in Filippone and Girolami, 2014 to sample the model hyperparameters from their posterior distribution. The PM approach allows us to approximately integrate out the latent variables whilst maintaining the correct invariant distribution, decoupling the samples and thus providing us with better mixing.

Moreover, adoption of GP models in practice has been hindered by the high computational cost of fitting and predicting for large datasets, given these tasks have cost $\mathcal{O}(n^3)$. One method to overcome this was proposed by Herbrich, Lawrence, and Seeger, 2003, which selects a much smaller subset of the training data while minimising the impact on predictive performance. We describe how this can be combined with our model in Section 2.3.

Importantly, we have found no current literature on sparse estimation for fully-Bayesian GPC models, and very little on making these models scalable. Hensman, Matthews, and Ghahramani, 2015 used a variation approach, while Hernández-Lobato and Hernández-Lobato, 2016 uses expectation propagation for sparsity, but is not fully Bayesian.

The rest of the report proceeds as follows:

- In Section 2, the entire modelling process is described, with detailed explanations of how and why quantities are obtained at each step. We present the sparse approximation for use of this model with large datasets, and we prove that the pseudo-marginal approach defines a valid Metropolis-Hastings algorithm. The procedure is presented as a schematic algorithm.
- In Section 3, the computational aspects of the model are discussed; we detail the ways in which runtime can be reduced, and explain the ‘tricks’ used to provide the method with numerical stability. We also compare the performance of the model to a novel Bayesian logistic regression on the spambase data set.

- We conclude in Section 4, collecting our findings and detailing areas for further work.

2 Model

We specify our model as follows: Let $X = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ be a set of n -many input vectors each described by d -many covariates, and each associated with a response $y_i \in \{-1, +1\}$, contained in the vector $\mathbf{y} = \{y_1, \dots, y_n\}$. Let $\mathbf{f} = \{f_1, \dots, f_n\}$ be the set of latent variables. Using a generative approach, the classifier assumes the class labels are Bernoulli distributed with success probability determined by a transformation of the class label and latent variable:

$$p(y_i|f_i) = \Phi(y_i f_i). \quad (1)$$

Here, Φ is the cumulative density function of the standard Gaussian distribution. By independence, the likelihood is the product over all observations

$$p(\mathbf{y}|\mathbf{f}) = \prod_{i=1}^n p(y_i|f_i). \quad (2)$$

We model $\mathbf{f}$ with a zero mean Gaussian process prior, with covariance matrix $\mathbf{K}$:

$$\mathbf{f} \sim \mathcal{N}(\mathbf{f}|\mathbf{0}, \mathbf{K}). \quad (3)$$

Each entry $K_{i,j}$ is determined by the evaluation

$$k(\mathbf{x}_i, \mathbf{x}_j|\boldsymbol{\theta}) = \sigma \exp \left[-\frac{1}{2}(\mathbf{x}_i - \mathbf{x}_j)^\top \mathbf{A}(\mathbf{x}_i - \mathbf{x}_j) \right] \quad (4)$$

of our chosen kernel, which is in turn defined by the hyperparameter vector $\boldsymbol{\theta} = (\sigma, \gamma)$. σ controls the variance of the marginal prior, while γ is the length-scale control such that $\mathbf{A} = \gamma^{-2} \mathbf{I}_d$, giving an isotropic covariance matrix. This specific kernel is known as the Gaussian kernel. We place Gamma priors on the hyperparameters $\boldsymbol{\theta}$ as follows:

$$\sigma \sim \text{Gamma}(\alpha_\sigma, \beta_\sigma), \quad (5)$$

$$\gamma \sim \text{Gamma}(\alpha_\gamma, \beta_\gamma), \quad (6)$$

where $\alpha_\sigma, \beta_\sigma, \alpha_\gamma, \beta_\gamma > 0$.

2.1 Predictive Distribution

With the aim of classifying new data, we want to obtain an approximate predictive distribution for a prediction y_* . Conditioning on a realisation of the hyperparameters $\boldsymbol{\theta}$, in the standard GPC setting we would compute

$$p(y_*|\mathbf{y}, \boldsymbol{\theta}) = \int p(y_*|f_*) p(f_*|\mathbf{f}, \boldsymbol{\theta}) q(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta}) d\mathbf{f}_* d\mathbf{f}. \quad (7)$$

Here $\boldsymbol{\theta}$ could be a ML estimate that maximizes the approximate likelihood or one sample from the approximate posterior $\hat{p}(\boldsymbol{\theta}|\mathbf{y})$. Note that we suppress notation

for dependence on the data X when writing posterior distributions, and treat this as implicit.

However, we want a fully Bayesian approach, so we integrate out the hyperparameters also, yielding the distribution

$$p(y_*|\mathbf{y}) = \int p(y_*|f_*) p(f_*|\mathbf{f}, \boldsymbol{\theta}) p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{y}) d\mathbf{f}_* d\mathbf{f} d\boldsymbol{\theta}. \quad (8)$$

It is not possible to analytically integrate equation (8) with respect to $\mathbf{f}$ or $\boldsymbol{\theta}$, so we draw samples from $p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{y})$ with a Monte Carlo method (described in Section 2.4) and use an unbiased estimate for the predictive distribution given by

$$p(y_*|\mathbf{y}) \simeq \frac{1}{N} \sum_{i=1}^N \int p(y_*|f_*) p(f_*|\mathbf{f}^{(i)}, \boldsymbol{\theta}^{(i)}) d\mathbf{f}_*, \quad (9)$$

where $\mathbf{f}^{(i)}, \boldsymbol{\theta}^{(i)}$ denotes the i -th sample from $p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{y})$.

For simplicity of notation, let $\mathbf{K}$ be the covariance matrix evaluated at $\boldsymbol{\theta}$, $\mathbf{k}_*$ be the vector whose i -th element is $k(\mathbf{x}_i, \mathbf{x}_*|\boldsymbol{\theta})$, and $k_{**} = k(\mathbf{x}_*, \mathbf{x}_*|\boldsymbol{\theta})$ be the kernel evaluated on the new set of observations $\mathbf{x}_*$.

Given the properties of multivariate Gaussian variables, and as per the specification of the GP model, f_* is distributed as $\mathcal{N}(f_*|\mu_*, \beta_*^2)$ with $\mu_* = \mathbf{k}_*^T \mathbf{K}^{-1} \mathbf{f}$ and $\beta_*^2 = k_{**} - \mathbf{k}_*^T \mathbf{K}^{-1} \mathbf{k}_*$.

Approximating $p(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta})$ with the Laplace approximation (2.2) makes it possible to perform analytic integration with respect to $\mathbf{f}$ in equation (8). In particular, the integration with respect to $\mathbf{f}$ yields $\mathcal{N}(f_*|m_*, s_*^2)$ with the predictive mean

$$m_* = \mathbf{k}_*^T \mathbf{K}^{-1} \boldsymbol{\mu}_q, \quad (10)$$

and predictive variance

$$s_*^2 = k_{**} - \mathbf{k}_*^T \mathbf{K}^{-1} \mathbf{k}_* + \mathbf{k}_*^T \mathbf{K}^{-1} \boldsymbol{\Sigma}_q \mathbf{K}^{-1} \mathbf{k}_*, \quad (11)$$

where $\boldsymbol{\mu}_q$ and $\boldsymbol{\Sigma}_q$ parameterise the mean and covariance respectively of the Laplace approximation.

The integration with respect to f_* has a closed form, as it is a convolution of a Gaussian and a cumulative Gaussian:

$$t_* = \int p(y_*|f_*) \mathcal{N}(f_*|m_*, s_*^2) df_* = \Phi\left(\frac{m_*}{\sqrt{1 + s_*^2}}\right). \quad (12)$$

We make a prediction y_* for $\mathbf{x}_*$ from the predictive distribution according to the $a_- - a_+$ decision rule $\delta : [0, 1] \rightarrow \{-1, +1\}$ such that

$$y_* = \delta(t_*) = \begin{cases} +1, & t_* > a_+/(a_- + a_+), \\ -1, & \text{otherwise.} \end{cases} \quad (13)$$

Here, $a_-, a_+ \geq 0$ are chosen to quantify the loss associated with wrongly predicting y_- and y_+ respectively.

Note the similarities of this classification rule to the hard-margin support vector machine (SVM) (see, for instance, Williams and Rasmussen, 2006, Section 6).

2.2 Laplace Approximation

By approximating the posterior of the latent variables given the data and hyperparameters $p(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta})$ (the target density) with a Gaussian $\mathcal{N}(\mathbf{f}|\boldsymbol{\mu}_q, \boldsymbol{\Sigma}_q)$, we obtain the **Laplace approximation**:

$$p(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta}) \approx q(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}(\mathbf{f}|\boldsymbol{\mu}_q, \boldsymbol{\Sigma}_q). \quad (14)$$

Substituting this into the predictive distribution (8) gives conjugacy, permitting analytical integration. For the approximation to be sufficiently close in location and curvature, we set $\boldsymbol{\Sigma}_q^{-1}$ equal to the negative Hessian of the target density.

Given a realisation of the hyperparameters $\boldsymbol{\theta}$, define

$$\Psi(\mathbf{f}) = \log[p(\mathbf{y}|\mathbf{f})] + \log[p(\mathbf{f}|\boldsymbol{\theta})] + C \quad (15)$$

as the logarithm of the target density, constant with respect to terms independent of $\mathbf{f}$ (represented by C).

We select a Gaussian $q(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}(\mathbf{f}|\hat{\mathbf{f}}, \hat{\boldsymbol{\Sigma}})$ where

$$\hat{\mathbf{f}} = \arg \max_{\mathbf{f}} \Psi(\mathbf{f}), \quad \hat{\boldsymbol{\Sigma}}^{-1} = -\nabla_{\mathbf{f}} \nabla_{\mathbf{f}} \Psi(\hat{\mathbf{f}}), \quad (16)$$

where the solution $\hat{\mathbf{f}}$ can be solved by iterating the Newton-Raphson formula

$$\mathbf{f}_{\text{new}} = \mathbf{f} - (\nabla_{\mathbf{f}} \nabla_{\mathbf{f}} \Psi(\mathbf{f}))^{-1} \nabla_{\mathbf{f}} \Psi(\mathbf{f}), \quad (17)$$

starting from $\mathbf{f} = \mathbf{0}$ until convergence. The gradient and the Hessian of the log-target density are:

$$\nabla_{\mathbf{f}} \Psi(\mathbf{f}) = \nabla_{\mathbf{f}} \log[p(\mathbf{y}|\mathbf{f})] - \mathbf{K}^{-1} \mathbf{f}, \quad (18)$$

$$\nabla_{\mathbf{f}} \nabla_{\mathbf{f}} \Psi(\mathbf{f}) = \nabla_{\mathbf{f}} \nabla_{\mathbf{f}} \log[p(\mathbf{y}|\mathbf{f})] - \mathbf{K}^{-1}. \quad (19)$$

The log-target density has a unique maximum by concavity.

2.3 Sparse Approximation

For large data sets, it is necessary to use some approximation to reduce the $\mathcal{O}(n^3)$ complexity of fitting the full model. For a discussion of possible methods, see Williams and Rasmussen, 2006, Section 8.5. We propose the use of a *subset of data* (SD) method as described in Herbrich, Lawrence, and Seeger, 2003, referred to as the *Informative Vector Machine*. This defines a sparse GPC model by greedily adding data entries to the active subset $I \subset \{1, \dots, n\}$ until $|I| = d$, where $d \ll n$ is chosen as the reduced subset size. Points are added according to an entropy heuristic; at each iteration, a score function is calculated for each point in the inactive set $J = \{1, \dots, n\} \setminus I$ and we select the maximiser as candidate.

Given a Laplace approximation $q(\mathbf{f}|\mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}(\mathbf{f}|\boldsymbol{\mu}, \boldsymbol{\Sigma})$ from (14), the chosen heuristic, *differential entropy score* (26), favours the point which most decreases the posterior variance when included in the active set.

We adapt the following from Herbrich, Lawrence, and Seeger, 2003, Section 3 into our notation. We

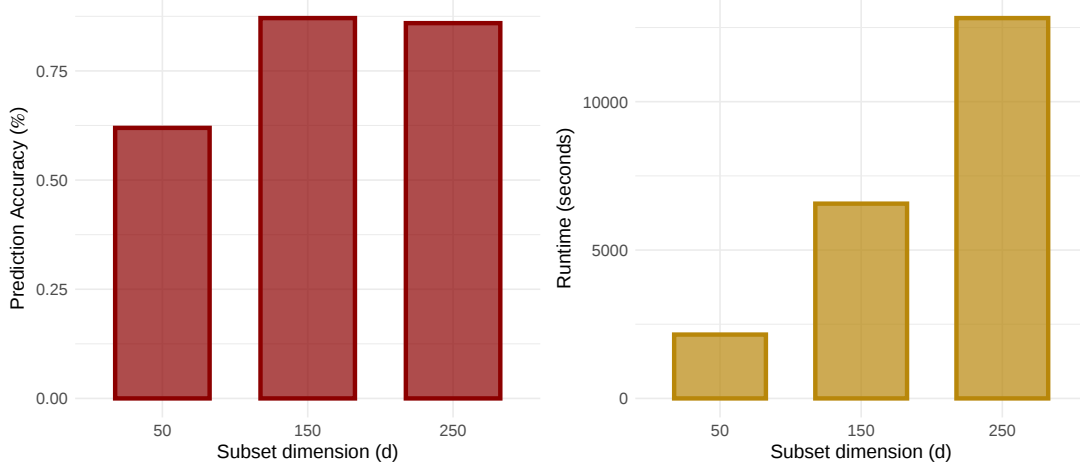


Figure 1: Predictive accuracy and runtime, by IVM subset size.

Algorithm 1: Informative vector machine algorithm (IVM)

input : Target sparsity $d \ll n$
initialise : $I = \emptyset, \mathbf{m} = \mathbf{0}$,
 $\mathbf{\Pi} = \text{diag}(\mathbf{0})$,
 $\text{diag}(\mathbf{\Sigma}) = \text{diag}(\mathbf{K})$,
 $\boldsymbol{\mu} = \mathbf{0}, J = \{1, \dots, n\}$
output : $I \subset \{1, \dots, n\}$ with $|I| = d$

```

1 repeat
2   for  $j \in J$  do
3     | Compute  $\Delta_j$  from (26)
4   end
5    $i = \arg \max_{j \in J} \Delta_j$ 
6   Do updates for  $p_i$  and  $m_i$  according to (22)
7   Update matrices  $\mathbf{L}, \mathbf{M}, \text{diag}(\mathbf{\Sigma})$  and  $\boldsymbol{\mu}$ 
   according to (23)
8    $I \leftarrow I \cup \{i\}, J \leftarrow J \setminus \{i\}$ .
9 until  $|I| = d$ 
    
```

have $\mathbf{\Sigma} := (\mathbf{K}^{-1} + \mathbf{\Pi})^{-1}, \boldsymbol{\mu} := \mathbf{\Sigma} \mathbf{\Pi} \mathbf{m}$ and $\mathbf{\Pi} := \nabla_{\mathbf{f}} \nabla_{\mathbf{f}} \log[p(\mathbf{y}|\mathbf{f})] = \text{diag}(\mathbf{p})$ for vectors $\mathbf{m}, \mathbf{p}$. If I is the current active set, then all of components of the $\mathbf{p}$ and $\mathbf{m}$ which are not in I are zero, and by the Woodbury identity we have

$$\mathbf{\Sigma} = \mathbf{K} - \mathbf{M}^T \mathbf{M}, \quad \mathbf{M} = \mathbf{L}^{-1} \mathbf{\Pi}_I^{1/2} \mathbf{K}_{I,\cdot} \in \mathbb{R}^{d,n}, \quad (20)$$

where $\mathbf{L}$ is the lower-triangular Cholesky factor of

$$\mathbf{B} = \mathbf{I} + \mathbf{\Pi}_I^{1/2} \mathbf{K}_I \mathbf{\Pi}_I^{1/2} \in \mathbb{R}^{d,d}. \quad (21)$$

In order to compute the differential entropy score for a point $j \notin I$, we have to know $\Sigma_{j,j}$ and μ_j . Thus, when including i into the active set I , we need to update $\text{diag}(\mathbf{\Sigma})$ and $\boldsymbol{\mu}$ accordingly, which in turn requires the matrices $\mathbf{L}$ and $\mathbf{M}$ to be kept up-to-date. The update

equations for p_i, m_i are

$$\begin{aligned}
 p_i &= \frac{\nu_i}{1 - \Sigma_{i,i} \nu_i}, & m_i &= \mu_i + \frac{\alpha_i}{\nu_i}, & \text{where} \\
 z_i &= \frac{y_i \cdot (\mu_i)}{\sqrt{1 + \Sigma_{i,i}}}, & \alpha_i &= \frac{y_i \cdot N(z_i|0,1)}{\Phi(z_i) \sqrt{1 + \Sigma_{i,i}}}, \\
 \nu_i &= \alpha_i \left(\alpha_i + \frac{\mu_i}{1 + \Sigma_{i,i}} \right).
 \end{aligned} \quad (22)$$

We then update $\mathbf{L} \rightarrow \mathbf{L}^{\text{new}}$ by appending the row $(\mathbf{l}^T, l)$, and $\mathbf{M} \rightarrow \mathbf{M}^{\text{new}}$ by appending the row $\mathbf{h}^T$, where

$$\begin{aligned}
 l &= \sqrt{p_i} M_{\cdot,i}, & l &= \sqrt{1 + p_i K_{i,i} - \mathbf{l}^T \mathbf{l}}, \\
 \mathbf{h} &= \mathbf{l}^{-1} (\sqrt{p_i} K_{\cdot,i} - \mathbf{M}^T \mathbf{l}).
 \end{aligned} \quad (23)$$

The final updates are

$$\text{diag}(\mathbf{\Sigma}^{\text{new}}) \leftarrow \text{diag}(\mathbf{\Sigma}) - (h_j^2)_j, \quad (24)$$

$$\boldsymbol{\mu}^{\text{new}} \leftarrow \boldsymbol{\mu} + \alpha_i l p_i^{-1/2} \mathbf{h}. \quad (25)$$

The differential entropy score for $j \notin I$ can be computed based on the variables in (22) (with $i \rightarrow j$) as

$$\Delta_j = \frac{1}{2} \log(1 - \Sigma_{j,j} \nu_j) \quad (26)$$

This has total complexity $\mathcal{O}(nd^2)$. Fitting the subsequent GPC model will have complexity $\mathcal{O}(d^3)$, which leads to significant runtime reductions in practice, and does so while minimising the trade-off in predictive performance, given we have selected an information-optimal subset of data.

For the rest of the paper, we continue to refer to the full dataset without loss of generality. Any calculations can be assumed to have been performed on a sparse subset instead.

2.4 MCMC Sampling from $p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{y})$

As we have pointed out, sampling from the posterior $p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{y})$ with joint proposals is not feasible, given the strong correlation between the two sampled quantities. Instead, we update $\mathbf{f}$ and $\boldsymbol{\theta}$ in turn with Gibbs samples.

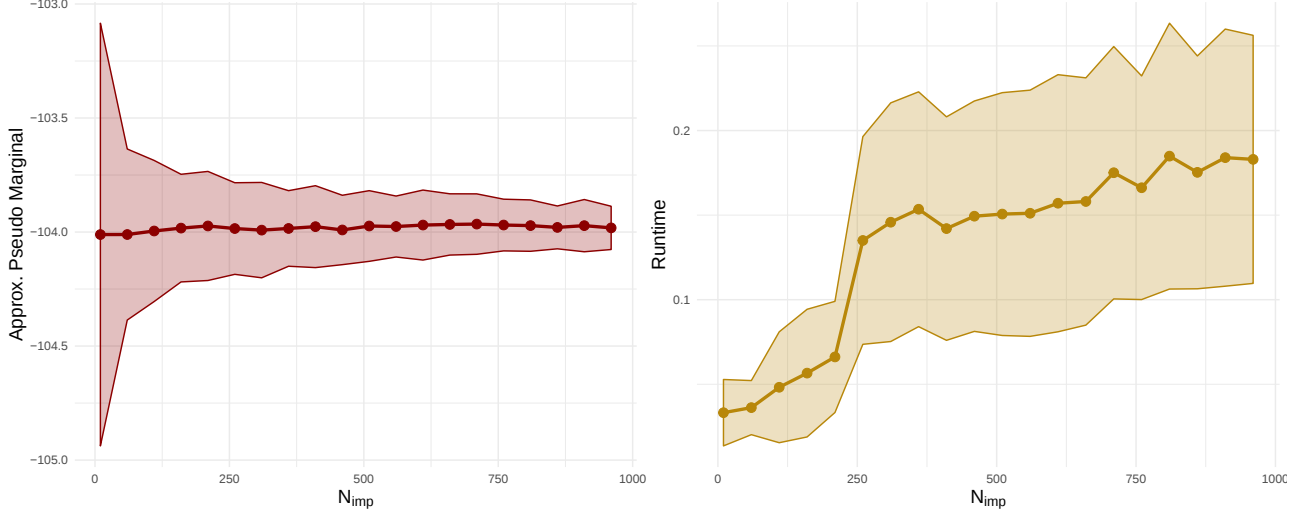


Figure 2: Number of importance samples N_{imp} against the mean and 95% interval of the approximate pseudo-marginal evaluation calculated from (30) (left), and against the mean and 95% interval of the computation time for each step (right). At each value of N_{imp} , the approximate marginal value was calculated 50 times.

2.4.1 Sampling from $p(\mathbf{f}|\mathbf{y}, \theta)$

We use the *elliptical slice sampling* (ELL-SS) (Murray, Prescott Adams, and MacKay, 2010) method to sample our latent variables $\mathbf{f}$. We express the posterior for $\mathbf{f}$ as

$$p(\mathbf{f}|\mathbf{y}, \theta) \propto \mathcal{N}(\mathbf{f}|\hat{\mathbf{f}}, \hat{\Sigma})p(\mathbf{y}|\mathbf{f}) \quad (27)$$

with the likelihood function $p(\mathbf{y}|\mathbf{f}) = L(\mathbf{f})$. This method samples according to the ELL-SS algorithm below, constrained to an ellipse centred at the current mean $\mathbf{f}$, with eccentricity determined by the covariance $\hat{\Sigma}$. This encourages steps of greater length in directions with greater covariance, and obviates the need for sample rejection.

The benefits of this are that no tuning is required, and each iteration has complexity $\mathcal{O}(n^2)$ (or $\mathcal{O}(d^2)$ under the reduction in Section 2.3), which is dominated by the complexity of fitting the entire model.

2.4.2 Sampling from $p(\theta|\mathbf{y})$ with the Pseudo-Marginal approach

It remains to sample from $p(\theta|\mathbf{y})$. While we can't do this by integrating out the hyperparameters $\mathbf{f}$, we can draw MCMC samples using an unbiased estimate of $p(\mathbf{y}|\theta)$.

Consider the Metropolis-Hastings algorithm with proposal $\pi(\theta'|\theta)$. We aim to sample from

$$p(\theta|\mathbf{y}) \propto p(\mathbf{y}|\theta)p(\theta). \quad (28)$$

The acceptance ratio would be

$$z = \frac{p(\mathbf{y}|\theta')p(\theta') \pi(\theta|\theta')}{p(\mathbf{y}|\theta)p(\theta) \pi(\theta'|\theta)}. \quad (29)$$

We are unable to do this due to the intractability of integrating out the latent variables for $p(\mathbf{y}|\theta)$. As

the following theorem demonstrates, however, we can sample from $p(\theta|\mathbf{y})$ using only an unbiased estimator for $p(\mathbf{y}|\theta)$.

One possible unbiased estimator $\tilde{p}(\mathbf{y}|\theta)$ can be found through importance sampling. Drawing N_{imp} -many samples $\mathbf{f}_i$ from $q(\mathbf{f}|\mathbf{y}, \theta)$ allows the approximation

$$\tilde{p}(\mathbf{y}|\theta) \simeq \frac{1}{N_{imp}} \sum_{i=1}^{N_{imp}} \frac{p(\mathbf{y}|\mathbf{f}_i)p(\mathbf{f}_i|\theta)}{q(\mathbf{f}_i|\mathbf{y}, \theta)}. \quad (30)$$

Lemma 1. $\tilde{p}(\mathbf{y}|\theta)$ is an unbiased estimator for $p(\mathbf{y}|\theta)$.

Proof. Each importance sample $\mathbf{f}_i$ is independent, so we only need to show $(p(\mathbf{y}|\mathbf{f}_1)p(\mathbf{f}_1|\theta))q(\mathbf{f}_1|\mathbf{y}, \theta)^{-1}$ has expectation equal to $E[p(\mathbf{y}|\theta)]$.

The Laplace approximation $q(\mathbf{f}_1|\mathbf{y}, \theta)$ has expectation $p(\mathbf{f}_1|\mathbf{y}, \theta)$, so by Bayes' theorem $(p(\mathbf{y}|\mathbf{f}_1)p(\mathbf{f}_1|\theta))q(\mathbf{f}_1|\mathbf{y}, \theta)^{-1}$ is unbiased, and hence $\tilde{p}(\mathbf{y}|\theta)$ is unbiased. $\square$

From the results of Beaumont, 2003, we can plug in an unbiased estimate $\tilde{p}(\mathbf{y}|\theta)$ and draw samples from the correct posterior $p(\theta|\mathbf{y})$.

Theorem 1. Using the unbiased estimator $\tilde{p}(\mathbf{y}|\theta)$ in place of $p(\mathbf{y}|\theta)$, we have a valid Metropolis-Hastings algorithm with $p(\theta|\mathbf{y})$ as the invariant distribution.

Proof. Here, we follow the argument from Beaumont, 2003, Appendix. With importance samples

$$\mathbf{f}'_{1:N_{imp}} = \{\mathbf{f}'_i, i = 1, \dots, N_{imp}\} \quad (31)$$

and acceptance ratio

$$\tilde{z} = \frac{\tilde{p}(\mathbf{y}|\theta')p(\theta') \pi(\theta|\theta')}{\tilde{p}(\mathbf{y}|\theta)p(\theta) \pi(\theta'|\theta)}, \quad (32)$$

Algorithm 2: Elliptical Slice Sampling (ELL-SS)

input : current state $\mathbf{f}$, a routine that samples from $\mathcal{N}(\hat{\mathbf{f}}, \hat{\Sigma})$, log-likelihood function $\log L$
output : a new state $\mathbf{f}'$. When $\mathbf{f}$ is drawn from $p^*(\mathbf{f}) \propto \mathcal{N}(\hat{\mathbf{f}}; \hat{\Sigma})L(\mathbf{f})$, the marginal distribution of $\mathbf{f}'$ is also p^*

- 1 Choose ellipse: $\nu \sim \mathcal{N}(\hat{\mathbf{f}}, \hat{\Sigma})$
- 2 Log-likelihood threshold:

$$u \sim \text{Uniform}[0, 1]$$

$$\log y \leftarrow \log L(\mathbf{f}) + \log u$$
- 3 Draw an initial proposal, also defining a bracket:

$$\theta \sim \text{Uniform}[0, 2\pi]$$

$$[\theta_{\min}, \theta_{\max}] \leftarrow [\theta - 2\pi, \theta]$$
- 4 $\mathbf{f}' \leftarrow \mathbf{f} \cos \theta + \nu \sin \theta$
- 5 **if** $\log L(\mathbf{f}') > \log y$ **then**
- 6 Accept: return $\mathbf{f}'$
- 7 **else**
- 8 Shrink the bracket and try a new point:
- 9 **if** $\theta < 0$ **then**
- 10 $\theta_{\min} \leftarrow \theta$
- 11 **else**
- 12 $\theta_{\max} \leftarrow \theta$
- 13 **end**
- 14 $\theta \sim \text{Uniform}[\theta_{\min}, \theta_{\max}]$
- 15 GoTo 4
- 16 **end**

we substitute our unbiased estimator $\tilde{p}(\mathbf{y}|\theta)$ for $p(\mathbf{y}|\theta)$ (as per Lemma 1) into (32). Expanding the summation, we may express this as follows:

$$\begin{aligned}
 \tilde{z} &= \frac{\sum_{j=1}^{N_{\text{imp}}} \left[p(\mathbf{y}|\mathbf{f}'_j) p(\mathbf{f}'_j|\theta') p(\theta') \prod_{i \neq j} q(\mathbf{f}'_i|\mathbf{y}, \theta') \right]}{\sum_{j=1}^{N_{\text{imp}}} \left[p(\mathbf{y}|\mathbf{f}_j) p(\mathbf{f}_j|\theta) p(\theta) \prod_{i \neq j} q(\mathbf{f}_i|\mathbf{y}, \theta) \right]} \times \\
 &\quad \frac{\prod_{i=1}^{N_{\text{imp}}} q(\mathbf{f}_i|\mathbf{y}, \theta) \pi(\theta|\theta')}{\prod_{i=1}^{N_{\text{imp}}} q(\mathbf{f}'_i|\mathbf{y}, \theta') \pi(\theta'|\theta)}.
 \end{aligned} \tag{33}$$

This is a Hastings ratio with the importance sample

$$\pi(\mathbf{f}'_{1:N_{\text{imp}}}, \theta'|\mathbf{f}, \theta) = \prod_{i=1}^{N_{\text{imp}}} q(\mathbf{f}'_i|\mathbf{y}, \theta') \pi(\theta'|\theta) \tag{34}$$

and the target joint distribution

$$p(\mathbf{y}|\mathbf{f}) p(\mathbf{f}|\theta) p(\theta) = p(\mathbf{y}, \mathbf{f}, \theta), \tag{35}$$

which is proportional to $p(\theta|\mathbf{y})$, and we hence have a valid MH algorithm. $\square$

Selecting the number N_{imp} of importance samples to draw is non-trivial. In Figure 2, we demonstrate how

the pseudo-marginal likelihood changes with respect to N_{imp} for the example spam dataset used in Section 3.2, with the computation time at each value of N_{imp} .

It is clear that all values of N_{imp} give a mean likelihood at approximately the same value. The lower values of N_{imp} , within the range of 1 to 100, have large but quickly decreasing variance, while the variance decreases at a far slower rate for values greater than 100. The computation time increases at a constant rate with respect to N_{imp} , other than a jump at $N_{\text{imp}} = 200$.

The calculation of the pseudo-marginal approximation is the dominant cost in every step of the Metropolis-Hastings algorithm, so the choice of N_{imp} is critical when attempting to keep computation time low. For this reason, we opt to use $N_{\text{imp}} = 100$ in further analysis to balance computation time and accuracy.

2.5 Algorithm

Algorithm 3 demonstrates the entire modelling procedure, to aid the reader in both following and implementing our work.

3 Implementation

Having outlined and justified our chosen GPC model mathematically, we now turn to the implementation of our model as software. Given the computationally-intensive nature of GP models, we should give serious consideration to the choices we make in the structure of the code. To evaluate the effectiveness of the model, we compare predictive and computational performance with the simpler logistic regression model on the spam-base dataset.

3.1 Computational Considerations

Examining Algorithm 3, we can see a loop for sampling from $q(\mathbf{f}|\mathbf{y}, \theta')$ embedded in a larger loop to sample θ . We should hence choose a compiled language over an interpreted language (in our case, Rcpp over R) when implementing to see serious performance gains. Indeed, the great number of linear algebraic operations means that using a package designed to access linear algebra subroutines would be preferable. We have opted to use RcppArmadillo (Eddelbuettel and Sanderson, 2014).

The algorithm can be run over multiple chains, each initiated with different initial samples but the same priors, and analysis of different chains allows us to check for good mixing and metastability.

Within each chain we parallelise the construction of the covariance matrix $\mathbf{K}$, sending each kernel evaluation (4) to a different worker, as well as the computation of the pseudo-marginal approximation, since neither of these operations are dependent on previous

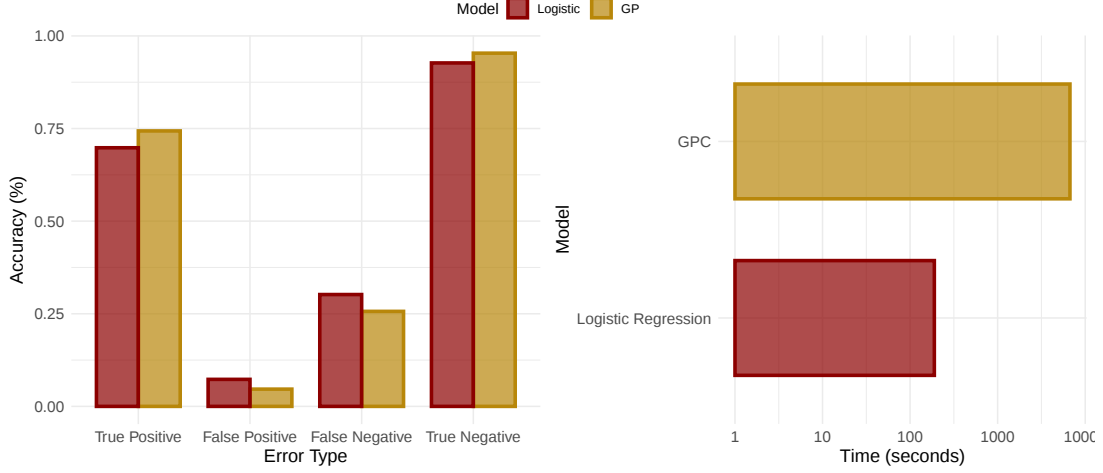


Figure 3: Comparative confusion plots of predictive accuracy for GPC and logistic regression (left) and time plots in seconds for GPC and logistic regression, on log scale (right)

Algorithm 3: Gaussian Process Classification via Pseudo-Marginal Approach

input : A prior for θ and $\mathbf{f}$, number of iterations N , data $(X, \mathbf{y})$, subset size $d < n$

```

1 Sample  $\theta^{(1)}$  from prior
2 Select a subset of size  $d$  from the data  $(X, \mathbf{y})$ 
  using IVM (Algorithm 1)
3 for  $i = 2 : N$  do
4   case Sampling  $\theta$  do
5     Draw  $\theta'$  from proposal  $\pi(\theta'|\theta^{(i-1)})$ 
6     Use Laplace approximation to
      approximate  $p(\mathbf{f}|\mathbf{y}, \theta')$  with  $q(\mathbf{f}|\mathbf{y}, \theta')$ 
      from (14)
7     Obtain  $N_{imp}$  samples from  $q(\mathbf{f}|\mathbf{y}, \theta')$  to
      get  $\tilde{p}(\mathbf{y}|\theta')$  from (30)
8     Compute acceptance probability
       $A = \min\{1, \tilde{z}\}$  with  $\tilde{z}$  from (32)
9     Draw  $u \sim \text{Unif}(0, 1)$ 
10    if  $A > u$  then
11      Next pair is  $(\theta', \tilde{p}(\mathbf{y}|\theta'))$ 
12       $\theta^{(i)} = \theta'$ 
13    else
14      Next pair is  $(\theta^{(i-1)}, \tilde{p}(\mathbf{y}|\theta^{(i-1)}))$ 
15       $\theta^{(i)} = \theta^{(i-1)}$ 
16    end
17  end
18  case Sampling  $\mathbf{f}|\theta$  do
19    Plug  $\theta^{(i)}$  into the prior for  $\mathbf{f}$ 
20    Sample  $\mathbf{f}$  from this prior using ELL-SS
      (Algorithm 2)
21  end
22 end
23 Use posteriors of  $\theta$  and  $\mathbf{f}|\theta$  to give predictions
  with (12)
```

iterations, and parallelisation is straight forward. Implementing this parallel approach sped up the model run time, making it approximately twice as fast.

In evaluating the Laplace approximation (Section 2.2), when estimating $\hat{\Sigma}^{-1}$ we are in effect adding low-rank updates to another inverted matrix $\mathbf{K}^{-1}$ (and similarly in evaluating $\hat{\mathbf{f}}$), meaning we can find computational savings with the Woodbury identity. We can also use this identity to avoid directly inverting the Gram matrix $\mathbf{K}$, an operation which is highly susceptible to numerical error in practice. Details on this are available in Williams and Rasmussen, 2006, Algorithm 3.1.

Then, in evaluating the approximate marginal distribution $\tilde{p}(\mathbf{y}|\theta)$ for a sample, we need to multiply many small probabilities. Given that these are so small, finding the product is likely to cause numerical underflow on a machine. We hence use the *LogSumExp* trick (*LogSumExp Trick*), instead summing the logarithms of the probabilities then taking the exponent. We also work with log probabilities in general to avoid numerical underflow and overflow. This includes calculating the log acceptance probability in the Metropolis-Hastings algorithm.

3.2 Model Comparison

To understand how well the GPC model performs relative to the simpler Bayesian logistic regression model, we compare the models by the loss that they induce (we use a $a_- - a_+$ loss with $a_- = 1, a_+ = 1$), and the time taken to produce a fit.

For a detailed explanation of the links between the two models, confer Williams and Rasmussen, 2006, Section 6.5.

We use the e-mail spam dataset provided by Dua and Graff, 2017 for comparison between the logistic regression and our GPC approach, and to test performance of our model. Briefly, this dataset contains $n = 4601$

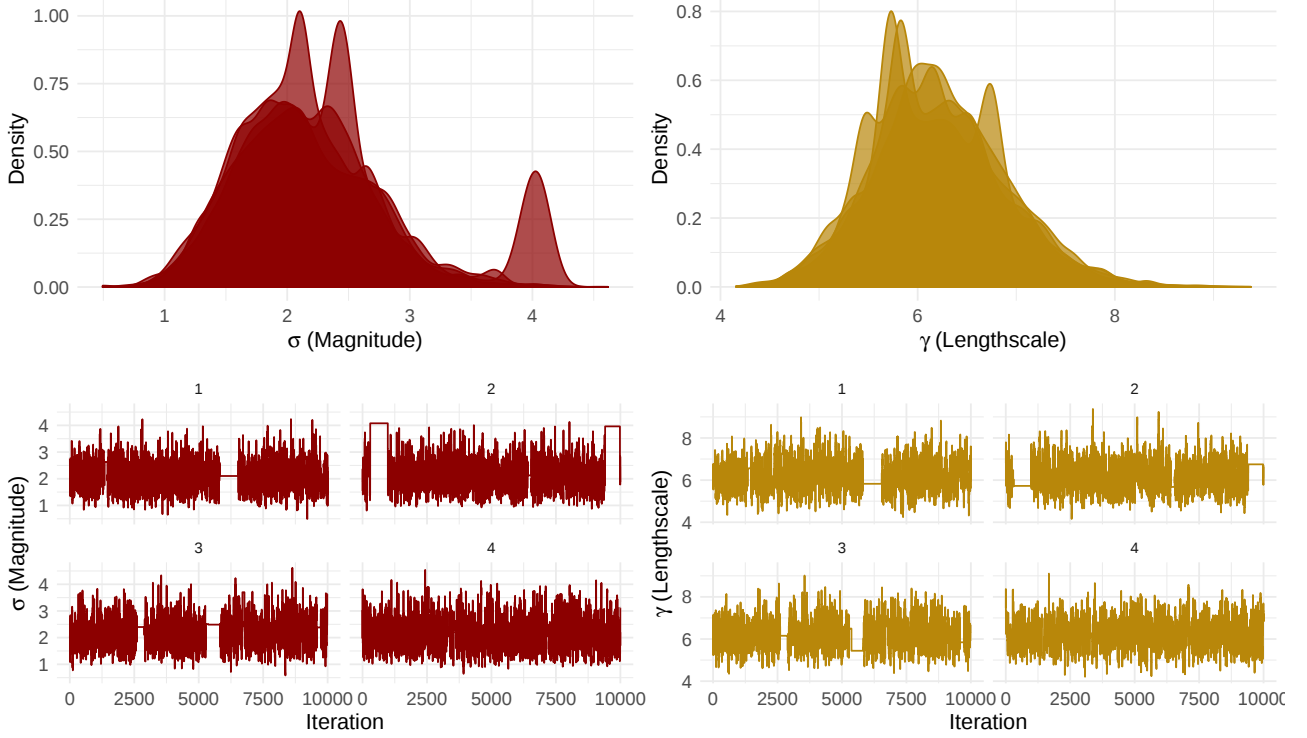


Figure 4: Posterior distributions and trace plots across four chains for the two kernel hyperparameters. Model run with 20,000 iterations, of which the first 10,000 iterations were discarded as ‘burn in’.

observations (emails) on $p = 57$ predictors (attributes of each email), each classified as spam (+) or not spam (−).

Figure 2 suggests that we use a value of 100 for N_{imp} , as we obtain reliable estimates of the marginal likelihood, for a manageable run time. We must also decide the value d corresponding to the size of the subset selected by IVM. In doing this, we are balancing predictive accuracy and run time.

In Figure 1 we can see that runtime increases significantly with d , but predictive accuracy seems to level out beyond $d = 150$, so we select this value. We fit a model with 10000 steps, and discard the first 5000 of these as burn-in. Hyperparameters are initialised with $\theta = (1, 1)$, and the variance of our random-walk proposal mechanism for the kernel parameters was tuned such that the algorithm’s acceptance ratio was consistently between 0.2 and 0.3; doing this ensures good mixing.

The fully-Bayesian logistic regression model was implemented in `rjags` (Plummer, 2019), using a subset of 13 predictors. These predictors were hand selected according to how informative we perceived they would be, and the subset was used as it would be computationally expensive to optimise the model over a 57-dimensional space, and it may not be possible given a finite set of data. The `rjags` model was used to produce 20,000 parameter samples, of which the initial 10,000 samples were discarded as burn-in. To compare the two models, a training set containing 70% of all observations was used to train the models, and their

performance was evaluated using a test set containing the remaining 30% of the observations.

In Figure 3, we can see that the GPC model consistently outperforms the logistic regression model on this data set, giving marginal but significant improvements for each error type.

Indeed, for datasets with a greater number of predictors p , our method should outperform logistic regression on both terms, since ours does not scale in complexity with p and has a more explicit control of regularity in the kernel space.

Finally, we can also see in Figure 3 that the fitting the GP model with $d = 150$ takes approximately 100 times as long as fitting logistic regression. This should be taken into account when deciding which model should be used in practice - if time is an important factor, a smaller value of d should be considered.

3.3 Fitting for a Large Number of Iterations

We now present the Gaussian process classification model, fit with a large number of iterations (20,000) and multiple chains (4) in the Metropolis-Hasting algorithm. This result is presented so that the reader may gain confidence in the convergence of the model.

Using the values of $N_{imp} = 100$ and $d = 150$ as before, Figure 4 shows the posterior distributions of the two hyperparameters used to define the Gaussian kernel (4), σ and γ : the magnitude and the lengthscale

respectively.

The posteriors produced by each chain are very similar, indicating that all four chains converged to approximately the same value for both σ and γ . We see that the posterior means have converged far from the initial values $\theta = (1, 1)$ proposed, suggesting that the chain has efficiently explored the state space. Given the successful convergence of all chains given rather poor initial conditions, the method proposed seems to be reliable.

4 Conclusion

In this report, we have explored the pseudo-marginal approach to a fully-Bayesian Gaussian process classification model described by Filippone and Girolami, 2014, and extended this to large datasets. Motivated by the desire to extend the logistic regression classifier to capture more complicated, non-linear relationships between observed variables, we outlined the flexible specification of the Gaussian process model.

Using the Laplace approximation of the posterior of the latent variables $\mathbf{f}$, and combining elliptical slice sampling and a pseudo-marginal approach to sample from the joint posterior of $\mathbf{f}, \theta$, we have shown a way to overcome the difficulties of classifying with a GP model, and to express our uncertainty in predictions. Using the sparse *Informative Vector Machine* representation of Herbrich, Lawrence, and Seeger, 2003, we have proposed a possible method for extending the model for large datasets.

By comparing results against the logistic regression classifier, we have found that our model is capable of making superior predictions while trading off speed in fitting. We have given details on how we have implemented the procedure efficiently in software, and have made our own implementation available as an R package at github.com/dannyjameswilliams/gpc.

Bibliography

- Beaumont, Mark A (2003). “Estimation of population growth or decline in genetically monitored populations”. In: *Genetics* 164.3, pp. 1139–1160.
- Dua, Dheeru and Casey Graff (2017). *UCI Machine Learning Repository*. URL: <http://archive.ics.uci.edu/ml>.
- Eddelbuettel, Dirk and Conrad Sanderson (2014). “RcppArmadillo: Accelerating R with high-performance C++ linear algebra”. In: *Computational Statistics & Data Analysis* 71, pp. 1054–1063.
- Filippone, Maurizio and Mark Girolami (2014). “Pseudo-marginal Bayesian inference for Gaussian processes”. In: *IEEE Transactions on Pattern Analysis and Machine Intelligence* 36.11, pp. 2214–2226.

- Hensman, James, Alexander Matthews, and Zoubin Ghahramani (2015). “Scalable variational Gaussian process classification”. In:
- Herbrich, Ralf, Neil D Lawrence, and Matthias Seeger (2003). “Fast sparse Gaussian process methods: The informative vector machine”. In: *Advances in neural information processing systems*, pp. 625–632.
- Hernández-Lobato, Daniel and José Miguel Hernández-Lobato (2016). “Scalable gaussian process classification via expectation propagation”. In: *Artificial Intelligence and Statistics*, pp. 168–176.
- LogSumExp Trick. URL: <https://en.wikipedia.org/wiki/LogSumExp>.
- Murray, Iain, Ryan Prescott Adams, and David JC MacKay (2010). “Elliptical slice sampling”. In: *JMLR: WI&CP*.
- Plummer, Martyn (2019). *rjags: Bayesian Graphical Models using MCMC*. R package version 4-10. URL: <https://CRAN.R-project.org/package=rjags>.
- Williams, Christopher KI and Carl Edward Rasmussen (2006). *Gaussian processes for machine learning*. Vol. 2. 3. MIT press Cambridge, MA.