

Dynamics with Gaussian Processes

Neil D. Lawrence

work with Magnus Rattray, Mauricio Alvarez, Pei Gao, Antti Honkela,
David Luengo, Guido Sanguinetti, Michalis Titsias, Jennifer Withers

Natural Computing Applications Forum, Sheffield

6th September 2008

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Dimensionality Reduction I

- Linear relationship between the data, $\mathbf{X} \in \Re^{N \times d}$, and a reduced dimensional representation, $\mathbf{F} \in \Re^{N \times q}$, where $q \ll d$.

$$\mathbf{X} = \mathbf{F}\mathbf{W} + \epsilon,$$

$$\epsilon \sim \mathcal{N}(\mathbf{0}, \Sigma)$$

- Integrate out $\mathbf{F}$, optimize with respect to $\mathbf{W}$.
- For temporal data and a particular Gaussian prior in the latent space: Kalman filter/smoothing.
- More generally consider a Gaussian process (GP) prior,

$$p(\mathbf{F}|\mathbf{t}) = \prod_{i=1}^q \mathcal{N}(\mathbf{f}_{:,i} | \mathbf{0}, \mathbf{K}_{\mathbf{f}_{:,i}, \mathbf{f}_{:,i}}).$$

- Given the covariance functions for $\{f_i(t)\}$ the implied covariance functions for $\{x_i(t)\}$ — semi-parametric latent factor model (Teh et al., 2005).
- Kalman filter/smoothing approach has been preferred
 - ▶ linear computational complexity in N .
 - ▶ Advances in sparse approximations have made the general GP framework practical. (Snelson and Ghahramani, 2006; Quiñero Candela and Rasmussen, 2005, also imminent work by Titsias).

Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Zero mean Gaussian distribution

- A multi-variate Gaussian distribution is defined by a mean and a covariance matrix.

$$N(\mathbf{f}|\mu, \mathbf{K}) = \frac{1}{(2\pi)^{\frac{N}{2}} |\mathbf{K}|^{\frac{1}{2}}} \exp\left(-\frac{(\mathbf{f} - \mu)^T \mathbf{K}^{-1} (\mathbf{f} - \mu)}{2}\right).$$

- We will consider the special case where the mean is zero,

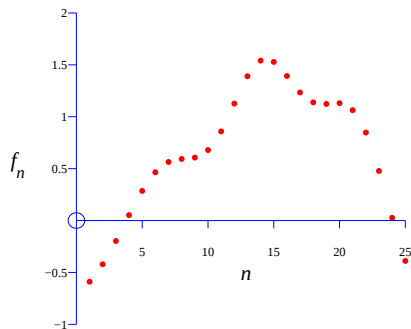
$$N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = \frac{1}{(2\pi)^{\frac{N}{2}} |\mathbf{K}|^{\frac{1}{2}}} \exp\left(-\frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}\right).$$

Multi-variate Gaussians

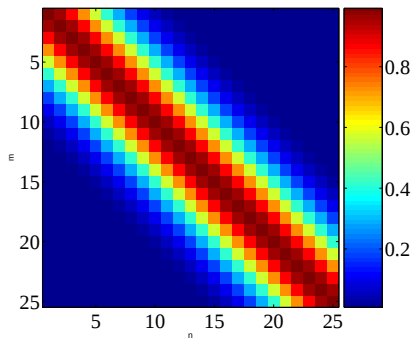
- We will consider a Gaussian with a particular structure of covariance matrix.
- Generate a single sample from this 25 dimensional Gaussian distribution, $\mathbf{f} = [f_1, f_2 \dots f_{25}]$.
- We will plot these points against their index.

Gaussian Distribution Sample

demGPSample



(a)



(b)

Figure: (a) 25 instantiations of a function, f_n , (b) colormap of covariance matrix.

The covariance matrix

- Covariance matrix shows correlation between points f_m and f_n if n is near to m .
- Less correlation if n is distant from m .
- Our ordering of points means that the *function appears smooth*.
- Let's focus on the joint distribution of two points from the 25.

The covariance matrix

- Covariance matrix shows correlation between points f_m and f_n if n is near to m .
- Less correlation if n is distant from m .
- Our ordering of points means that the *function appears smooth*.
- Let's focus on the joint distribution of two points from the 25.

The covariance matrix

- Covariance matrix shows correlation between points f_m and f_n if n is near to m .
- Less correlation if n is distant from m .
- Our ordering of points means that the *function appears smooth*.
- Let's focus on the joint distribution of two points from the 25.

The covariance matrix

- Covariance matrix shows correlation between points f_m and f_n if n is near to m .
- Less correlation if n is distant from m .
- Our ordering of points means that the *function appears smooth*.
- Let's focus on the joint distribution of two points from the 25.

Prediction of f_2 from f_1

demGPCov2D([1 2])

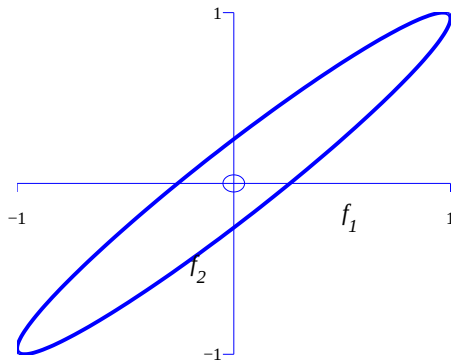


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_2 \end{bmatrix}$ is $\mathbf{K}_{12} = \begin{bmatrix} 1 & 0.966 \\ 0.966 & 1 \end{bmatrix}$.

Prediction of f_2 from f_1

demGPCov2D([1 2])

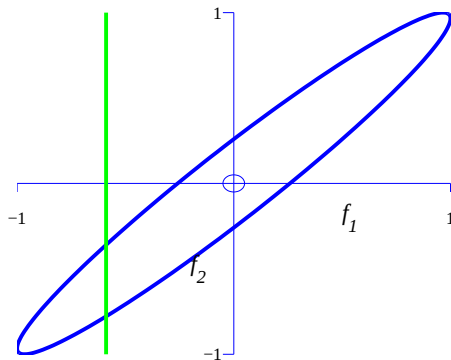


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_2 \end{bmatrix}$ is $\mathbf{K}_{12} = \begin{bmatrix} 1 & 0.966 \\ 0.966 & 1 \end{bmatrix}$.

Prediction of f_2 from f_1

demGPCov2D([1 2])

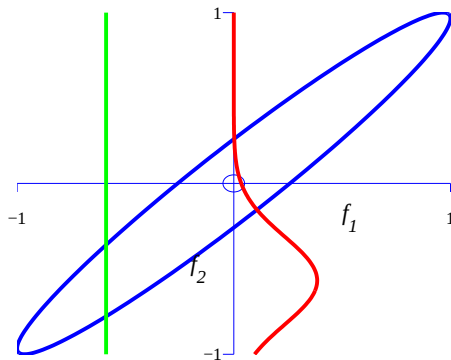


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_2 \end{bmatrix}$ is $\mathbf{K}_{12} = \begin{bmatrix} 1 & 0.966 \\ 0.966 & 1 \end{bmatrix}$.

Prediction of f_5 from f_1

demGPCov2D([1 5])

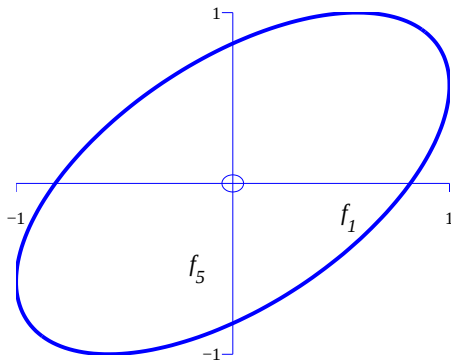


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_5 \end{bmatrix}$ is $\mathbf{K}_{15} = \begin{bmatrix} 1 & 0.574 \\ 0.574 & 1 \end{bmatrix}$.

Prediction of f_5 from f_1

demGPCov2D([1 5])

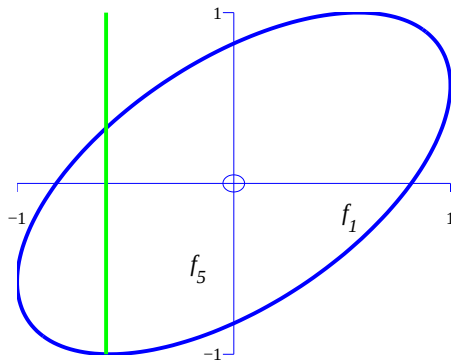


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_5 \end{bmatrix}$ is $\mathbf{K}_{15} = \begin{bmatrix} 1 & 0.574 \\ 0.574 & 1 \end{bmatrix}$.

Prediction of f_5 from f_1

demGPCov2D([1 5])

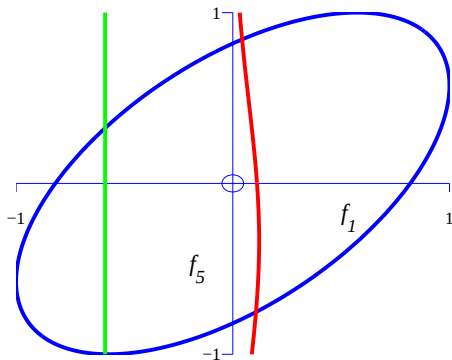


Figure: Covariance for $\begin{bmatrix} f_1 \\ f_5 \end{bmatrix}$ is $\mathbf{K}_{15} = \begin{bmatrix} 1 & 0.574 \\ 0.574 & 1 \end{bmatrix}$.

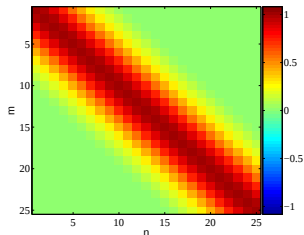
Covariance Functions

Where did this covariance matrix come from?

RBF Kernel Function

$$k(t, t') = \alpha \exp \left(-\frac{\|t - t'\|^2}{2l^2} \right)$$

- Covariance matrix is built using the *inputs* to the function t .
- For the example above it was based on Euclidean distance.
- The covariance function is also known as a kernel.



demCovFuncSample

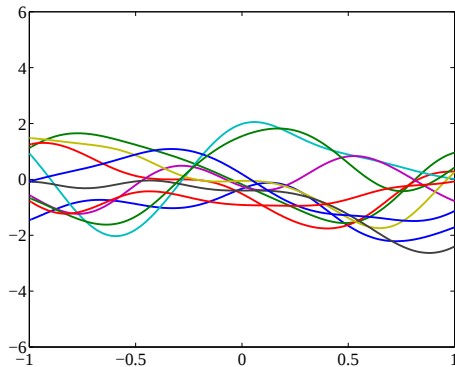


Figure: RBF kernel with $l = 10^{-\frac{1}{2}}$, $\alpha = 1$

demCovFuncSample

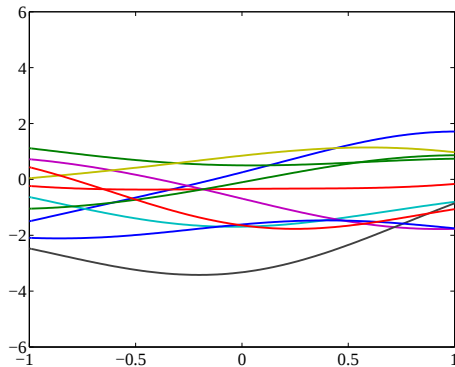


Figure: RBF kernel with $l = 1$, $\alpha = 1$

demCovFuncSample

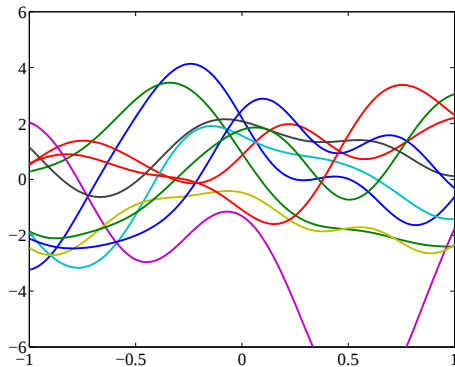


Figure: RBF kernel with $l = 0.3$, $\alpha = 4$

Gaussian Process Regression

demRegression

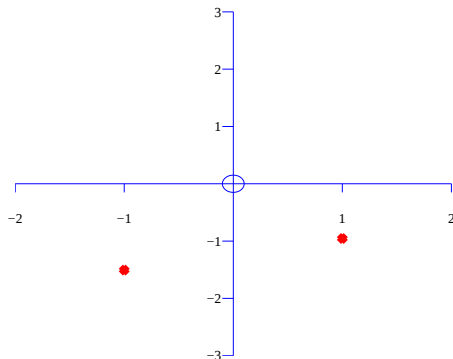


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

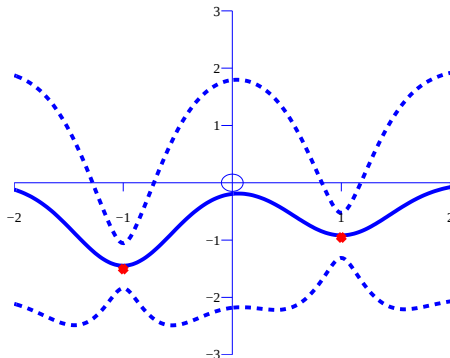


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

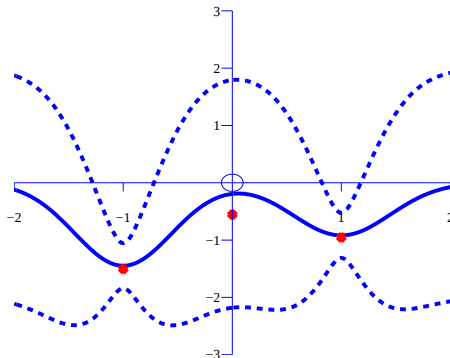


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

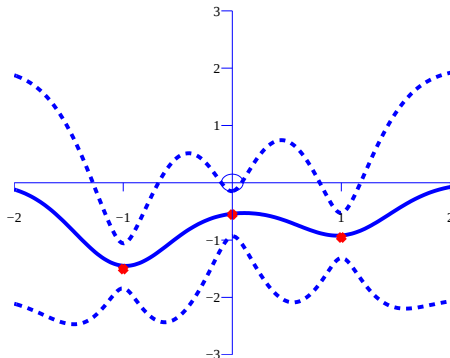


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

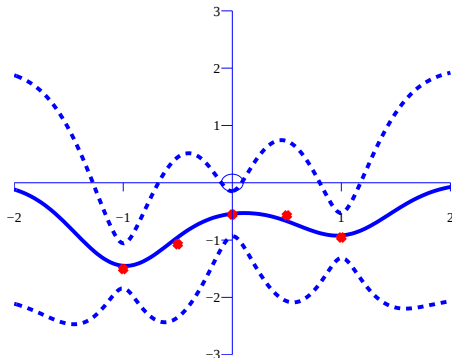


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

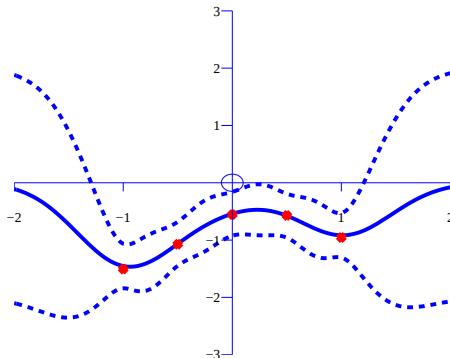


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

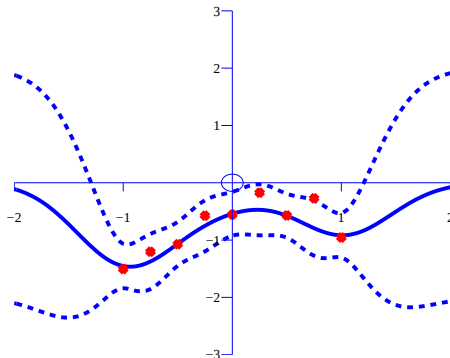


Figure: Examples include WiFi localization, C14 calibration curve.

Gaussian Process Regression

demRegression

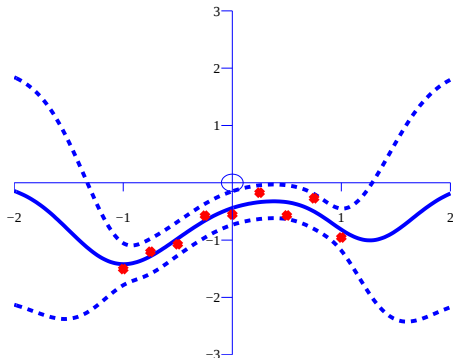
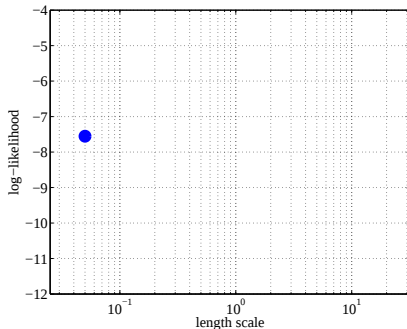
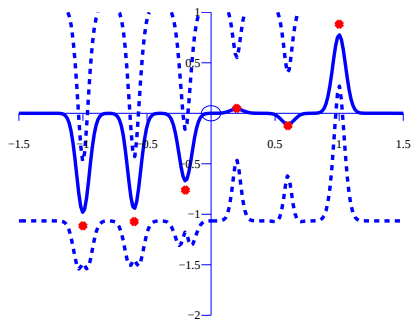


Figure: Examples include WiFi localization, C14 calibration curve.

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

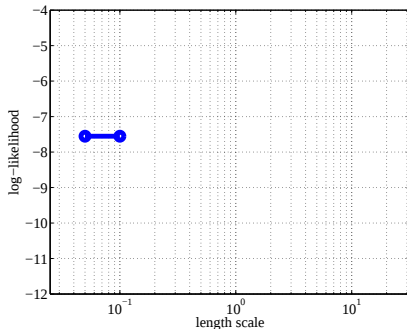
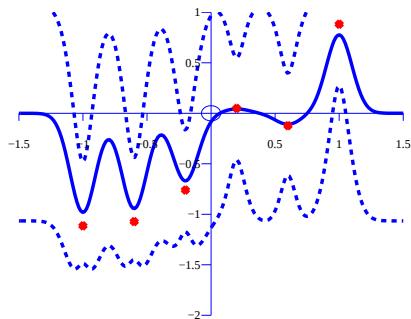


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

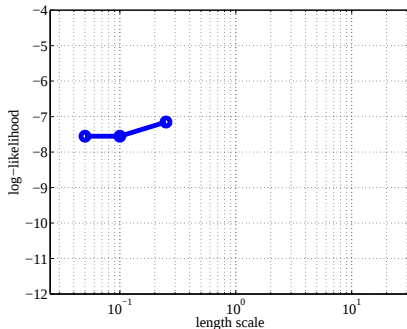
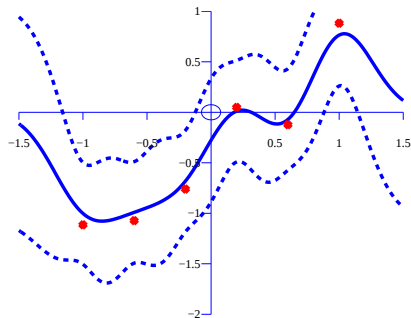


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

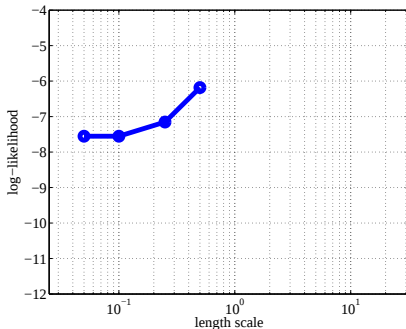
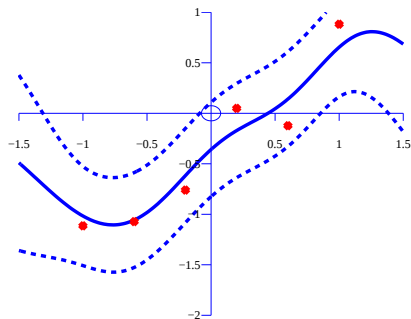


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

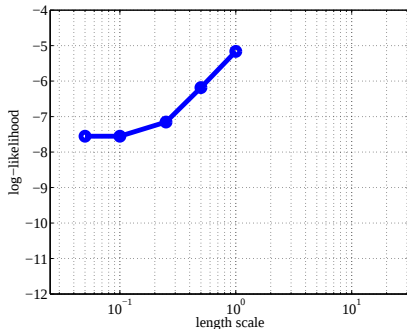
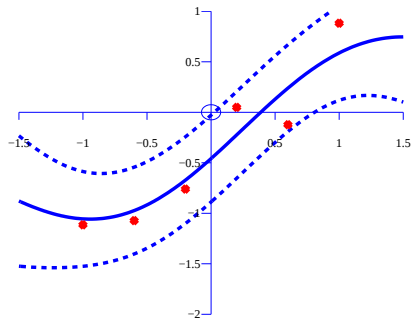


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

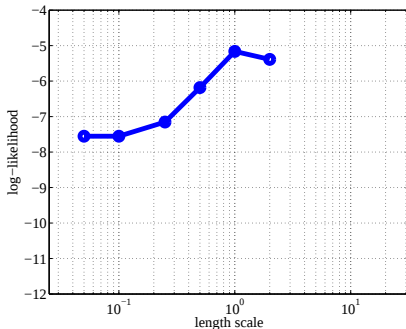
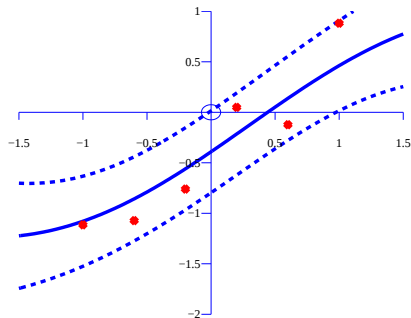


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

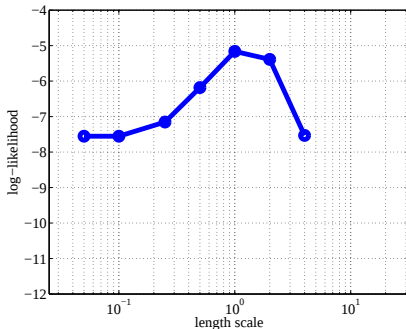
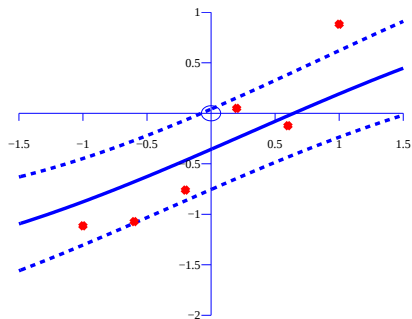


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

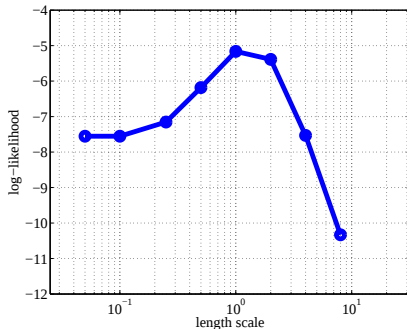
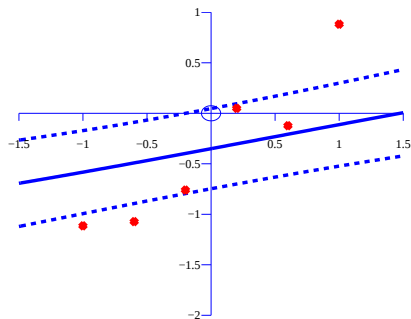


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern

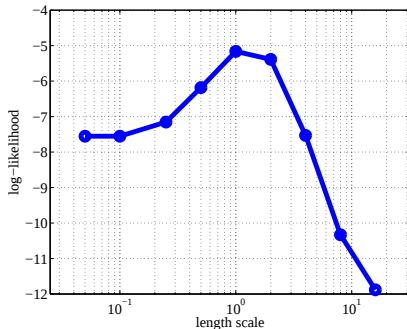
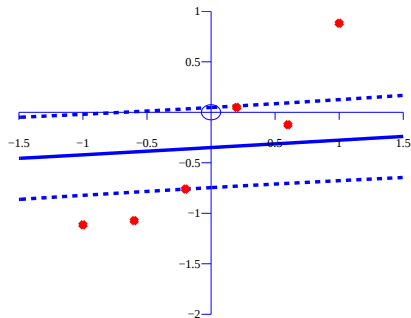


$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

Learning Kernel Parameters

Can we determine length scales and noise levels from the data?

demOptimiseKern



$$\log N(\mathbf{f}|\mathbf{0}, \mathbf{K}) = -\frac{N}{2} \log 2\pi - \frac{1}{2} \log |\mathbf{K}| - \frac{\mathbf{f}^T \mathbf{K}^{-1} \mathbf{f}}{2}$$

- These models rely on the latent variables to provide the dynamic information.
- We now introduce a further dynamical system with a *mechanistic* inspiration.
- Physical Interpretation:
 - ▶ the latent functions, $f_i(t)$ are q forces.
 - ▶ We observe the displacement of d springs to the forces.,
 - ▶ Interpret system as the force balance equation, $\mathbf{X}\mathbf{D} = \mathbf{F}\mathbf{S} + \epsilon$.
 - ▶ Forces act, e.g. through levers — a matrix of sensitivities, $\mathbf{S} \in \mathbb{R}^{q \times d}$.
 - ▶ Diagonal matrix of spring constants, $\mathbf{D} \in \mathbb{R}^{d \times d}$.
 - ▶ Original System: $\mathbf{W} = \mathbf{S}\mathbf{D}^{-1}$.

- Add a damper and give the system mass.

$$\mathbf{F}\mathbf{S} = \ddot{\mathbf{X}}\mathbf{M} + \dot{\mathbf{X}}\mathbf{C} + \mathbf{X}\mathbf{D} + \epsilon.$$

- Now have a second order mechanical system.
- It will exhibit inertia and resonance.
- There are many systems that can also be represented by differential equations.
 - ▶ When being forced by latent function(s), $\{f_i(t)\}_{i=1}^q$, we call this a *latent force model*.

Gaussian Process priors and Latent Force Models

- For Gaussian process we can compute the covariance matrices for the output displacements.
- For one displace the model is

$$m_k \ddot{x}_k(t) + c_k \dot{x}_k(t) + d_k x_k(t) = b_k + \sum_{i=0}^M s_{ik} f_i(t), \quad (1)$$

where, m_k is the k th diagonal element from $\mathbf{M}$ and similarly for c_k and d_k . s_{ik} is the i , k th element of $\mathbf{S}$.

- Model the latent forces as q independent, GPs with RBF covariances

$$k_{f_i f_l}(t, t') = \exp \left(-\frac{(t - t')^2}{\sigma_i^2} \right) \delta_{il}.$$

Covariance for ODE Model

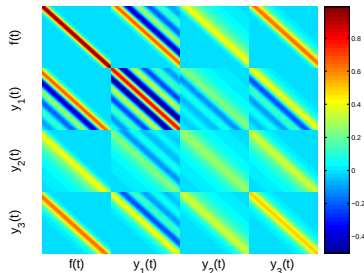
- RBF Kernel function for $f(t)$

$$x_j(t) = \frac{1}{m_j \omega_j} \sum_{i=1}^q S_{ji} \exp(-\alpha_j t) \int_0^t f_i(u) \exp(\alpha_j u) \sin(\omega_j(t-u)) du$$

- Joint distribution for $x_1(t)$, $x_2(t)$, $x_3(t)$ and $f(t)$.

Damping ratios:

ζ_1	ζ_2	ζ_3
0.125	2	1



Joint Sampling of $x(t)$ and $f(t)$

- demLfmSample

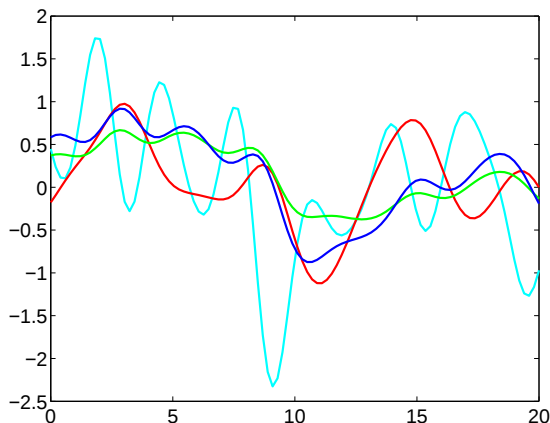


Figure: Joint samples from the ODE covariance, *cyan*: $f(t)$, *red*: $x_1(t)$ (underdamped) and *green*: $x_2(t)$ (overdamped) and *blue*: $x_3(t)$ (critically damped).

Joint Sampling of $x(t)$ and $f(t)$

- demLfmSample

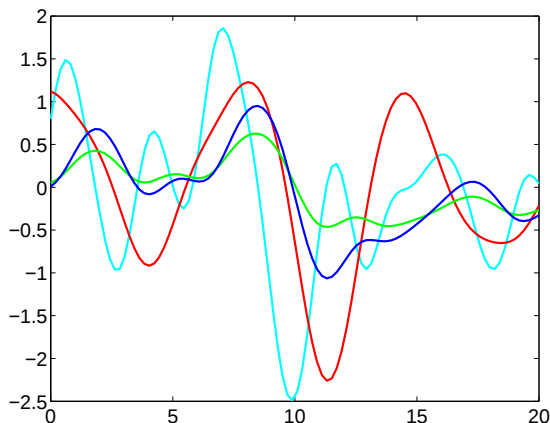


Figure: Joint samples from the ODE covariance, *cyan*: $f(t)$, *red*: $x_1(t)$ (underdamped) and *green*: $x_2(t)$ (overdamped) and *blue*: $x_3(t)$ (critically damped).

Joint Sampling of $x(t)$ and $f(t)$

- demLfmSample

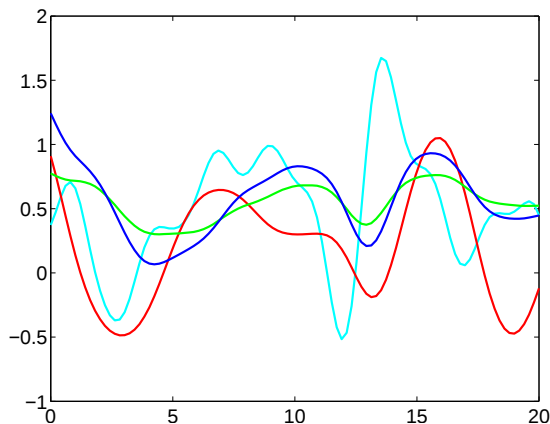


Figure: Joint samples from the ODE covariance, *cyan*: $f(t)$, *red*: $x_1(t)$ (underdamped) and *green*: $x_2(t)$ (overdamped) and *blue*: $x_3(t)$ (critically damped).

Joint Sampling of $x(t)$ and $f(t)$

- demLfmSample

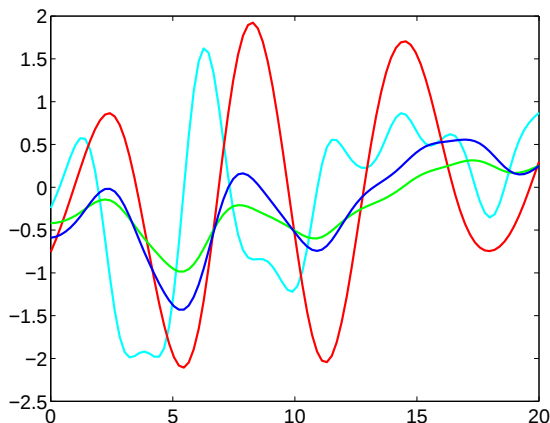


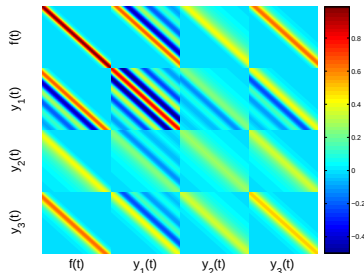
Figure: Joint samples from the ODE covariance, *cyan*: $f(t)$, *red*: $x_1(t)$ (underdamped) and *green*: $x_2(t)$ (overdamped) and *blue*: $x_3(t)$ (critically damped).

- RBF Kernel function for $f(t)$

$$x_j(t) = \frac{1}{m_j \omega_j} \sum_{i=1}^q S_{ji} \exp(-\alpha_j t) \int_0^t f_i(u) \exp(\alpha_j u) \sin(\omega_j(t-u)) du$$

- Joint distribution for $x_1(t)$, $x_2(t)$, $x_3(t)$ and $f(t)$.
- Damping ratios:

ζ_1	ζ_2	ζ_3
0.125	2	1



Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

Example: Transcriptional Regulation

- First Order Differential Equation

$$\frac{dx_j(t)}{dt} = B_j + S_j f(t) - D_j x_j(t)$$

- Can be used as a model of gene transcription: Barenco et al., 2006.
- $x_j(t)$ – concentration of gene j 's mRNA
- $f(t)$ – concentration of active transcription factor
- Model parameters: baseline B_j , sensitivity S_j and decay D_j
- Application: identifying co-regulated genes (targets)
- Problem: how do we fit the model when $f(t)$ is not observed?

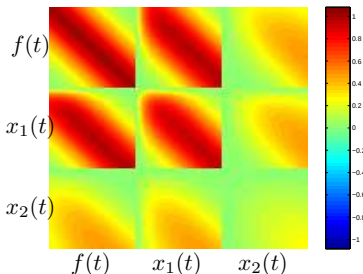
RBF covariance function for $f(t)$

$$x_i(t) = \frac{B_i}{D_i} + S_i \exp(-D_i t) \int_0^t f(u) \exp(D_i u) du.$$

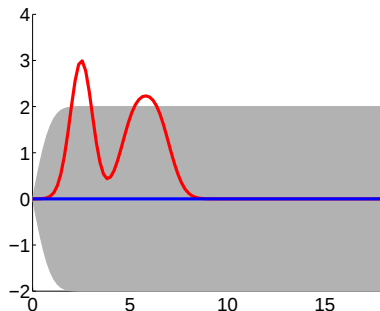
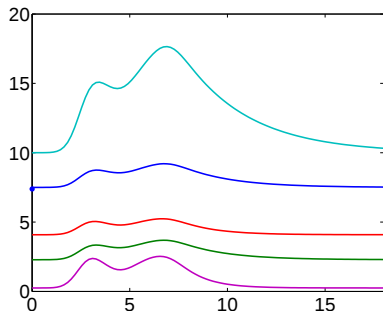
- Joint distribution for $x_1(t)$, $x_2(t)$ and $f(t)$.

► Here:

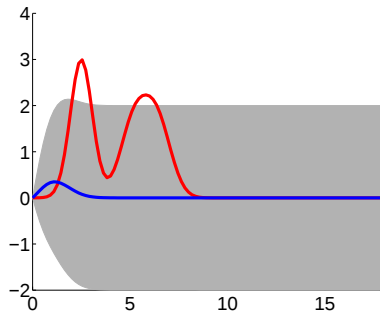
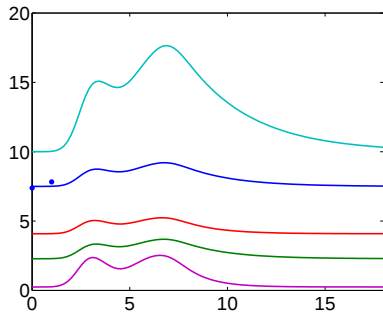
D_1	S_1	D_2	S_2
5	5	0.5	0.5



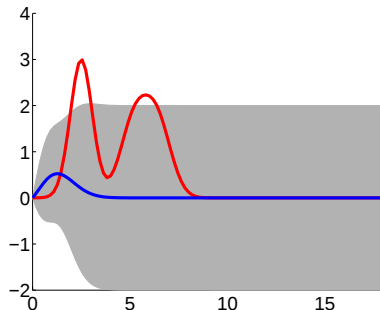
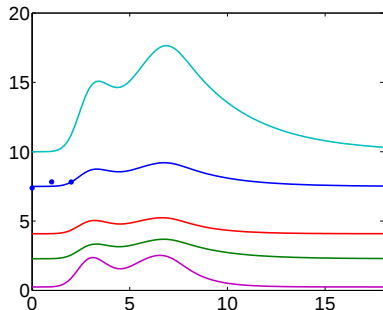
Artificial Example: Inferring $f(t)$



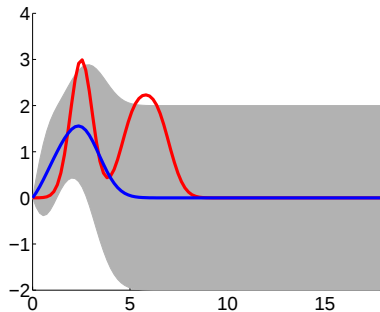
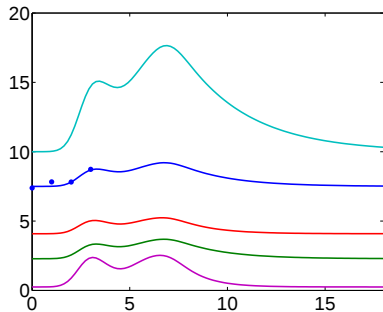
Artificial Example: Inferring $f(t)$



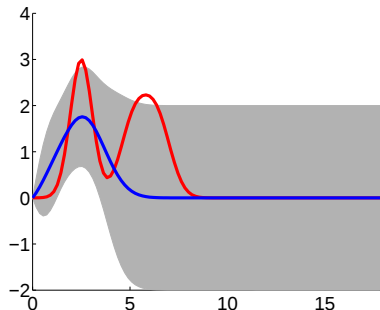
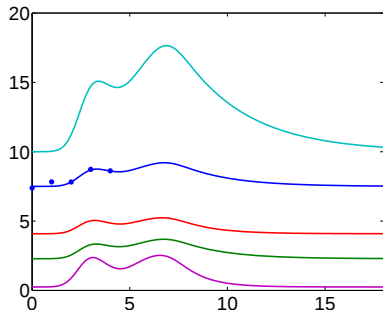
Artificial Example: Inferring $f(t)$



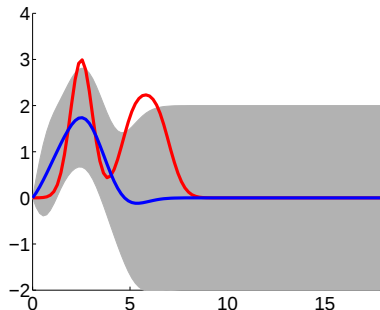
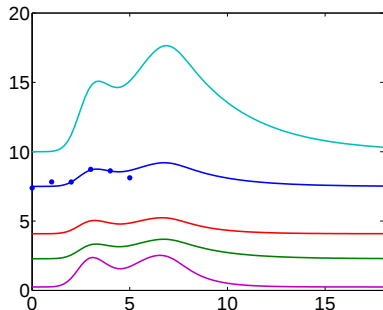
Artificial Example: Inferring $f(t)$



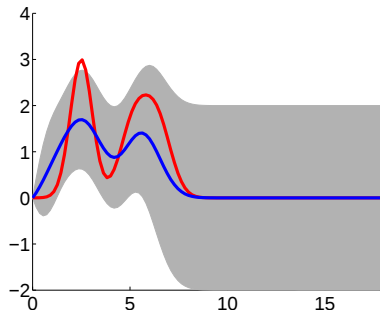
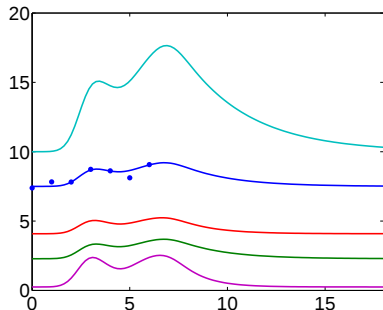
Artificial Example: Inferring $f(t)$



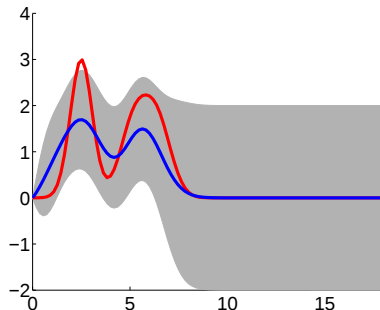
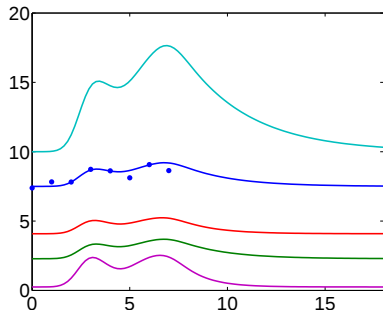
Artificial Example: Inferring $f(t)$



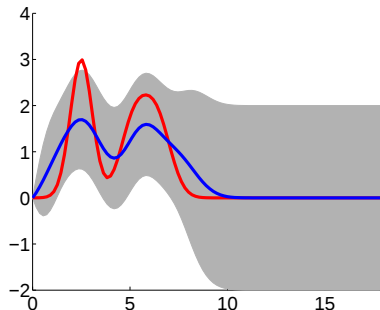
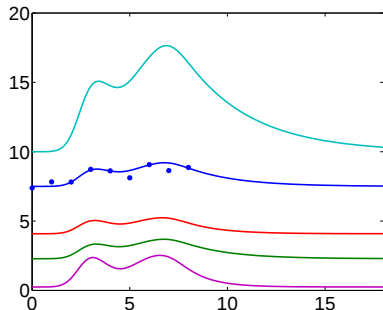
Artificial Example: Inferring $f(t)$



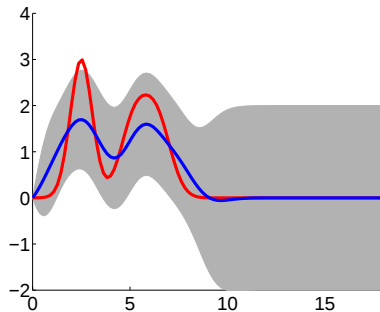
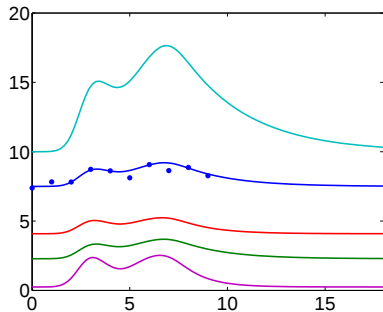
Artificial Example: Inferring $f(t)$



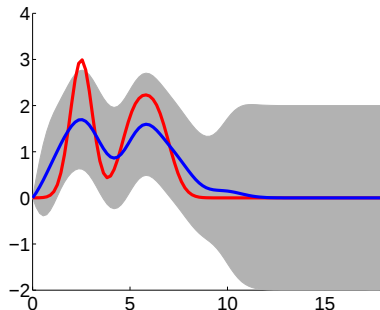
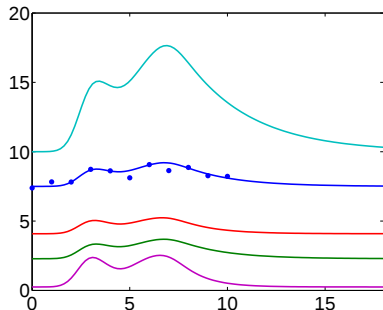
Artificial Example: Inferring $f(t)$



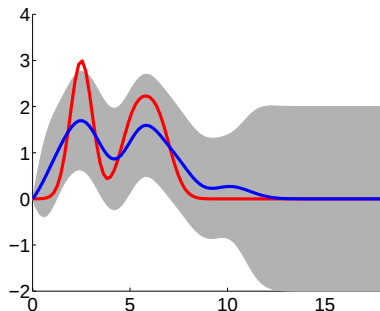
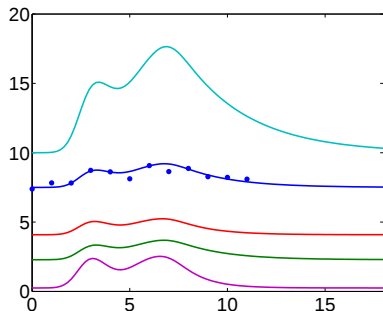
Artificial Example: Inferring $f(t)$



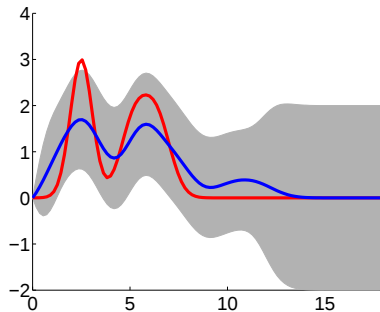
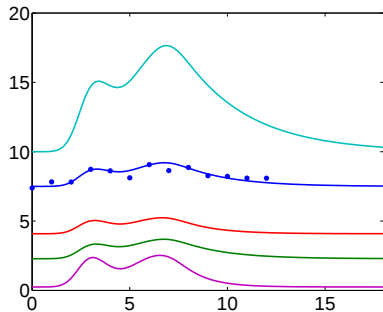
Artificial Example: Inferring $f(t)$



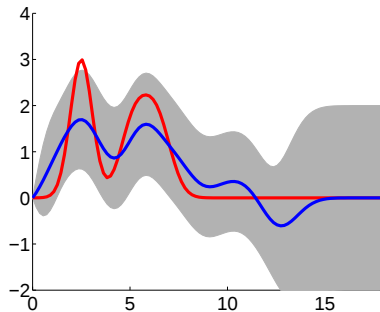
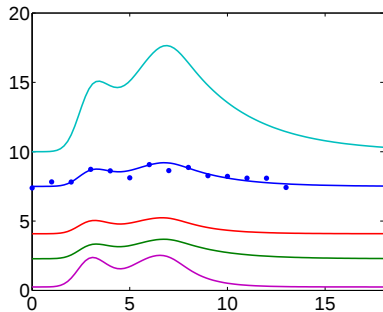
Artificial Example: Inferring $f(t)$



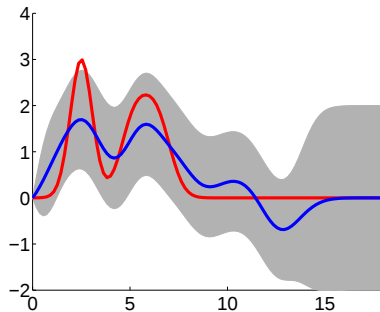
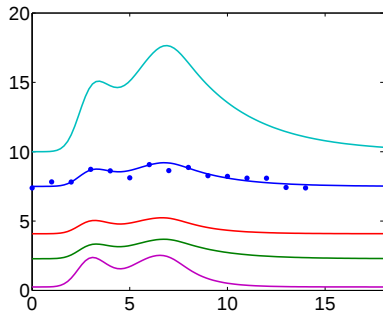
Artificial Example: Inferring $f(t)$



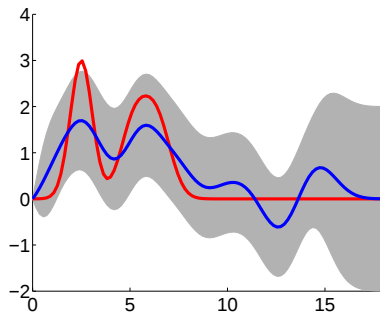
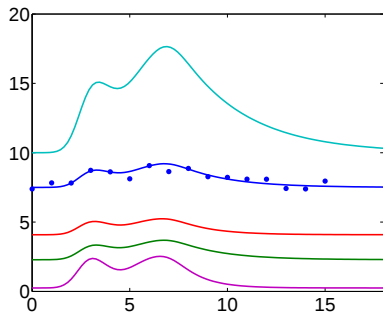
Artificial Example: Inferring $f(t)$



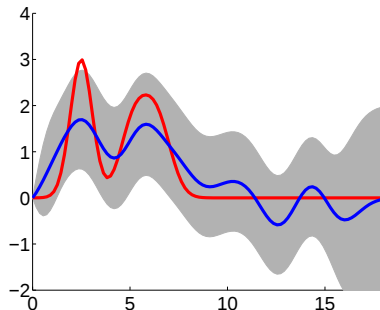
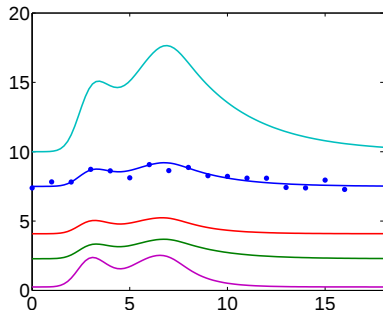
Artificial Example: Inferring $f(t)$



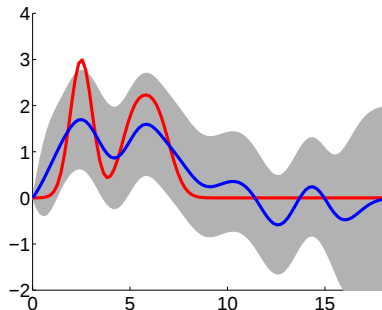
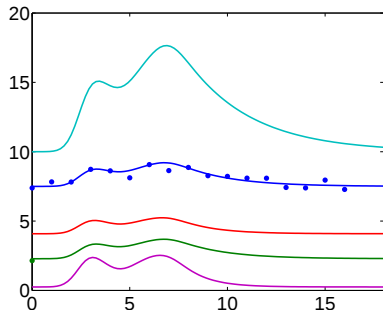
Artificial Example: Inferring $f(t)$



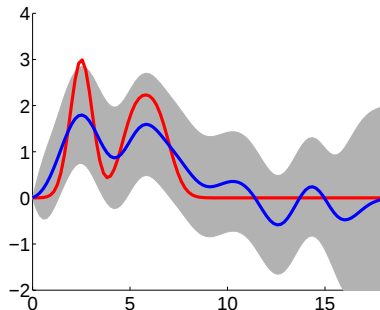
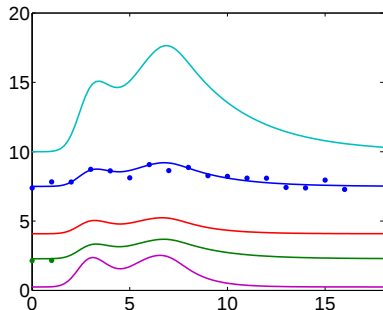
Artificial Example: Inferring $f(t)$



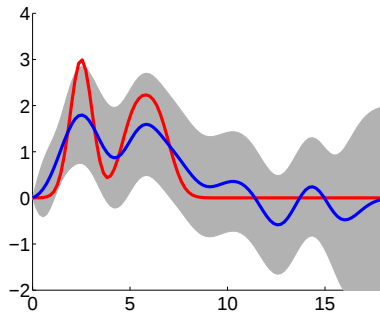
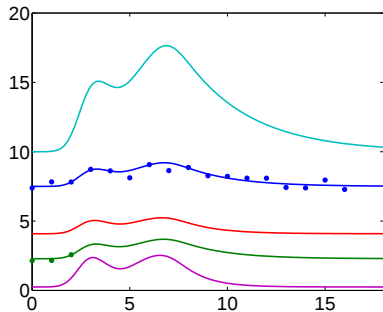
Artificial Example: Inferring $f(t)$



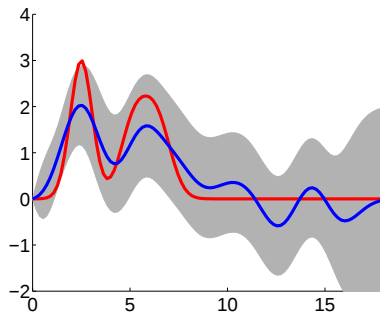
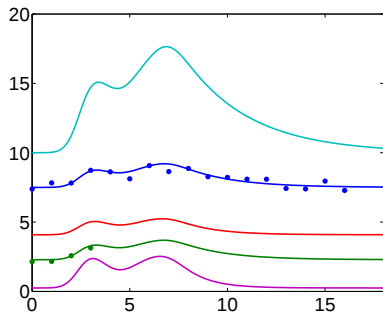
Artificial Example: Inferring $f(t)$



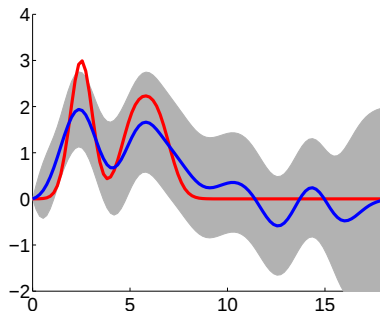
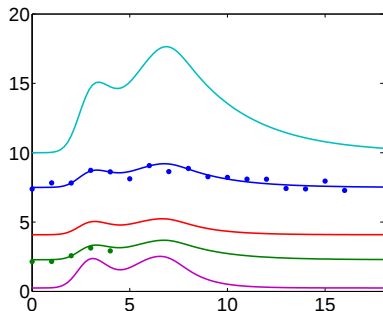
Artificial Example: Inferring $f(t)$



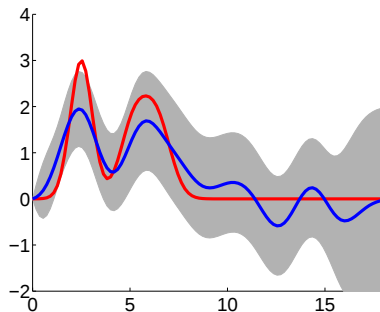
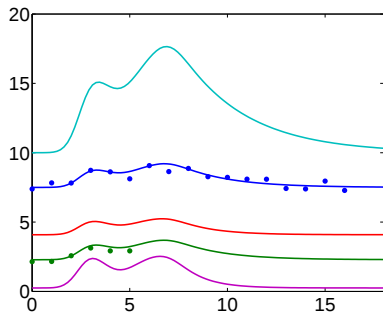
Artificial Example: Inferring $f(t)$



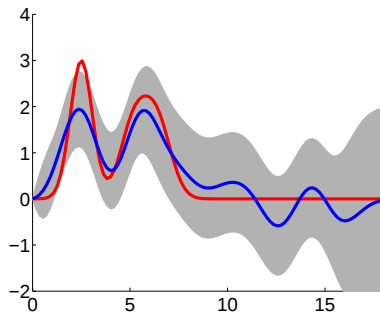
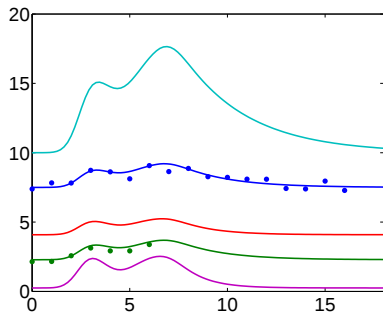
Artificial Example: Inferring $f(t)$



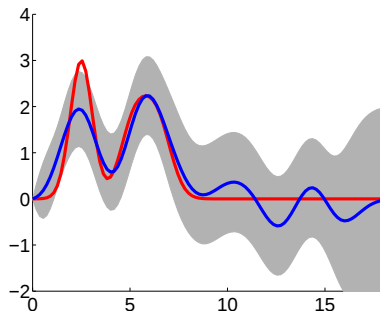
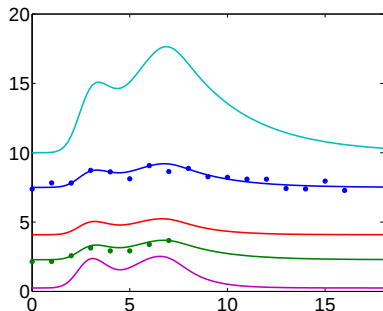
Artificial Example: Inferring $f(t)$



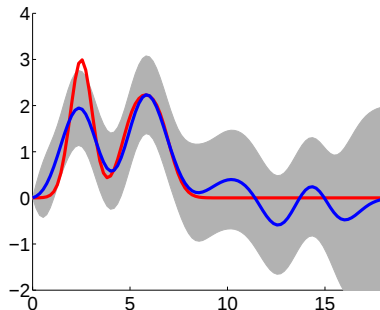
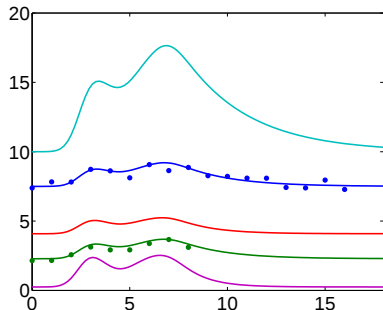
Artificial Example: Inferring $f(t)$



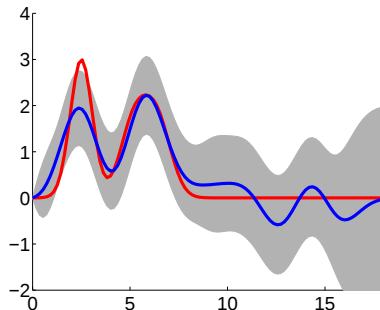
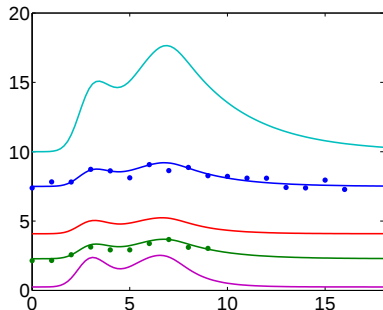
Artificial Example: Inferring $f(t)$



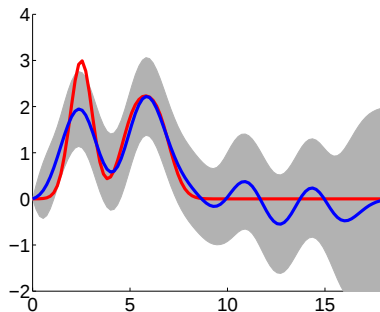
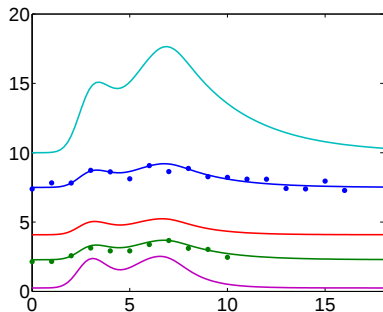
Artificial Example: Inferring $f(t)$



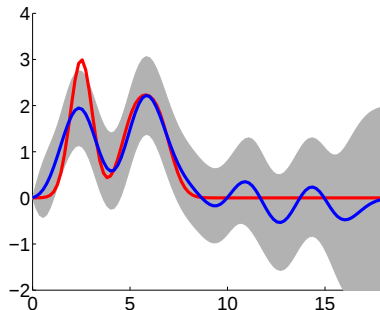
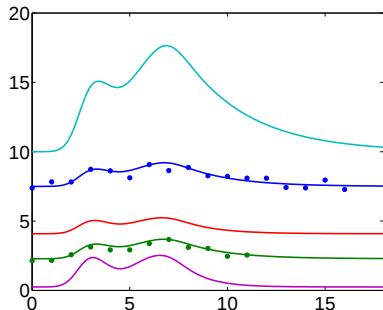
Artificial Example: Inferring $f(t)$



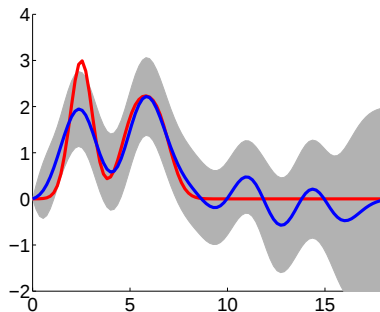
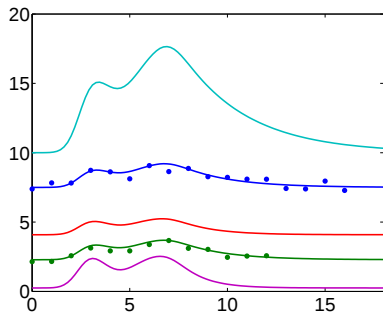
Artificial Example: Inferring $f(t)$



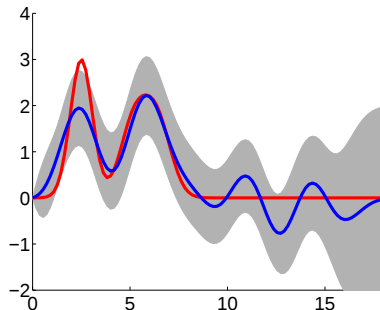
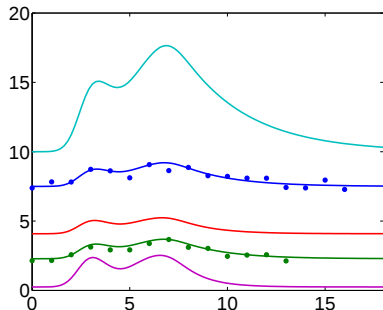
Artificial Example: Inferring $f(t)$



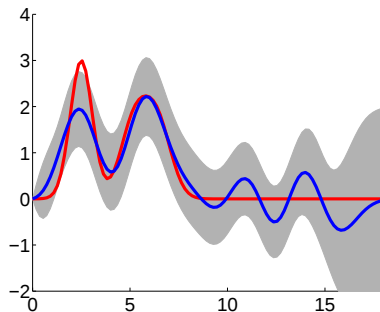
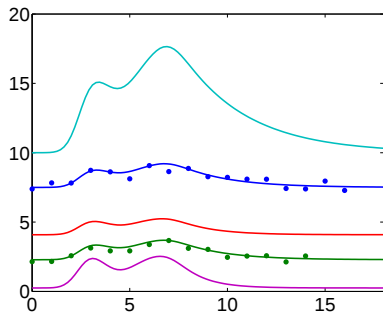
Artificial Example: Inferring $f(t)$



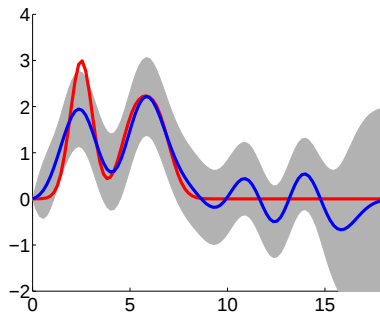
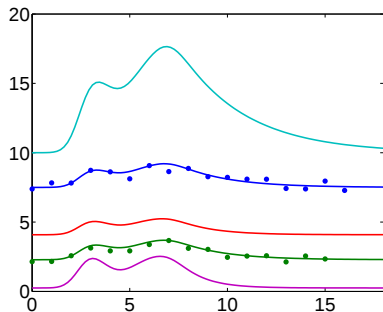
Artificial Example: Inferring $f(t)$



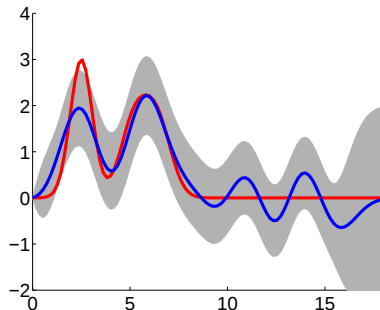
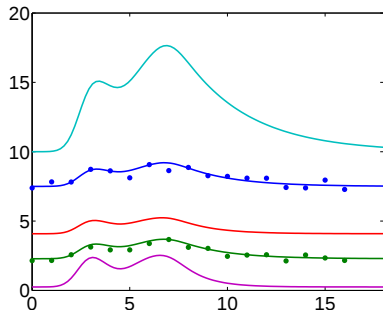
Artificial Example: Inferring $f(t)$



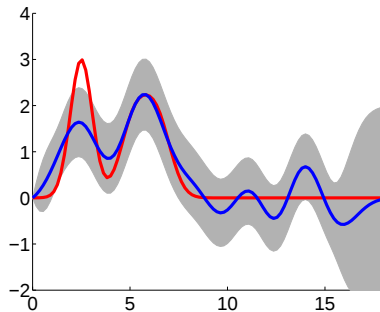
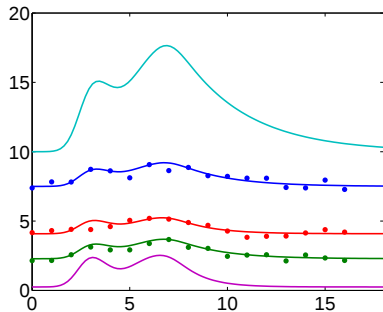
Artificial Example: Inferring $f(t)$



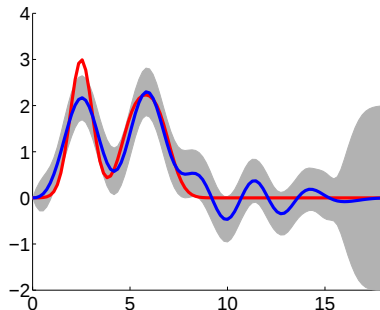
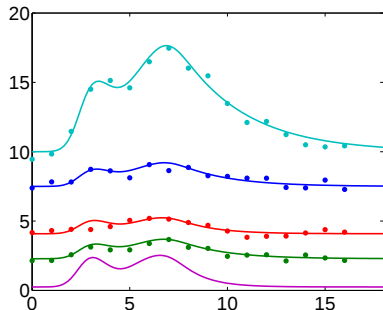
Artificial Example: Inferring $f(t)$



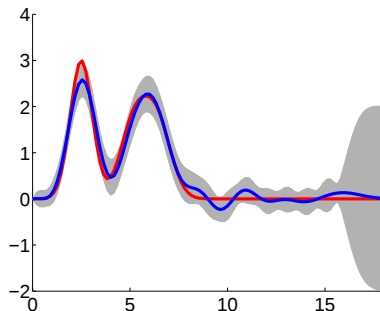
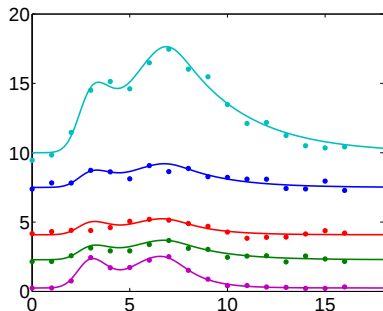
Artificial Example: Inferring $f(t)$



Artificial Example: Inferring $f(t)$



Artificial Example: Inferring $f(t)$



- Responsible for Repairing DNA damage
- Activates DNA Repair proteins
- Pauses the Cell Cycle (prevents replication of damage DNA)
- Initiates *apoptosis* (cell death) in the case where damage can't be repaired.
- Large scale feedback loop with NF- κ B.

p53 DNA Damage Repair

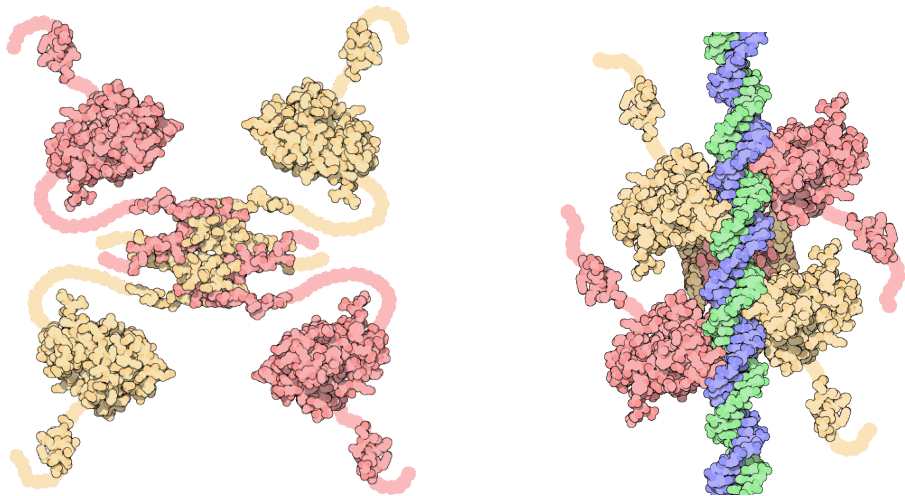


Figure: p53. *Left* unbound, *Right* bound to DNA. Images by David S. Goodsell from <http://www.rcsb.org/> (see the “Molecule of the Month” feature).

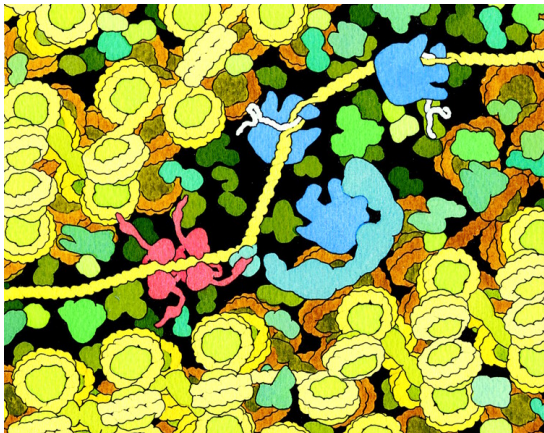


Figure: Repair of DNA damage by p53. Image from Goodsell (1999).

Modelling Assumption

- Assume p53 affects targets as a single input module network motif (SIM).

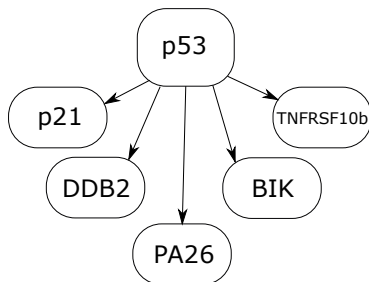
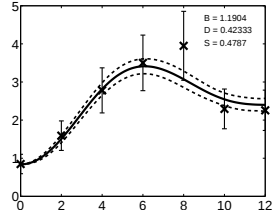
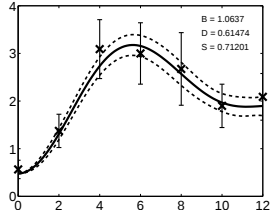
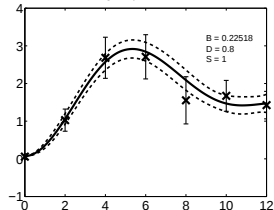
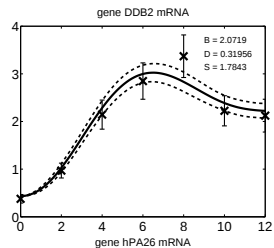
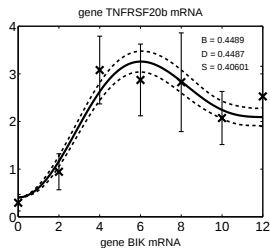
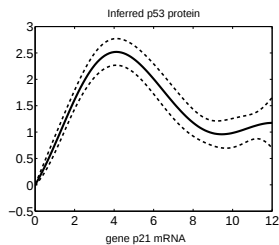
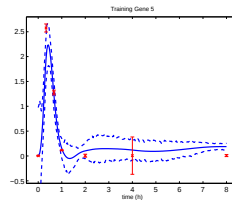
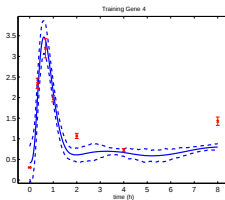
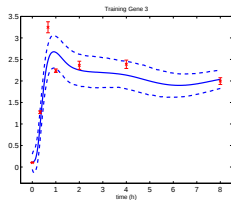
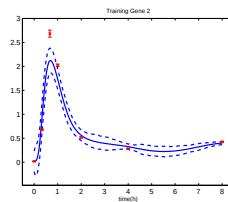
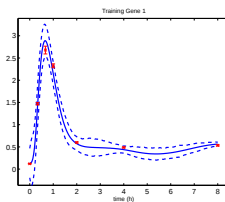
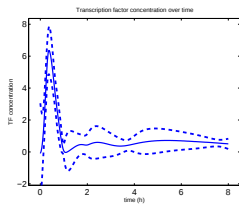


Figure: p53 SIM network motif as modelled by Barenco et al. 2006.

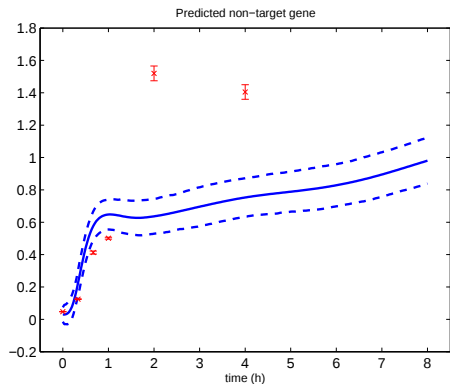
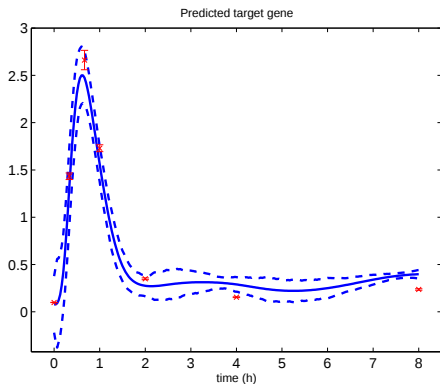


- Target Ranking for Elk-1.
- Elk-1 is phosphorylated by ERK from the EGF signalling pathway.
- Predict concentration of Elk-1 from known targets.
- Rank other targets of Elk-1.



Elk-1 target selection

Fitted model used to rank potential targets of Elk-1



Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Mauricio Alvarez

- Solutions to these differential equations is normally as a convolution.

$$x_i(t) = \int f(u) k_i(u-t) du + h_i(t)$$

$$x_i(t) = \int_0^t f(u) g_i(u) du + h_i(t)$$

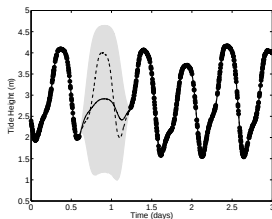
- Convolution Processes (Higdon, 2002; Boyle and Freaan, 2005).
- Convolutions lead to $N \times d$ size covariance matrices $O(N^3 d^3)$ complexity, $O(N^2 d^2)$ storage.
- Model is conditionally independent over $\{x_i(t)\}_{i=1}^d$ given $f(t)$.

Mauricio Alvarez

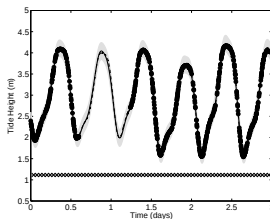
- Can assume conditional independence given $\{f(t_i)\}_{i=1}^k$.
 - ▶ Result is very similar to PITC approximation (Quiñonero Candela and Rasmussen, 2005).
 - ▶ Reduces to $O(N^3 dk^2)$ complexity, $O(N^2 dk)$ storage.
 - ▶ Can also do a FITC style approximation (Snelson and Ghahramani, 2006).
 - ▶ Reduces to $O(Ndk^2)$ complexity, $O(Ndk)$ storage.

Mauricio Alvarez

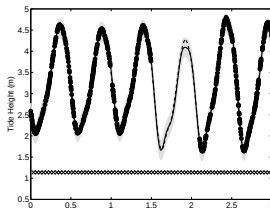
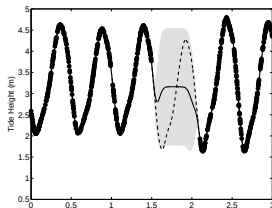
- Network of tide height sensors in the solent — tide heights are correlated.
- Data kindly provided by Alex Rogers (see Rogers et al., 2008).
- $d = 3$ and $N = 1000$ of the 4320 for the training set.
- Simulate sensor failure by knocking out one sensor for a given time.
- For the other two sensors we used all 1000 training observations.
- Take $k = 100$.



(a) Bramblemet Independent



(b) Bramblemet PITC



Mauricio Alvarez

- Jura dataset — concentrations of several heavy metals.
- Prediction 259 data, validation 100 data points.
- Predict *primary variables* (cadmium and copper) at prediction locations in conjunction with some *secondary variables* (nickel and zinc for cadmium; lead, nickel and zinc for copper) (Goovaerts, 1997, p. 248,249).

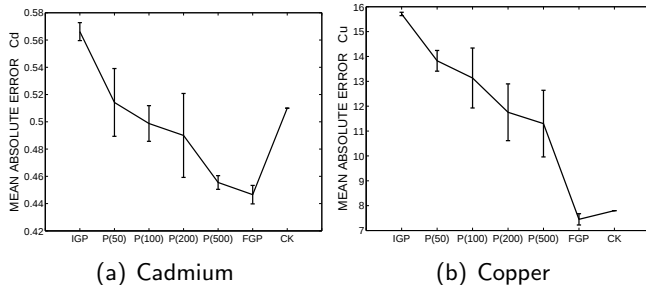


Figure: Mean absolute error. IGP stands for independent GP, $P(M)$ stands for PITC with M inducing values, FGP stands for full GP and CK stands for ordinary co-kriging.

Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models**
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Models of non-linear regulation

- Non-linear Activation: Michaelis-Menten Kinetics

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i f(t)}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Rogers and Girolami (2006)

- Non-linear Repression

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Khanin et al., 2006, PNAS 103

Models of non-linear regulation

- Non-linear Activation: Michaelis-Menten Kinetics

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i f(t)}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Rogers and Girolami (2006)

- Non-linear Repression

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Khanin et al., 2006, PNAS 103

- Non-linear Activation: Michaelis-Menten Kinetics

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i f(t)}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Rogers and Girolami (2006)

- Non-linear Repression

$$\frac{dx_i(t)}{dt} = B_i + \frac{S_i}{\gamma_i + f(t)} - D_i x_i(t)$$

used by Khanin et al., 2006, PNAS 103

MAP Laplace Approximation

Consider the following modification to the model,

$$\frac{dx_j(t)}{dt} = B_j + S_j g(f(t)) - D_j x_j(t),$$

where $g(\cdot)$ is a non-linear function. The differential equation can still be solved,

$$x_j(t) = \frac{B_j}{D_j} + S_j \int_0^t e^{-D_j(t-u)} g_j(f(u)) du$$

Use Laplace's method (Laplace, 1774),

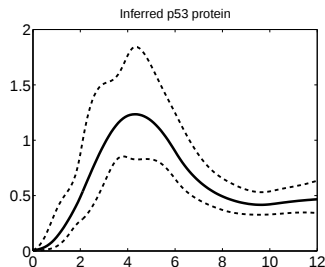
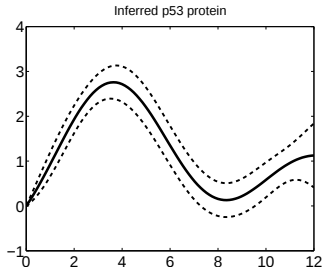
$$p(\mathbf{f} | \mathbf{x}) = N(\hat{\mathbf{f}}, \mathbf{A}^{-1}) \propto \exp\left(-\frac{1}{2} (\mathbf{f} - \hat{\mathbf{f}})^T \mathbf{A} (\mathbf{f} - \hat{\mathbf{f}})\right)$$

where $\hat{\mathbf{f}} = \operatorname{argmax}_{\mathbf{f}} p(\mathbf{f} | \mathbf{x})$ and $\mathbf{A} = -\nabla \nabla \log p(\mathbf{f} | \mathbf{y})|_{\mathbf{f}=\hat{\mathbf{f}}}$ is the Hessian of the negative posterior at that point.

- The Michaelis-Menten activation model uses the following non-linearity

$$g_j(f(t)) = \frac{e^{f(t)}}{\gamma_j + e^{f(t)}},$$

where we are using a GP $f(t)$ to model the log of the TF activity.



(a)

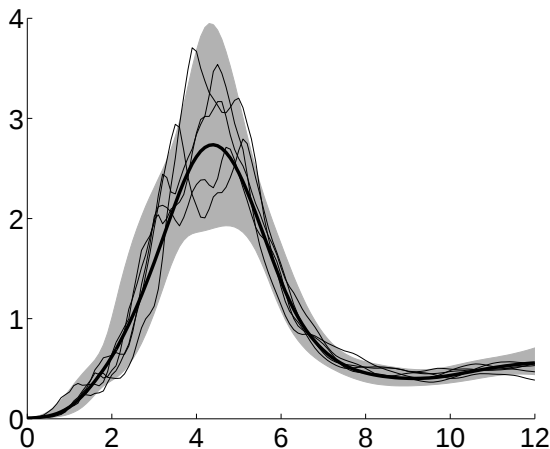


Figure: Laplace approximation error bars along with samples from the true posterior distribution.

- Sample in Gaussian processes

$$p(\mathbf{f}|\mathbf{x}) \propto p(\mathbf{x}|\mathbf{f}) p(\mathbf{f})$$

- Likelihood relates GP to data through

$$x_j(t) = \alpha_j e^{-D_j t} + \frac{B_j}{D_j} + S_j \int_0^t e^{-D_j(t-u)} g_j(f(u)) du$$

- We use *control points* for fast sampling.



Sampling using control points

- Separate the points in $\mathbf{f}$ into two groups:
 - ▶ few control points $\mathbf{f}_c$
 - ▶ and the large majority of the remaining points $\mathbf{f}_\rho = \mathbf{f} \setminus \mathbf{f}_c$
- Sample the control points $\mathbf{f}_c$ using a proposal $q\left(\mathbf{f}_c^{(t+1)}|\mathbf{f}_c^{(t)}\right)$
- Sample the remaining points $\mathbf{f}_\rho$ using the conditional GP prior $p\left(\mathbf{f}_\rho^{(t+1)}|\mathbf{f}_c^{(t+1)}\right)$
- The whole proposal is

$$Q\left(\mathbf{f}^{(t+1)}|\mathbf{f}^{(t)}\right) = p\left(\mathbf{f}_\rho^{(t+1)}|\mathbf{f}_c^{(t+1)}\right) q\left(\mathbf{f}_c^{(t+1)}|\mathbf{f}_c^{(t)}\right)$$

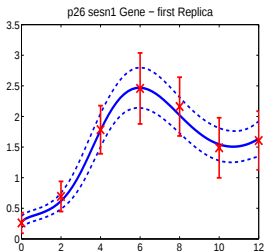
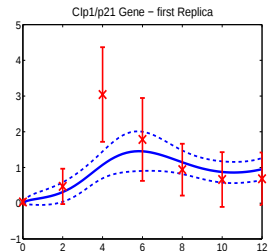
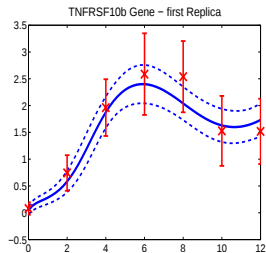
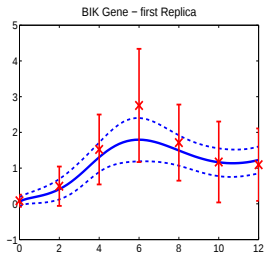
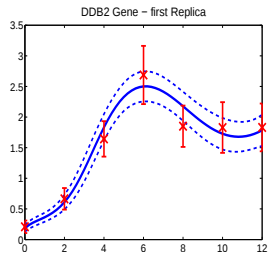
- Its like sampling from the prior $p(\mathbf{f})$ but imposing random walk behaviour through the control points.

- One transcription factor (p53) that acts as an activator. We consider the Michaelis-Menten kinetic equation

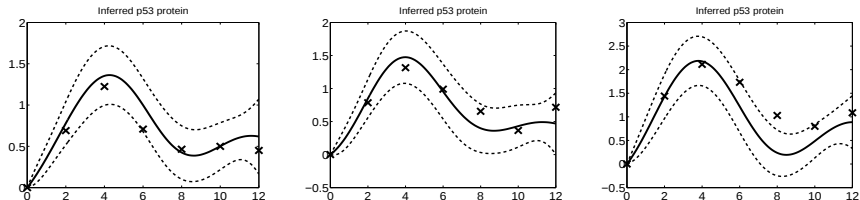
$$\frac{dx_j(t)}{dt} = B_j + S_j \frac{\exp(f(t))}{\exp(f(t)) + \gamma_j} - D_j x_j(t)$$

- MCMC details:
 - ▶ 7 control points are used (placed in a equally spaced grid)
 - ▶ Running time 4/5 hours for 2 million sampling iterations plus burn in
 - ▶ Acceptance rate for $\mathbf{f}$ after burn in was between 15% – 25%

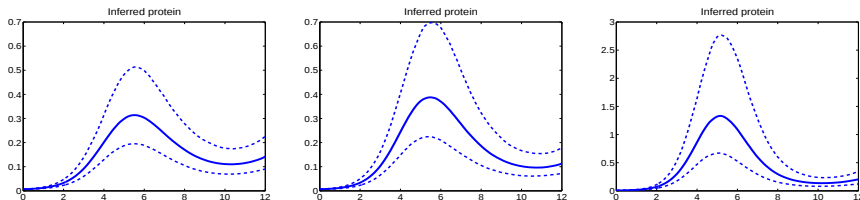
Data used by Barenco et al. (2006): Predicted gene expressions for the 1st replica



Data used by Barenco et al. (2006): Protein concentrations

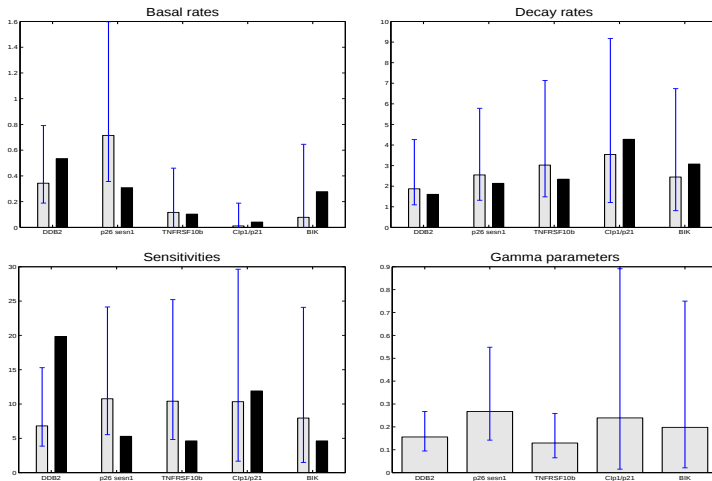


Linear model (Barenco et al. predictions are shown as crosses)



Nonlinear (Michaelis-Menten kinetic equation)

p53 Data Kinetic parameters



Our results (grey) compared with Barenco et al. (2006) (black). Note that Barenco et al. use a linear model

Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations**
- 7 Discussion and Future Work
- 8 Acknowledgements

Antti Honkela

- Transcription factor protein also has governing mRNA.
- This mRNA can be measured.
- In signalling systems this measurement can be misleading because it is activated (phosphorylated) transcription factor that counts.
- In development phosphorylation plays less of a role.

Data from Furlong Lab in EMBL Heidelberg.

- Describe mesoderm development.

Antti Honkela

We take the production rate of active transcription factor to be given by

$$\begin{aligned}\frac{df(t)}{dt} &= \sigma y(t) - \delta f(t) \\ \frac{dx_j(t)}{dt} &= B_j + S_j f(t) - D_j x_j(t)\end{aligned}$$

The solution for $f(t)$, setting transient terms to zero, is

$$f(t) = \sigma \exp(-\delta t) \int_0^t y(u) \exp(\delta u) du .$$

RBF covariance function for $y(t)$

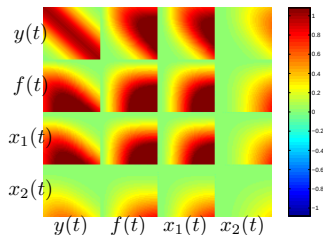
$$f(t) = \sigma \exp(-\delta t) \int_0^t y(u) \exp(\delta u) du$$

$$x_i(t) = \frac{B_i}{D_i} + S_i \exp(-D_i t) \int_0^t f(u) \exp(D_i u) du.$$

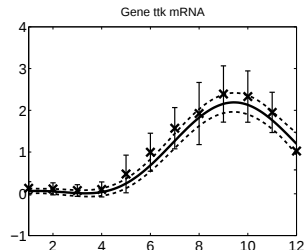
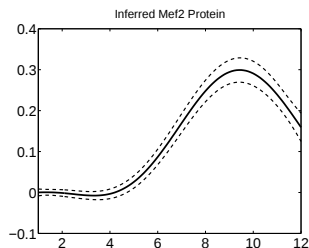
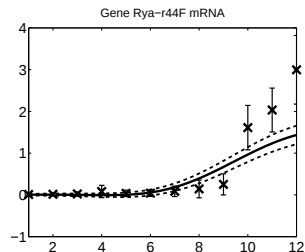
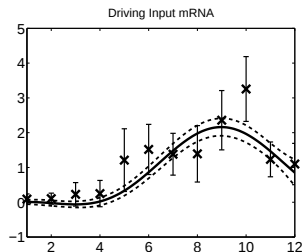
- Joint distribution for $x_1(t)$, $x_2(t)$, $f(t)$ and $y(t)$.

- Here:

δ	D_1	S_1	D_2	S_2
0.1	5	5	0.5	0.5



Results for Mef2 using the Cascade model



Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

- Integration of probabilistic inference with mechanistic models.
- These results are small simple systems.
- Ongoing work:
 - ▶ Scaling up to larger systems
 - ▶ Applications to other types of system, e.g. non-steady-state metabolomics, spatial systems etc.
 - ▶ Improved approximations.
 - ▶ Stochastic differential equations

Outline

- 1 Introduction
- 2 Gaussian Process Review
- 3 Covariance Functions
- 4 Convolutions and Computational Complexity
- 5 Non-linear Response Models
- 6 Cascaded Differential Equations
- 7 Discussion and Future Work
- 8 Acknowledgements

Acknowledgements

- Investigators: Neil Lawrence and Magnus Rattray
- Researchers: Peo Gao, Antti Honkela, Michalis Titsias, Mauricio Alvarez, David Luengo and Jennifer Withers
- Charles Girardot and Eileen Furlong of EMBL in Heidelberg (mesoderm development in *D. Melanogaster*).
- Martino Barenco and Mike Hubank at the Institute of Child Health in UCL (p53 pathway).

Funded by the BBSRC award “Improved Processing of microarray data using probabilistic models” and EPSRC award “Gaussian Processes for Systems Identification with applications in Systems Biology”

References I

- M. Barenco, D. Tomescu, D. Brewer, R. Callard, J. Stark, and M. Hubank. Ranked prediction of p53 targets using hidden variable dynamic modeling. *Genome Biology*, 7(3):R25, 2006. [PDF].
- P. Boyle and M. Frea. Dependent Gaussian processes. In L. Saul, Y. Weiss, and L. Bouttou, editors, *Advances in Neural Information Processing Systems*, volume 17, pages 217–224, Cambridge, MA, 2005. MIT Press. [PDF].
- D. S. Goodsell. The molecular perspective: p53 tumor suppressor. *The Oncologist*, Vol. 4, No. 2, 138–139, April 1999, 4(2): 138–139, 1999.
- P. Goovaerts. *Geostatistics For Natural Resources Evaluation*. Oxford University Press, 1997. [Google Books] .
- D. M. Higdon. Space and space-time modelling using process convolutions. In C. Anderson, V. Barnett, P. Chatwin, and A. El-Shaarawi, editors, *Quantitative methods for current environmental issues*, pages 37–56. Springer-Verlag, 2002.
- R. Khanin, V. Viciotti, and E. Wit. Reconstructing repressor protein levels from expression of gene targets in *E. Coli*. *Proc. Natl. Acad. Sci. USA*, 103(49):18592–18596, 2006. [PDF]. [DOI].
- P. S. Laplace. Mémoire sur la probabilité des causes par les événements. In *Mémoires de mathématique et de physique, présentés à l'Académie Royale des Sciences, par divers savans, & lû dans ses assemblées* 6, pages 621–656, 1774. Translated in ?.
- J. Quiñero Candela and C. E. Rasmussen. A unifying view of sparse approximate Gaussian process regression. *Journal of Machine Learning Research*, 6:1939–1959, 2005. [PDF].
- A. Rogers, M. A. Osborne, S. D. Ramchurn, S. J. Roberts, and N. R. Jennings. Towards real-time information processing of sensor network data using computationally efficient multi-output Gaussian processes. In *Proceedings of the International Conference on Information Processing in Sensor Networks (IPSN 2008)*, 2008. In press.
- S. Rogers and M. Girolami. Model based identification of transcription factor regulatory activity via Markov chain Monte Carlo. Presentation at MASAMB '06, 2006.
- E. Snelson and Z. Ghahramani. Sparse Gaussian processes using pseudo-inputs. In Y. Weiss, B. Schölkopf, and J. C. Platt, editors, *Advances in Neural Information Processing Systems*, volume 18, Cambridge, MA, 2006. MIT Press. [PDF].
- Y. W. Teh, M. Seeger, and M. I. Jordan. Semiparametric latent factor models. In R. G. Cowell and Z. Ghahramani, editors, *Proceedings of the Tenth International Workshop on Artificial Intelligence and Statistics*, pages 333–340, Barbados, 6–8 January 2005. Society for Artificial Intelligence and Statistics. [PDF].