

Probabilistic Dimensional Reduction with the Gaussian Process Latent Variable Model

An Overview of the GP-LVM and its extensions with some state
of the art applications

Neil Lawrence
Machine Learning Group
School of Computer Science
University of Manchester, U.K.

8th February 2007

Outline

- 1 Motivation
 - High Dimensional Data
 - Smooth Low Dimensional Embedded Spaces
 - Existing Methodologies
- 2 GP-LVM
 - Mathematical Foundations
 - Examples
- 3 Extensions
 - Back Constraints
 - Dynamics
 - Hierarchical GP-LVM
- 4 Conclusions
 - Summary



Online Resources

All source code and slides are available online

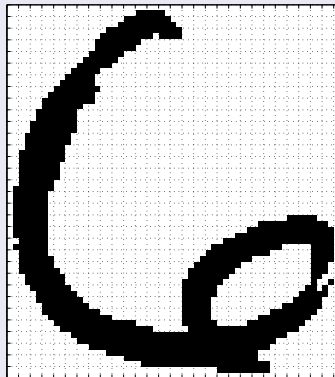
- This talk available from my home page (see talks link on left hand side).
- MATLAB examples in the 'oxford' toolbox (vrs 0.131), `demGplvmTalk`.
 - <http://www.dcs.shef.ac.uk/~neil/oxford/>.
- And the 'fgplvm' toolbox (vrs 0.15).
 - <http://www.dcs.shef.ac.uk/~neil/fgplvm/>.
- MATLAB commands used for examples given in typewriter font.



High Dimensional Data

USPS Data Set Handwritten Digit

- 3648 Dimensions
 - 64 rows by 57 columns
- Space contains more than just this digit.
- Even if we sample every nanosecond from now until the end of the universe, you won't see the original six!



High Dimensional Data

USPS Data Set Handwritten Digit

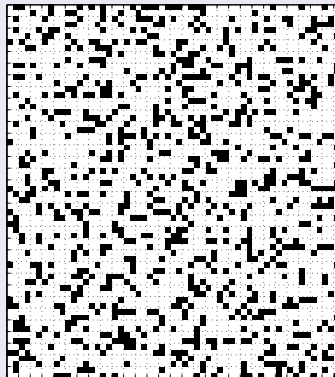
- 3648 Dimensions
 - 64 rows by 57 columns
- Space contains more than just this digit.
- Even if we sample every nanosecond from now until the end of the universe, you won't see the original six!



High Dimensional Data

USPS Data Set Handwritten Digit

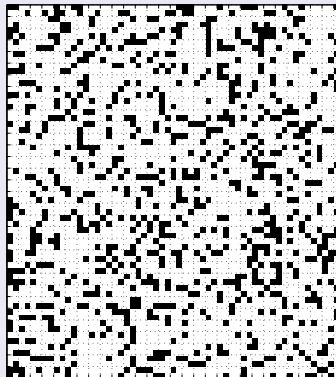
- 3648 Dimensions
 - 64 rows by 57 columns
- Space contains more than just this digit.
- Even if we sample every nanosecond from now until the end of the universe, you won't see the original six!



High Dimensional Data

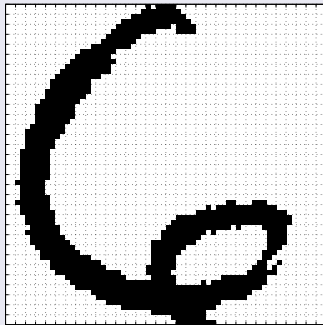
USPS Data Set Handwritten Digit

- 3648 Dimensions
 - 64 rows by 57 columns
- Space contains more than just this digit.
- Even if we sample every nanosecond from now until the end of the universe, you won't see the original six!



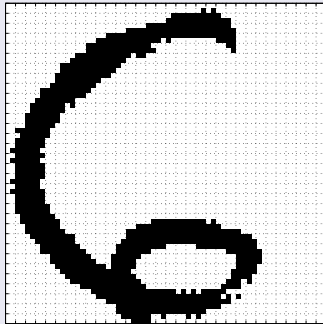
Simple Model of Digit

Rotate a 'Prototype'



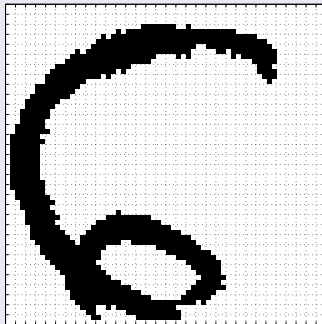
Simple Model of Digit

Rotate a 'Prototype'



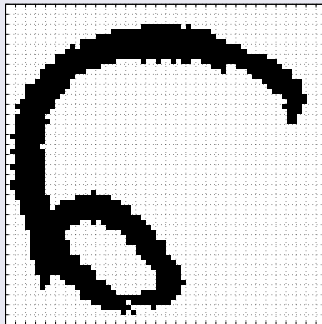
Simple Model of Digit

Rotate a 'Prototype'



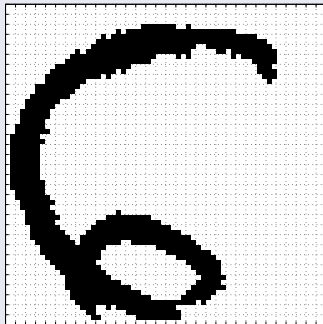
Simple Model of Digit

Rotate a 'Prototype'



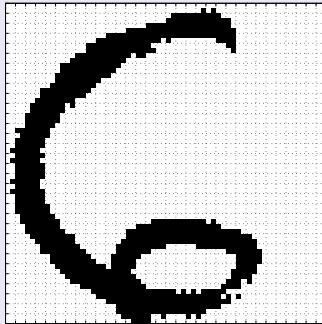
Simple Model of Digit

Rotate a 'Prototype'



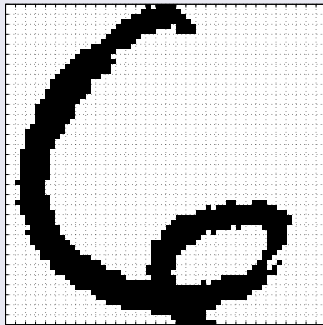
Simple Model of Digit

Rotate a 'Prototype'



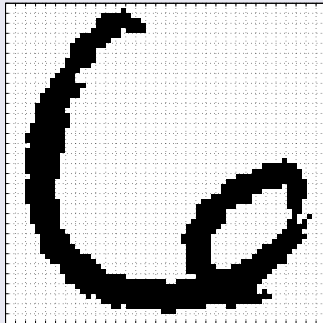
Simple Model of Digit

Rotate a 'Prototype'



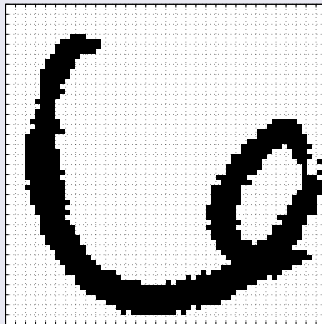
Simple Model of Digit

Rotate a 'Prototype'



Simple Model of Digit

Rotate a 'Prototype'



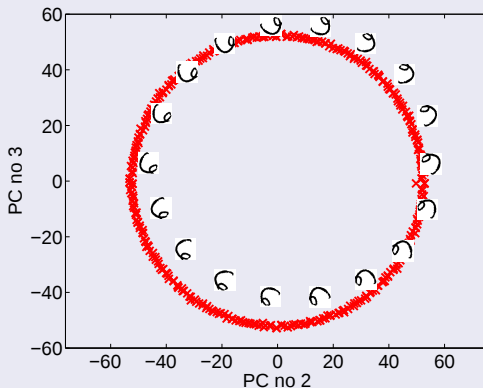
MATLAB Demo

```
demDigitsManifold[2 3], 'all')
```



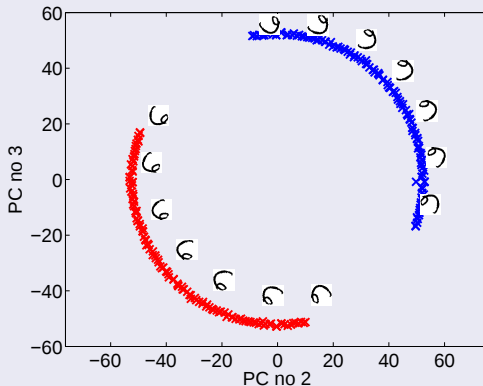
MATLAB Demo

```
demDigitsManifold[2 3], 'all')
```



MATLAB Demo

```
demDigitsManifold([2 3], 'sixnine')
```



Low Dimensional Manifolds

Pure Rotation is too Simple

- In practice the data may undergo several distortions.
 - e.g. digits undergo 'thinning', translation and rotation.
- For data with 'structure':
 - we expect fewer distortions than dimensions;
 - we therefore expect the data to live on a lower dimensional manifold.
- Conclusion: deal with high dimensional data by looking for lower dimensional non-linear embedding.



Existing Methods

Spectral Approaches

- Classical Multidimensional Scaling (MDS) [Mardia et al., 1979].
 - Uses eigenvectors of similarity matrix.
 - Isomap [Tenenbaum et al., 2000] is MDS with a particular proximity measure.
 - Kernel PCA [Schölkopf et al., 1998]
 - Provides an low dimensional representation and a mapping.
 - Mapping is implied through the use of a kernel function as a similarity matrix.
 - Locally Linear Embedding [Roweis and Saul, 2000].
 - Looks to preserve locally linear relationships in a low dimensional space.



Existing Methods II

Iterative Methods

- Multidimensional Scaling (MDS)
 - Iterative optimisation of a stress function [Kruskal, 1964].
 - Sammon Mappings [Sammon, 1969].
 - Strictly speaking not a mapping — similar to iterative MDS.
- NeuroScale [Lowe and Tipping, 1997]
 - Augmentation of iterative MDS methods with a mapping.



Existing Methods III

Probabilistic Approaches

- Probabilistic PCA [Tipping and Bishop, 1999, Roweis, 1998]
 - A linear method.
- Density Networks [MacKay, 1995]
 - Use importance sampling and a multi-layer perceptron.
- Generative Topographic Mapping (GTM) [Bishop et al., 1998]
 - Uses a grid based sample and an RBF network.



Existing Methods III

Probabilistic Approaches

- Probabilistic PCA [Tipping and Bishop, 1999, Roweis, 1998]
 - A linear method.
- Density Networks [MacKay, 1995]
 - Use importance sampling and a multi-layer perceptron.
- Generative Topographic Mapping (GTM) [Bishop et al., 1998]
 - Uses a grid based sample and an RBF network.



Existing Methods III

Probabilistic Approaches

- Probabilistic PCA [Tipping and Bishop, 1999, Roweis, 1998]
 - A linear method.
- Density Networks [MacKay, 1995]
 - Use importance sampling and a multi-layer perceptron.
- Generative Topographic Mapping (GTM) [Bishop et al., 1998]
 - Uses a grid based sample and an RBF network.



Existing Methods III

Probabilistic Approaches

- Probabilistic PCA [Tipping and Bishop, 1999, Roweis, 1998]
 - A linear method.
- Density Networks [MacKay, 1995]
 - Use importance sampling and a multi-layer perceptron.
- Generative Topographic Mapping (GTM) [Bishop et al., 1998]
 - Uses a grid based sample and an RBF network.

Difficulty for Probabilistic Approaches

Propagate a probability distribution through a non-linear mapping.



The New Model

A Probabilistic Non-linear PCA

- PCA has a probabilistic interpretation [Tipping and Bishop, 1999].
- It is difficult to 'non-linearise'.

Dual Probabilistic PCA

- We present a new probabilistic interpretation of PCA [Lawrence, 2005].
- This interpretation can be made non-linear.
- The result is non-linear probabilistic PCA.



Notation

q — dimension of latent/embedded space

d — dimension of data space

n — number of data points

centred data, $\mathbf{Y} = [\mathbf{y}_{1,:}, \dots, \mathbf{y}_{n,:}]^T = [\mathbf{y}_{:,1}, \dots, \mathbf{y}_{:,d}] \in \mathbb{R}^{n \times d}$

latent variables, $\mathbf{X} = [\mathbf{x}_{1,:}, \dots, \mathbf{x}_{n,:}]^T = [\mathbf{x}_{:,1}, \dots, \mathbf{x}_{:,q}] \in \mathbb{R}^{n \times q}$

mapping matrix, $\mathbf{W} \in \mathbb{R}^{d \times q}$

$\mathbf{a}_{i,:}$ is a vector from the i th row of a given matrix $\mathbf{A}$

$\mathbf{a}_{:,j}$ is a vector from the j th row of a given matrix $\mathbf{A}$



Reading Notation

X and **Y** are *design matrices*

- Covariance given by $n^{-1}\mathbf{Y}^T\mathbf{Y}$.
- Inner product matrix given by $\mathbf{Y}\mathbf{Y}^T$.



Linear Dimensionality Reduction

Linear Latent Variable Model

- Represent data, $\mathbf{Y}$, with a lower dimensional set of latent variables $\mathbf{X}$.
- Assume a linear relationship of the form

$$\mathbf{y}_{i,:} = \mathbf{W}\mathbf{x}_{i,:} + \boldsymbol{\eta}_{i,:},$$

where

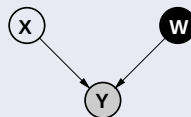
$$\boldsymbol{\eta}_{i,:} \sim N(\mathbf{0}, \sigma^2 \mathbf{I}).$$



Linear Latent Variable Model

Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Standard** Latent variable approach:
 - Define Gaussian prior over *latent space*, $\mathbf{X}$.
 - Integrate out *latent variables*.



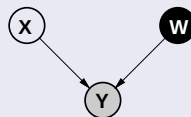
$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$



Linear Latent Variable Model

Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Standard** Latent variable approach:
 - Define Gaussian prior over *latent space*, $\mathbf{X}$.
 - Integrate out *latent variables*.



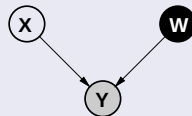
$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$



Linear Latent Variable Model

Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Standard** Latent variable approach:
 - Define Gaussian prior over *latent space*, $\mathbf{X}$.
 - Integrate out *latent variables*.



$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$

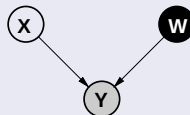
$$p(\mathbf{X}) = \prod_{i=1}^n N(\mathbf{x}_{i,:} | \mathbf{0}, \mathbf{I})$$



Linear Latent Variable Model

Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Standard** Latent variable approach:
 - Define Gaussian prior over *latent space*, $\mathbf{X}$.
 - Integrate out *latent variables*.



$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$

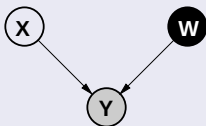
$$p(\mathbf{X}) = \prod_{i=1}^n N(\mathbf{x}_{i,:} | \mathbf{0}, \mathbf{I})$$

$$p(\mathbf{Y}|\mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{0}, \mathbf{W}\mathbf{W}^T + \sigma^2 \mathbf{I})$$



Linear Latent Variable Model II

Probabilistic PCA Max. Likelihood Soln [Tipping and Bishop, 1999]



$$p(\mathbf{Y}|\mathbf{W}) = \prod_{i=1}^n \mathcal{N}(\mathbf{y}_{i,:} | \mathbf{0}, \mathbf{W}\mathbf{W}^T + \sigma^2\mathbf{I})$$



Linear Latent Variable Model II

Probabilistic PCA Max. Likelihood Soln [Tipping and Bishop, 1999]

$$p(\mathbf{Y}|\mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:}|\mathbf{0}, \mathbf{C}), \quad \mathbf{C} = \mathbf{W}\mathbf{W}^T + \sigma^2\mathbf{I}$$

$$\log p(\mathbf{Y}|\mathbf{W}) = -\frac{n}{2} \log |\mathbf{C}| - \frac{1}{2} \text{tr}(\mathbf{C}^{-1}\mathbf{Y}^T\mathbf{Y}) + \text{const.}$$

If $\mathbf{U}_q$ are first q principal eigenvectors of $n^{-1}\mathbf{Y}^T\mathbf{Y}$ and the corresponding eigenvalues are Λ_q ,

$$\mathbf{W} = \mathbf{U}_q \mathbf{L} \mathbf{V}^T, \quad \mathbf{L} = (\Lambda_q - \sigma^2 \mathbf{I})^{\frac{1}{2}}$$

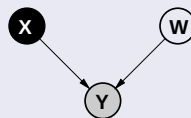
where $\mathbf{V}$ is an arbitrary rotation matrix.



Linear Latent Variable Model III

Dual Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Novel** Latent variable approach:
 - Define Gaussian prior over *parameters*, $\mathbf{W}$.
 - Integrate out *parameters*.



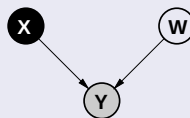
$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$



Linear Latent Variable Model III

Dual Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Novel** Latent variable approach:
 - Define Gaussian prior over *parameters*, $\mathbf{W}$.
 - Integrate out *parameters*.



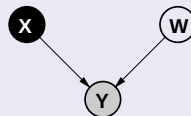
$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$



Linear Latent Variable Model III

Dual Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Novel** Latent variable approach:
 - Define Gaussian prior over *parameters*, $\mathbf{W}$.
 - Integrate out *parameters*.



$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$

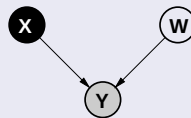
$$p(\mathbf{W}) = \prod_{i=1}^d N(\mathbf{w}_{i,:} | \mathbf{0}, \mathbf{I})$$



Linear Latent Variable Model III

Dual Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Novel** Latent variable approach:
 - Define Gaussian prior over *parameters*, $\mathbf{W}$.
 - Integrate out *parameters*.



$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$

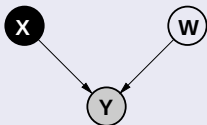
$$p(\mathbf{W}) = \prod_{i=1}^d N(\mathbf{w}_{i,:} | \mathbf{0}, \mathbf{I})$$

$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{X}\mathbf{X}^T + \sigma^2 \mathbf{I})$$



Linear Latent Variable Model IV

Dual Probabilistic PCA Max. Likelihood Soln [Lawrence, 2004]



$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{X}\mathbf{X}^T + \sigma^2 \mathbf{I})$$



Linear Latent Variable Model IV

Dual Probabilistic PCA Max. Likelihood Soln [Lawrence, 2004]

$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{K}), \quad \mathbf{K} = \mathbf{X}\mathbf{X}^T + \sigma^2 \mathbf{I}$$

$$\log p(\mathbf{Y}|\mathbf{X}) = -\frac{d}{2} \log |\mathbf{K}| - \frac{1}{2} \text{tr}(\mathbf{K}^{-1} \mathbf{Y}\mathbf{Y}^T) + \text{const.}$$

If $\mathbf{U}'_q$ are first q principal eigenvectors of $d^{-1} \mathbf{Y}\mathbf{Y}^T$ and the corresponding eigenvalues are Λ_q ,

$$\mathbf{X} = \mathbf{U}'_q \mathbf{L} \mathbf{V}^T, \quad \mathbf{L} = (\Lambda_q - \sigma^2 \mathbf{I})^{\frac{1}{2}}$$

where $\mathbf{V}$ is an arbitrary rotation matrix.



Linear Latent Variable Model IV

Probabilistic PCA Max. Likelihood Soln [Tipping and Bishop, 1999]

$$p(\mathbf{Y}|\mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:}|\mathbf{0}, \mathbf{C}), \quad \mathbf{C} = \mathbf{W}\mathbf{W}^T + \sigma^2\mathbf{I}$$

$$\log p(\mathbf{Y}|\mathbf{W}) = -\frac{n}{2} \log |\mathbf{C}| - \frac{1}{2} \text{tr}(\mathbf{C}^{-1} \mathbf{Y}^T \mathbf{Y}) + \text{const.}$$

If $\mathbf{U}_q$ are first q principal eigenvectors of $n^{-1} \mathbf{Y}^T \mathbf{Y}$ and the corresponding eigenvalues are Λ_q ,

$$\mathbf{W} = \mathbf{U}_q \mathbf{L} \mathbf{V}^T, \quad \mathbf{L} = (\Lambda_q - \sigma^2 \mathbf{I})^{\frac{1}{2}}$$

where $\mathbf{V}$ is an arbitrary rotation matrix.



Equivalence of Formulations

The Eigenvalue Problems are equivalent

- Solution for Probabilistic PCA (solves for the mapping)

$$\mathbf{Y}^T \mathbf{Y} \mathbf{U}_q = \mathbf{U}_q \Lambda_q \quad \mathbf{W} = \mathbf{U}_q \mathbf{L} \mathbf{V}^T$$

- Solution for Dual Probabilistic PCA (solves for the latent positions)

$$\mathbf{Y} \mathbf{Y}^T \mathbf{U}'_q = \mathbf{U}'_q \Lambda_q \quad \mathbf{X} = \mathbf{U}'_q \mathbf{L} \mathbf{V}^T$$

- Equivalence is from

$$\mathbf{U}_q = \mathbf{Y}^T \mathbf{U}'_q \Lambda_q^{-\frac{1}{2}}$$



Gaussian Process (GP)

Prior for Functions

- Probability Distribution over Functions
 - Functions are infinite dimensional.
 - Prior distribution over *instantiations* of the function: finite dimensional objects.
- Can prove by induction that GP is 'consistent'.
- Mean and Covariance Functions
 - Instead of mean and covariance matrix, GP is defined by mean function and covariance function.
 - Mean function often taken to be zero or constant.
 - Covariance function must be *positive definite*.
 - Class of valid covariance functions is the same as the class of *Mercer kernels*.



Gaussian Processes II

Zero mean Gaussian Process

- A (zero mean) Gaussian process likelihood is of the form

$$p(\mathbf{y}|\mathbf{X}) = N(\mathbf{y}|\mathbf{0}, \mathbf{K}),$$

where $\mathbf{K}$ is the covariance function or *kernel*.

- The *linear kernel* with noise has the form

$$\mathbf{K} = \mathbf{X}\mathbf{X}^T + \sigma^2\mathbf{I}$$

- Priors over non-linear functions are also possible.
 - To see what functions look like, we can sample from the prior process.



Covariance Samples

demCovFuncSample

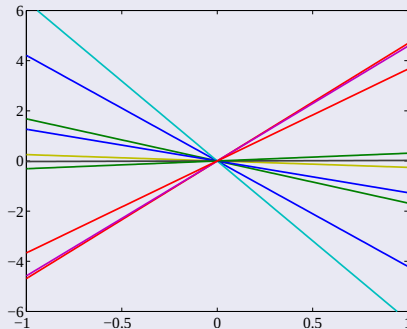


Figure: linear kernel, $\mathbf{K} = \mathbf{xx}^T$



Covariance Samples

demCovFuncSample

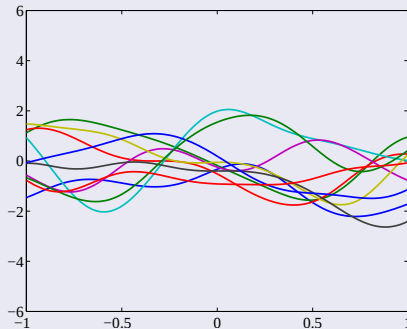


Figure: RBF kernel with $\gamma = 10$, $\alpha = 1$



Covariance Samples

demCovFuncSample

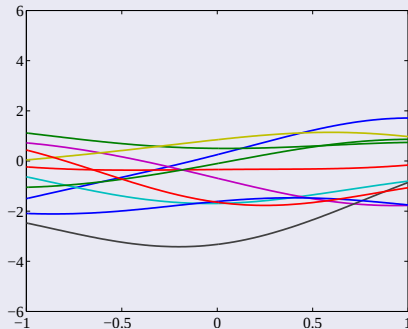


Figure: RBF kernel with $l = 1$, $\alpha = 1$



Covariance Samples

demCovFuncSample

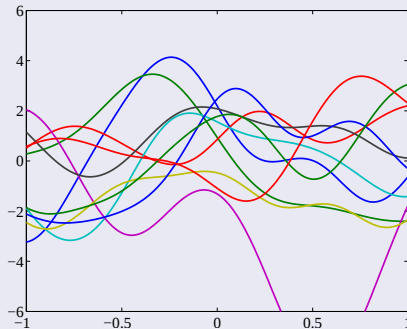


Figure: RBF kernel with $l = 0.3$, $\alpha = 4$



Covariance Samples

demCovFuncSample

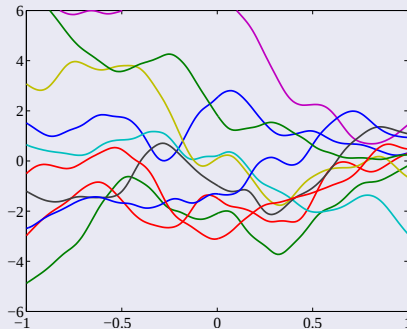


Figure: MLP kernel with $\alpha = 8$, $w = 100$ and $b = 100$



Covariance Samples

demCovFuncSample

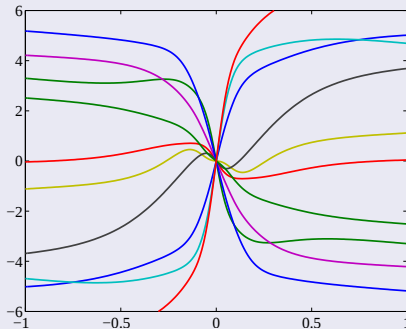


Figure: MLP kernel with $\alpha = 8$, $b = 0$ and $w = 100$



Covariance Samples

demCovFuncSample

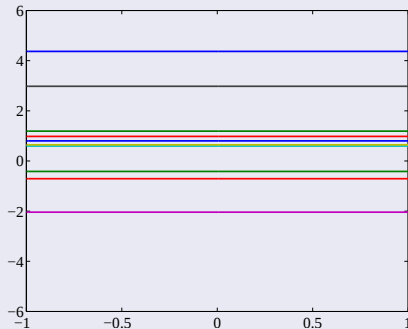


Figure: bias kernel with $\alpha = 1$ and



Covariance Samples

demCovFuncSample

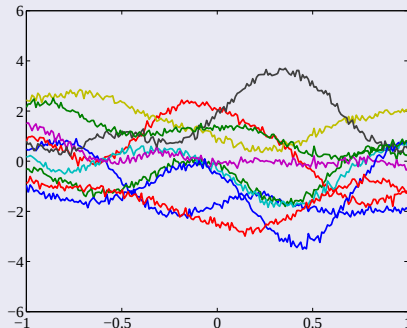


Figure: summed combination of: RBF kernel, $\alpha = 1$, $l = 0.3$; bias kernel, $\alpha = 1$; and white noise kernel, $\beta = 100$



Gaussian Process Regression

Posterior Distribution over Functions

- Gaussian processes are often used for regression.
- We are given a known inputs $\mathbf{X}$ and targets $\mathbf{Y}$.
- We assume a prior distribution over functions by selecting a kernel.
- Combine the prior with data to get a *posterior* distribution over functions.



Gaussian Process Regression

demRegression

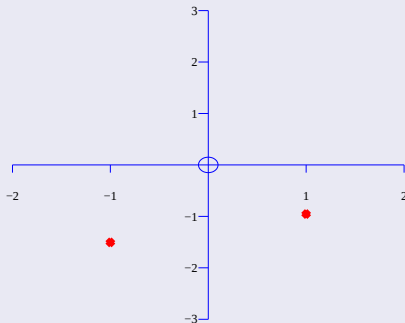


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

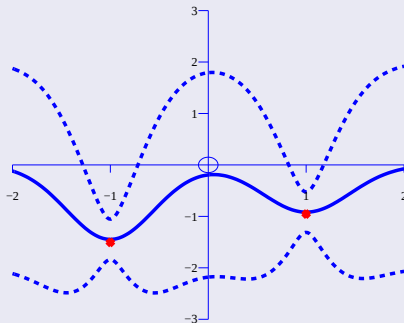


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

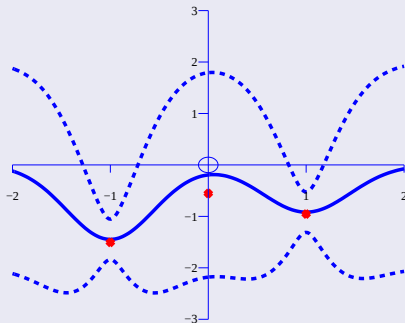


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

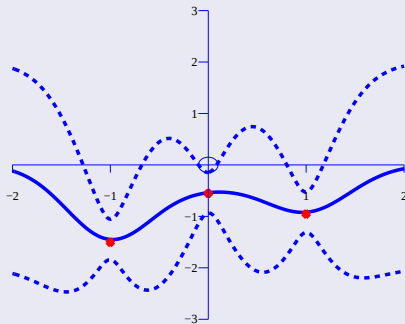


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

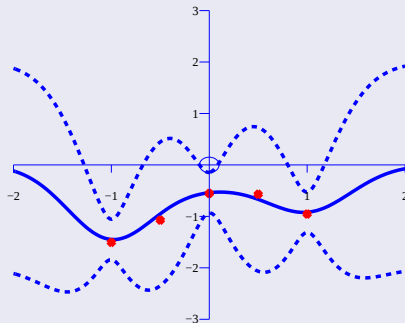


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

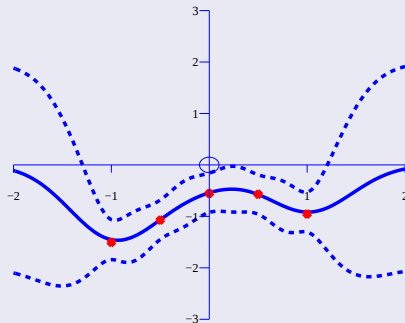


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

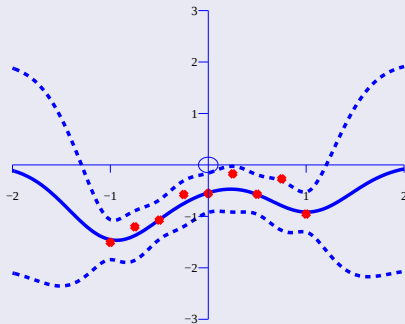


Figure: Examples include WiFi localization, C14 calibration curve.



Gaussian Process Regression

demRegression

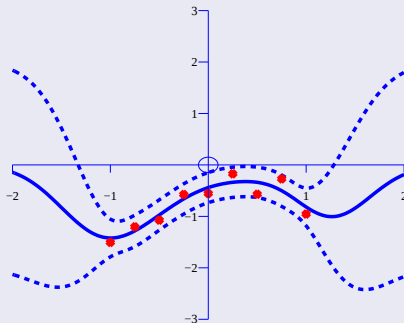


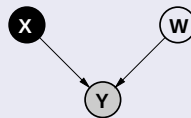
Figure: Examples include WiFi localization, C14 calibration curve.



Non-Linear Latent Variable Model

Dual Probabilistic PCA

- Define *linear-Gaussian relationship* between latent variables and data.
- **Novel** Latent variable approach:
 - Define Gaussian prior over *parameters*, $\mathbf{W}$.
 - Integrate out *parameters*.



$$p(\mathbf{Y}|\mathbf{X}, \mathbf{W}) = \prod_{i=1}^n N(\mathbf{y}_{i,:} | \mathbf{W}\mathbf{x}_{i,:}, \sigma^2 \mathbf{I})$$

$$p(\mathbf{W}) = \prod_{i=1}^d N(\mathbf{w}_{i,:} | \mathbf{0}, \mathbf{I})$$

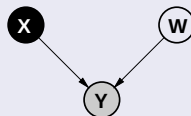
$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{X}\mathbf{X}^T + \sigma^2 \mathbf{I})$$



Non-Linear Latent Variable Model

Dual Probabilistic PCA

- Inspection of the marginal likelihood shows ...
 - The covariance matrix is a covariance function.
 - We recognise it as the 'linear kernel'.
 - We call this the Gaussian Process Latent Variable model (GP-LVM).



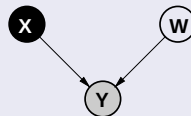
$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{X}\mathbf{X}^T + \sigma^2 \mathbf{I})$$



Non-Linear Latent Variable Model

Dual Probabilistic PCA

- Inspection of the marginal likelihood shows ...
 - The covariance matrix is a covariance function.
 - We recognise it as the 'linear kernel'.
 - We call this the Gaussian Process Latent Variable model (GP-LVM).



$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{K})$$

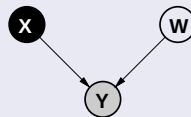
$$\mathbf{K} = \mathbf{X}\mathbf{X}^T + \sigma^2\mathbf{I}$$



Non-Linear Latent Variable Model

Dual Probabilistic PCA

- Inspection of the marginal likelihood shows ...
 - The covariance matrix is a covariance function.
 - We recognise it as the 'linear kernel'.
 - We call this the Gaussian Process Latent Variable model (GP-LVM).



$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{K})$$

$$\mathbf{K} = \mathbf{X}\mathbf{X}^T + \sigma^2\mathbf{I}$$

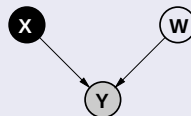
This is a product of Gaussian processes with linear kernels.



Non-Linear Latent Variable Model

Dual Probabilistic PCA

- Inspection of the marginal likelihood shows ...
 - The covariance matrix is a covariance function.
 - We recognise it as the 'linear kernel'.
 - We call this the Gaussian Process Latent Variable model (GP-LVM).



$$p(\mathbf{Y}|\mathbf{X}) = \prod_{j=1}^d N(\mathbf{y}_{:,j} | \mathbf{0}, \mathbf{K})$$

$$\mathbf{K} = ?$$

Replace linear kernel with non-linear kernel for non-linear model.



Non-linear Latent Variable Models

RBF Kernel

- The RBF kernel has the form $k_{ij} = k(\mathbf{x}_{i,:}, \mathbf{x}_{j,:})$, where

$$k(\mathbf{x}_{i,:}, \mathbf{x}_{j,:}) = \alpha \exp \left(-\frac{(\mathbf{x}_{i,:} - \mathbf{x}_{j,:})^T (\mathbf{x}_{i,:} - \mathbf{x}_{j,:})}{2l^2} \right).$$

- No longer possible to optimise wrt $\mathbf{X}$ via an eigenvalue problem.
- Instead find gradients with respect to $\mathbf{X}$, α , l and σ^2 and optimise using conjugate gradients.



Applications

Style Based Inverse Kinematics

Facilitating animation through modelling human motion with the GP-LVM [Grochow et al., 2004]

Tracking

Tracking using models of human motion learnt with the GP-LVM [Urtasun et al., 2005, 2006]

.



Stick Man

Generalization with less Data than Dimensions

- Powerful uncertainty handling of GPs leads to surprising properties.
- Non-linear models can be used where there are fewer data points than dimensions *without overfitting*.
- Example: Modelling a stick man in 102 dimensions with 55 data points!



Stick Man II

demStick1

Figure: The latent space for the stick man motion capture data.



Stick Man II

demStick1

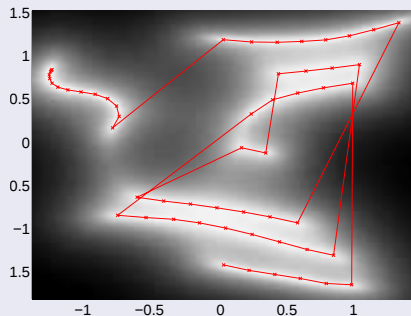


Figure: The latent space for the stick man motion capture data.

Back Constraints I

Local Distance Preservation [Lawrence and Quiñero Candela, 2006]

- Most dimensional reduction techniques preserve local distances.
- The GP-LVM does not.
- GP-LVM maps smoothly from latent to data space.
 - Points close in latent space are close in data space.
 - This does not imply points close in data space are close in latent space.
- Kernel PCA maps smoothly from data to latent space.
 - Points close in data space are close in latent space.
 - This does not imply points close in latent space are close in data space.

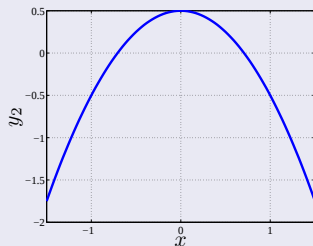
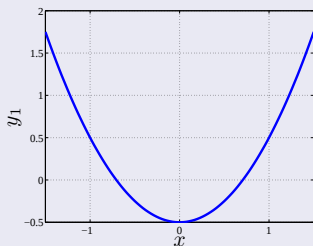


Back Constraints II

Forward Mapping (demBackMapping in oxford toolbox)

- Mapping from 1-D latent space to 2-D data space.

$$y_1 = x^2 - 0.5, \quad y_2 = -x^2 + 0.5$$

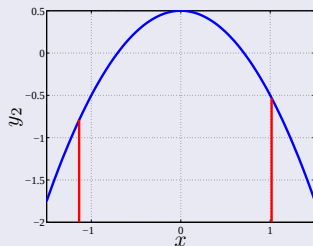
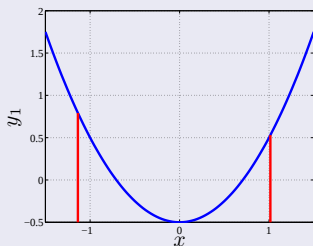


Back Constraints II

Forward Mapping (demBackMapping in oxford toolbox)

- Mapping from 1-D latent space to 2-D data space.

$$y_1 = x^2 - 0.5, \quad y_2 = -x^2 + 0.5$$

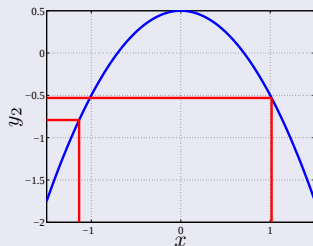
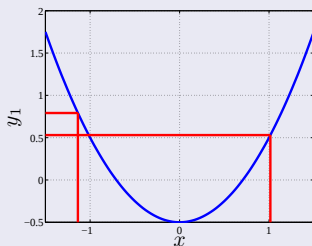


Back Constraints II

Forward Mapping (demBackMapping in oxford toolbox)

- Mapping from 1-D latent space to 2-D data space.

$$y_1 = x^2 - 0.5, \quad y_2 = -x^2 + 0.5$$

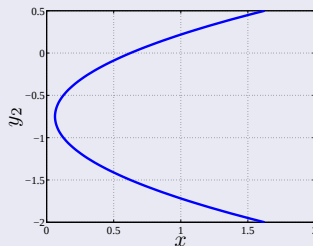
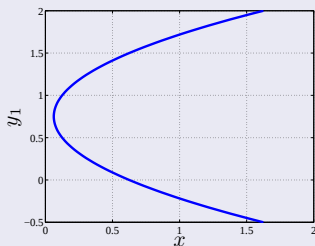


Back Constraints II

Backward Mapping (demBackMapping in oxford toolbox)

- Mapping from 2-D data space to 1-D latent.

$$x = 0.5 (y_1^2 + y_2^2 + 1)$$

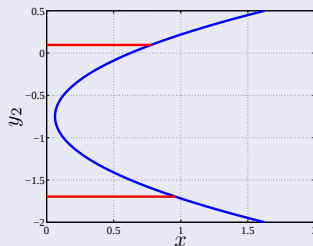
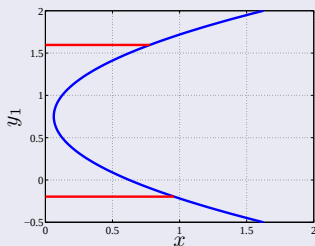


Back Constraints II

Backward Mapping (demBackMapping in oxford toolbox)

- Mapping from 2-D data space to 1-D latent.

$$x = 0.5 (y_1^2 + y_2^2 + 1)$$

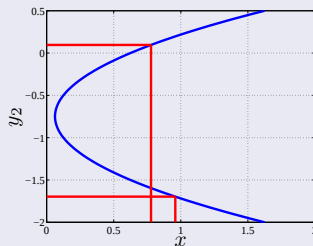
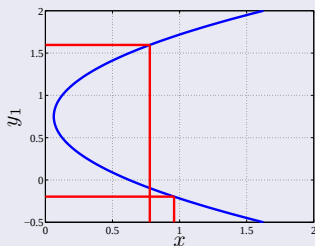


Back Constraints II

Backward Mapping (demBackMapping in oxford toolbox)

- Mapping from 2-D data space to 1-D latent.

$$x = 0.5 (y_1^2 + y_2^2 + 1)$$



NeuroScale

Multi-Dimensional Scaling with a Mapping

- Lowe and Tipping [1997] made latent positions a function of the data.

$$x_{ij} = f_j(\mathbf{y}_i; \mathbf{w})$$

- Function was either multi-layer perceptron or a radial basis function network.
- Their motivation was different from ours:
 - They wanted to add the advantages of a true mapping to multi-dimensional scaling.



Back Constraints in the GP-LVM

Back Constraints

- We can use the same idea to force the GP-LVM to respect local distances.[Lawrence and Quiñero Candela, 2006]
 - By constraining each $\mathbf{x}_i$ to be a 'smooth' mapping from $\mathbf{y}_i$ local distances can be respected.
- This works because in the GP-LVM we maximise wrt latent variables, we don't integrate out.
- Can use any 'smooth' function:
 - 1 Neural network.
 - 2 RBF Network.
 - 3 Kernel based mapping.



Optimising BC-GPLVM

Computing Gradients

- GP-LVM normally proceeds by optimising

$$L(\mathbf{X}) = \log p(\mathbf{Y}|\mathbf{X})$$

with respect to $\mathbf{X}$ using $\frac{dL}{d\mathbf{X}}$.

- The back constraints are of the form

$$x_{ij} = f_j(\mathbf{y}_{i,:}; \mathbf{B})$$

where $\mathbf{B}$ are parameters.

- We can compute $\frac{dL}{d\mathbf{B}}$ via chain rule and optimise parameters of mapping.



Motion Capture Results

demStick1 and demStick3

Figure: The latent space for the motion capture data with (*right*) and without (*left*) dynamics. The dynamics use a Gaussian process with an RBF kernel.



Motion Capture Results

demStick1 and demStick3

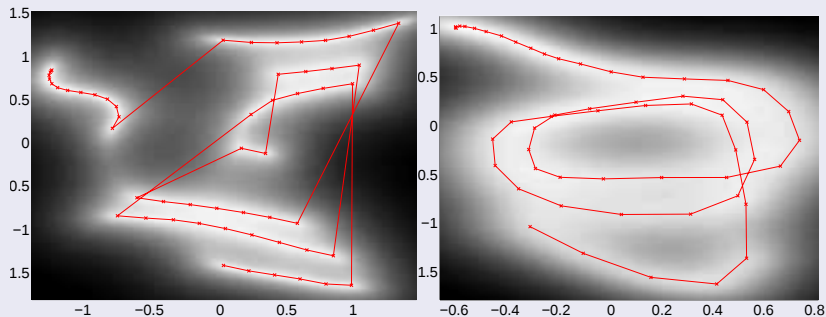
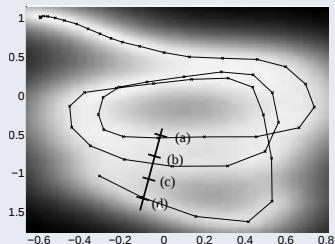


Figure: The latent space for the motion capture data with (*right*) and without (*left*) dynamics. The dynamics use a Gaussian process with an RBF kernel.



Stick Man Results

demStickResults



Projection into data space from four points in the latent space. The

Adding Dynamics

MAP Solutions for Dynamics Models

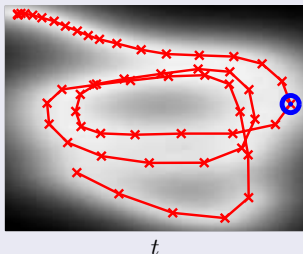
- Data often has a temporal ordering.
- Markov-based dynamics are often used.
- For the GP-LVM
 - Marginalising such dynamics is intractable.
 - But: MAP solutions are trivial to implement.
- Many choices: Kalman filter, Markov chains *etc.*.
- Wang et al. [2006] suggest using a Gaussian Process.



Gaussian Process Dynamics

GP-LVM with Dynamics

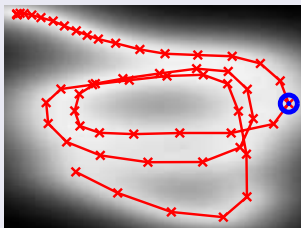
- Autoregressive Gaussian process mapping in latent space between time points.



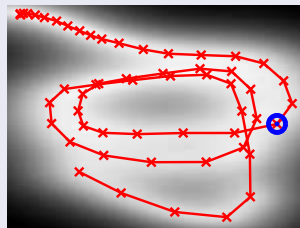
Gaussian Process Dynamics

GP-LVM with Dynamics

- Autoregressive Gaussian process mapping in latent space between time points.



t



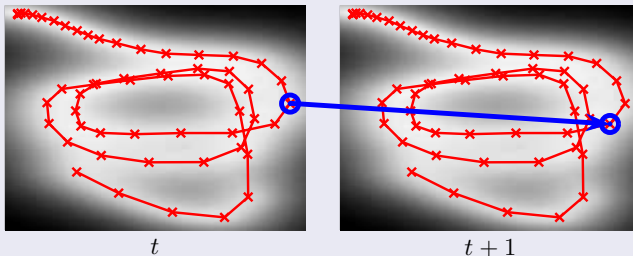
$t + 1$



Gaussian Process Dynamics

GP-LVM with Dynamics

- Autoregressive Gaussian process mapping in latent space between time points.



Motion Capture Results

demStick1 and demStick2

Figure: The latent space for the motion capture data without dynamics (*left*), with auto-regressive dynamics (*right*) based on an RBF kernel.



Motion Capture Results

demStick1 and demStick2

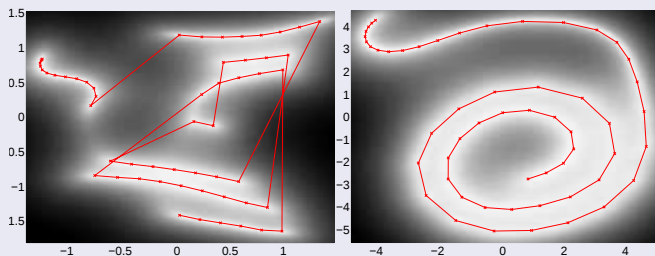


Figure: The latent space for the motion capture data without dynamics (*left*), with auto-regressive dynamics (*right*) based on an RBF kernel.

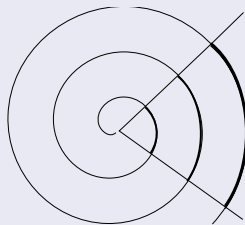




Regressive Dynamics

Inner Groove Distortion

- Autoregressive unimodal dynamics, $p(\mathbf{x}_t | \mathbf{x}_{t-1})$.
- Forces spiral visualisation.
- Poorer model due to inner groove distortion.



Regressive Dynamics

Direct use of Time Variable

- Instead of auto-regressive dynamics, consider regressive dynamics.
- Take $\mathbf{t}$ as an input, use a prior $p(\mathbf{X}|\mathbf{t})$.
- User a Gaussian process prior for $p(\mathbf{X}|\mathbf{t})$.
- Also allows us to consider variable sample rate data.



Motion Capture Results

demStick1, demStick2 and demStick5

Figure: The latent space for the motion capture data without dynamics (*left*), with auto-regressive dynamics (*middle*) and with regressive dynamics (*right*) based on an RBF kernel.



Motion Capture Results

demStick1, demStick2 and demStick5

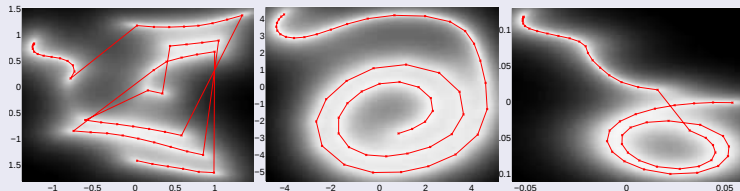


Figure: The latent space for the motion capture data without dynamics (*left*), with auto-regressive dynamics (*middle*) and with regressive dynamics (*right*) based on an RBF kernel.



Hierarchical GP-LVM

Stacking Gaussian Processes

- Regressive dynamics provides a simple hierarchy.
 - The input space of the GP is governed by another GP.
- By stacking GPs we can consider more complex hierarchies.
- Ideally we should marginalise latent spaces
 - In practice we seek MAP solutions.



Two Correlated Subjects

demHighFive1

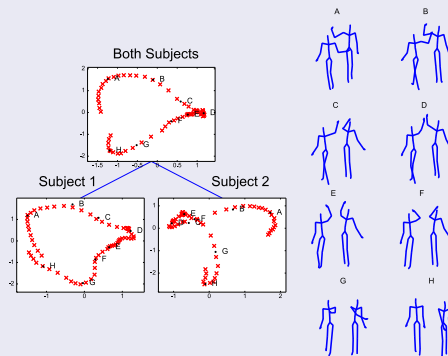


Figure: Hierarchical model of a 'high five'.



Within Subject Hierarchy

Decomposition of Body

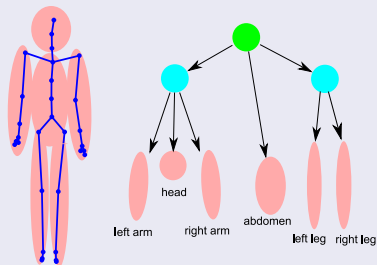


Figure: Decomposition of a subject.



Single Subject Run/Walk

demRunWalk1

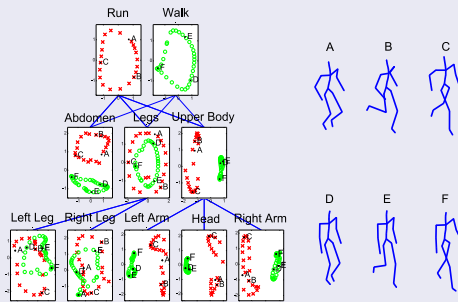


Figure: Hierarchical model of a walk and a run.



Summary

- GP-LVM is a Probabilistic Non-Linear Generalisation of PCA.
- Works Effectively as a Probabilistic Model in High Dimensional Spaces.
- Back constraints can be introduced to force local distance preservation.
- Dynamics can be introduced for modelling data with a temporal structure.



Initialisation

'Swiss Roll'

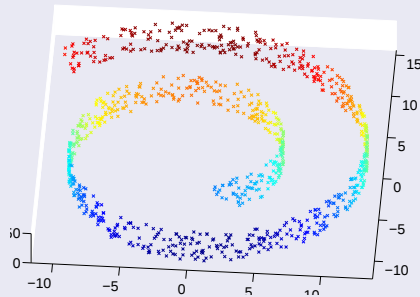


Figure: The 'Swiss Roll' data set is data in three dimensions that is inherently two dimensional.

Initialisation II

Quality of solution is Initialisation Dependent

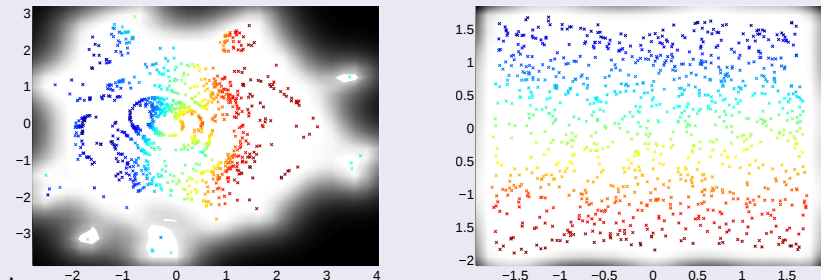


Figure: *Left:* Swiss roll solution initialised by PCA. *Right:* Swiss roll solution initialised by Isomap.



Linear Back Constraints I

- Special case of back constraints is a *linear* back constraint.

$$\mathbf{X} = \mathbf{YB}$$

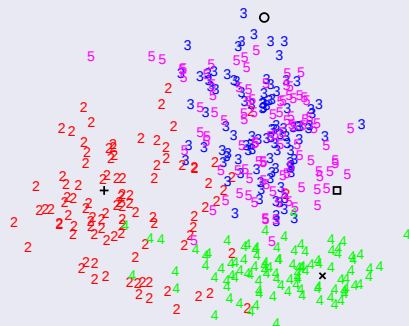
where $\mathbf{B} \in \mathbb{R}^{d \times q}$.

- Maximise the likelihood with respect to the projection matrix $\mathbf{B}$.
- Seems strange to sacrifice the 'non-linearity' of the model in this way.
- Motivate this through a digits data set.



Linear Back Constraints II

Digits Model with Linear Back Constraints



Linear Back Constraints III

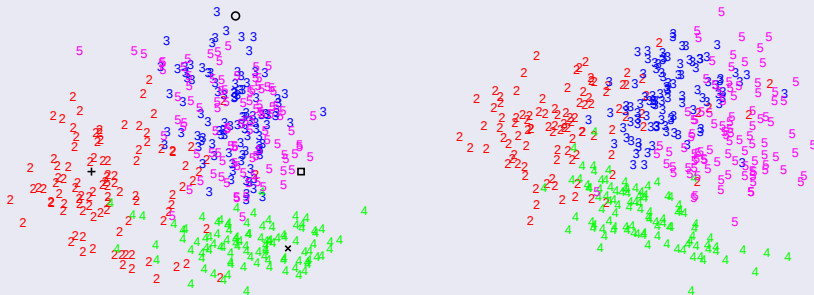


Figure: Linear projections from PCA (*left*) and linear back constrained GP-LVM (*right*)



1-Nearest Neighbour in $\mathbf{X}$

Comparison for increasing latent dimensionality

q	2	3	4
PCA Errors	131	115	47
Linear constrained GP-LVM Errors	79	60	39

c.f. 24 errors in data space



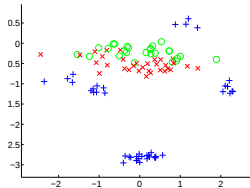
Oil Data I

Example Data set

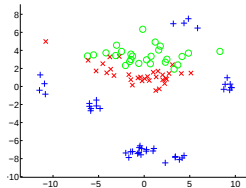
- Oil flow data [Bishop and James, 1993].
- Three phases of flow (stratified, annular, homogenous).
- Twelve measurement probes.
- 1000 data points.
- We sub-sampled to 100 data points
- Compare, with KPCA, MDS, Sammon mappings, PCA and GTM.



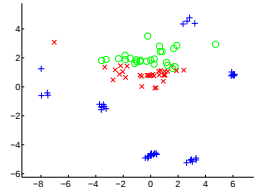
Oil Data II



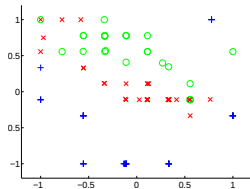
(a) PCA



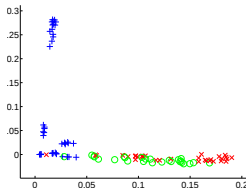
(b) Non-metric MDS



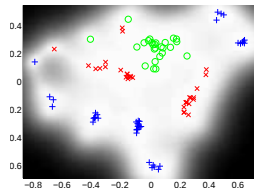
(c) Sammon Mapping



(d) GTM



(e) Kernel PCA



(f) GP-LVM



Oil Data III

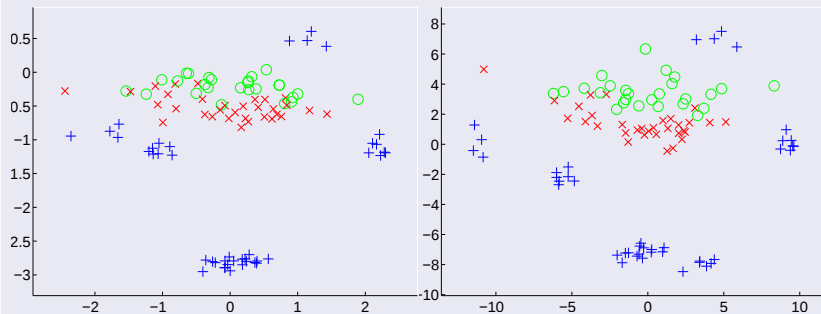


Figure: *Left* PCA, *right* Non-metric MDS



Oil Data III

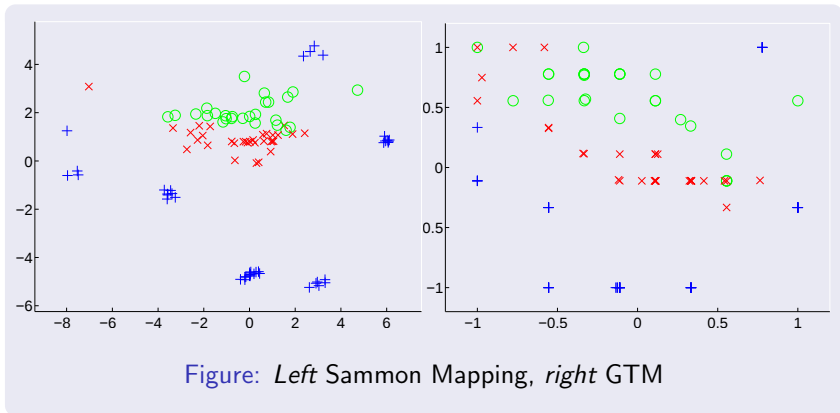


Figure: *Left* Sammon Mapping, *right* GTM



Oil Data III

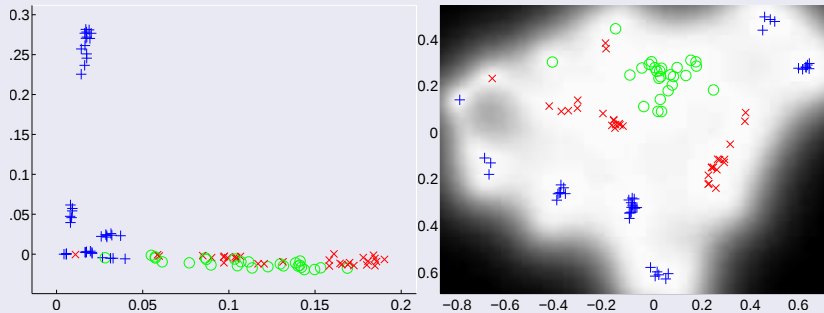


Figure: *Left* Kernel PCA, *right* GP-LVM



Oil Data IV

Nearest neighbour errors in $\mathbf{X}$ space

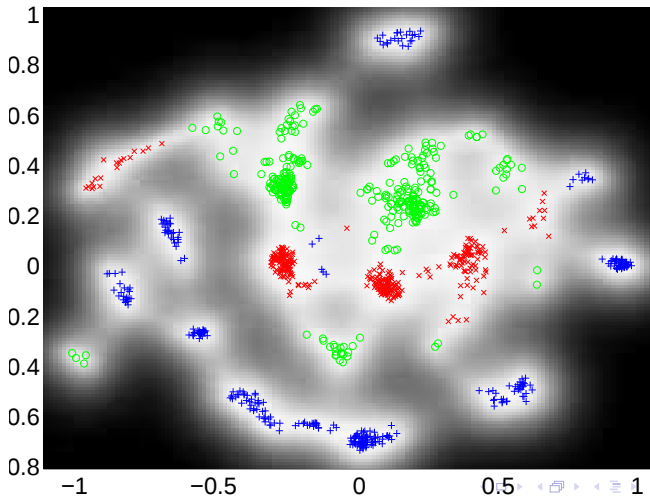
- Nearest neighbour classification in latent space.

Method	PCA	Non-metric MDS	Sammon Mapping
Errors	20	13	6
Method	GTM*	Kernel PCA*	GP-LVM
Errors	7	13	4

* These models require parameter selection.



Full Oil Data Set I



Full Oil Data Set II

Nearest Neighbour error in $\mathbf{X}$

- Nearest neighbour classification in latent space.

Method	PCA	GTM	GP-LVM
Errors	162	11	1

cf 2 errors in data space.



Vowel Data

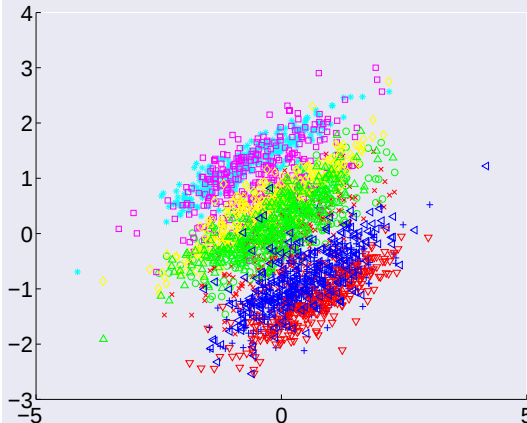
Vocal Joystick Data

- Vowel sounds from a vocal joystick system [Bilmes et al., 2006].
 - <http://ssli.ee.washington.edu/vj>
- Vowels are from a single speaker and represented as:
 - cepstral coefficients (12 dimensions) and
 - 'deltas' (further 12 dimensions).
- 2700 data points in total (300 for each vowel).



PCA Results

PCA (used as initialisation for GP-LVM)

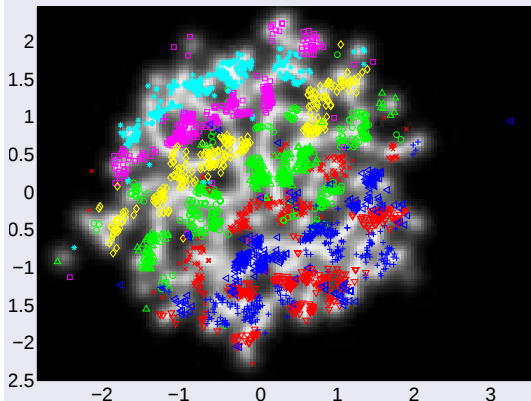


The different vowels are shown as follows: /a/ red cross /ae/ green circle /ao/ blue plus /e/ cyan asterisk /i/ pink square /ibar/ yellow diamond /o/ red down triangle /schwa/ green up triangle and /u/ blue left triangle.



GP-LVM Results

demVowels2

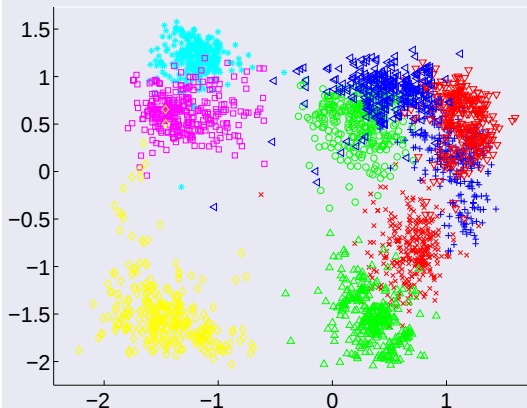


The different vowels are shown as follows: /a/ red cross /ae/ green circle /ao/ blue plus /e/ cyan asterix /i/ pink square /ibar/ yellow diamond /o/ red down triangle /schwa/ green up triangle and /u/ blue left triangle.



Isomap Results

demVowelsIsomap

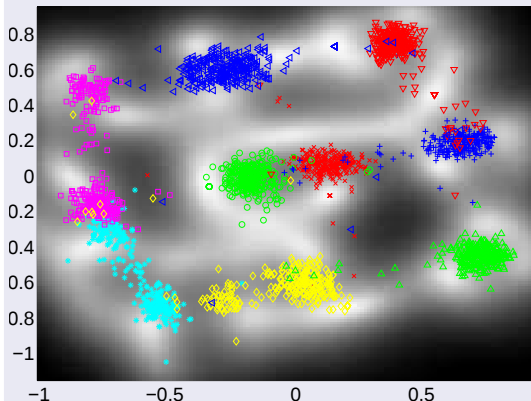


The different vowels are shown as follows: /a/ red cross /ae/ green circle /ao/ blue plus /e/ cyan asterix /i/ pink square /ibar/ yellow diamond /o/ red down triangle /schwa/ green up triangle and /u/ blue left triangle.



BC-GPLVM Results

demVowels3



The different vowels are shown as follows: /a/ red cross /ae/ green circle /ao/ blue plus /e/ cyan asterix /i/ pink square /ibar/ yellow diamond /o/ red down triangle /schwa/ green up triangle and /u/ blue left triangle.



1-Nearest Neighbour in $\mathbf{X}$

Comparison of the Approaches

- Nearest neighbour classification in latent space.

Method	GP-LVM	Isomap	BC-GP-LVM
Errors	226	458	155

cf 24 errors in data space.



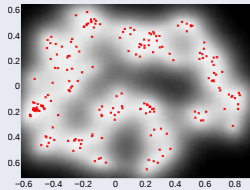
Robot SLAM I

Navigating by WiFi

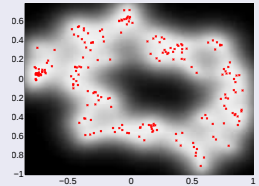
- Wireless access point signal strengths measured by robot moving around building.
 - 215 separate signal strength readings.
 - 30 separate access points.
- Robot moves in two dimensions so we expect data to be inherently 2-D.
- Learn GP-LVM, GP-LVM with Dynamics, back constrained GP-LVM and back constrained GP-LVM with dynamics.[Ferris et al., 2007]



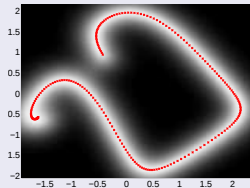
Robot SLAM II



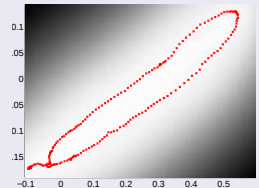
(a) Standard GP-LVM



(b) Standard GP-LVM



(c) Standard GP-LVM



(d) Standard GP-LVM

Non-Gaussian Data

Modelling Binary Data

- A common form of non-Gaussian data is *binary data*.
 - Can use Assumed Density Filtering to model binary data.
 - This can also easily be extended to the Expectation Propagation Algorithm Minka [2001].
- Practical consequences:
 - d times slower.
 - requires d times more storage.



Modelling Binary Twos

Cedar CD ROM digits

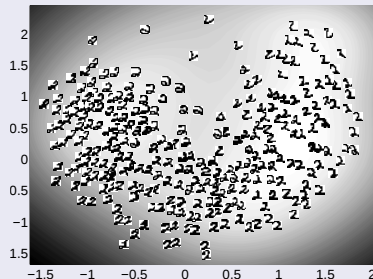
- We model 700 examples of binary 8×8 handwritten twos.
- Use a standard GP-LVM (a Gaussian noise assumption).
- Compare with ADF approximation for the Bernoulli noise model.



Two Results I



(a) Gaussian Noise Model



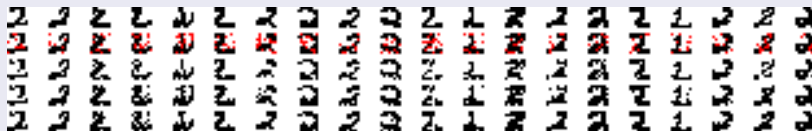
(b) Bernoulli Noise Model

Figure:



Two Results II

Reconstruction of Deleted Pixels



Reconstruction Method	Pixel Error Rate
GP-LVM Bernoulli noise	23.5%
GP-LVM Gaussian noise	35.9%
Missing pixels 'not ink'	51.5%



References

- J. Bilmes, J. Malkin, X. Li, S. Harada, K. Kilanski, K. Kirchhoff, R. Wright, A. Subramanya, J. Landay, P. Dowden, and H. Chizeck. The vocal joystick. In *Proceedings of the IEEE Conference on Acoustics, Speech and Signal Processing*. IEEE, May 2006. To appear.
- C. M. Bishop and G. D. James. Analysis of multiphase flows using dual-energy gamma densitometry and neural networks. *Nuclear Instruments and Methods in Physics Research*, A327:580–593, 1993.
- C. M. Bishop, M. Svensén, and C. K. I. Williams. GTM: the Generative Topographic Mapping. *Neural Computation*, 10(1):215–234, 1998.
- B. D. Ferris, D. Fox, and N. D. Lawrence. WiFi-SLAM using Gaussian process latent variable models. In M. M. Veloso, editor, *Proceedings of the 20th International Joint Conference on Artificial Intelligence (IJCAI 2007)*, pages 2480–2485, 2007.
- K. Grochow, S. L. Martin, A. Hertzmann, and Z. Popovic. Style-based inverse kinematics. In *ACM Transactions on Graphics (SIGGRAPH 2004)*, 2004.
- J. B. Kruskal. Multidimensional scaling by optimizing goodness-of-fit to a nonmetric hypothesis. *Psychometrika*, 29(1):1–28, 1964.
- N. D. Lawrence. Gaussian process models for visualisation of high dimensional data. In S. Thrun, L. Saul, and B. Schölkopf, editors, *Advances in Neural Information Processing Systems*, volume 16, pages 329–336, Cambridge, MA, 2004. MIT Press.
- N. D. Lawrence. Probabilistic non-linear principal component analysis with Gaussian process latent variable models. *Journal of Machine Learning Research*, 6:1783–1816, Nov 2005.
- N. D. Lawrence and J. Quiñero Candela. Local distance preservation in the GP-LVM through back constraints. In W. Cohen and A. Moore, editors, *Proceedings of the International Conference in Machine Learning*, volume 23, pages 513–520. Omnipress, 2006. ISBN 1-59593-383-2.
- D. Lowe and M. E. Tipping. Neuroscale: Novel topographic feature extraction with radial basis function networks. In M. C. Mozer, M. I. Jordan, and T. Petsche, editors, *Advances in Neural Information Processing Systems*, volume 9, pages 543–549, Cambridge, MA, 1997. MIT Press.
- D. J. C. MacKay. Bayesian neural networks and density networks. *Nuclear Instruments and Methods in Physics Research*, A, 354(1):73–80, 1995.
- K. V. Mardia, J. T. Kent, and J. M. Bibby. *Multivariate analysis*. Academic Press, London, 1979.

