

Correlated prior and time series

Václav Šmíd

May 4, 2020

Transformation of Gaussian Distribution

Consider normal distribution

$$e \sim \mathcal{N}(0, 1) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}e^2\right),$$

the distribution of

$$x = f(e) = \mu + \sigma e \quad \Leftrightarrow \quad e = g(x) = \frac{x - \mu}{\sigma}$$

is:

$$p(x) = p_e\left(\frac{x - \mu}{\sigma}\right) |g'(x)| = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{x - \mu}{\sigma}\right)^2\right) \frac{1}{\sigma} = \mathcal{N}(\mu, \sigma)$$

Multivariate Normal

In multivariate case:

$$\begin{aligned} \mathbf{e} &\sim \mathcal{N}(\mathbf{0}, I), & \Leftrightarrow & & e_i &\sim \mathcal{N}(0, 1), \quad i = 1, \dots, n \\ &\propto \exp\left(-\frac{1}{2}\mathbf{e}^T \mathbf{e}\right) & & & \prod_{i=1}^n p(e_i) &\propto \prod_{i=1}^n \exp\left(-\frac{1}{2}e_i^2\right) \end{aligned}$$

either with transformation

$$\begin{aligned} \mathbf{x} &= \boldsymbol{\mu} + S\mathbf{e} & \Leftrightarrow & & \mathbf{e} &= S^{-1}(\mathbf{x} - \boldsymbol{\mu}) \\ & & & & \mathbf{e}^T \mathbf{e} &= (\mathbf{x} - \boldsymbol{\mu})^T S^{-T} S^{-1}(\mathbf{x} - \boldsymbol{\mu}) \end{aligned}$$

$$\mathbf{x} \sim \frac{1}{\sqrt{2\pi}^n} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T S^{-T} S^{-1}(\mathbf{x} - \boldsymbol{\mu})\right) |S^{-1}|$$

with $\Sigma = S S^T$ and $\det(\Sigma) = |S|^2$

Least squares – different way I

Consider regression problem:

$$\begin{aligned} \mathbf{y} &= X\boldsymbol{\theta} + \sigma \mathbf{e} & \mathbf{e} &\sim \mathcal{N}(\mathbf{0}, I), \\ p(\mathbf{y}|X, \mathbf{x}) &= \mathcal{N}(X\boldsymbol{\theta}, \sigma^2 I) \\ \boldsymbol{\theta} &= \alpha \mathbf{e}_{\theta} & \mathbf{e}_{\theta} &\sim \mathcal{N}(\mathbf{0}, I) \\ p(\boldsymbol{\theta}) &= \mathcal{N}(\mathbf{0}, I), \end{aligned}$$

Build joint distribution

$$\begin{aligned} \begin{bmatrix} \mathbf{y} \\ \boldsymbol{\theta} \end{bmatrix} &= \begin{bmatrix} \sigma I & X \\ 0 & I \end{bmatrix} \begin{bmatrix} \mathbf{e} \\ \mathbf{e}_{\theta} \end{bmatrix} \\ p(\mathbf{y}, \boldsymbol{\theta}) &= \mathcal{N}\left(\mathbf{0}, \begin{bmatrix} \sigma I & X \\ 0 & I \end{bmatrix} \begin{bmatrix} \sigma I & 0 \\ X^T & I \end{bmatrix}\right) = \mathcal{N}\left(\mathbf{0}, \begin{bmatrix} \sigma^2 I + XX^T & X^T \\ X & I \end{bmatrix}\right) \end{aligned}$$

Least squares – different way II

Joint

$$p(\mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}\left(\mathbf{0}, \begin{bmatrix} \sigma^2 I + \mathbf{X}\mathbf{X}^T & \mathbf{X} \\ \mathbf{X}^T & I \end{bmatrix}\right)$$

with decomposition

$$\begin{aligned} p(\mathbf{y}) &= \mathcal{N}(\mathbf{0}, \sigma^2 I + \mathbf{X}\mathbf{X}^T), \\ p(\boldsymbol{\theta}|\mathbf{y}) &= \mathcal{N}(\hat{\boldsymbol{\theta}}, \Sigma_{\boldsymbol{\theta}}) \\ \hat{\boldsymbol{\theta}} &= \mathbf{X}^T(\sigma^2 I + \mathbf{X}\mathbf{X}^T)^{-1}\mathbf{y}, \\ \Sigma_{\boldsymbol{\theta}} &= I - \mathbf{X}^T(\sigma^2 I + \mathbf{X}\mathbf{X}^T)^{-1}\mathbf{X} \end{aligned}$$

decomposition into:

1. generative model,
2. different form of least squares. (inverse in dimension of $\mathbf{y}$!)

Generative models of 2nd order polynomials

we have a least squares model:

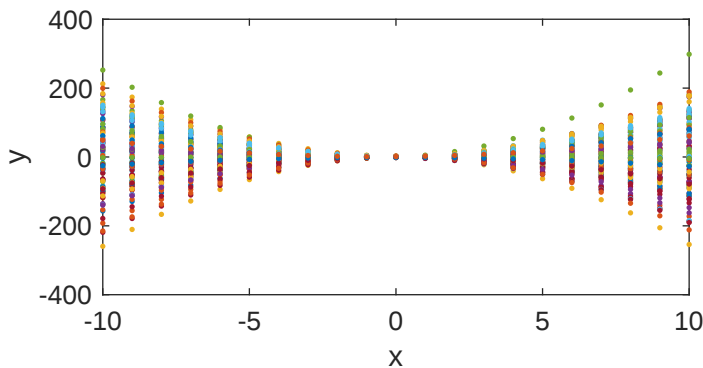
$$p(\mathbf{y}|\mathbf{M}, \boldsymbol{\theta}) = \mathcal{N}(\mathbf{M}\boldsymbol{\theta}, \sigma^2 \mathbf{I}),$$

$$p(\boldsymbol{\theta}) = \mathcal{N}(\mathbf{0}, \mathbf{I}),$$

$$p(\mathbf{y}) = \int p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta} \quad p(\mathbf{y}) = \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I} + \mathbf{M}\mathbf{M}^T),$$

Consider samples:

- ▶ $\boldsymbol{\theta}^{(i)}$ from $p(\boldsymbol{\theta})$,
- ▶ $\mathbf{y}^{(i)}$ from $p(\mathbf{y}|\mathbf{M}, \boldsymbol{\theta}^{(i)})$ to generate samples from the joint $(\mathbf{y}^{(i)}, \mathbf{x}^{(i)})$



Least squares – prediction

Provide predictions $\bar{\mathbf{y}}$ at given $\bar{\mathbf{x}}$ and known $(\mathbf{x}, \mathbf{y})$

$$p(\bar{\mathbf{y}}|\bar{\mathbf{x}}, \mathbf{x}, \mathbf{y}) = \int p(\bar{\mathbf{y}}|\bar{\mathbf{x}}, \theta) p(\theta|\mathbf{x}, \mathbf{y}) d\theta$$

Alternative:

$$\begin{bmatrix} \mathbf{y} \\ \bar{\mathbf{y}} \\ \boldsymbol{\theta} \end{bmatrix} = \begin{bmatrix} \sigma I & 0 & X \\ 0 & \sigma I & \bar{X} \\ 0 & 0 & \alpha I \end{bmatrix} \begin{bmatrix} \mathbf{e} \\ \bar{\mathbf{e}} \\ \mathbf{e}_\theta \end{bmatrix}$$

$$p(\mathbf{y}, \bar{\mathbf{y}}, \boldsymbol{\theta}) = \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \sigma^2 I + XX^T & X\bar{X}^T & X \\ \bar{X}X^T & \sigma^2 I + \bar{X}\bar{X}^T & \bar{X} \\ X^T & \bar{X}^T & I \end{bmatrix} \right)$$

$$p(\mathbf{y}, \bar{\mathbf{y}}) = \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \sigma^2 I + XX^T & X\bar{X}^T \\ \bar{X}X^T & \sigma^2 I + \bar{X}\bar{X}^T \end{bmatrix} \right)$$

We do not need parameters to predict! (we need all data!)

All we need is kernel $K_{\mathbf{x}\mathbf{x}} = XX^T$, $K_{\mathbf{x}\bar{\mathbf{x}}}$, $K_{\bar{\mathbf{x}}\bar{\mathbf{x}}}$

Gaussian process (non-parametric model)

Definition

[Rasmussen, 2006] A Gaussian process is a collection of random variables, any finite number of which have a joint Gaussian distribution. Value of a process at location $\mathbf{x}$ is

$$f(\mathbf{x}) \sim \mathcal{GP}(m(\mathbf{x}), k(\mathbf{x}, \mathbf{x}'))$$

$$m(\mathbf{x}) = \mathbb{E}(f(\mathbf{x})),$$

$$k(\mathbf{x}, \mathbf{x}') = \mathbb{E}((f(\mathbf{x}) - m(\mathbf{x}))(f(\mathbf{x}') - m(\mathbf{x}'))),$$

where $m(\mathbf{x})$ is the mean and $k(\mathbf{x}, \mathbf{x}')$ is a correlation function. $\mathbf{x}'$ is any other point in the collection.

- ▶ in polynomial regression $k(\mathbf{x}, \mathbf{x}') = \mathbf{x}^T \mathbf{x}'$
- ▶ “better” kernels:

$$k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{1}{2l}|\mathbf{x} - \mathbf{x}'|^2\right)$$

$$k(\mathbf{x}, \mathbf{x}') = (1 + |\mathbf{x} - \mathbf{x}'|^2)^{-\alpha}$$

with hyper-parameters $l, \alpha \dots$

Example

Joint distribution of two values $y = f(x) + e$

$$p(f(x), \overline{f(x)}) = \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} k(x, x) & k(x, \bar{x}) \\ k(\bar{x}, x) & k(\bar{x}, \bar{x}) \end{bmatrix} \right)$$

$$= \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} 1 & k(x, \bar{x}) \\ k(\bar{x}, x) & 1 \end{bmatrix} \right)$$

$$p(y, \bar{y}) = \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \sigma + 1 & k(x, \bar{x}) \\ k(\bar{x}, x) & 1 + \sigma \end{bmatrix} \right)$$

$$p(\bar{y}|y) = \mathcal{N} (k(x, x')f(x), (1 + \sigma - k(x, x')^2/(1 + \sigma))) ,$$

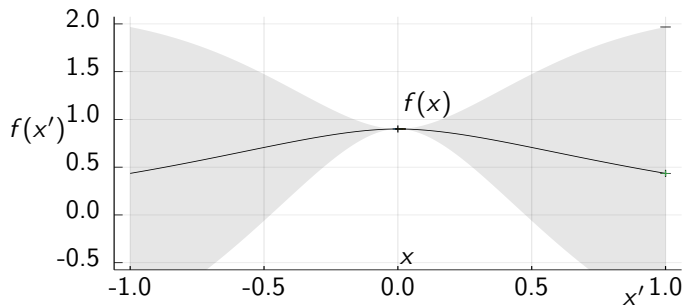
$$p(\bar{y}) = \mathcal{N} (0, 1 + \sigma) ,$$

1D Example

Joint distribution of two values $y = f(x) + e$,

$$p(f(\bar{x})|y) = \mathcal{N}(k(x, \bar{x})y, (1 + \sigma - k(x, \bar{x})^2)),$$

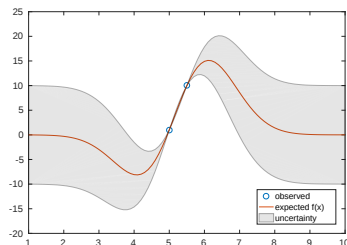
$$p(\bar{y}|y) = \mathcal{N}(k(x, \bar{x})y, (1 + \sigma - k(x, \bar{x})^2/(1 + \sigma))),$$



1D Example – adding points

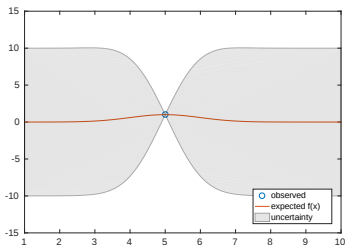
Joint distribution of two values $y = f(x) + e$

$$p(\bar{y}|\bar{y}) = \mathcal{N} \left(K(X, \bar{X})y, K(\bar{X}, \bar{X}) + \sigma I - K(\bar{X}, X) (\sigma I + K(X, X))^{-1} K(X, \bar{X}) \right),$$

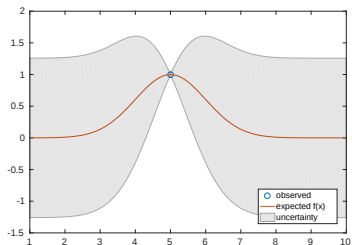


Sensitivity to kernel parameters

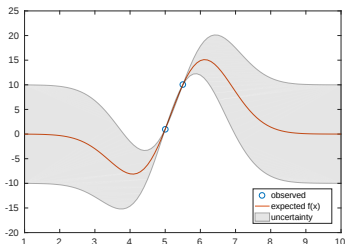
$$\ell = 1, \sigma = 0.1$$



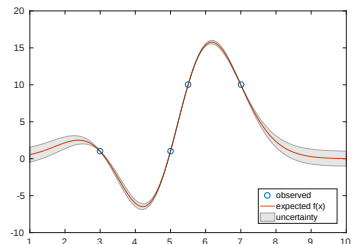
$$\ell = 0.1, \sigma = 0.1$$



$$\ell = 1, \sigma = 0.1$$

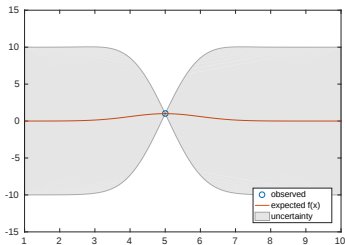


$$\ell = 1, \sigma = 0.1$$

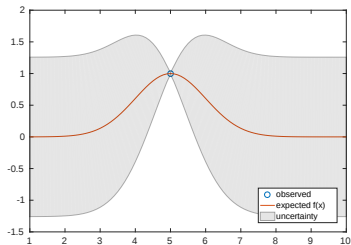


Sensitivity to kernel parameters (function smoothness)

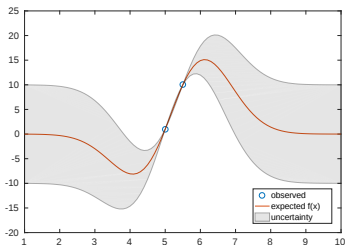
$\ell = 1, \sigma = 0.1$



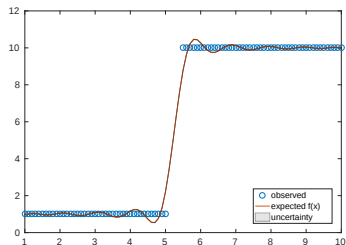
$\ell = 0.1, \sigma = 0.1$



$\ell = 1, \sigma = 0.1$



$\ell = 1, \sigma = 0.1$



Learning kernel parameters

Technically we have:

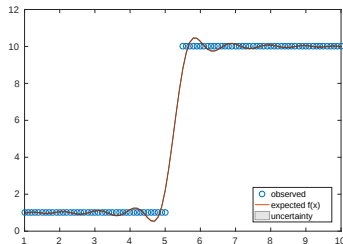
$$p(\mathbf{y}|\theta) = \mathcal{N}(\mathbf{0}, \sigma^2 I + K_\theta(X, X^T))$$

we can learn

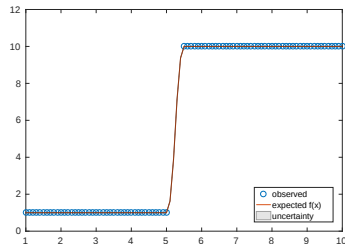
$$\theta = \arg \max_{\theta} p(\mathbf{y}|\theta)$$

which implies minimization of $|\sigma^2 I + K_\theta(X, X^T)|$. Or VB, EM, MCMC...

$$\ell = 1, \sigma = 0.1$$



$$\hat{\ell} = 0.55, \hat{\sigma} = 1e-6$$



Non-parametric models

- ▶ Data interpolators
- ▶ Nice analytical properties (Gaussian)
- ▶ Algebra on kernels: $K(X, \bar{X}) = K_1(X, \bar{X}) + K_2(X, \bar{X})$
- ▶ Extensions
 - ▶ Student-t processes, etc.
 - ▶ Non-stationary processes
 - ▶ Hierarchical processes
- ▶ Book:
 - ▶ Carl Edward Rasmussen and Christopher K.I. Williams. Gaussian processes for machine learning. Vol. 2. No. 3. Cambridge, MA: MIT press, 2006.

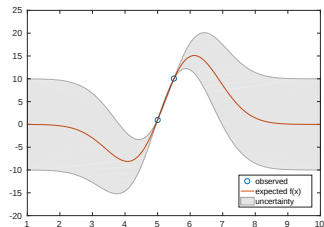
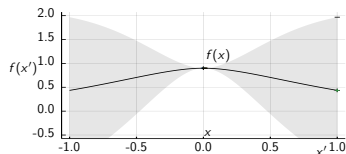
Assignment

Derive a Gaussian process approximation for

- ▶ 1 point of observation
- ▶ 2 points of observation

Choose freely

- ▶ kernel,
- ▶ prediction range...
- ▶ y



Back to State-space model

Consider state space

$$x_n = ax_{n-1} + \sqrt{q}v_n,$$

$$y_n = cx_n + \sqrt{r}w_n$$

as a regression problem: $\mathbf{y} = [y_1, y_2 \dots y_n]$ and $\mathbf{x} = [x_1, x_2 \dots, x_n]$

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \end{bmatrix} = \begin{bmatrix} c & & \\ & c & \\ & & c \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix} + \sqrt{r} \begin{bmatrix} e_1 \\ e_2 \\ \vdots \end{bmatrix} \quad p(\mathbf{y}|\mathbf{x}) = \mathcal{N}(\mathbb{C}\mathbf{x}, rI)$$

$$\begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix} = \begin{bmatrix} 0 & a & 0 & \dots \\ 0 & 0 & a & \dots \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix} + \sqrt{q} \begin{bmatrix} e_1 \\ e_2 \\ \vdots \end{bmatrix} \quad p(\mathbf{x}) = ?$$

Correlated model

Prior on $\mathbf{x}$:

$$\begin{aligned} p(x_1) &= \mathcal{N}(\hat{x}_1, q_1) \Leftrightarrow x_1 = \hat{x}_1 + \sqrt{p_1}e_1, & e_1 &\sim \mathcal{N}(0, 1) \\ p(x_i|x_{i-1}) &= \mathcal{N}(ax_{i-1}, q) & x_i &= ax_{i-1} + \sqrt{q}e_i, & i &= 2, \dots, n \end{aligned}$$

$$\begin{aligned} p_1 e_1 &= x_1 - \hat{x}_1 \\ q e_2 &= x_2 - ax_1 \\ q e_3 &= x_3 - ax_2 \\ &\vdots \end{aligned} \Leftrightarrow \begin{bmatrix} \sqrt{p_1} & & & \\ & \sqrt{q} & & \\ & & \sqrt{q} & \\ & & & \ddots \end{bmatrix} \begin{bmatrix} e_1 \\ e_2 \\ e_3 \\ \vdots \end{bmatrix} = \begin{bmatrix} 1 & & & \\ -a & 1 & & \\ & -a & 1 & \\ & & & \ddots \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \end{bmatrix} -$$

Implies

$$\Sigma_0 \mathbf{e} = \mathbb{A} \mathbf{x} - \mu_0 \Rightarrow \mathbb{A}^{-1}(\Sigma_0 \mathbf{e} + \mu_0) = \mathbf{x} \Rightarrow p(\mathbf{x}) = \mathcal{N}(\mathbb{A}^{-1} \hat{\mu}_0, \mathbb{A}^{-1} \Sigma_0 \mathbb{A}^{-1})$$

Higher order correlation

Consider vector

$$\mathbf{x} = [x_1, x_2, \dots, x_n]$$

with prior probabilities

$$x_1 = \hat{x}_1 + \sigma_1 \mathbf{e}_1, \quad \mathbf{e}_1 \sim \mathcal{N}(0, 1)$$

$$x_2 = ax_1 + \hat{x}_2 + \sigma_2 \mathbf{e}_1, \quad \mathbf{e}_1 \sim \mathcal{N}(0, 1)$$

$$x_i = ax_{i-1} + bx_{i-2} + \sigma \mathbf{e}_i, \quad i = 2, \dots, n$$

$$\begin{bmatrix} \sigma_1 & & & \\ & \sigma_2 & & \\ & & \sigma & \\ & & & \ddots \end{bmatrix} \begin{bmatrix} \mathbf{e}_1 \\ \mathbf{e}_2 \\ \mathbf{e}_3 \\ \vdots \end{bmatrix} = \begin{bmatrix} 1 & & & \\ -a & 1 & & \\ -b & -a & 1 & \\ & & & \ddots \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \end{bmatrix} - \begin{bmatrix} \hat{x}_1 \\ \hat{x}_2 \\ 0 \\ 0 \end{bmatrix}$$

Implies

$$\Sigma_0 \mathbf{e} = A\mathbf{x} - \mu_0 \quad \Rightarrow \quad A^{-1}(\Sigma_0 \mathbf{e} + \mu_0) = \mathbf{x} \quad \Rightarrow \quad \mathbf{x} = \mathcal{N}(A^{-1}\hat{\mu}_0, A^{-1}\Sigma_0 A^{-T})$$

Special case #1

State Space model for $a = 0, c = 1$

$$\begin{aligned}x_n &= \sqrt{q}v_n, \\y_n &= x_n + \sqrt{r}w_n\end{aligned}$$

We have

$$\begin{aligned}p(\mathbf{y}|\mathbf{x}) &= \mathcal{N}(\mathbf{x}, rI), \\p(\mathbf{x}) &= \mathcal{N}(qI)\end{aligned} \quad \Leftrightarrow \quad \begin{aligned}p(y_i|x_i) &= \mathcal{N}(x_i, r), \\p(x_i) &= \mathcal{N}(q)\end{aligned}$$

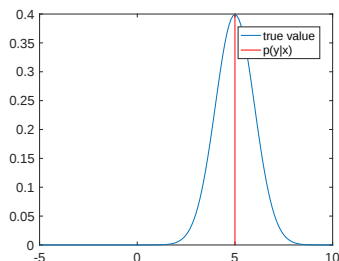
Lecture #2. (Normal observation y of constant parameter m)

Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

With observation y .



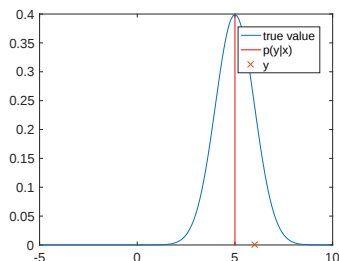
$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$

Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

With observation y .



$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$

Bayes rule for Normal

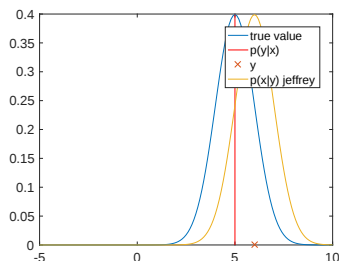
Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

With observation y .

Posterior with Jeffrey's prior

$$p(x|y) \propto p(x|y) \\ = \mathcal{N}(y, r)$$



$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$

Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

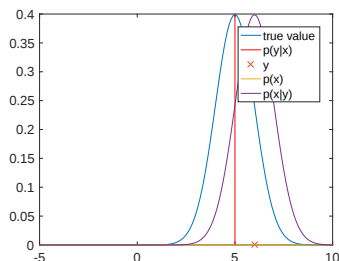
With observation y .

Posterior with Jeffrey's prior

$$p(x|y) \propto p(x|y) \\ = \mathcal{N}(y, r)$$

with normal prior $p(x) = \mathcal{N}(0, q)$

$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$



Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

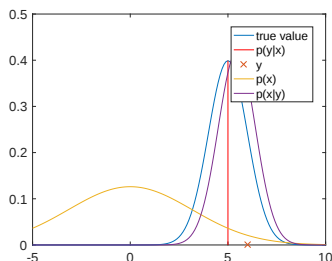
With observation y .

Posterior with Jeffrey's prior

$$p(x|y) \propto p(x|y) \\ = \mathcal{N}(y, r)$$

with normal prior $p(x) = \mathcal{N}(0, q)$

$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$



Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

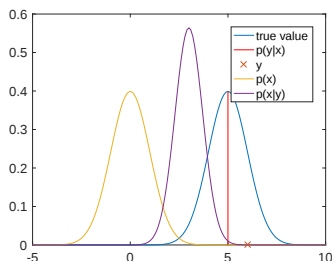
With observation y .

Posterior with Jeffrey's prior

$$p(x|y) \propto p(x|y) \\ = \mathcal{N}(y, r)$$

with normal prior $p(x) = \mathcal{N}(0, q)$

$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$



Bayes rule for Normal

Likelihood of unknown x :

$$p(y|x) = \mathcal{N}(x, r) \\ \propto \exp\left(-\frac{1}{2}\left(\frac{x-y}{r}\right)^2\right)$$

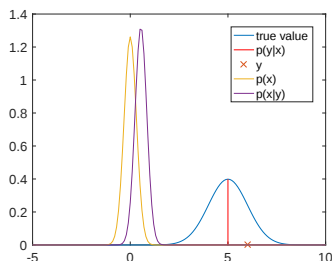
With observation y .

Posterior with Jeffrey's prior

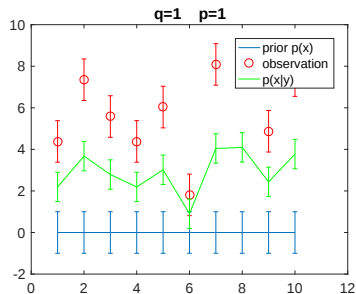
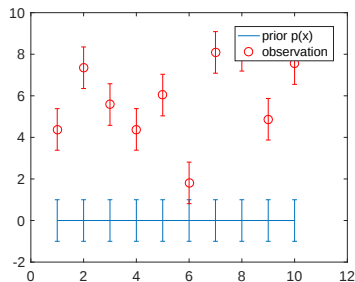
$$p(x|y) \propto p(x|y) \\ = \mathcal{N}(y, r)$$

with normal prior $p(x) = \mathcal{N}(0, q)$

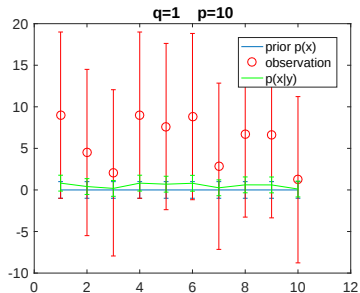
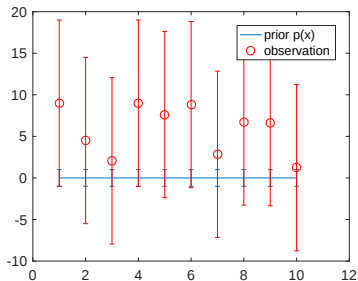
$$p(x|y, \phi) = \mathcal{N}\left(\left(\frac{r+q}{qr}\right)^{-1} \frac{0r+yq}{qr}, \left(\frac{r+q}{qr}\right)^{-1}\right) = \mathcal{N}\left(\left(\frac{qy}{r+q}\right), \frac{qr}{q+r}\right)$$



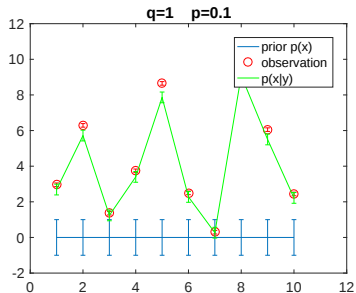
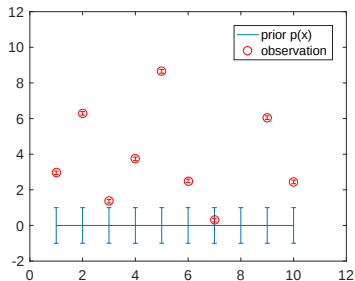
Independent Kalman



Independent Kalman



Independent Kalman



Special case #2

State Space model for $a = 1, q = 0; c = 1$

$$x_n = x_{n-1}$$

$$y_n = x_n + \sqrt{r}w_n$$

We have

$$\begin{aligned} p(\mathbf{x}|\mathbf{y}) &\propto p(y_1|x_1)p(y_2|x_2)\dots p(x_1)p(x_2)\dots \\ &\propto p(y_1|x_1)p(y_2|x_2)\dots p(x_1)\delta(x_2 - x_1)\dots \\ &\propto p(y_1|x_1)p(y_2|x_1)\dots p(x_1) \end{aligned}$$

Multiple observation of a single unknown x_1 .

Random walk prior

independent prior

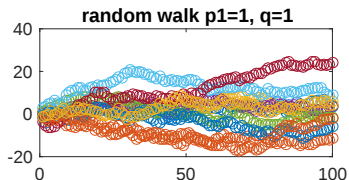
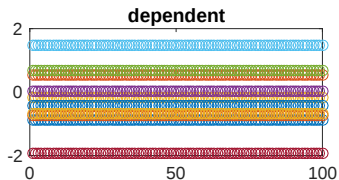
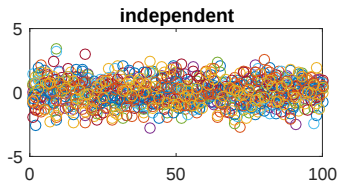
$$p(x_i) = \mathcal{N}(0, q)$$

dependent prior

$$p(x_i) = \delta(x_{i-1})$$

Random walk

$$p(x_i) = x_{i-1} + \mathcal{N}(0, q)$$



Inference with random walk

