

Plan for today

We talk about:

- Markov property, conditional independence, precision matrix
- Gaussian Markov random fields
- Integrated nested Laplace approximations

Numerical algorithms for sparse matrices: scaling properties

Cholesky factorization of “sparse” $\mathbf{Q}$

- Time: $\mathcal{O}(n)$

This is to be compared with general $\mathcal{O}(n^3)$ algorithm for the Cholesky factorization of a dense matrix.

Numerical algorithms for sparse matrices: scaling properties

Cholesky factorization of “sparse” $\mathbf{Q}$

- Time: $\mathcal{O}(n)$
- Space: $\mathcal{O}(n^{3/2})$

This is to be compared with general $\mathcal{O}(n^3)$ algorithm for the Cholesky factorization of a dense matrix.

Numerical algorithms for sparse matrices: scaling properties

Cholesky factorization of “sparse” $\mathbf{Q}$

- Time: $\mathcal{O}(n)$
- Space: $\mathcal{O}(n^{3/2})$
- Space-time: $\mathcal{O}(n^2)$

This is to be compared with general $\mathcal{O}(n^3)$ algorithm for the Cholesky factorization of a dense matrix.

Gaussian Markov random field (GMRF)

A **GMRF** is a simple construct

- A normal distributed random vector

$$\mathbf{x} = (x_1, \dots, x_n)^T$$

- Additional Markov properties:

$$x_i \perp x_j \mid \mathbf{x}_{-ij}$$

x_i and x_j are conditional independent (CI).

Conditional independence

If $x_i \perp x_j \mid \mathbf{x}_{-ij}$ for a set of $\{i, j\}$, then we need to constrain the parametrisation of the GMRF.

Conditional independence

If $x_i \perp x_j \mid \mathbf{x}_{-ij}$ for a set of $\{i, j\}$, then we need to constrain the parametrisation of the GMRF.

- Covariance matrix: difficult
- Precision matrix: easy

Conditional independence and the precision matrix

The density of a zero mean Gaussian

$$\pi(\mathbf{x}) \propto |\mathbf{Q}|^{1/2} \exp \left(-\frac{1}{2} \mathbf{x}^T \mathbf{Q} \mathbf{x} \right)$$

Conditional independence and the precision matrix

The density of a zero mean Gaussian

$$\pi(\mathbf{x}) \propto |\mathbf{Q}|^{1/2} \exp\left(-\frac{1}{2}\mathbf{x}^T \mathbf{Q} \mathbf{x}\right)$$

Constraining the parametrisation to obey CI properties

Theorem

$$x_i \perp x_j \mid \mathbf{x}_{-ij} \iff Q_{ij} = 0$$

Use a (undirected) graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ to represent the CI properties,

$\mathcal{V}$ Vertices: $1, 2, \dots, n$.

$\mathcal{E}$ Edges $\{i, j\}$

Use a (undirected) graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ to represent the CI properties,

$\mathcal{V}$ Vertices: $1, 2, \dots, n$.

$\mathcal{E}$ Edges $\{i, j\}$

- No edge between i and j if $x_i \perp x_j \mid \mathbf{x}_{-ij}$.
- An edge between i and j if $x_i \not\perp x_j \mid \mathbf{x}_{-ij}$.

Definition of a GMRF

A random vector $\mathbf{x} = (x_1, \dots, x_n)^T$ is called a GMRF wrt the graph $\mathcal{G} = (\mathcal{V} = \{1, \dots, n\}, \mathcal{E})$ with mean $\boldsymbol{\mu}$ and precision matrix $\mathbf{Q} > 0$, iff its density has the form

$$\pi(\mathbf{x}) = (2\pi)^{-n/2} |\mathbf{Q}|^{1/2} \exp \left(-\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})^T \mathbf{Q} (\mathbf{x} - \boldsymbol{\mu}) \right)$$

and

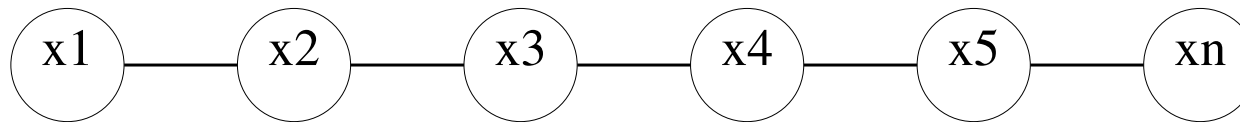
$$Q_{ij} \neq 0 \iff \{i, j\} \in \mathcal{E} \quad \text{for all } i \neq j.$$

Simple example of a GMRF

Auto-regressive process of order 1

$$x_t \mid x_{t-1}, \dots, x_1 \sim \mathcal{N}(\phi x_{t-1}, 1), \quad t = 2, \dots, n$$

and $x_1 \sim \mathcal{N}(0, (1 - \phi^2)^{-1})$.



Tridiagonal precision matrix

$$\mathbf{Q} = \begin{pmatrix} 1 & -\phi & & & \\ -\phi & 1 + \phi^2 & -\phi & & \\ & \ddots & \ddots & \ddots & \\ & & -\phi & 1 + \phi^2 & -\phi \\ & & & -\phi & 1 \end{pmatrix}$$

Interpretation of elements of $\mathbf{Q}$

Let $\mathbf{x}$ be a GMRF wrt $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ with mean $\boldsymbol{\mu}$ and precision matrix $\mathbf{Q} \succ 0$, then

$$\begin{aligned} \mathbb{E}(x_i \mid \mathbf{x}_{-i}) &= \mu_i - \frac{1}{Q_{ii}} \sum_{j:j \sim i} Q_{ij}(x_j - \mu_j), \\ \text{Var}(x_i \mid \mathbf{x}_{-i}) &= 1/Q_{ii} \quad \text{and} \\ \text{Corr}(x_i, x_j \mid \mathbf{x}_{-ij}) &= -\frac{Q_{ij}}{\sqrt{Q_{ii} Q_{jj}}}, \quad i \neq j. \end{aligned}$$

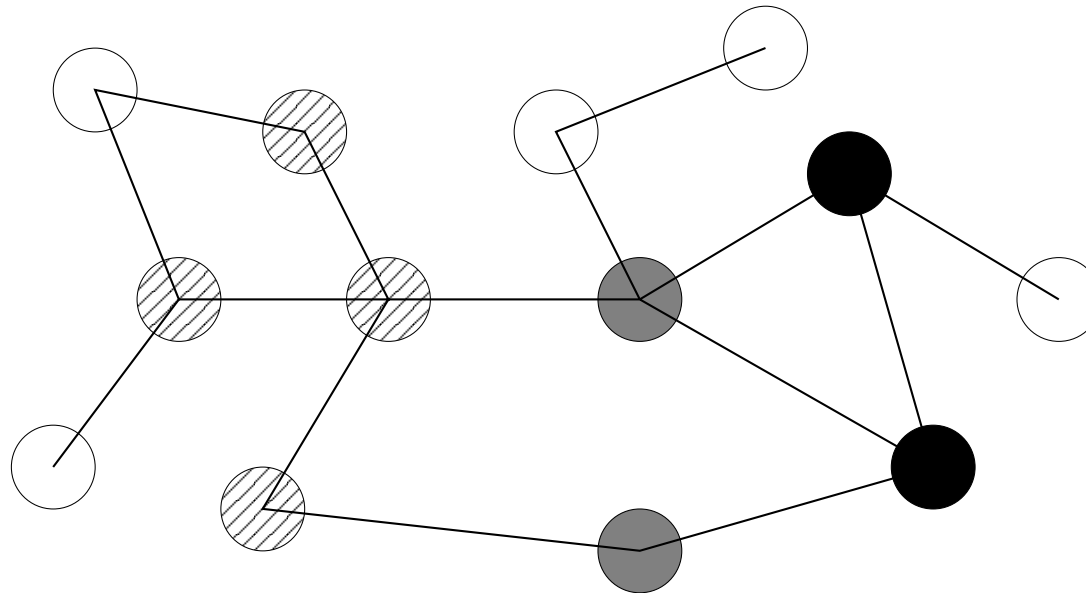
Markov properties

Let $\mathbf{x}$ be a GMRF wrt $\mathcal{G} = (\mathcal{V}, \mathcal{E})$. Then the following are equivalent.

The global Markov property:

$$\mathbf{x}_A \perp \mathbf{x}_B \mid \mathbf{x}_C$$

for all disjoint sets A , B and C where C separates A and B , and A and B are non-empty.



Hierarchical Bayesian models

Hierarchical models are an extremely useful tool in Bayesian model building.

Three parts:

- **Observation model $y|x, \theta$** : Encodes information about observed data.
- **The latent model $x|\theta$** : The unobserved process.
- **Hyperpriors for θ** : Models for all of the parameters in the observation and latent processes.

Bayesian computing

The posterior distribution is given by

$$p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y}) \propto p(\mathbf{y} | \mathbf{x}, \boldsymbol{\theta}) p(\mathbf{x} | \boldsymbol{\theta}) p(\boldsymbol{\theta})$$

We are interested in the **posterior marginal quantities** like $p(x_i | \mathbf{y})$ and $p(\theta_j | \mathbf{y})$. We have discussed how to do Bayesian inference using MCMC, for example Stan, and discussed pros and cons.

Approximate inference

Except for the (trivial) cases when everything can be computed exactly (maybe up to very small integration error), we can never do exact inference.

- **Integrated nested Laplace approximations** (INLA) are a promising alternative to inference via MCMC in latent Gaussian models (Rue et al, 2009, JRSS-B).
- The methodology is particularly attractive if the latent Gaussian model is a **Gaussian Markov random field** (GMRF) (Rue and Held, 2005).

Model framework

Latent Gaussian models

$$\mathbf{y}|\mathbf{x}, \boldsymbol{\theta} \sim \prod_i \pi(y_i|\eta_i, \boldsymbol{\theta})$$

$$\mathbf{x}|\boldsymbol{\theta} \sim \pi(\mathbf{x}|\boldsymbol{\theta}) \sim \mathcal{N}(\mathbf{0}, \mathbf{Q}(\boldsymbol{\theta})^{-1})$$

Gaussian!

$$\boldsymbol{\theta} \sim \pi(\boldsymbol{\theta})$$

Not Gaussian

where the precision matrix $\mathbf{Q}(\boldsymbol{\theta})$ is sparse, coming from a **Gaussian Markov random fields** (GMRFs).

The sparseness can be exploited for very quick computations for the Gaussian part of the model through numerical algorithms for sparse matrices.

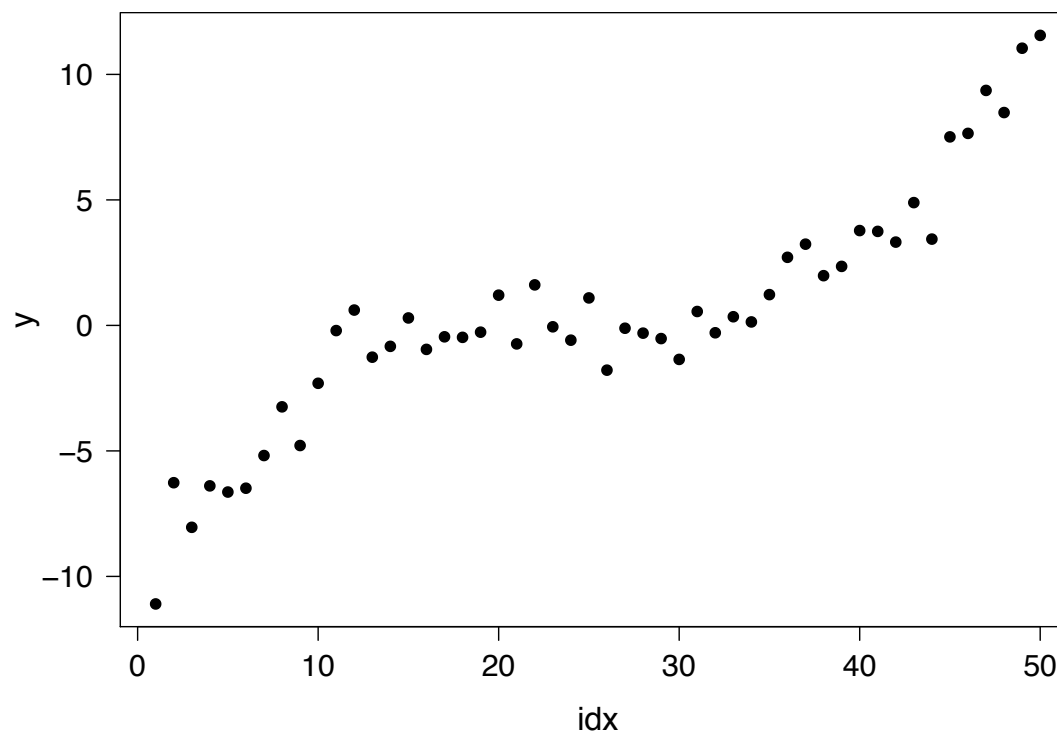
How does INLA work?

Observations

$$y_i = m(i) + \epsilon_i, \quad i = 1, \dots, n$$

Here, we assume that $m(i)$ is a smooth function wrt i and $\epsilon_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, \tau_0)$ with *known* precision τ_0 .

```
1 n = 50
2 idx = 1:n
3 # generate something smooth
  representing m
4 fun = 100*((idx-n/2)/n)^3
5 # add some noise
6 y = fun + rnorm(n, mean=0,
  sd=1)
7 plot(idx, y)
```



Assumed hierarchical model

1. **Data:** Gaussian observations with known precision

$$y_i \mid x_i, \theta \sim \mathcal{N}(x_i, \tau_0)$$

2. **Latent model:** A Gaussian model for the smooth function⁶

$$\pi(\mathbf{x} \mid \theta) \propto \theta^{(n-2)/2} \exp \left(-\frac{\theta}{2} \sum_{i=3}^n (x_i - 2x_{i-1} + x_{i-2})^2 \right)$$

3. **Hyperparameter:** The smoothing parameter θ which we assign a $\Gamma(a, b)$ prior

$$\pi(\theta) \propto \theta^{a-1} \exp(-b\theta), \quad \theta > 0$$

⁶`model="rw2"`

Derivation of posterior marginals (I)

Since

$$\mathbf{x}, \mathbf{y} \mid \theta \sim \mathcal{N}(\cdot, \cdot)$$

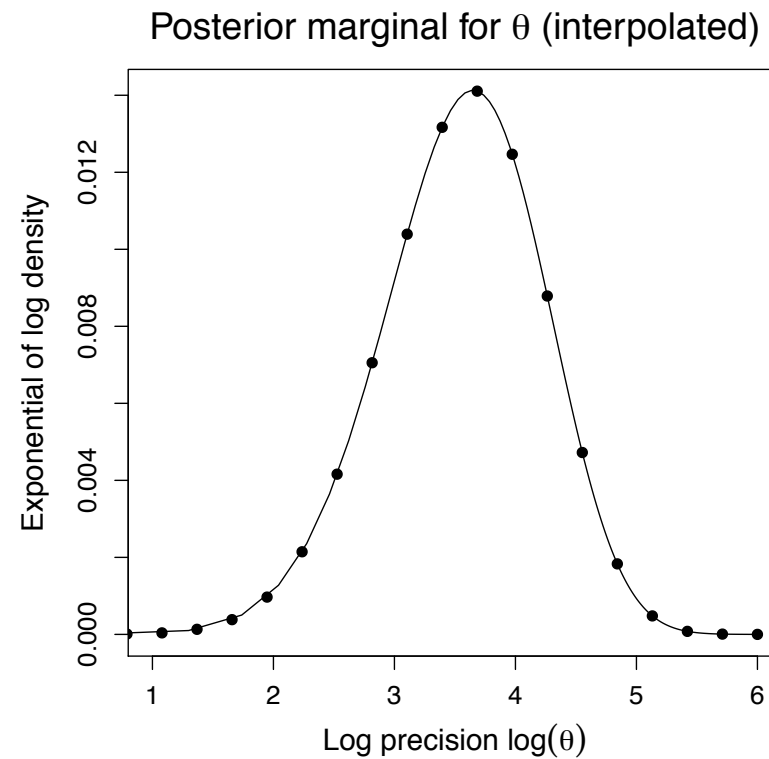
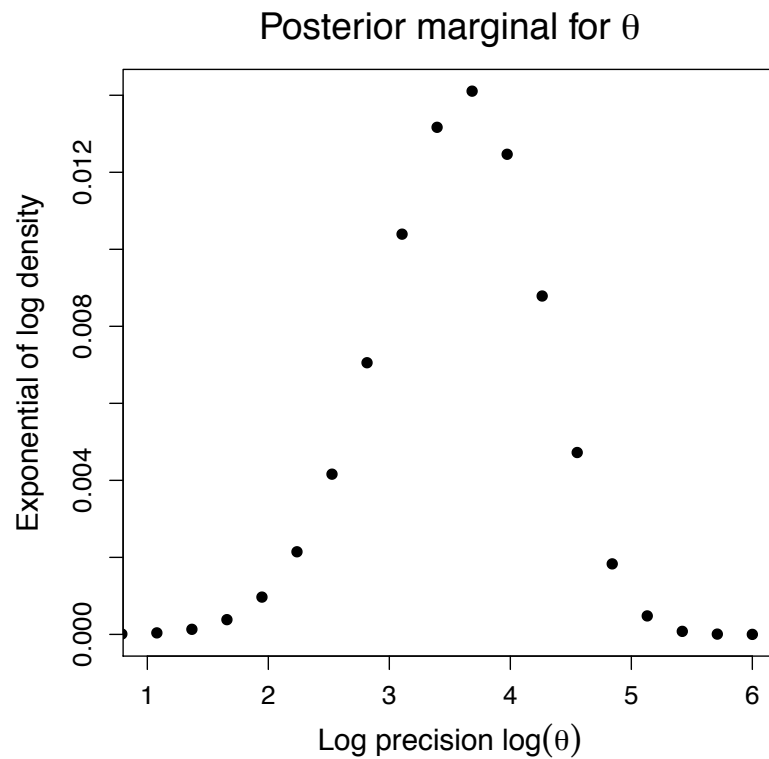
(derived using $\pi(\mathbf{x}, \mathbf{y} \mid \theta) \propto \pi(\mathbf{y} \mid \mathbf{x}, \theta) \pi(\mathbf{x} \mid \theta)$),

we can compute (numerically) all marginals, using that

$$\pi(\theta \mid \mathbf{y}) \propto \frac{\overbrace{\pi(\mathbf{x}, \mathbf{y} \mid \theta)}^{\text{Gaussian}} \pi(\theta)}{\underbrace{\pi(\mathbf{x} \mid \mathbf{y}, \theta)}_{\text{Gaussian}}}$$

Posterior marginal for hyperparameter

Select a grid of points to represent the density $\theta \mid \mathbf{y}$. (Here, the points are chosen to be equi-distant).



Derivation of posterior marginals (II)

From

$$\mathbf{x} \mid \mathbf{y}, \theta \sim \mathcal{N}(\cdot, \cdot)$$

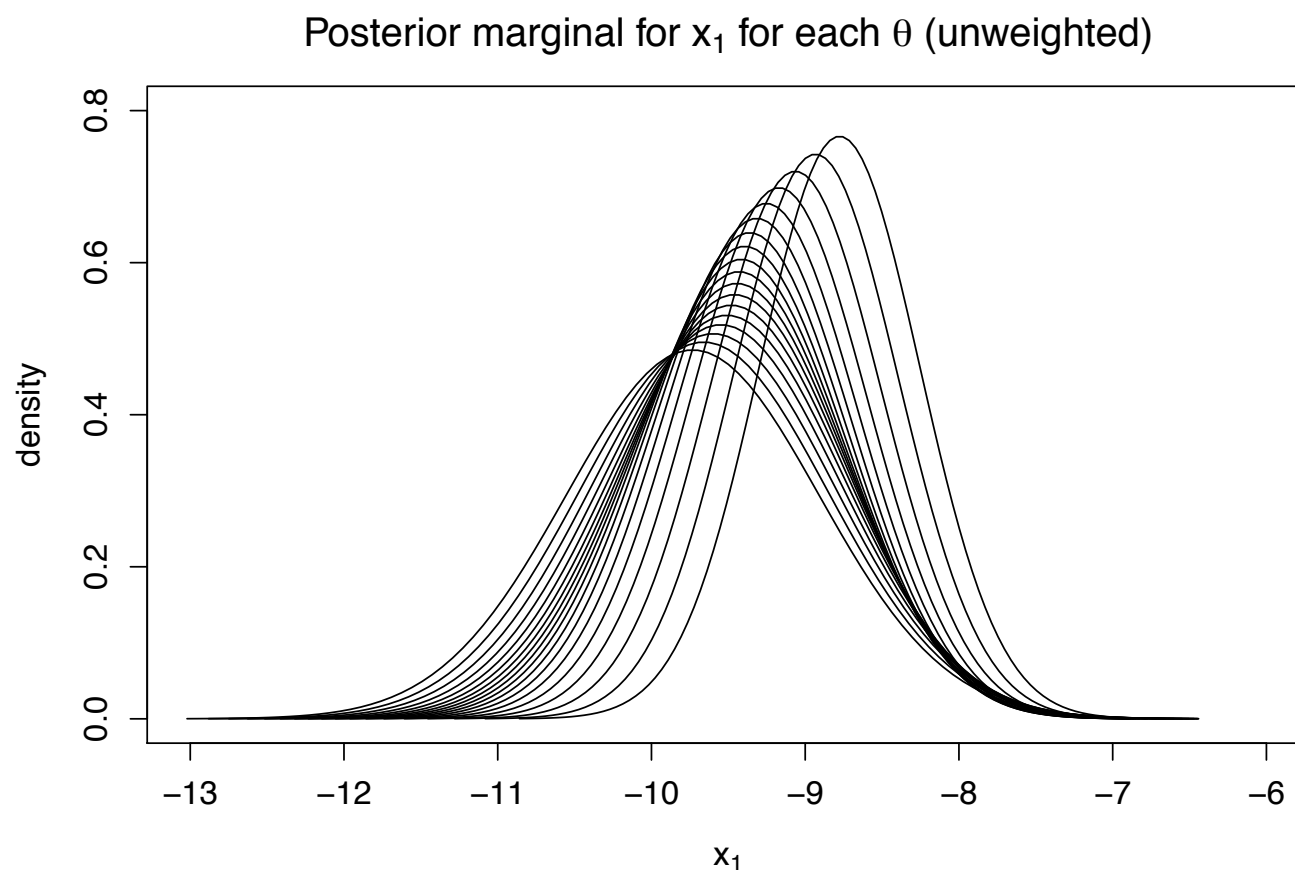
we can compute

$$\begin{aligned} \pi(x_i \mid \mathbf{y}) &= \int \underbrace{\pi(x_i \mid \theta, \mathbf{y})}_{\text{Gaussian}} \pi(\theta \mid \mathbf{y}) d\theta \\ &\approx \sum_k \pi(x_i \mid \theta_k, \mathbf{y}) \pi(\theta_k \mid \mathbf{y}) \Delta_k \end{aligned}$$

where θ_k , $k = 1, \dots, K$, correspond to representative points of $\theta \mid \mathbf{y}$ and Δ_k are the corresponding weights.

Posterior marginal for latent parameters

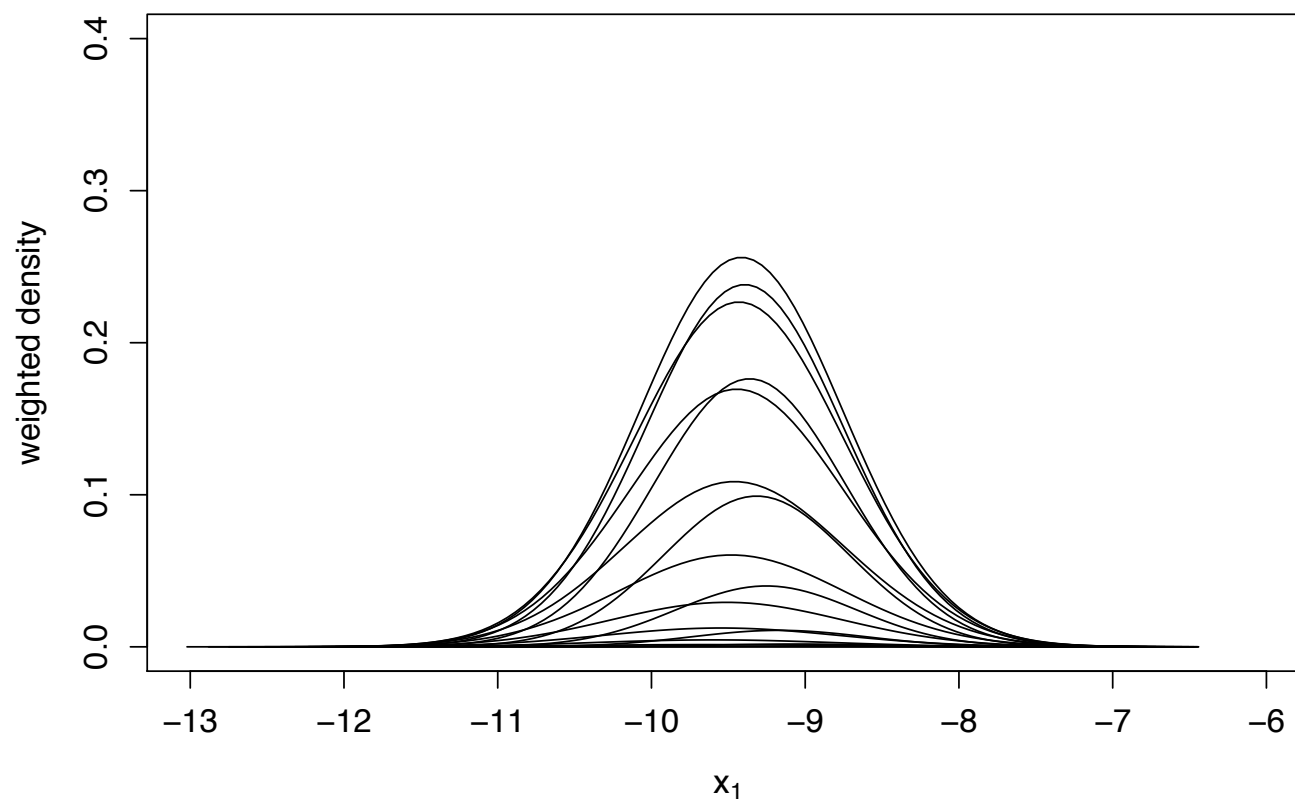
Compute the conditional marginal posterior for each x_i given θ_k . Here, shown for x_1 .



Posterior marginal for latent parameters

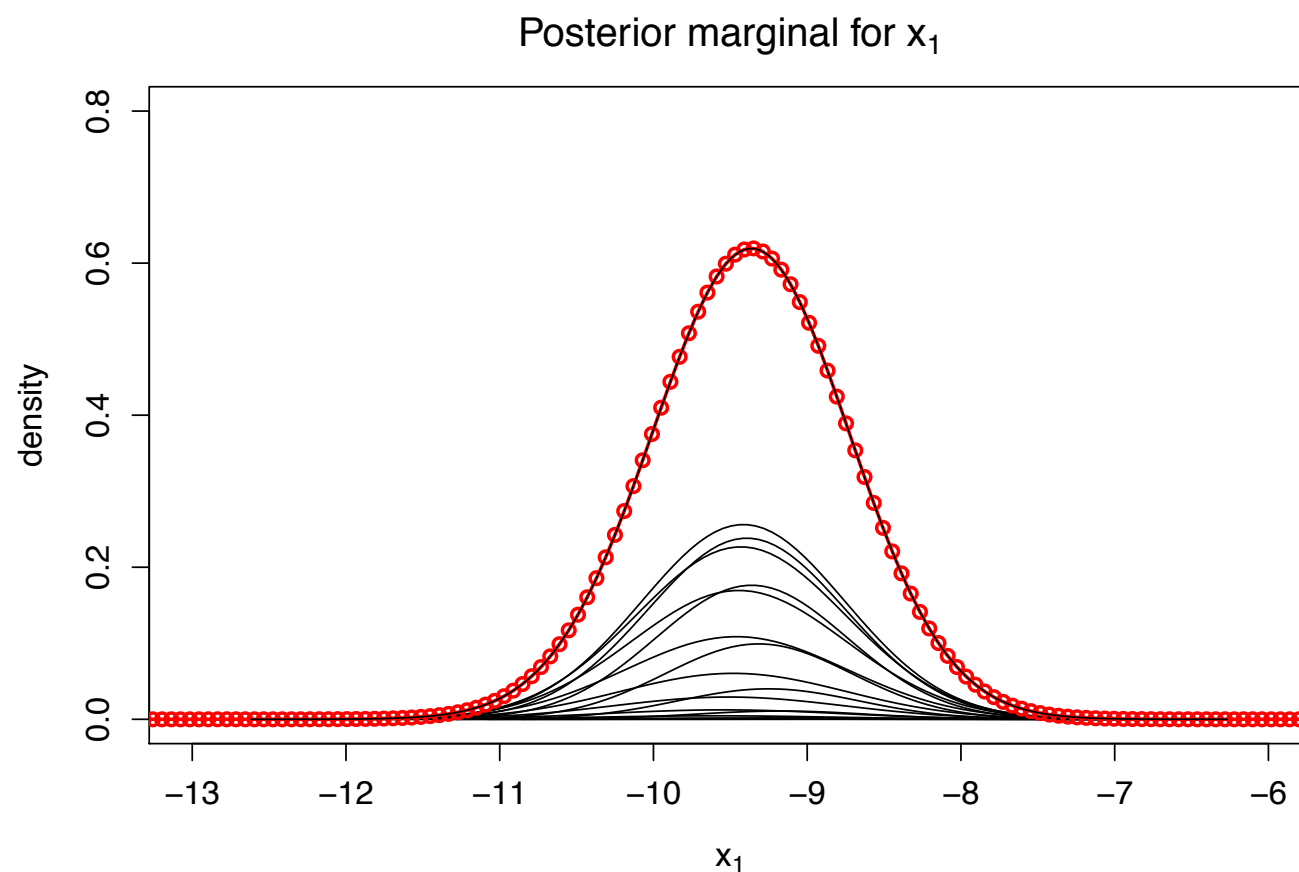
Weigh the resulting (conditional) marginal posterior by the density associated with each θ_k .

Posterior marginal for x_1 for each θ (weighted)



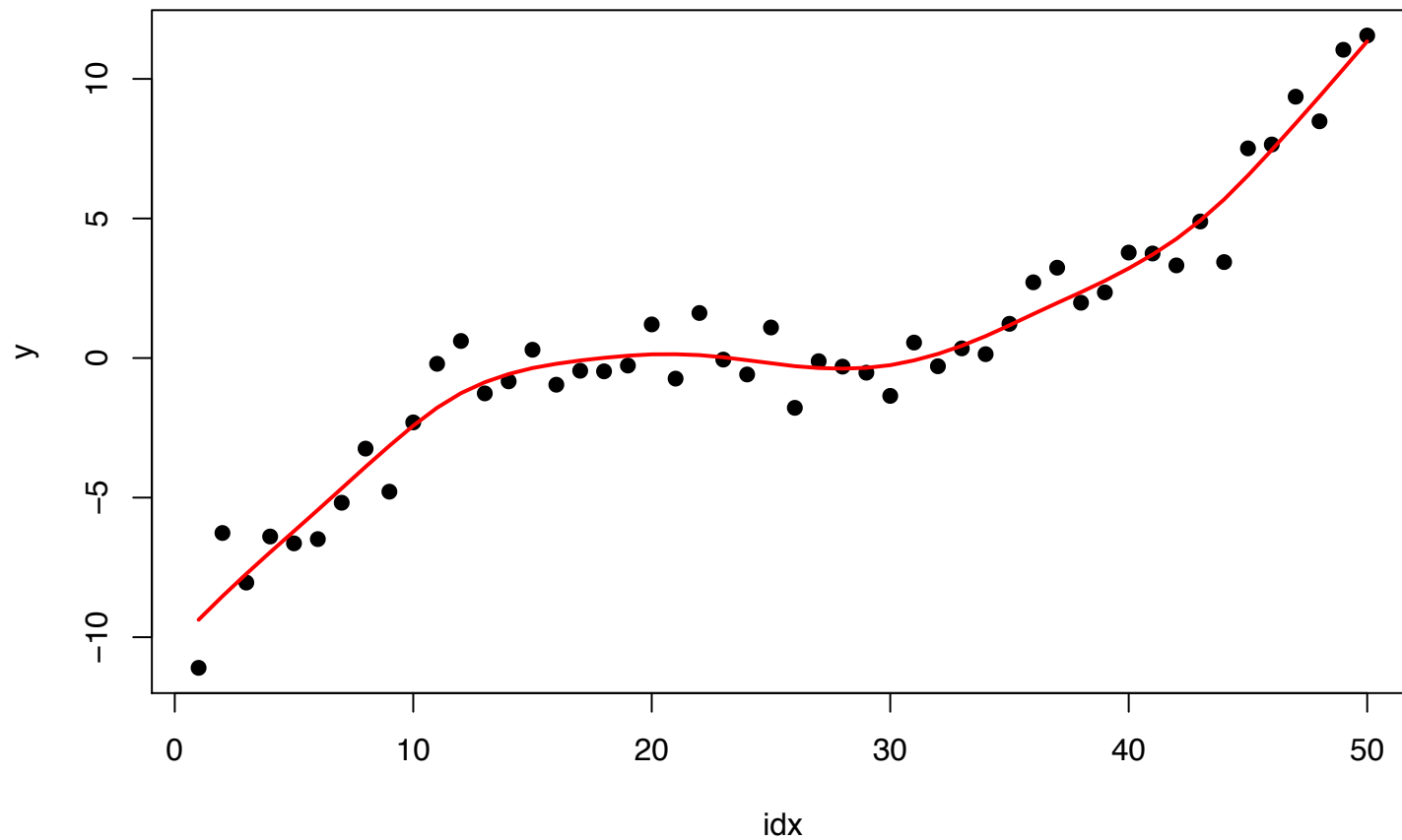
Posterior marginal for latent parameters

Numerically sum over all conditional densities to obtain the posterior marginal for each x_i .



Fitted spline

The posterior marginals are used to calculate summary statistics, like means, variances and credible intervals:



Extensions

This is the basic idea behind INLA. It is quite simple.

However, we need to extend this basic idea so we can deal with

- More than one hyperparameter
- Non-Gaussian observations

How, do things change?

$$\pi(\boldsymbol{\theta} \mid \mathbf{y}) \propto \frac{\overbrace{\pi(\mathbf{x}, \mathbf{y} \mid \boldsymbol{\theta})}^{\text{Non-Gaussian, BUT KNOWN}} \pi(\boldsymbol{\theta})}{\underbrace{\pi(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta})}_{\text{Non-Gaussian and UNKNOWN}}}$$

Extensions

This is the basic idea behind INLA. It is quite simple.

However, we need to extend this basic idea so we can deal with

- More than one hyperparameter
- Non-Gaussian observations

How, do things change?

$$\pi(\boldsymbol{\theta} \mid \mathbf{y}) \propto \frac{\overbrace{\pi(\mathbf{x}, \mathbf{y} \mid \boldsymbol{\theta})}^{\text{Non-Gaussian, BUT KNOWN}} \pi(\boldsymbol{\theta})}{\underbrace{\pi(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta})}_{\text{Non-Gaussian and UNKNOWN}}}$$

Complications... Mostly practical

The non-Gaussian part of the model

- In many cases $\pi(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta})$ is very close to a Gaussian distribution, and can be replaced with a Laplace approximation
- This means that all the really hard, high-dimensional integrals with respect to the latent field are easy, and only the integrals with respect to the hyperparameters remain
- If the number of hyperparameters is low, these integrals can be done efficiently numerically

The GMRF (Laplace) approximation

Let $\mathbf{x}$ denote a GMRF with precision matrix $\mathbf{Q}$ and mean $\boldsymbol{\mu}$. Approximate

$$\pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}) \propto \exp \left(-\frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \sum_{i=1}^n \log \pi(y_i|x_i) \right)$$

by using a second-order Taylor expansion of $\log \pi(y_i|x_i)$ around $\boldsymbol{\mu}_0$, say.

Recall

$$f(x) \approx f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2} f''(x_0)(x - x_0)^2 = a + bx - \frac{1}{2} cx^2$$

with $b = f'(x_0) - f''(x_0)x_0$ and $c = -f''(x_0)$.

The GMRF approximation (II)

Thus,

$$\begin{aligned}\tilde{\pi}(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}) &\propto \exp \left(-\frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \sum_{i=1}^n (a_i + b_i x_i - 0.5 c_i x_i^2) \right) \\ &\propto \exp \left(-\frac{1}{2} \mathbf{x}^\top (\mathbf{Q} + \text{diag}(\mathbf{c})) \mathbf{x} + \mathbf{b}^\top \mathbf{x} \right)\end{aligned}$$

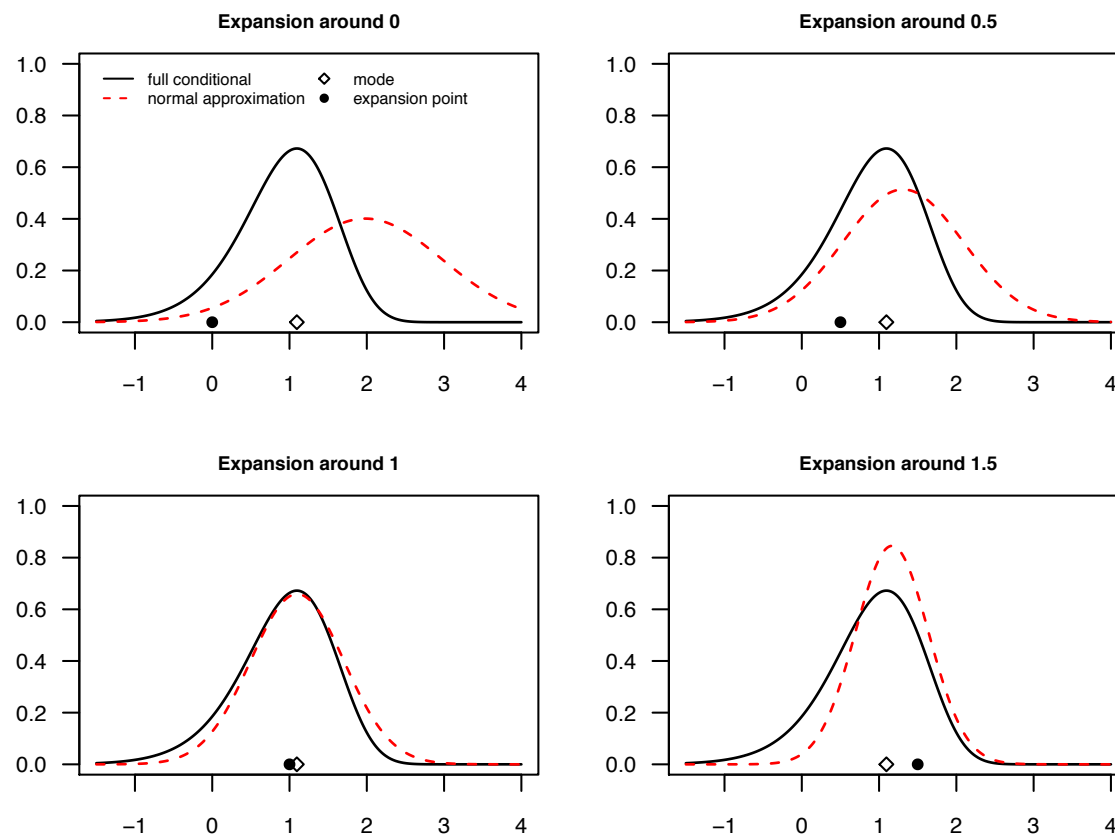
to get a Gaussian approximation with precision matrix $\mathbf{Q} + \text{diag}(\mathbf{c})$ and mean given by the solution of $(\mathbf{Q} + \text{diag}(\mathbf{c}))\boldsymbol{\mu} = \mathbf{b}$. The canonical parameterisation is

$$\mathcal{N}_c(\mathbf{b}, \mathbf{Q} + \text{diag}(\mathbf{c}))$$

which corresponds to

$$\mathcal{N}((\mathbf{Q} + \text{diag}(\mathbf{c}))^{-1} \mathbf{b}, (\mathbf{Q} + \text{diag}(\mathbf{c}))^{-1}).$$

Illustration



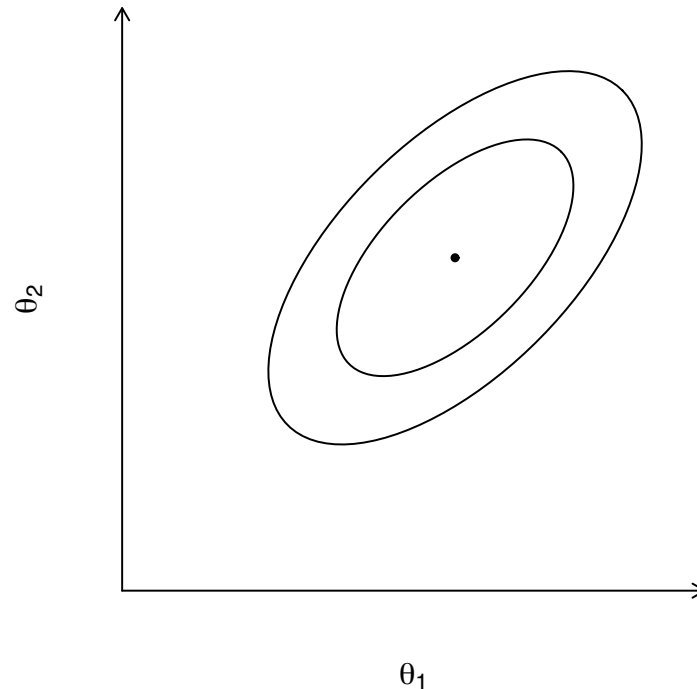
If $y \mid x, \theta$ is Gaussian "the approximation" is exact!

Exploring $\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$

$\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$ is “numerically explored” to find suitable support points $\boldsymbol{\theta}_k$.

Main use: Select good evaluation points $\boldsymbol{\theta}_k$ for the numerical integration when approximating $\tilde{\pi}(x_i|\mathbf{y})$

- Locate the mode

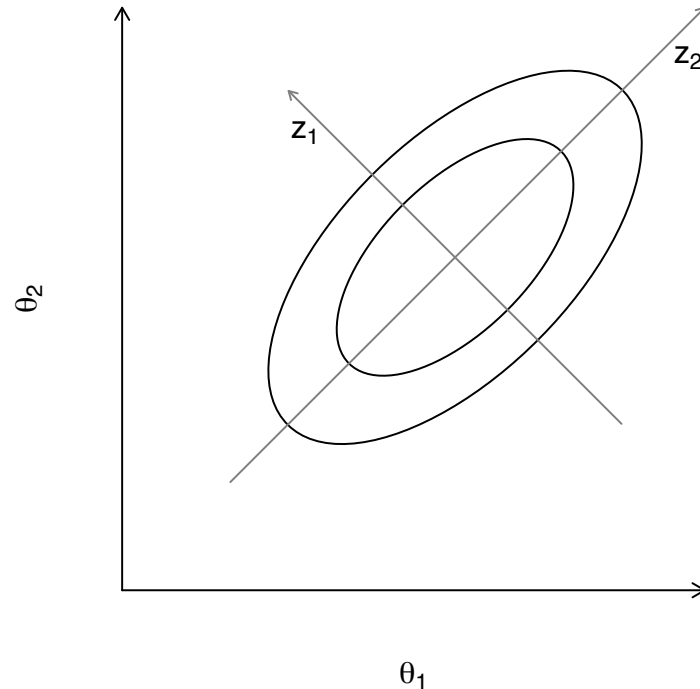


Exploring $\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$

$\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$ is “numerically explored” to find suitable support points $\boldsymbol{\theta}_k$.

Main use: Select good evaluation points $\boldsymbol{\theta}_k$ for the numerical integration when approximating $\tilde{\pi}(x_i|\mathbf{y})$

- Locate the mode
- Compute the Hessian to construct principal components

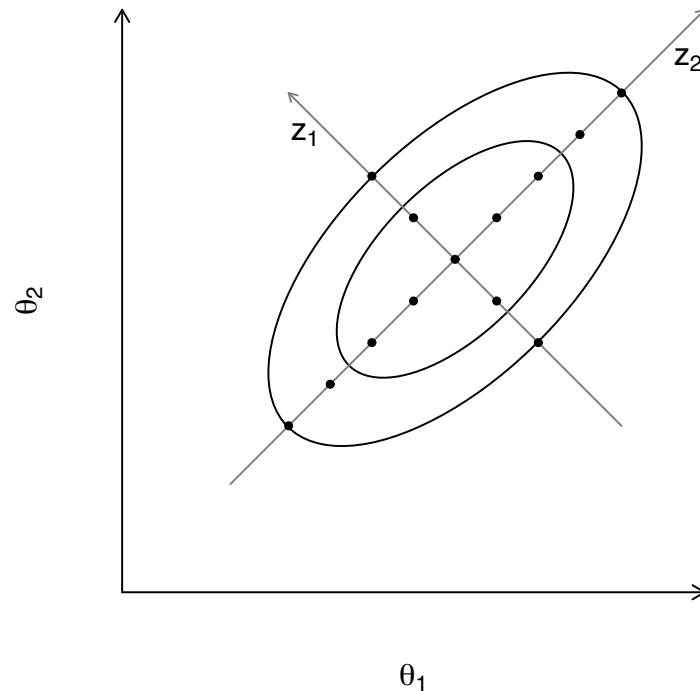


Exploring $\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$

$\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$ is “numerically explored” to find suitable support points $\boldsymbol{\theta}_k$.

Main use: Select good evaluation points $\boldsymbol{\theta}_k$ for the numerical integration when approximating $\tilde{\pi}(x_i|\mathbf{y})$

- Locate the mode
- Compute the Hessian to construct principal components
- Grid-search to locate bulk of the probability mass

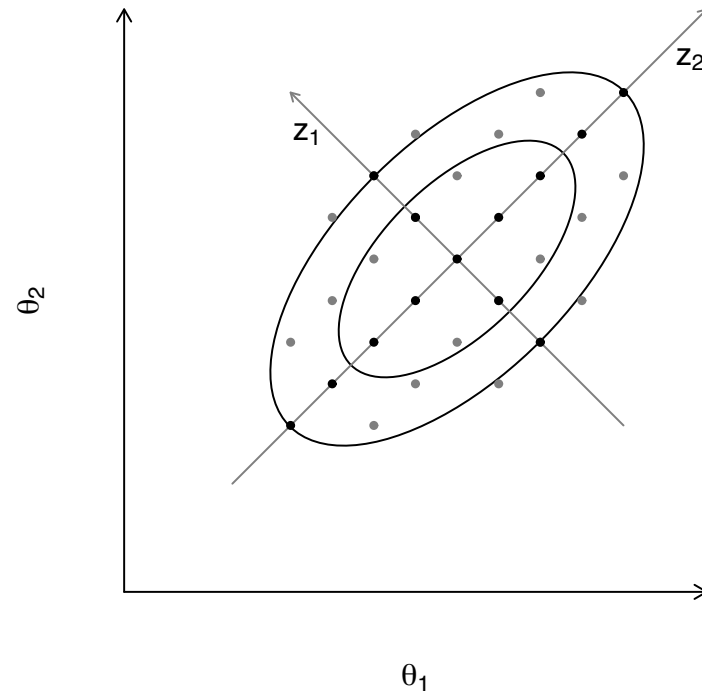


Exploring $\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$

$\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y})$ is “numerically explored” to find suitable support points $\boldsymbol{\theta}_k$.

Main use: Select good evaluation points $\boldsymbol{\theta}_k$ for the numerical integration when approximating $\tilde{\pi}(x_i|\mathbf{y})$

- Locate the mode
- Compute the Hessian to construct principal components
- Grid-search to locate bulk of the probability mass



All points found have equal area weight Δ_k .

Approximating $\pi(x_i|\boldsymbol{\theta}, \mathbf{y})$

For approximating the first component $\pi(x_i|\boldsymbol{\theta}, \mathbf{y})$ we can use

- a **Gaussian approximation**, easily extractable from $\tilde{\pi}_G(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})$.
However, **errors in location and/or lack of skewness** possible.
- a **Laplace approximation**

$$\tilde{\pi}_{LA}(x_i|\boldsymbol{\theta}, \mathbf{y}) \propto \frac{\pi(\mathbf{x}, \boldsymbol{\theta}, \mathbf{y})}{\tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \boldsymbol{\theta}, \mathbf{y})} \bigg|_{\mathbf{x}_{-i}=\mathbf{x}_{-i}^*(x_i, \boldsymbol{\theta})}.$$

The approximation is very accurate but very expensive.

- a **simplified Laplace approximation** based on fitting a skew-normal distribution to a series expansion of $\tilde{\pi}_{LA}$.

INLA: Overview

Step I Approximate $\pi(\boldsymbol{\theta}|\mathbf{y})$ using the Laplace approximation and select good evaluation points $\boldsymbol{\theta}_k$.

Step II For each $\boldsymbol{\theta}_k$ and i approximate $\pi(x_i|\boldsymbol{\theta}_k, \mathbf{y})$ using the Laplace or simplified Laplace approximation for selected values of x_i .

Step III For each i , sum out $\boldsymbol{\theta}_k$

$$\tilde{\pi}(x_i|\mathbf{y}) = \sum_k \tilde{\pi}(x_i|\boldsymbol{\theta}_k, \mathbf{y}) \times \tilde{\pi}(\boldsymbol{\theta}_k|\mathbf{y}) \times \Delta_k.$$

Build a log spline corrected Gaussian to represent $\tilde{\pi}(x_i|\mathbf{y})$.

Limitations

- The dimension of the latent field $\mathbf{x}$ can be large (10^2 – 10^6)
- But the dimension of the hyperparameters θ must be small (≤ 9)

In other words, each random effect can be big, but there cannot be too many random effects unless they share parameters.