

Stephan Bark
University of Hohenheim

Optimal Design for Multi-Environment Trials in Plant Breeding: A Mixed Model and Bayesian Updating Framework

Supervised by Hans-Peter Piepho and Volker Schmid

Thirteenth Working Seminar on Statistical Methods in Variety Testing, at Poznań,
Poland, on 07 - 12 September 2025

- 1 Introduction
- 2 Frequentist Framework
 - Mixed Model
 - Variance Components
 - Optimal Design
- 3 Bayesian Updating Framework
 - Mixed Model
 - Cycle Posterior Sampling Algorithm
 - Inverse Gamma Variance Components
 - Inverse Wishart Variance Components
 - Optimal Design
- 4 Discussion
- 5 References
- 6 Appendix

Multi-Environment Trials in Plant Breeding

- New crop variety development depends on performance across environments.
- Multi-environment trials (MET) evaluate varieties over multiple locations and years.
- METs help to assess genotype performance and genotype-environment interactions.
- The objective is to identify robust, high-yielding varieties for specific target populations of environments (TPE) distributed e.g across a country over multiple years.



Figure: Field experiment of rice.

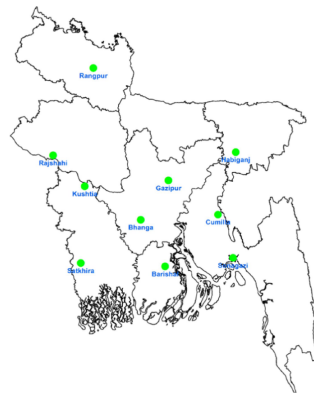


Figure: Location of MET of rice varieties from 2001 till 2020 provided by Rahman et al. (2023).

Problem Setting: Mixed Models for METs

1) Frequentist perspective using restricted maximum likelihood (REML)



2) Fully Bayesian perspective using Markov chain Monte Carlo (MCMC)

Objectives:

- Describe, discuss and compare both approaches in the context of the MET data set provided by Rahman et al. (2023)
- Explain implementation of fully Bayesian approach as Bayesian updating approach
- Demonstrate benefits of fully Bayesian approach in context of the optimal design approach by Prus and Piepho (2024)

Application of Optimal Design Calculation

- Prus and Piepho (2024) proposed a method to optimally distribute a fixed total number of trials across climatic zones e.g. of a country.
- Optimality should be achieved with respect to a criterion $\Phi^*(\xi)$.
- This criterion is influenced by both the design ξ and a fixed set of variance components from a mixed model fitted to MET data.
- ξ corresponds to either $J_1, \dots, J_Z$ integer numbers of trials per zone with $\sum_{z=1}^Z J_z = J$ or $w_1, \dots, w_Z$ zone weights per zone with $\sum_{z=1}^Z w_z = 1$ and $w_z \geq 0$.

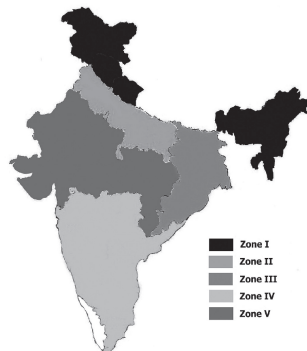


Figure: Indian maize breeding zones used by Kleinknecht et al. (2013).

Frequentist Framework

- (1) Estimate a Frequentist mixed model via restricted maximum likelihood (REML) method to the MET of Rahman et al. (2023), which outputs variance component point estimates.
- (2) Plug in variance components into design criterion $\Phi(\xi)$.
- (3) Optimize $\Phi(\xi)$ to generate optimal design ξ^* .

Frequentist Mixed Model Formulation

Definition

The mixed model in a conditional perspective is defined as

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \sum_{k=1}^{K+1} \mathbf{Z}_k \mathbf{b}_k + \boldsymbol{\epsilon},$$

where $\mathbf{b}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{G}_k)$ with $\mathbf{G}_k = \sigma_k^2 \mathbf{I}_{q_k}$,
and $\mathbf{b}_{K+1} \sim \mathcal{N}(\mathbf{0}, \mathbf{G}_{K+1})$ with $\mathbf{G}_{K+1} = \bigoplus_{g=1}^G \boldsymbol{\Sigma}_Z$,
and $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{R})$ with $\mathbf{R} = \bigoplus_{e=1}^E \mathbf{R}_e$ with $\mathbf{R}_e = \sigma_e^2 \mathbf{I}_{N_e}$,
and $\mathbf{y} \mid \mathbf{X}, \boldsymbol{\beta}, \mathbf{b}, \mathbf{R} \sim \mathcal{N}\left(\mathbf{X}\boldsymbol{\beta} + \sum_{k=1}^{K+1} \mathbf{Z}_k \mathbf{b}_k, \mathbf{R}\right)$.

- $\mathbf{b}_{K+1}$ corresponded to the nested effect of genotypes within zones.
- A dependence of the same genotype between the zones is assumed for the random effect $\mathbf{b}_{K+1}$.
- The effect of a genotype in each of the zones is assumed to be unstructured: In the variance covariance matrix $\boldsymbol{\Sigma}_Z$ each component is estimated freely.

Results of Variance Components

Table: REML estimated variance components and their corresponding standard errors of the fitted genotype-zone effect unstructured mixed model. The underlying data was provided by Rahman et al. (2023).

Variance (Short Form)	REML Estimate	Standard Error
σ_1^2 (var_year)	0.0287	0.0317
... (var_zone_year)	0	-
... (var_gen_year)	0.0246	0.0071
... (var_zone_loc_year)	0.561	0.066
... (var_gen_zone_year)	0	-
... (var_zone_loc_rep_year)	0.0046	0.0012
σ_K^2 (var_gen_zone_loc_year)	0.2504	0.0131
Σ_Z (Sigma_gen_zone)	see $\hat{\Sigma}_Z$ below	not provided here
$\frac{1}{E} \sum_{e=1}^E \sigma_e^2$ (env_mean_var_resid)	0.3249	0.3725

- The σ_e^2 's have been estimated individually and their mean has been calculated afterwards!

Results of Variance Components

For the nested effect of genotypes within zones:

$$\hat{\Sigma}_Z = \begin{bmatrix} 0.3121 & 0.2478 & 0.2943 & 0.1713 \\ 0.2478 & 0.2157 & 0.2388 & 0.1285 \\ 0.2943 & 0.2388 & 0.2955 & 0.1559 \\ 0.1713 & 0.1285 & 0.1559 & 0.1031 \end{bmatrix}.$$

→ High variance and low covariance estimates in $\hat{\Sigma}_Z$ lead to a larger number of trials!

Frequentist Optimal Design Results

Considering the estimated variance components above based on the MET data from

- **Zone 1** of Kushtia, Rajshahi and Rangpur;
- **Zone 2** of Barishal and Satkhira;
- **Zone 3** of Cumilla, Sonagazi and Habiganj;
- **Zone 4** of Bhanga and Gazipur;

e.g. for a MET with 3 years and 20 locations, the optimal zone weights are

$$w_1 = 0.49 / w_2 = 0.17 / w_3 = 0.29 / w_4 = 0.05$$

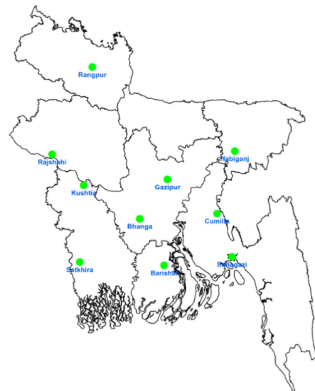


Figure: Location of MET of rice varieties from 2001 till 2020 provided by Rahman et al. (2023).

Bayesian Updating Framework

- (1) Introduce prior distributional assumptions for the variance components in the mixed model.
- (2) Specify a Bayesian mixed model for the MET of Rahman et al. (2023), which outputs marginal posterior distributions of variance components.
- (3) Implement a posterior sampling algorithm using Markov chain Monte Carlo (MCMC) methods to objectively inform mixed model variance components through historical MET data cycles.
- (4) Use fully informed marginal posteriors after the last cycle to draw variance component samples.
- (5) Optimize $\Phi(\xi)$ for each drawn set of variance components to generate optimal design posterior sample.
→ Allows for additional uncertainty measures in the design.

Bayesian Mixed Model Reformulation

Motivation: Uncertainty measures of the optimal design later in the pipeline.

Definition

The Bayesian model joint posterior based on the conditional mixed model perspective has the property

$$\underbrace{p(\beta, \mathbf{b}_1, \dots, \mathbf{b}_{K+1}, \sigma_1^2, \dots, \sigma_K^2, \boldsymbol{\Sigma}_Z, \sigma_1^2, \dots, \sigma_E^2 \mid \mathbf{y}, \mathbf{X})}_{\text{Mixed model parameter vector } \boldsymbol{\Psi}} \stackrel{\text{[Bayes Theorem]}}{\propto} \underbrace{\underbrace{p(\mathbf{y} \mid \mathbf{X}, \boldsymbol{\Psi})}_{\text{Likelihood}} \times p(\beta) \times \prod_{k=1}^{K+1} p(\mathbf{b}_k \mid \mathbf{G}_k)}_{\text{Priors}} \times \prod_{k=1}^K p(\sigma_k^2) \times p(\boldsymbol{\Sigma}_Z) \times \prod_{e=1}^E p(\sigma_e^2),$$

where $\mathbf{b}_{K+1} \sim \mathcal{N}(\mathbf{0}, \mathbf{G}_{K+1})$,
and $\mathbf{b}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{G}_k)$,
and $\boldsymbol{\Sigma}_Z \sim \text{Inv-Wishart}(\nu, \mathbf{S})$,
and $\sigma_k^2 \sim \text{Inv-Gamma}(\alpha_k, \beta_k)$,
and $\sigma_e^2 \sim \text{Inv-Gamma}(\alpha_e, \beta_e)$.

Cycle Procedure to Objectively Inform Priors

A 22-year historical dataset is split up with respect to L multi-year cycles.

Important notes:

- All full conditional posteriors of all mixed model parameters have been proven to be conjugated to their priors.
→ The MCMC algorithm can be based on efficient Gibbs sampling.
- $\prod_{k=1}^{K+1} p(\mathbf{b}_k | \mathbf{G}_k) \stackrel{ind.}{\sim} \prod_{k=1}^{K+1} \mathcal{N}(\mathbf{0}, \mathbf{G}_k)$ are conditional prior distributions of the random effects.
- The structural assumptions of the priors should not be violated during the cycling process to stay in the Bayesian mixed model framework!
- The idea is to only pass the estimated posterior parameters of the variance component posteriors through the cycles.

Cycle Procedure via Gibbs Sampling

For each multi-year cycle data set $\mathbf{y}_l, \mathbf{X}_l$ with $l = 1, \dots, L$:

(1) Initialize $\boldsymbol{\Psi}_{(l)}^{(0)} = (\boldsymbol{\beta}_{(l)}^{(0)}, (\mathbf{b}_1)_{(l)}^{(0)}, \dots, (\mathbf{b}_{K+1})_{(l)}^{(0)}, (\sigma_1^2)_{(l)}^{(0)}, \dots, (\sigma_K^2)_{(l)}^{(0)}, (\boldsymbol{\Sigma}_Z)_{(l)}^{(0)}, (\sigma_1^2)_{(l)}^{(0)}, \dots, (\sigma_E^2)_{(l)}^{(0)})^\top$.

(2) For each posterior sample component $t = 1, \dots, B, \dots, T$:

(2.1) sample $\boldsymbol{\beta}_{(l)}^{(t)} \sim p_{(l)}(\boldsymbol{\beta} | (\mathbf{b})_{(l)}^{(t-1)}, \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$,

(2.2) sample for $k = 1, \dots, K + 1$: $(\mathbf{b}_k)_{(l)}^{(t)} \sim p_{(l)}(\mathbf{b}_k | \boldsymbol{\beta}_{(l)}^{(t)}, \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$,

(2.3) sample for $k = 1, \dots, K$: $(\sigma_k^2)_{(l)}^{(t)} \sim p_{(l)}(\sigma_k^2 | \boldsymbol{\beta}_{(l)}^{(t)}, \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$,

(2.4) sample $(\boldsymbol{\Sigma}_Z)_{(l)}^{(t)} \sim p_{(l)}(\boldsymbol{\Sigma}_Z | \boldsymbol{\beta}_{(l)}^{(t)}, \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$,

(2.5) sample for $e = 1, \dots, E$: $(\sigma_e^2)_{(l)}^{(t)} \sim p_{(l)}(\sigma_e^2 | \boldsymbol{\beta}_{(l)}^{(t)}, \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$,

→ All $p_{(l)}(\cdot | \dots, \mathbf{X}_{(l)}, \mathbf{y}_{(l)})$ are known full conditional normal, inverse Gamma or inverse Wishart distributions that implicitly depends on $\mathbf{X}_{(1)}, \mathbf{y}_{(1)}, \dots, \mathbf{X}_{(l-1)}, \mathbf{y}_{(l-1)}$

Cycle Procedure via Gibbs Sampling

(2) **Note:** During all cycles full conditionals of $\mathbf{b}_k$ involves the prior $\mathbf{b}_k | \sigma_k^2 \sim \mathcal{N}(\mathbf{0}, \sigma_k^2 \mathbf{I}_{Q_k})$ or $\mathbf{b}_{K+1} | \Sigma_Z \sim \mathcal{N}(\mathbf{0}, \bigoplus_{g=1}^G \Sigma_Z)$,

and after the first cycle full conditionals of σ_k^2 involves the updated prior $\sigma_k^2 \sim p_{(l)}(\sigma_k^2 | \mathbf{X}_{(l-1)}, \mathbf{y}_{(l-1)}) = \text{Inv-Gamma}\left((\hat{\alpha}_k)_{(l-1)}^{(\text{MLE})}, (\hat{\beta}_k)_{(l-1)}^{(\text{MLE})}\right)$,

and full conditional of Σ_Z involves the updated prior $\Sigma_Z \sim p_{(l)}(\Sigma_Z | \mathbf{X}_{(l-1)}, \mathbf{y}_{(l-1)}) = \text{Inv-Wishart}\left((\hat{\nu})_{(l-1)}^{(\text{MLE})}, (\hat{\mathbf{S}})_{(l-1)}^{(\text{MLE})}\right)$,

and full conditional of σ_e^2 involves the updated prior $\sigma_e^2 \sim p_{(l)}(\sigma_e^2 | \mathbf{X}_{(l-1)}, \mathbf{y}_{(l-1)}) = \text{Inv-Gamma}\left((\hat{\alpha}_e)_{(l-1)}^{(\text{MLE})}, (\hat{\beta}_e)_{(l-1)}^{(\text{MLE})}\right)$.

Cycle Procedure via Gibbs Sampling

(3) Compute and store

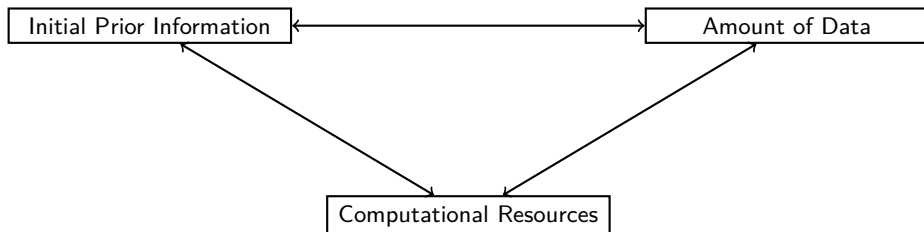
$(\hat{\alpha}_k)_{(l)}^{(\text{MLE})}, (\hat{\beta}_k)_{(l)}^{(\text{MLE})}, (\hat{\nu})_{(l)}^{(\text{MLE})}, (\hat{\mathbf{S}})_{(l)}^{(\text{MLE})}, (\hat{\alpha}_e)_{(l)}^{(\text{MLE})}, (\hat{\beta}_e)_{(l)}^{(\text{MLE})} \forall k$ and $\forall e$ via maximum likelihood estimation based on the current cycle l posterior sample:

$$\left[(\hat{\alpha}_k)_{(l)}^{(\text{MLE})}, (\hat{\beta}_k)_{(l)}^{(\text{MLE})} \right]^\top = \underset{[\alpha_k, \beta_k > 0]}{\operatorname{argmax}} \mathcal{L} \left(\alpha_k, \beta_k \mid (\sigma_k^2)_{(l)}^{(B+1)}, \dots, (\sigma_k^2)_{(l)}^{(T)} \right)$$

$$\left[(\hat{\nu})_{(l)}^{(\text{MLE})}, (\hat{\mathbf{S}})_{(l)}^{(\text{MLE})} \right]^\top = \underset{[\nu > Z+1 \text{ and } \mathbf{S} \text{ pos. def.}]}{\operatorname{argmax}} \mathcal{L} \left(\nu, \mathbf{S} \mid (\boldsymbol{\Sigma}_Z)_{(l)}^{(B+1)}, \dots, (\boldsymbol{\Sigma}_Z)_{(l)}^{(T)} \right)$$

$$\left[(\hat{\alpha}_e)_{(l)}^{(\text{MLE})}, (\hat{\beta}_e)_{(l)}^{(\text{MLE})} \right]^\top = \underset{[\alpha_e, \beta_e > 0]}{\operatorname{argmax}} \mathcal{L} \left(\alpha_e, \beta_e \mid (\sigma_e^2)_{(l)}^{(B+1)}, \dots, (\sigma_e^2)_{(l)}^{(T)} \right)$$

Trade Off for Proper Convergence with MCMC



Prior Type	Convergence Risk	Stability	Bias Risk
Strongly Informative	Low	High	High (can over-shrink estimates)
Weakly Informative	Medium	Medium	Low (balances data and prior influence)
Uninformative (flat, improper)	High	Low	None (tends to divergence)

Table: Comparison of prior types and their effects on convergence, stability, and bias risk inspired by Gelman and Yao (2020).

Trade Off for Proper Convergence with MCMC

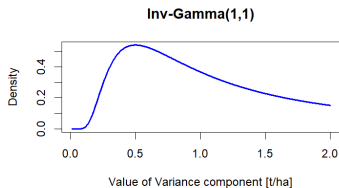
For demonstration purposes, comparison of two extreme estimation settings:

- (1) Low initial prior info and number of iterations
- (2) High initial prior info and number of iterations

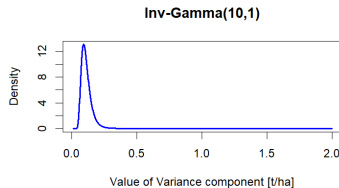
Initial Prior Information in the First Cycle

First, consider the scalar valued variance components $\sigma_1^2, \dots, \sigma_K^2$ and $\sigma_1^2, \dots, \sigma_E^2$.

Inv-Gamma Prior:



Low informative



High informative



Low Informative Initial Prior Parametrization

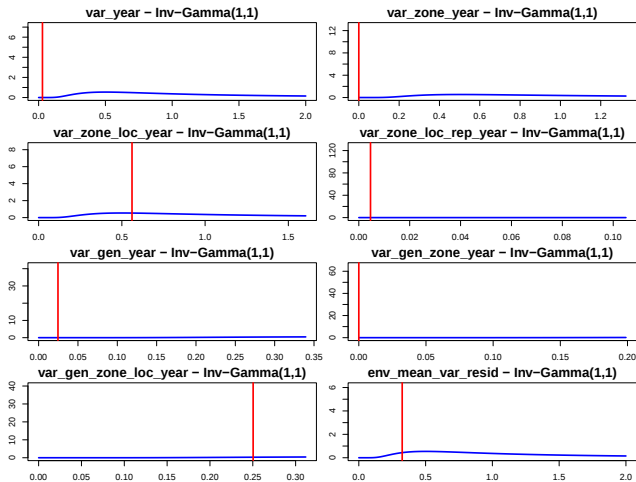


Figure: Inverse Gamma with shape and scale parameter equal to one [in blue]; Original MLE of the mixed model analysis [in red]; Note that the chosen x- and y-axis scales will match the posterior results on the next five slides.

Low Initial Prior Info and Number of Iterations

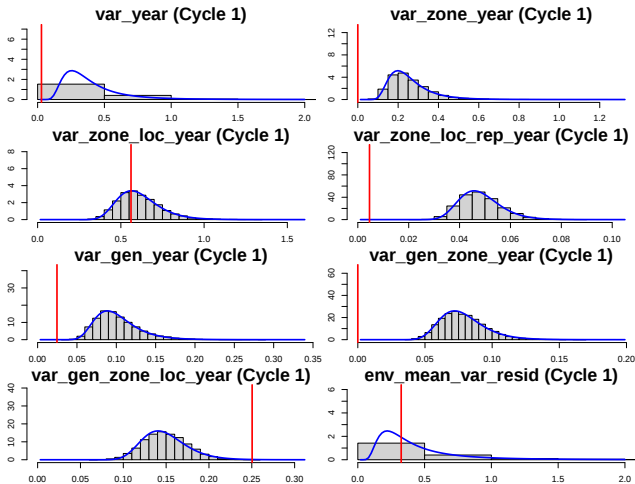


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

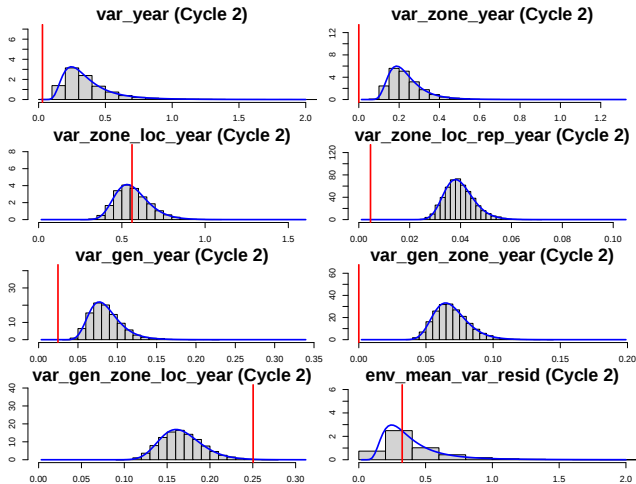


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

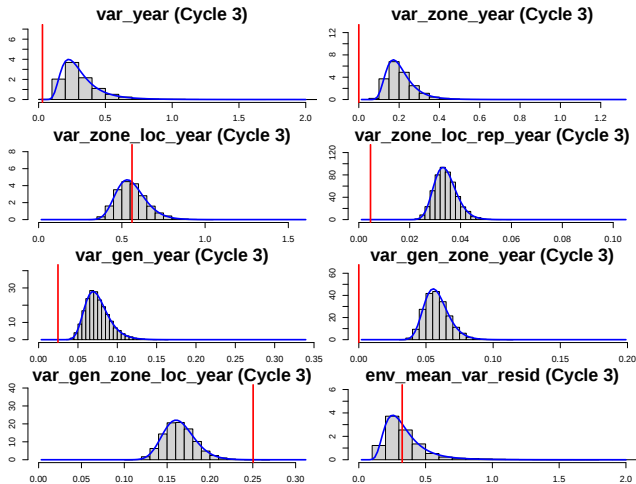


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

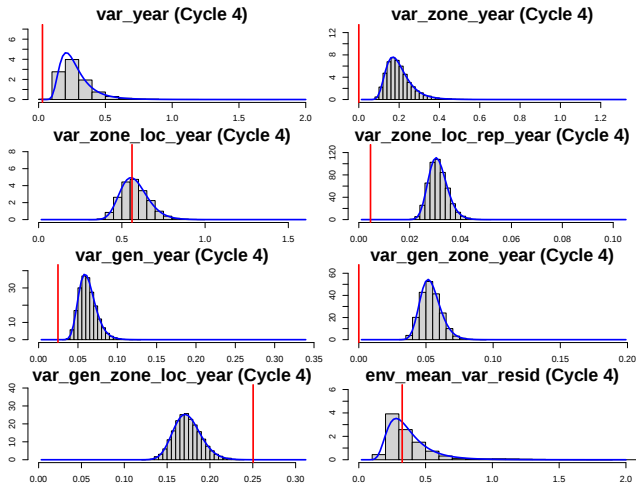


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

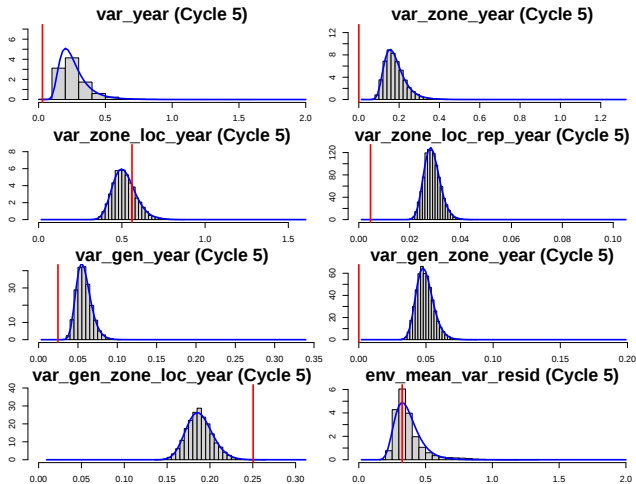


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

No proper convergence achieved after first sampled cycle!

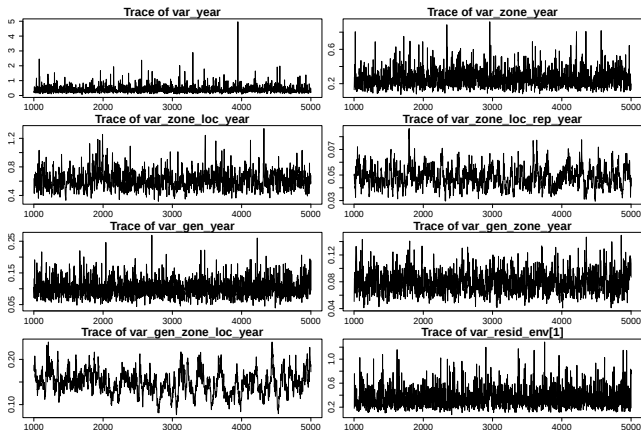


Figure: Trace plots corresponding to the 2000 posterior sampled values considering the thinning of two of chain one **of the first cycle**. On the x-axis is the number of the sample element of the Markov chain; On the y-axis is the value of the sampled variance component as yield in tons per hectare.

High Informative Initial Prior Parametrization

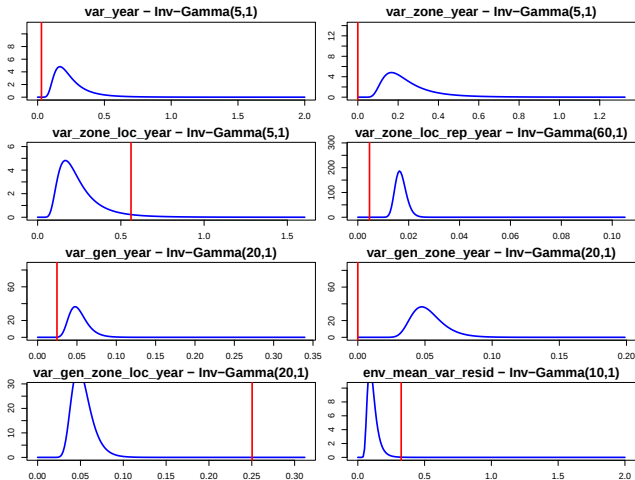


Figure: Inverse Gamma with shape and scale parameter individually chosen for each variance component [in blue]; Original MLE of the mixed model analysis [in red].

High Initial Prior Info and Number of Iterations

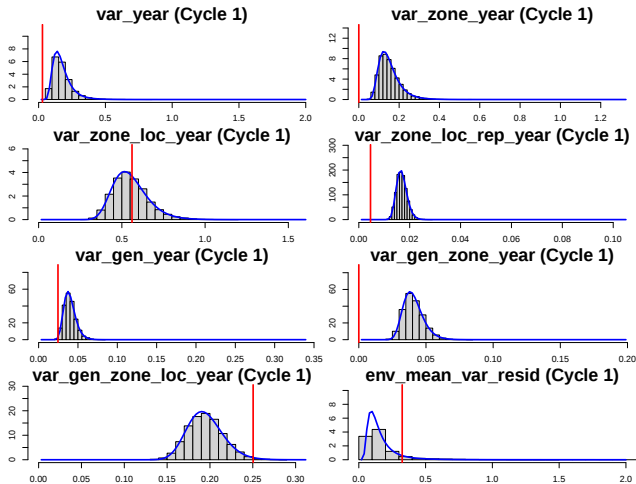


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

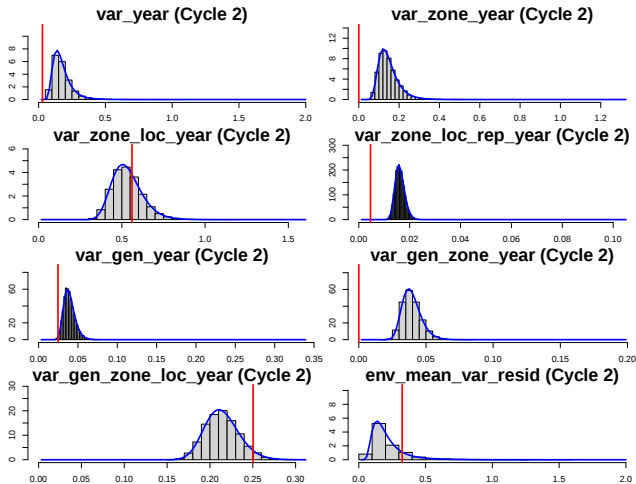


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

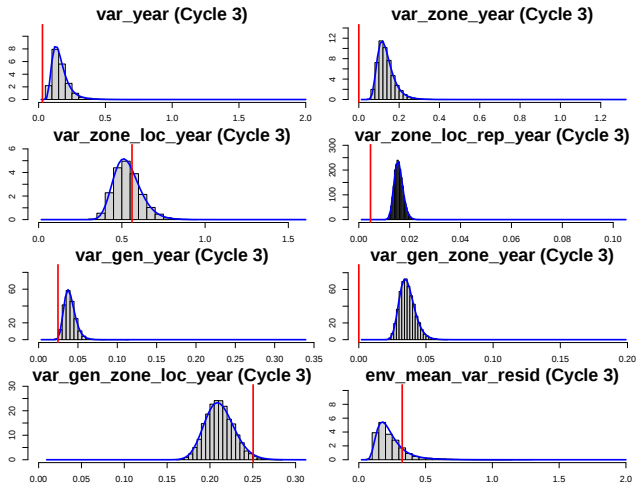


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

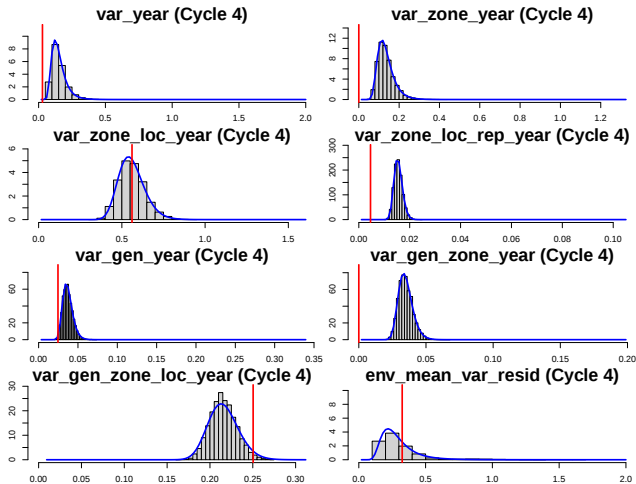


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

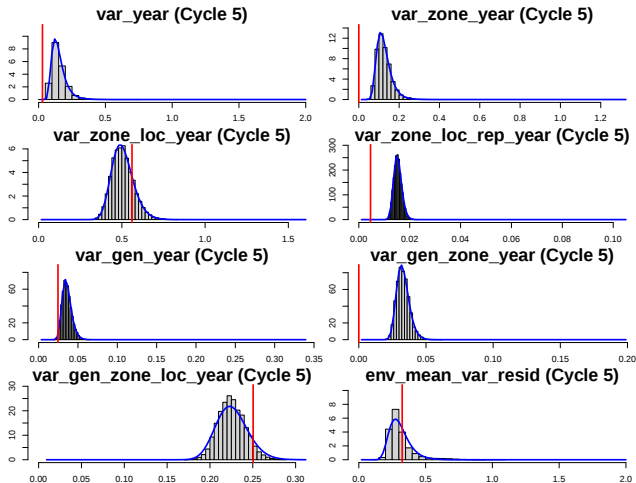


Figure: Histograms of the drawn Gibbs sample [in grey]; Fitted posterior density under inverse gamma approximation assumption [in blue]; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

Proper convergence achieved after first sampled cycle!

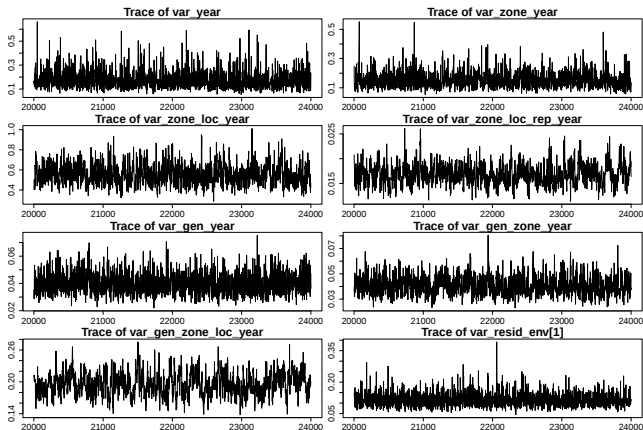


Figure: Trace plots corresponding to the 2000 posterior sampled values considering the thinning of two of chain one of the first cycle. On the x-axis is the number of the sample element of the Markov chain; On the y-axis is the value of the sampled variance component as yield in tons per hectare.

Initial Prior Information in the First Cycle

Second, consider the matrix-valued variance-covariance component Σ_Z .

Inv-Wishart Prior:

- Increasing information with increasing degree of freedom parameter ν .
→ Similar as with increasing the shape parameter of an inverse Gamma.
- Informing idea of scale matrix parameter S :

$$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 0.1 & 0.1 & 0.1 \\ 0.1 & 1 & 0.1 & 0.1 \\ 0.1 & 0.1 & 1 & 0.1 \\ 0.1 & 0.1 & 0.1 & 1 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 0.5 & 0.5 & 0.5 \\ 0.5 & 1 & 0.5 & 0.5 \\ 0.5 & 0.5 & 1 & 0.5 \\ 0.5 & 0.5 & 0.5 & 1 \end{bmatrix} \rightarrow \begin{bmatrix} 1 & 0.9 & 0.9 & 0.9 \\ 0.9 & 1 & 0.9 & 0.9 \\ 0.9 & 0.9 & 1 & 0.9 \\ 0.9 & 0.9 & 0.9 & 1 \end{bmatrix}$$

Low Informative Initial Prior Parametrization

$$\text{Inverse-Wishart}(\nu, \mathbf{S}) \rightarrow \text{Inverse-Wishart}\left(5, \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}\right)$$

Low Initial Prior Info and Number of Iterations

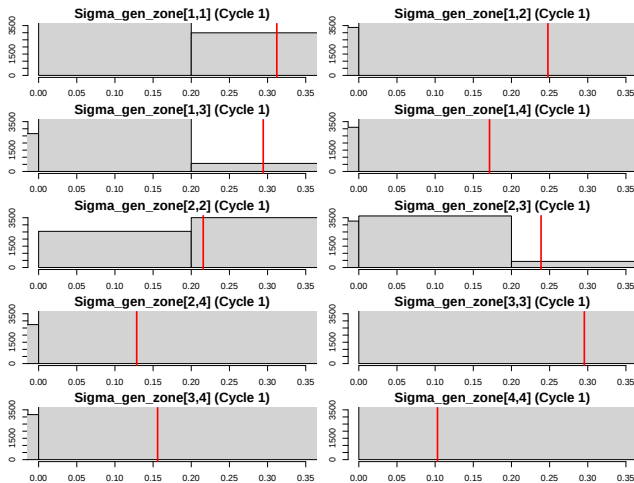


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

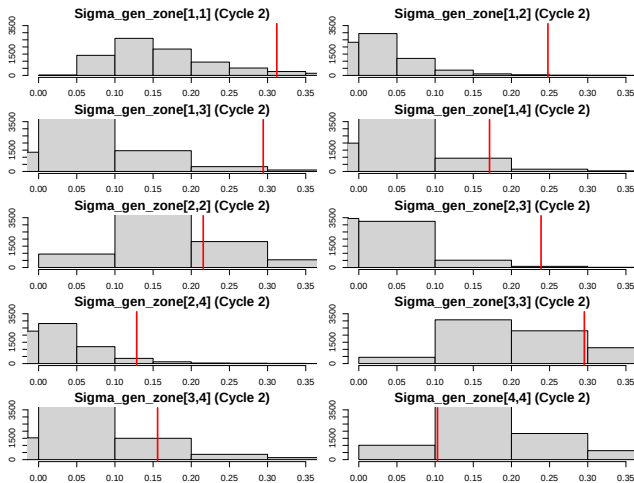


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

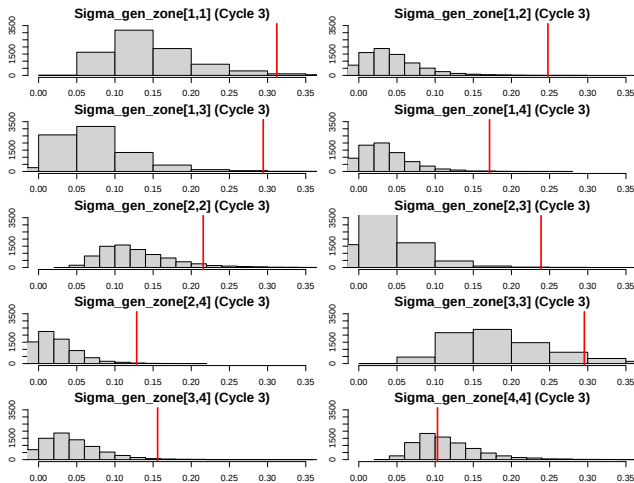


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

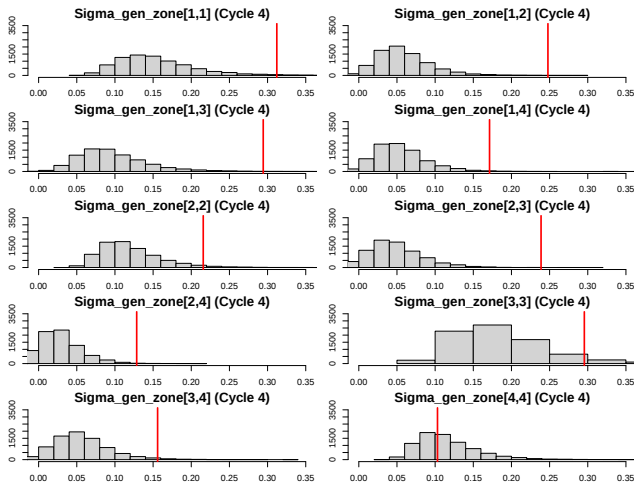


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

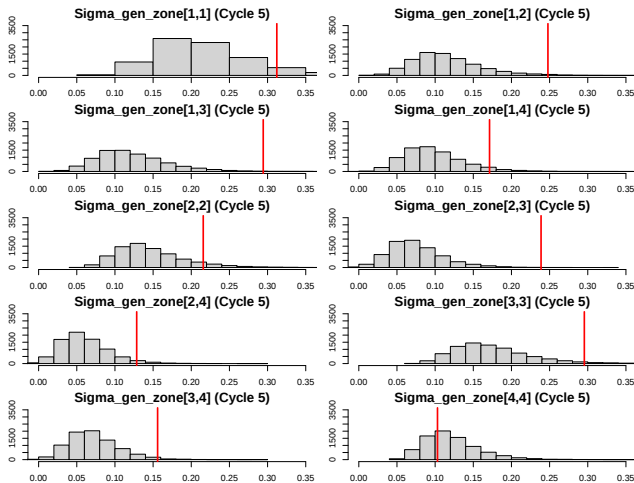


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

Low Initial Prior Info and Number of Iterations

No proper convergence achieved after first sampled cycle!

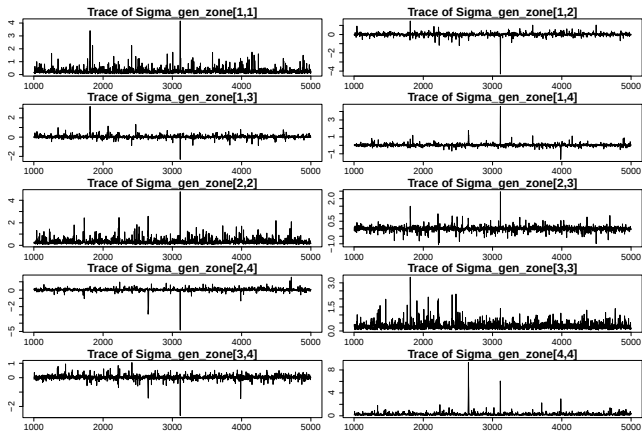


Figure: Trace plots corresponding to the 2000 posterior sampled values considering the thinning of two of chain one **of the first cycle**. On the x-axis is the number of the sample element of the Markov chain; On the y-axis is the value of the sampled variance component as yield in tons per hectare.

Low Initial Prior Info and Number of Iterations

Minor correlation issue between chains with inverse Wishart!

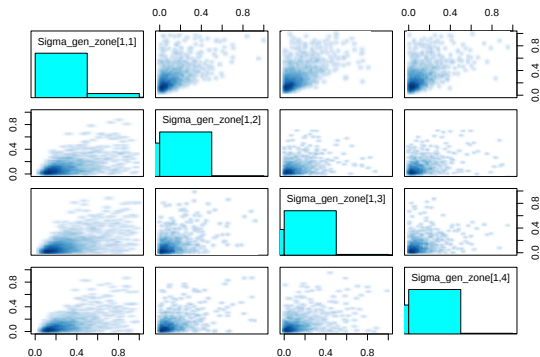


Figure: Pairs plots for diagonal variance parameters of `Sigma_gen_zone` corresponding to the 2000 posterior sampled values considering the thinning of two of chain one **of the first cycle**. Of the x- and y-axis of the off diagonal scatter plots are the sampled values of the corresponding variance component as yield in tons per hectare.

High Informative Initial Prior Parametrization

$$\text{Inverse-Wishart}(\nu, \mathbf{S}) \rightarrow \text{Inverse-Wishart}\left(20, \begin{bmatrix} 1 & 0.9 & 0.9 & 0.9 \\ 0.9 & 1 & 0.9 & 0.9 \\ 0.9 & 0.9 & 1 & 0.9 \\ 0.9 & 0.9 & 0.9 & 1 \end{bmatrix}\right)$$

High Initial Prior Info and Number of Iterations

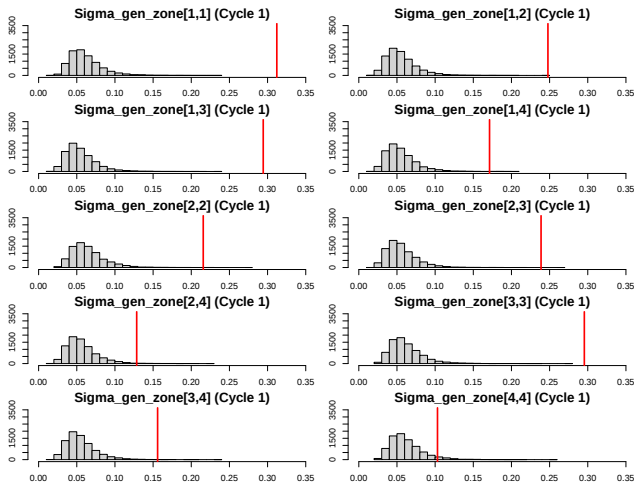


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

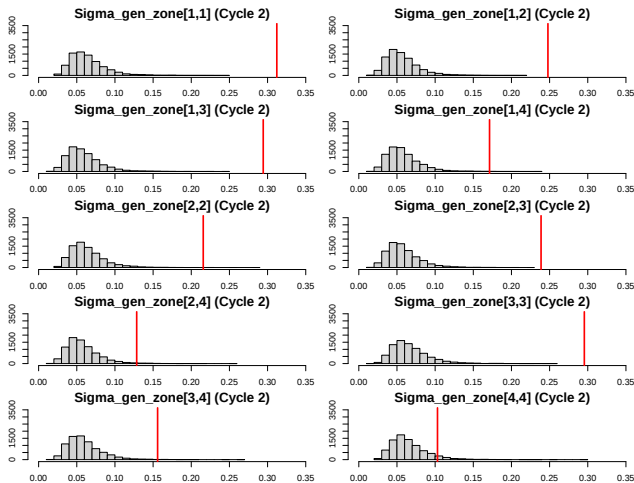


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

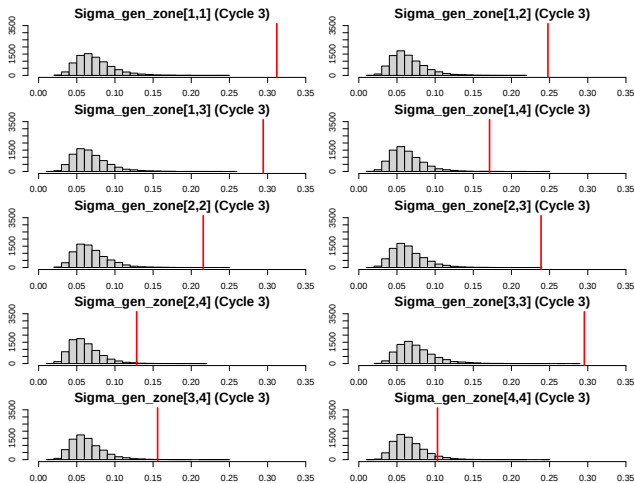


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

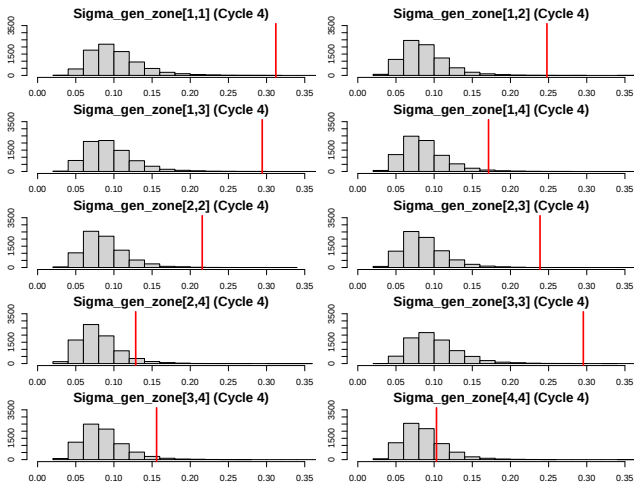


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

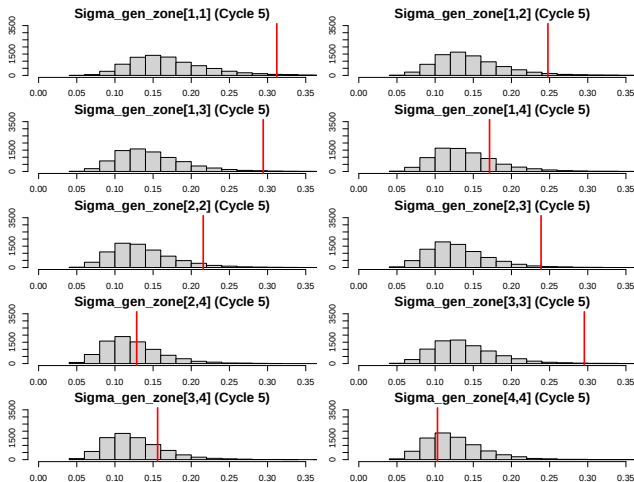


Figure: Histograms of the drawn Gibbs sample [in grey] with initial identity scale matrix; Original MLE of the mixed model analysis [in red]; Note that in all histograms the breaks are fixed to 20, where wider bars hint to more values out of plotted x-axis range.

High Initial Prior Info and Number of Iterations

Proper convergence achieved after first sampled cycle!

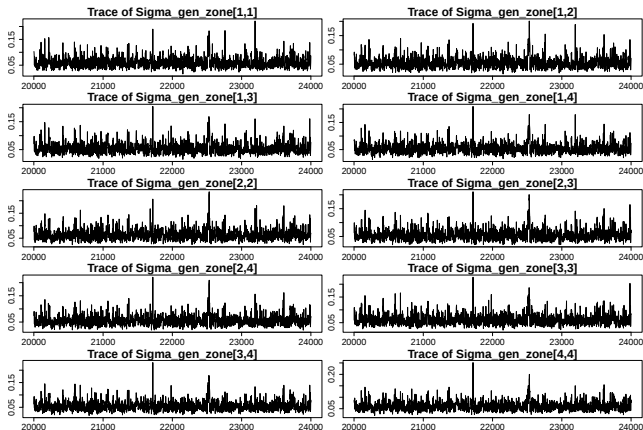


Figure: Trace plots corresponding to the 2000 posterior sampled values considering the thinning of two of chain one **of the first cycle**. On the x-axis is the number of the sample element of the Markov chain; On the y-axis is the value of the sampled variance component as yield in tons per hectare.

High Initial Prior Info and Number of Iterations

Major correlation issue between chains with inverse Wishart!

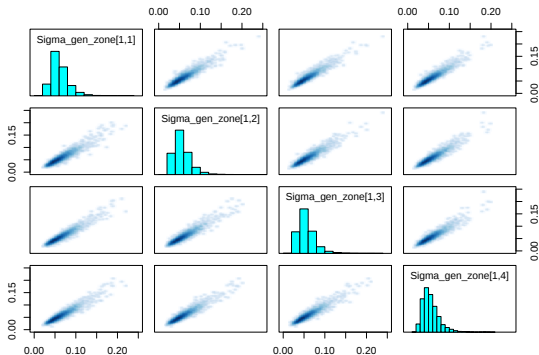
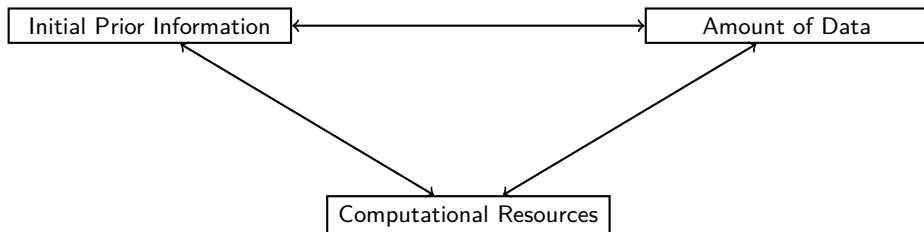


Figure: Pairs plots for diagonal variance parameters of `Sigma_gen_zone` corresponding to the 2000 posterior sampled values considering the thinning of two of chain one **of the first cycle**. Of the x- and y-axis of the off diagonal scatter plots are the sampled values of the corresponding variance component as yield in tons per hectare.

Trade Off for Proper Convergence with MCMC



Prior Type	Convergence Risk	Stability	Bias Risk
Strongly Informative	Low	High	High (can over-shrink estimates)
Weakly Informative	Medium	Medium	Low (balances data and prior influence)
Uninformative (flat, improper)	High	Low	None (tends to divergence)

Table: Comparison of prior types and their effects on convergence, stability, and bias risk inspired by Gelman and Yao (2020).

Bayesian Optimal Design

Considering the posterior sample of variance components above based on the MET data from

- **Zone 1** of Kushtia, Rajshahi and Rangpur;
- **Zone 2** of Barishal and Satkhira;
- **Zone 3** of Cumilla, Sonagazi and Habiganj;
- **Zone 4** of Bhanga and Gazipur;

e.g. for a MET with 3 years and 20 locations, the posterior mean and standard deviation of optimal zone weights are

$$\bar{w}_1 = 0.33 (0.04) / \bar{w}_2 = 0.23 (0.04) / \\ \bar{w}_3 = 0.26 (0.04) / \bar{w}_4 = 0.19 (0.05)$$

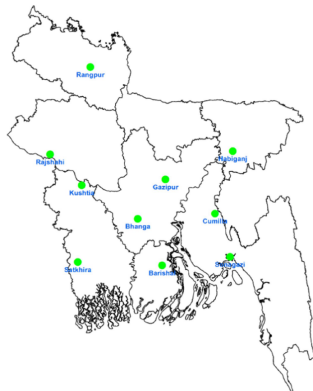


Figure: Location of MET of rice varieties from 2001 till 2020 provided by Rahman et al. (2023).

Frequentist and Bayesian Updating Framework

Frequentist Mixed Model

- + Uses the whole MET data in one go
- Does not consider historical prior information
- Variance component point estimates
- REML can estimate variance components to exactly zero (Studnicki and Piepho, 2024) assumed to be non zero by model assumption
- + Less computer intensive
- Less robust and stable estimation process

Bayesian Mixed Model

- Uses the MET data iteratively split into historical subsequent data cycles
- + Does objectively inform priors
- + Variance component posterior distribution estimates
- + Delivers more accurate estimation, as all variance components by model assumptions are considered and non zero
- MCMC might introduce bias, if priors are not truly objectively informed
- More computer intensive
- + More robust and stable estimation process

Further work on Bayesian Updating for MET data

The objective prior informing process:

- Inverse gamma and Wishart prior assumptions on variance components should be improved:
 - **Inverse gamma alternatives:** half-t distributions, half-Cauchy distributions, improper uniform priors, Jeffreys prior or log-normal priors (Gelman, 2006); penalized complexity (PC) priors (Simpson et al., 2017)
 - **Inverse Wishart alternatives:** Lewandowski-Kurowicka-Joe (LKJ) for correlation matrix and independent standard deviation (SD) prior (Lewandowski et al., 2009); scaled inverse Wishart (Alvarez et al., 2014)
- Alternative sampling strategies should be exploited.

⇒ Maximize efficiency and minimize bias in the analysis!

Further work on Bayesian Updating for MET data

So far, only MCMC as standard to sample from the posterior distribution.

Alternative estimation strategies:

- Variational inference (VI) with scalability in large/high-dimensional data and high model complexity context.
→ Loss of accuracy
- Integrated Nested Laplace Approximation (INLA) with high efficiency in spatial statistics and complex random effect structure context.
→ Small loss if accuracy

Further work on Bayesian Updating for MET data

In general: Increase stability and robustness of the estimation process for large MET data sets with complex experimental design structures and increase accuracy in corresponding variety testing analyses.

→ Improve genomic selection!

Interesting application situation:

- Optimal design with respect to climatic zones (Prus and Piepho, 2024)
- Genomic breeding values (GBV) in context of complex genotype-environment interaction structures (e.g. Buntaran et al., 2022)
- Factor analytic models (e.g. Piepho and Williams, 2024)
- Random intercept and slope model with covariates (e.g. Hrachov et al., 2025)
- Integrate spatial statistics methods

Thank you for your attention!



Figure: Rice Field in Bangladesh.

References

- Gelman, A., & Yao, Y. (2020). Holes in bayesian statistics. *Journal of Physics G*.
<https://doi.org/10.48550/arXiv.2002.06467>
- Kleinknecht, K., Möhring, J., Singh, K.-P., Zaidi, P.-H., Atlin, G.-N., & Piepho, H.-P. (2013). Comparison of the performance of best linear unbiased estimation and best linear unbiased prediction of genotype effects from zoned indian maize data. *Crop Science*, 53(4), 1384–1391.
<https://doi.org/10.2135/cropsci2013.02.0073>
- Prus, M., & Piepho, H.-P. (2024). Optimizing the allocation of trials to sub-regions in crop variety testing with multiple years and locations. *Journal of Agricultural, Biological and Environmental Statistics*.
<https://doi.org/10.1007/s13253-024-00659-1>
- Rahman, N. M. F., Malik, W. A., Kabir, M. S., Baten, M. A., Hossain, M. I., & et al., D. N. R. P. (2023). 50 years of rice breeding in bangladesh: Genetic yield trends. *TAG. Theoretical and applied genetics. Theoretische und angewandte Genetik*, 136(1), 18. <https://doi.org/10.1007/s00122-023-04260-x>
- Studnicki, M., & Piepho, H.-P. (2024). Hierarchical modelling of variance components makes analysis of resolvable incomplete block designs more efficient. *TAG. Theoretical and applied genetics. Theoretische und angewandte Genetik*, 137(6), 134. <https://doi.org/10.1007/s00122-024-04639-4>

Appendix: Rahman et al. (2023)'s data structure

- Randomized complete block design (RCBD) in each trial of the MET data.
- Increasing number of rice genotypes over the years.
- Rice genotypes are labeled with respect to the groups 'Short', 'Medium' and 'Long' in the data.

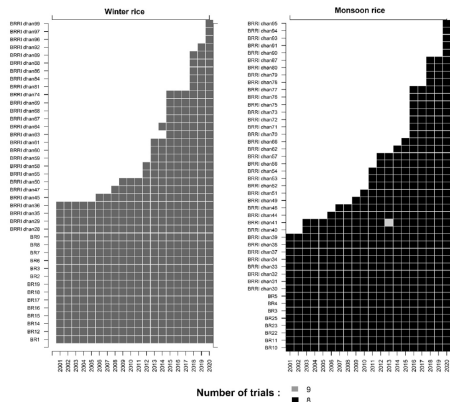


Figure: Classification of genotype-year combinations for winter rice series (left) and monsoon rice series (right). Cell shades of gray indicate number of trials in a year from Rahman et al. (2023).

Appendix: Descriptive Analysis

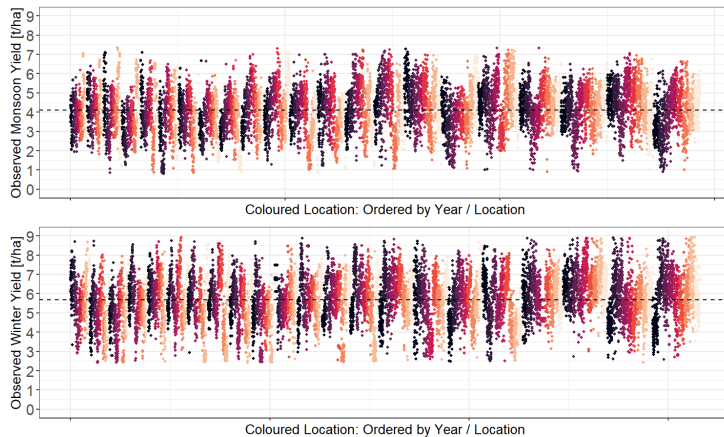


Figure: Observed yield in tons per hectare as the target variable of interest plotted and sorted by Environment (Year and Location combination); above for monsoon genotypes; below for Winter genotypes.

Frequentist Optimal Design Results

Table: Optimal numbers of trials per zone and efficiency of balanced design compared to optimal designs in case of unstructured variance-covariance structural assumption of Sigma_gen_zone in the model specified for different values of the total number of years H and locations J (Prus and Piepho, 2024).

H	J	Approximate design ξ_a^*				Eff_a	Exact design ξ_{ex}^*				Eff_{ex}
		w_1	w_2	w_3	w_4		J_1	J_2	J_3	J_4	
2	10	0.55	0.1	0.25	0.1	0.88	5	1	3	1	0.88
2	20	0.56	0.09	0.3	0.05	0.87	11	2	6	1	0.87
2	40	0.5	0.18	0.29	0.03	0.9	20	7	12	1	0.9
2	100	0.42	0.2	0.26	0.13	0.95	42	20	25	13	0.95
2	200	0.4	0.19	0.22	0.18	0.95	81	38	45	36	0.95
3	10	0.52	0.1	0.28	0.1	0.89	5	1	3	1	0.89
3	20	0.49	0.17	0.29	0.05	0.9	10	3	6	1	0.9
3	40	0.44	0.22	0.29	0.05	0.94	18	9	11	2	0.94
3	100	0.38	0.22	0.26	0.14	0.96	37	22	26	15	0.96
3	200	0.38	0.21	0.24	0.18	0.95	76	42	47	35	0.95
22	10	0.27	0.26	0.28	0.18	0.99	2	3	3	2	1
22	20	0.25	0.27	0.28	0.21	0.99	5	5	6	4	1
22	40	0.25	0.26	0.27	0.22	1	10	10	11	9	1
22	100	0.27	0.25	0.26	0.22	1	27	25	26	22	1
22	200	0.27	0.25	0.25	0.22	1	55	49	51	45	1

Bayesian Optimal Design

Table: Posterior mean and standard deviation in brackets of optimal approximate design ξ_a^* and efficiency of balanced design compared to optimal designs in case of unstructured variance-covariance structural assumption of Sigma_gen_zone in the underlying mixed model as specified for different values of the total number of years H and locations J (Prus and Piepho, 2024).

H	J	Posterior mean of approximate design ξ_a^*				$\overline{Eff}_a$
		$\bar{w}_1$	$\bar{w}_2$	$\bar{w}_3$	$\bar{w}_4$	
2	10	0.42 (0.07)	0.19 (0.07)	0.24 (0.07)	0.14 (0.05)	0.94 (0.031)
2	20	0.36 (0.05)	0.22 (0.05)	0.25 (0.05)	0.16 (0.06)	0.96 (0.027)
2	40	0.33 (0.03)	0.23 (0.04)	0.25 (0.03)	0.19 (0.04)	0.97 (0.019)
2	100	0.3 (0.02)	0.23 (0.03)	0.25 (0.02)	0.21 (0.03)	0.98 (0.014)
2	200	0.3 (0.02)	0.24 (0.02)	0.25 (0.02)	0.22 (0.02)	0.98 (0.012)
3	10	0.38 (0.05)	0.21 (0.06)	0.25 (0.06)	0.16 (0.05)	0.96 (0.025)
3	20	0.33 (0.04)	0.23 (0.04)	0.26 (0.04)	0.19 (0.05)	0.97 (0.018)
3	40	0.3 (0.02)	0.23 (0.03)	0.25 (0.03)	0.21 (0.03)	0.98 (0.012)
3	100	0.29 (0.02)	0.24 (0.02)	0.25 (0.02)	0.22 (0.02)	0.99 (0.008)
3	200	0.28 (0.02)	0.24 (0.02)	0.25 (0.02)	0.22 (0.02)	0.99 (0.008)
22	10	0.26 (0.01)	0.25 (0.01)	0.25 (0.01)	0.24 (0.02)	1 (0.004)
22	20	0.26 (0.01)	0.25 (0.01)	0.25 (0.01)	0.24 (0.01)	1 (0)
22	40	0.26 (0.01)	0.25 (0.01)	0.25 (0.01)	0.24 (0.01)	1 (0)
22	100	0.26 (0)	0.25 (0)	0.25 (0)	0.25 (0.01)	1 (0)
22	200	0.26 (0.01)	0.25 (0)	0.25 (0)	0.25 (0)	1 (0)

Appendix: Gibbs sampling with placeholder θ

For each multi-year cycle data set $\mathbf{y}_l, \mathbf{X}_l$ with $l = 1, \dots, L$:

① Initialize $\theta_{(l)}^{(0)} = ((\theta_1)_{(l)}^{(0)}, \dots, (\theta_D)_{(l)}^{(0)})^\top$ as draw from the priors.

② For iteration $t = 1, \dots, B, \dots, T$:

(1) Sample $(\theta_1)_{(l)}^{(t)} \sim p_{(l)}(\theta_1 | (\theta_2)_{(l)}^{(t-1)}, \dots, (\theta_D)_{(l)}^{(t-1)}, \mathbf{y}_{(l)}, \mathbf{X}_{(l)})$

(2) Sample $(\theta_2)_{(l)}^{(t)} \sim p_{(l)}(\theta_2 | (\theta_1)_{(l)}^{(t)}, (\theta_3)_{(l)}^{(t-1)}, \dots, (\theta_D)_{(l)}^{(t-1)}, \mathbf{y}_{(l)}, \mathbf{X}_{(l)})$

$\vdots$

(D) Sample $(\theta_D)_{(l)}^{(t)} \sim p_{(l)}(\theta_D | (\theta_1)_{(l)}^{(t)}, \dots, (\theta_{D-1})_{(l)}^{(t)}, \mathbf{y}_{(l)}, \mathbf{X}_{(l)})$

Full cond. posterior implicitly depends on $\mathbf{y}_{(1)}, \mathbf{X}_{(1)}, \dots, \mathbf{y}_{(l-1)}, \mathbf{X}_{(l-1)}$

③ Discard the B burn-in draws.

④ Compute and store posterior parameters w.r.t. variance components for the priors of the next cycle.

Appendix: Sensitivity Analysis of Genotype-Zone Effect

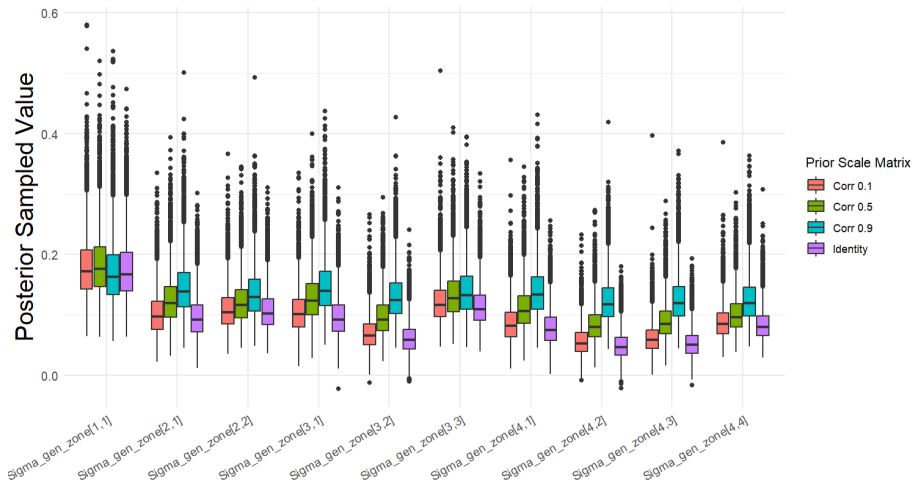


Figure: Box plots of posterior sampled individual variance and covariance components of $\Sigma_{\text{gen_zone}}$ matrix-valued posterior sample.