

Article

A Bayesian Hierarchical Semiparametric Approach to Modeling Interval-Valued Panel Data

Dengke Xu ^{1,*} , Mengen Qin ¹ and Ruiqin Tian ²¹ School of Economics, Hangzhou Dianzi University, Hangzhou 310018, China; 241150080@hdu.edu.cn² School of Mathematics, Hangzhou Normal University, Hangzhou 311121, China; ruiqintian@163.com

* Correspondence: xudengke@hdu.edu.cn

Abstract

This study proposes a Bayesian Hierarchical Semiparametric Center and Range (BHS-CRM) model to address nonlinearity and heterogeneity in interval-valued panel data. By incorporating Bayesian penalized B-splines (P-splines) into a hierarchical random effects structure, the proposed method utilizes a parallel dual-component framework to separately capture dynamic central tendencies and variability. Simulation studies demonstrate the model's robust performance, particularly in capturing nonlinear dynamics under high-noise conditions. Empirically, the model characterizes nonlinear temperature patterns in China's Carbon Emission Trading Scheme (ETS) and achieves competitive out-of-sample predictive accuracy.

Keywords: interval-valued panel data; Bayesian hierarchical model; semiparametric regression; P-splines; carbon price volatility

MSC: 62F15; 62G05; 62J05

1. Introduction

In fields such as finance, healthcare, and energy management, data complexity and uncertainty have become increasingly prominent, particularly when observations are characterized as intervals rather than single points. In financial markets, for instance, daily high and low prices serve as proxies for market volatility, while in healthcare, daily blood pressure ranges reflect physiological variability. These interval-valued data, by capturing both upper and lower bounds, offer a unique analytical framework for studying uncertainty—distinct from traditional point-valued data. However, modeling such data poses non-trivial challenges, particularly in reconciling the internal correlation between bounds, a task further complicated in panel data settings where temporal and cross-sectional dimensions interact. Conventional linear regression techniques, such as ordinary least squares (OLS), often fail to fully capture these dynamics, especially when underlying patterns exhibit nonlinearities manifested in variables such as trading volume or medication effects. This limitation necessitates advanced methodologies capable of addressing subject-specific random effects and nonlinear dynamics in interval-valued panel data, driven by the dual imperatives of theoretical advancement and practical applications, such as refining volatility forecasts or optimizing energy consumption predictions.

The analysis of interval-valued data began with the interval arithmetic of Moore [1], a foundational framework for imprecise measurements that spurred statistical innovation. Billard and Diday [2] advanced this with symbolic data analysis, treating intervals as



Academic Editor: Ioan M. Stancu-Minasian

Received: 28 July 2026

Revised: 13 September 2026

Accepted: 14 September 2026

Published: 17 September 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

structured objects rather than point collections, reshaping uncertainty modeling. Their Center and Min-Max Methods offered a simple start but missed bound interdependencies. Lima Neto and De Carvalho [3] introduced the Centre and Range Method (CRM), decomposing intervals into centers and ranges for a fuller representation. Building on this, Hao and Guo [4] proposed a constrained center and range joint model to prevent negative range predictions, improving the logical consistency of linear fits. More recently, Guan and Li [5] expanded this constrained CRM framework by incorporating an asymmetric Laplace distribution to handle residual uncertainties. Moving beyond linearity, Lim [6] introduced nonparametric additive models to capture more complex dependencies, while Kong et al. [7] developed a local linear smoothing approach to estimate interval bounds with greater flexibility. Sun et al. [8] addressed data heterogeneity via interval-valued functional clustering based on the Wasserstein distance, identifying distinct patterns within complex interval datasets. Furthermore, in applied research, extensive optimization algorithms have been developed to efficiently quantify and resolve interval uncertainties within complex engineering systems [9–11]. Despite these advances, such methods often fail to capture the complex trends, cycles, and nonlinear rhythms of time-varying data. This limitation is particularly evident when temporal and individual dimensions intersect, necessitating a framework to address these complexities.

In economic and financial applications, however, interval data analysis faces an additional layer of complexity: the dynamic interplay between time and individual heterogeneity. Unlike pure cross-sectional or time-series data, panel data addresses this multi-dimensional challenge by tracking multiple entities over time, as outlined by Baltagi [12]. However, introducing interval-valued observations into this framework complicates matters, as both the bounds vary dynamically across time and entities, demanding more than traditional linear tools. Ji et al. [13] advanced this field with nonlinear support vector regression for interval panel data, capturing complex temporal patterns missed by linear models. Yet, nonlinear relationships, as seen in stock volatility [14] or energy demand [15], often require greater flexibility. Liu et al. [16] introduced varying coefficient models, enhancing adaptability to shifting bounds, but limitations persist in fully addressing nonlinearity and heterogeneity in panel interval-valued data.

Bayesian methods provide a robust framework for addressing uncertainty, leveraging prior distributions and posterior inference, as detailed by Gelman et al. [17]. Zhang et al. [18] applied Bayesian quantile semiparametric regression, utilizing asymmetric Laplace distributions to flexibly model complex longitudinal data. In the context of interval-valued data, Xu and Qin [19] proposed a bivariate Bayesian framework to account for the internal correlation of bounds, which they further advanced with a parametrized method [20] to mitigate multicollinearity under high random error conditions. While these foundational works successfully address cross-sectional interval data, incorporating subject-specific random effects for panel settings requires drawing insights from broader longitudinal models. The growing vitality of this paradigm is evidenced by recent studies: Ferede et al. [21] and Langat et al. [22] demonstrated the efficacy of flexible Bayesian semiparametric mixed-effects models in handling complex longitudinal data, while Kowal and Wu [23] advanced efficient Monte Carlo inference strategies for such frameworks. In a related vein, Cui et al. [24] applied semiparametric Bayesian inference to interval-censored data, addressing complex patterns in settings like time-to-event analysis. Despite these strengths, Bayesian approaches for interval-valued panel data remain underexplored, particularly in handling nonlinear effects where linear assumptions fall short, necessitating more flexible modeling strategies.

Semiparametric methods offer a balanced approach to nonlinearity, combining parametric structure with nonparametric flexibility, as shown by Ruppert et al. [25] and Hastie

and Tibshirani [26] with generalized additive models (GAMs). Su and Ullah [27] extended these to panel data, capturing nonlinear effects with fixed effects, while Wood [28] refined GAMs for broader applications. In the context of environmental economics, Tian et al. [29] successfully utilized B-spline approximations in semiparametric varying-coefficient panel models to analyze carbon emission datasets. Additionally, Xu et al. [30] emphasized the importance of jointly modeling the mean and covariance structures in longitudinal data analysis to ensure estimation efficiency. However, the application of such semiparametric techniques to interval-valued panel data remains limited, leaving untapped potential for modeling nonlinearity and heterogeneity. Motivated by this untapped potential, this paper proposes a Bayesian semiparametric framework with random effects, enhancing precision in capturing the nonlinear dynamics and heterogeneity of interval-valued panel data.

Specifically, the main contributions of this study lie in the structural improvement of the modeling framework, which can be summarized in three aspects. First, we extend the Bayesian hierarchical framework to the analysis of interval-valued panel data. By employing a parallel dual-component (center and range) structure, the proposed model effectively accommodates both individual heterogeneity and temporal dynamics inherent in complex panel datasets. Second, we introduce a nonlinear non-parametric component into Bayesian interval data analysis. This structural advancement provides the essential flexibility required to model nonlinear relationships within interval bounds, overcoming the strict and often unrealistic linearity assumptions of traditional interval models. Third, from an algorithmic perspective, we handle the non-parametric component by integrating penalized B-splines (P-splines) with random walk priors. This specific design yields sparse precision matrices and enables a fully conjugate MCMC sampling framework. It not only avoids the convergence pitfalls of constrained optimization but also rigorously propagates structural uncertainty into the final predictions, yielding geometrically more robust interval reconstructions under high-noise environments.

The rest of this paper is structured as follows. Section 2 formulates the Bayesian Hierarchical Semiparametric Center and Range Model (BHS-CRM). Section 3 outlines the posterior inference procedure and provides the specific Gibbs sampling algorithm for parameter estimation. To validate the proposed method, Section 4 conducts Monte Carlo simulations to assess its finite sample performance and robustness against competing benchmarks. Section 5 applies the framework to empirical data from China's Carbon ETS to investigate the nonlinear effects of temperature on carbon price volatility. Finally, Section 6 presents the conclusion and discussion.

2. Methodology

The core contribution of this paper is the development of a Bayesian hierarchical framework that employs a parallel dual-component structure to analyze the center and range of longitudinal interval-valued responses. By integrating a semiparametric structure with random effects, the proposed approach simultaneously captures linear covariate effects, time-varying nonlinear dynamics, and subject-specific heterogeneity within a unified probabilistic model.

2.1. Model Specification

This paper proposes a Bayesian Hierarchical Semiparametric Center and Range Model (BHS-CRM). Let the observed outcome for individual i at time t be an interval-valued variable denoted by $y_{it} = [y_{it}^L, y_{it}^U]$, where y_{it}^L and y_{it}^U represent the lower and upper bounds, respectively. To capture the underlying dynamics, we decompose the interval outcome into two real-valued components: the center y_{it}^c and the range y_{it}^r . These are defined as:

$$y_{it}^c = \frac{y_{it}^L + y_{it}^U}{2}, \quad y_{it}^r = y_{it}^U - y_{it}^L. \quad (1)$$

Subsequently, for each component $s \in \{c, r\}$, the response y_{it}^s for individual i ($i = 1, \dots, n$) at time point t ($t = 1, \dots, T$) is formulated as a semiparametric random effects regression:

$$y_{it}^s = (\mathbf{X}_{it}^s)^\top \boldsymbol{\beta}^s + f^s(Z_{it}) + u_i^s + \epsilon_{it}^s, \quad (2)$$

where $\mathbf{X}_{it}^s$ denotes the vector of covariates associated with the linear fixed effects $\boldsymbol{\beta}^s$. The terms u_i^s and ϵ_{it}^s represent the subject-specific random intercepts and the idiosyncratic errors, respectively. We assume that they follow independent Gaussian distributions: $u_i^s \sim N(0, \sigma_{u,s}^2)$ and $\epsilon_{it}^s \sim N(0, \sigma_{\epsilon,s}^2)$. The term $f^s(\cdot)$ is an unknown smooth function capturing the nonlinear influence of the covariate Z_{it} .

To flexibly estimate the unknown nonlinear function $f^s(\cdot)$, we adopt the Bayesian penalized B-splines (P-splines) approach. This method builds upon the foundational framework introduced by Eilers and Marx [31], while adopting the Bayesian inferential paradigm developed by Lang and Brezger [32]. Unlike unpenalized B-splines, whose performance is notoriously sensitive to the number and placement of knots, P-splines avoid the knot-selection dilemma by utilizing a large number of equally spaced knots combined with a discrete difference penalty on adjacent coefficients. Specifically, the smooth function $f^s(Z_{it})$ is approximated as a linear combination of K B-spline basis functions defined on a set of equally spaced knots:

$$f^s(Z_{it}) = \sum_{\rho=1}^K B_\rho(Z_{it}) \theta_\rho^s = \mathbf{b}^s(Z_{it}) \boldsymbol{\theta}^s, \quad (3)$$

where $\mathbf{b}^s(Z_{it})$ denotes the $1 \times K$ row vector of basis functions evaluated at Z_{it} . These basis functions are constructed using cubic B-splines of degree $q = 3$. $\boldsymbol{\theta}^s$ represents the $K \times 1$ vector of unknown spline coefficients with elements $\theta_1^s, \dots, \theta_K^s$. In matrix notation for the full sample, this formulation aggregates to $\mathbf{f}^s = \mathbf{B}^s \boldsymbol{\theta}^s$, where $\mathbf{B}^s$ is the $N \times K$ design matrix with rows corresponding to $\mathbf{b}^s(Z_{it})$.

In accordance with the P-splines framework, the smoothness of the fitted curves is achieved by imposing discrete difference penalties on adjacent B-spline coefficients. Specifically, we employ second-order ($k = 2$) stochastic difference constraints on $\boldsymbol{\theta}^s$:

$$\Delta^k \theta_\rho^s \sim N(0, \tau_s^2), \quad (4)$$

which is equivalent to specifying a random walk process for the coefficients. The variance parameter τ_s^2 acts as the inverse smoothing parameter, controlling the trade-off between fidelity to the data and the roughness of the function estimate. This hierarchical formulation allows for the simultaneous estimation of the smooth functions and complexity parameters. Crucially, unlike continuous derivative-based penalties that yield dense matrices via integration, the discrete difference operator produces a highly sparse, banded precision matrix. This sparse structure significantly accelerates the matrix inversions required during Gibbs sampling, ensuring computational scalability and numerical stability for our high-dimensional hierarchical model.

The use of second-order difference penalties introduces a rank deficiency in the associated precision matrix $\mathbf{P}$, where $\text{rank}(\mathbf{P}) = K - 2$. This rank deficiency causes an inherent unidentifiability between the model intercept and the mean level of the nonlinear functions $f^c(Z)$ and $f^r(Z)$. To resolve this issue, this paper adopts a static basis-centering strategy to ensure that the parameters can be uniquely estimated. Specifically, let $\mathbf{b}^s(Z_{it})$ denote

the $1 \times K$ row vector of basis functions evaluated at observation Z_{it} . We first compute the sample mean vector of these basis evaluations across all N observations:

$$\bar{\mathbf{b}}^s = \frac{1}{N} \sum_{i=1}^n \sum_{t=1}^T \mathbf{b}^s(Z_{it}),$$

which is a $1 \times K$ row vector. Then, the centered basis matrix $\mathbf{B}^{s*}$ (of dimension $N \times K$) is constructed by subtracting these means from the original basis matrix:

$$\mathbf{B}^{s*} = \mathbf{B}^s - \mathbf{1}_N \bar{\mathbf{b}}^s,$$

with $\mathbf{1}_N$ denoting an $N \times 1$ vector of ones. This transformation implies the constraint $\sum_{i,t} f^s(Z_{it}) = 0$, restricting the nonlinear component to represent functional deviations from the global mean. Consequently, the covariate vector $\mathbf{X}_{it}^s$ must explicitly include an intercept term to capture the baseline level.

Substituting the centered basis into the original framework yields the final hierarchical formulation for the center and range submodels:

$$y_{it}^c = (\mathbf{X}_{it}^c)^\top \boldsymbol{\beta}^c + \mathbf{B}_{it}^{s*} \boldsymbol{\theta}^c + u_i^c + \epsilon_{it}^c, \quad (5)$$

$$y_{it}^r = (\mathbf{X}_{it}^r)^\top \boldsymbol{\beta}^r + \mathbf{B}_{it}^{r*} \boldsymbol{\theta}^r + u_i^r + \epsilon_{it}^r. \quad (6)$$

Here, $\mathbf{X}^c$ and $\mathbf{X}^r$ denote the full $N \times p_s$ design matrices for the fixed effects, explicitly including an intercept term. The term $\mathbf{B}_{it}^{s*}$ represents the $1 \times K$ row vector of centered spline basis evaluations for the observation at time t of subject i . Regarding the range model (y^r), while strictly non-negative distributions are theoretically preferable to enforce the non-negativity constraint, adopting them in our complex hierarchical framework would destroy the conjugacy of the Gibbs sampler. Therefore, a Gaussian specification is adopted in this study as a pragmatic, computationally efficient approximation. To rigorously enforce the non-negativity of the predicted range, as emphasized by Hao and Guo [4], our algorithm incorporates a standard post-hoc truncation safeguard, $\hat{y}^r = \max(0, \hat{y}^{r, \text{Gaussian}})$. Furthermore, following the CRM paradigm [3], the center and range components are modeled independently. This conditional independence preserves the conjugacy of the Gibbs sampler, thereby ensuring computational efficiency. Incorporating the potential correlation between these components is left as a direction for future research.

2.2. Prior Distributions

Within the Bayesian inferential framework, posterior sampling is implemented via Markov chain Monte Carlo (MCMC). For the collection of model parameters

$$\boldsymbol{\Phi} = \{\boldsymbol{\beta}^c, \boldsymbol{\beta}^r, \boldsymbol{\theta}^c, \boldsymbol{\theta}^r, \sigma_{\epsilon,c}^2, \sigma_{\epsilon,r}^2, \sigma_{u,c}^2, \sigma_{u,r}^2, \tau_c^2, \tau_r^2\},$$

the following prior specifications are adopted to reflect limited prior knowledge while ensuring posterior propriety.

2.2.1. Fixed Effects $\boldsymbol{\beta}$

The fixed-effect coefficient vectors $\boldsymbol{\beta}^c$ and $\boldsymbol{\beta}^r$ are assigned weakly informative multivariate normal priors:

$$\boldsymbol{\beta}^c \sim N(\boldsymbol{\mu}_0^c, \boldsymbol{\Sigma}_0^c), \quad \boldsymbol{\beta}^r \sim N(\boldsymbol{\mu}_0^r, \boldsymbol{\Sigma}_0^r),$$

where we set $\boldsymbol{\mu}_0^c = \mathbf{0}$ and $\boldsymbol{\mu}_0^r = \mathbf{0}$. The covariance matrices $\boldsymbol{\Sigma}_0^c$ and $\boldsymbol{\Sigma}_0^r$ are set to be diagonal with large variance entries (e.g., 10^3) to ensure the priors are sufficiently flat over the parameter space.

2.2.2. Variance Parameters

The observation error variances ($\sigma_{\epsilon,c}^2, \sigma_{\epsilon,r}^2$) and the random effects variances ($\sigma_{u,c}^2, \sigma_{u,r}^2$) are assigned inverse-gamma priors, which are conjugate with the normal likelihoods and facilitate efficient Gibbs updates:

$$\sigma_{\epsilon,s}^2 \sim \text{IG}(a, b), \quad \sigma_{u,s}^2 \sim \text{IG}(a_u, b_u), \quad s \in \{c, r\},$$

where $\text{IG}(\cdot, \cdot)$ denotes the inverse-gamma distribution. Here, $\sigma_{\epsilon,s}^2$ denotes the variance of the idiosyncratic errors ϵ_{it}^s , while $\sigma_{u,s}^2$ denotes the variance of the subject-specific random intercepts u_i^s . To maintain numerical stability and avoid improper posteriors in finite samples, these hyperparameters are consistently set to weakly informative values, specifically $a = b = a_u = b_u = 0.01$.

2.2.3. Spline Smoothing Parameters τ_s^2

As detailed in Section 2.1, the spline coefficients θ^s are assigned second-order random walk priors conditional on the smoothing parameters τ_s^2 . To complete the Bayesian hierarchy, we assign conjugate inverse-gamma hyperpriors to these smoothing variances:

$$\tau_s^2 \sim \text{IG}(a_\tau, b_\tau), \quad s \in \{c, r\}.$$

For the smoothing variance hyperparameters (a_τ, b_τ), we adopt values that balance flexibility with regularization, ensuring the algorithm converges robustly even in scenarios with high noise or small sample sizes. Specifically, we set $a_\tau = 1$ and $b_\tau = 0.005$ in our overall implementation.

To intuitively summarize the entire model specification discussed above, Figure 1 illustrates the overarching hierarchical structure of the proposed BHS-CRM framework, from the data transformation level up to the Bayesian priors.

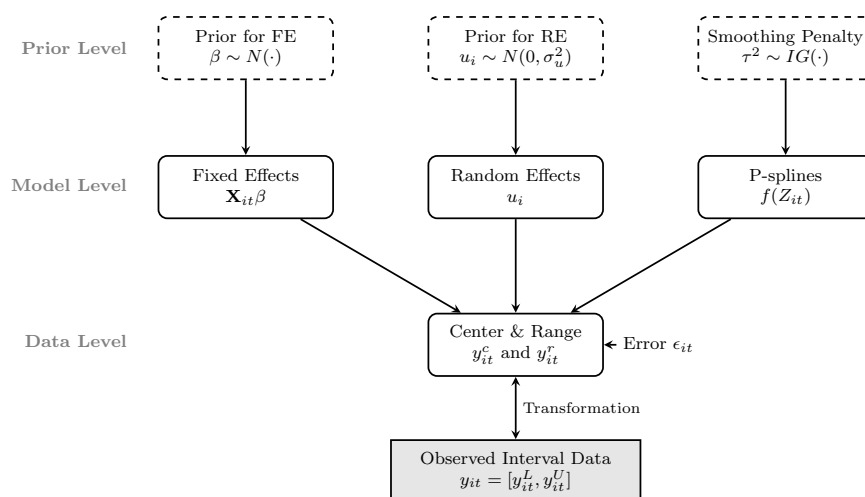


Figure 1. The hierarchical structure diagram of the proposed BHS-CRM framework.

3. Model Estimation

Posterior inference is conducted via Gibbs sampling. The Gibbs sampler leverages the fact that the full conditional posterior distributions of the primary model blocks belong to standard families, which avoids the need for Metropolis–Hastings steps and improves computational efficiency. Because the center and range components share the same structural form, a unified notation $s \in \{c, r\}$ is used below to present full conditional distributions concisely. For matrix notation, let $\mathbf{y}^s = (y_{11}^s, \dots, y_{1T}^s, \dots, y_{n1}^s, \dots, y_{nT}^s)^\top$ and $\mathbf{u}^s = (u_1^s, \dots, u_n^s)^\top$

denote the stacked response vector and random effects vector, respectively. Unlike standard frequentist methods that plug in deterministic point estimates for variance components, this fully Bayesian MCMC approach marginalizes over their joint posterior distribution. This fully propagates parameter uncertainty into the final estimates, effectively preventing overconfident credible intervals.

3.1. Full Conditional Distribution of the Fixed Effects β^s

Conditional on the remaining parameters Φ_{rest} and data $\mathbf{y}^s$, the full conditional distribution of β^s is Gaussian:

$$\beta^s \mid \Phi_{\text{rest}}, \mathbf{y}^s \sim N(\tilde{\mu}_{\beta^s}, \tilde{\Sigma}_{\beta^s}),$$

with

$$\begin{aligned} \tilde{\Sigma}_{\beta^s} &= \left(\frac{1}{\sigma_{\epsilon,s}^2} (\mathbf{X}^s)^\top \mathbf{X}^s + (\Sigma_0^s)^{-1} \right)^{-1}, \\ \tilde{\mu}_{\beta^s} &= \tilde{\Sigma}_{\beta^s} \left(\frac{1}{\sigma_{\epsilon,s}^2} (\mathbf{X}^s)^\top \tilde{\mathbf{y}}_\beta^s + (\Sigma_0^s)^{-1} \mu_0^s \right), \end{aligned}$$

where $\tilde{\mathbf{y}}_\beta^s = \mathbf{y}^s - \mathbf{B}^{s*} \theta^s - \mathbf{D} \mathbf{u}^s$ represents the partial residuals for the fixed effects. Here, $\mathbf{D}$ is the $N \times n$ incidence matrix that maps the subject-specific random effects $\mathbf{u}^s$ to the corresponding observations in the response vector.

3.2. Full Conditional Distribution of the Spline Coefficients θ^s

The full conditional distribution of the spline coefficient vector is also Gaussian:

$$\theta^s \mid \Phi_{\text{rest}}, \mathbf{y}^s \sim N(\tilde{\mu}_{\theta^s}, \tilde{\Sigma}_{\theta^s}),$$

with

$$\begin{aligned} \tilde{\Sigma}_{\theta^s} &= \left(\frac{1}{\sigma_{\epsilon,s}^2} (\mathbf{B}^{s*})^\top \mathbf{B}^{s*} + \frac{1}{\tau_s^2} \mathbf{P} \right)^{-1}, \\ \tilde{\mu}_{\theta^s} &= \tilde{\Sigma}_{\theta^s} \left(\frac{1}{\sigma_{\epsilon,s}^2} (\mathbf{B}^{s*})^\top \tilde{\mathbf{y}}_\theta^s \right), \end{aligned}$$

where $\tilde{\mathbf{y}}_\theta^s = \mathbf{y}^s - \mathbf{X}^s \beta^s - \mathbf{D} \mathbf{u}^s$. In practical computation, a minute ridge term $\delta \mathbf{I}$ (e.g., $\delta = 10^{-6}$) is typically added to the precision matrix prior to inversion to ensure strict positive definiteness and numerical stability. This standard computational safeguard does not affect the theoretical posterior inference.

3.3. Full Conditional Distribution of the Random Effects u_i^s

Subject-specific random effects u_i^s are conditionally independent and follow Gaussian full conditional distributions:

$$u_i^s \mid \Phi_{\text{rest}}, \mathbf{y}^s \sim N(\tilde{\mu}_{u_i^s}, (\tilde{\sigma}_{u_i^s})^2).$$

For subject i , the posterior variance and mean are

$$\begin{aligned} (\tilde{\sigma}_{u_i^s})^2 &= \left(\frac{T}{\sigma_{\epsilon,s}^2} + \frac{1}{\sigma_{u,s}^2} \right)^{-1}, \\ \tilde{\mu}_{u_i^s} &= (\tilde{\sigma}_{u_i^s})^2 \left(\frac{1}{\sigma_{\epsilon,s}^2} \sum_{t=1}^T e_{it}^s \right), \end{aligned}$$

where $e_{it}^s = y_{it}^s - (\mathbf{X}_{it}^s)^\top \beta^s - \mathbf{B}_{it}^{s*} \theta^s$ denotes the partial residual for individual i at time t excluding the random effect contribution.

3.4. Full Conditional Distributions of Variance and Smoothing Parameters

Due to conjugacy, the reciprocals of variance parameters follow Gamma distributions. Throughout this section, the notation $x \sim \text{Gamma}(a, b)$ implies a density function $p(x) \propto x^{a-1} \exp(-bx)$, where a is the shape parameter and b is the rate parameter.

3.4.1. Residual Variances $\sigma_{\epsilon,s}^2$

The full conditional distribution of the precision parameter $\frac{1}{\sigma_{\epsilon,s}^2}$ is

$$\frac{1}{\sigma_{\epsilon,s}^2} \sim \text{Gamma}(\tilde{c}^s, \tilde{d}^s),$$

where

$$\begin{aligned} \tilde{c}^s &= a + \frac{N}{2}, \\ \tilde{d}^s &= b + \frac{1}{2}(\mathbf{R}^s)^\top \mathbf{R}^s, \end{aligned}$$

and $\mathbf{R}^s = \mathbf{y}^s - \mathbf{X}^s \boldsymbol{\beta}^s - \mathbf{B}^{s*} \boldsymbol{\theta}^s - \mathbf{D} \mathbf{u}^s$ is the full residual vector. Note that a and b are the hyperparameters defined in Section 2.2.

3.4.2. Random Effects Variances $\sigma_{u,s}^2$

The full conditional distribution of $\frac{1}{\sigma_{u,s}^2}$ is

$$\frac{1}{\sigma_{u,s}^2} \sim \text{Gamma}(\tilde{a}_u^s, \tilde{b}_u^s),$$

with

$$\tilde{a}_u^s = a_u + \frac{n}{2}, \quad \tilde{b}_u^s = b_u + \frac{1}{2} \sum_{i=1}^n (u_i^s)^2.$$

Here, a_u and b_u denote the hyperparameters for the random effects variance defined in Section 2.2.

3.4.3. Smoothing Parameters τ_s^2

The full conditional distribution of the inverse smoothing parameter $\frac{1}{\tau_s^2}$ is

$$\frac{1}{\tau_s^2} \sim \text{Gamma}(\tilde{a}_\tau^s, \tilde{b}_\tau^s),$$

where

$$\tilde{a}_\tau^s = a_\tau + \frac{K-2}{2}, \quad \tilde{b}_\tau^s = b_\tau + \frac{1}{2}(\boldsymbol{\theta}^s)^\top \mathbf{P} \boldsymbol{\theta}^s.$$

Here, the penalty matrix $\mathbf{P}$ is constructed as $\mathbf{P} = \Delta_2^\top \Delta_2$, where Δ_2 denotes the $(K-2) \times K$ second-order difference matrix:

$$\Delta_2 = \begin{pmatrix} 1 & -2 & 1 & 0 & \cdots & 0 \\ 0 & 1 & -2 & 1 & \cdots & 0 \\ \vdots & \ddots & \ddots & \ddots & \ddots & \vdots \\ 0 & \cdots & 0 & 1 & -2 & 1 \end{pmatrix}.$$

The shape parameter increment $(K-2)/2$ reflects the rank deficiency of $\mathbf{P}$, determined by the order of the random walk prior, given that $\text{rank}(\mathbf{P}) = K-2$.

3.5. Computational Algorithm

Algorithm 1 summarizes the systematic Gibbs sampling procedure, where each parameter block is updated conditionally based on the derived distributions. In practice, the Gibbs sampler is initialized using OLS estimates for the fixed effects and zeros for the spline coefficients. Upon convergence, the resulting posterior samples enable statistical inference and forecasting. Although the MCMC algorithm inherently takes longer to run than standard frequentist estimation, the overall computation time remains highly tractable and is a reasonable trade-off for rigorous uncertainty quantification.

All computational procedures were implemented using R software (version 4.4.3; R Core Team, Vienna, Austria) and RStudio (version 2025.05.0; Posit Software, PBC, Boston, MA, USA).

Algorithm 1 Gibbs sampling algorithm for the BHS-CRM.

Gibbs sampling strategy for the parameter set Φ

Input: Observed data $\mathbf{y}^s, \mathbf{X}^s, \mathbf{D}$ for $s \in \{c, r\}$; hyperparameters $a, b, a_u, b_u, \dots$; initial values $\Phi^{(0)}$.

Output: Posterior MCMC sample sequence $\Phi^{(1)}, \Phi^{(2)}, \dots, \Phi^{(S)}$.

Procedure:

For $g = 1$ to S do:

For each component $s \in \{c, r\}$ do:

1. Sample fixed effects: $\beta^{s(g)} \sim N(\tilde{\mu}_{\beta^s}, \tilde{\Sigma}_{\beta^s})$
2. Sample spline coefficients: $\theta^{s(g)} \sim N(\tilde{\mu}_{\theta^s}, \tilde{\Sigma}_{\theta^s})$
3. Sample random effects: $\mathbf{u}^{s(g)} \sim N(\tilde{\mu}_{u^s}, \tilde{\Sigma}_{u^s})$
4. Sample residual error variance: $1/\sigma_s^{2(g)} \sim \text{Gamma}(\tilde{c}^s, \tilde{d}^s)$
5. Sample random effects variance: $1/\sigma_{u^s}^{2(g)} \sim \text{Gamma}(\tilde{a}_u^s, \tilde{b}_u^s)$
6. Sample smoothing parameter: $1/\tau_s^{2(g)} \sim \text{Gamma}(\tilde{a}_\tau^s, \tilde{b}_\tau^s)$

End For

End For

4. Simulation Study

In this section, we investigate the performance of the proposed BHS-CRM via Monte Carlo simulation. The analysis is designed in two parts. First, we conduct an internal validation of the BHS-CRM to assess its parameter estimation accuracy, convergence properties, and non-parametric goodness-of-fit.

Second, we perform a comprehensive comparative analysis, benchmarking the BHS-CRM against two standard frequentist models built upon the CRM framework: the LME-CRM and the PLME-CRM. Crucially, the PLME-CRM shares the exact same structural specification as our proposed BHS-CRM, as defined in Equation (2). This ensures that both models are evaluated on an identical functional basis, thereby isolating the performance gains resulting specifically from the Bayesian hierarchical estimation framework in handling parameter uncertainty and complex variance structures within interval-valued panel data.

4.1. Measurements

To evaluate model performance, we use two categories of metrics. First, we define interval prediction accuracy metrics to assess the accuracy of the reconstructed interval $\hat{y}_{it} = [\hat{y}_{it}^L, \hat{y}_{it}^U]$. Note that the predicted lower and upper bounds are derived from the estimated center and range components via $\hat{y}_{it}^L = \hat{y}_{it}^c - \hat{y}_{it}^r/2$ and $\hat{y}_{it}^U = \hat{y}_{it}^c + \hat{y}_{it}^r/2$, respectively. Second, we define sub-model diagnostic metrics to assess the individual performance of the center ($s = c$) and range ($s = r$) estimations.

- (1) The interval prediction accuracy metrics include the Root Mean Square Error for the lower bound (RMSE_L), upper bound (RMSE_U), and the combined interval bounds (RMSE_H), as well as the Relative Intersection score (RI):

$$\begin{aligned}\text{RMSE}_L &= \sqrt{\frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} (\hat{y}_{it}^L - y_{it}^L)^2}, \\ \text{RMSE}_U &= \sqrt{\frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} (\hat{y}_{it}^U - y_{it}^U)^2}, \\ \text{RMSE}_H &= \sqrt{\frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} (|\hat{y}_{it}^U - y_{it}^U| + |\hat{y}_{it}^L - y_{it}^L|)^2}, \\ \text{RI} &= \frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} \frac{\omega(y_{it} \cap \hat{y}_{it})}{\omega(y_{it} \cup \hat{y}_{it})},\end{aligned}$$

where $\omega(\cdot)$ denotes the interval width. Conceptually representing the intersection over union, the RI strictly penalizes both misaligned bounds and excessive widths, reaching a maximum of 1 only upon perfect alignment.

- (2) The sub-model diagnostic metrics include the Coefficient of Determination (R_s^2), the Root Mean Square Error (RMSE_s), and the Root Average Squared Error (RASE_s) for the non-parametric component, defined for $s \in \{c, r\}$:

$$\begin{aligned}R_s^2 &= 1 - \frac{\sum_{i,t \in \text{Test}} (y_{it}^s - \hat{y}_{it}^s)^2}{\sum_{i,t \in \text{Test}} (y_{it}^s - \bar{y}_{\text{test}}^s)^2}, \\ \text{RMSE}_s &= \sqrt{\frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} (y_{it}^s - \hat{y}_{it}^s)^2}, \\ \text{RASE}_s &= \sqrt{\frac{1}{M} \sum_{k=1}^M (\hat{f}^s(z_k) - f_{\text{true}}^s(z_k))^2},\end{aligned}$$

where z_k are M grid points uniformly distributed over the domain $[0, 1]$ of the covariate Z , and $f_{\text{true}}^s(\cdot)$ denotes the true generating function used in the simulation design.

4.2. Data Generating Process

To evaluate the proposed BHS-CRM, four distinct simulation configurations (S1 through S4) are designed to capture varying conditions of linearity and noise levels, as summarized in Table 1. The data are generated based on the semiparametric random effects specification for each component $s \in \{c, r\}$:

$$y_{it}^s = (\mathbf{X}_{it}^s)^\top \boldsymbol{\beta}^s + f^s(Z_{it}) + u_i^s + \epsilon_{it}^s.$$

The true fixed-effect regression coefficients are set to $\boldsymbol{\beta}^c = (8, 7, 6)^\top$ and $\boldsymbol{\beta}^r = (4, 6, 5)^\top$. The design vector is $\mathbf{X}_{it}^s = (1, X_{it2}, X_{it3})^\top$, where the covariates X_{itj} (for $j = 2, 3$) are drawn from a uniform distribution $U(0, 2)$. The time-varying covariate Z_{it} is drawn from $U(0, 1)$. Subject-specific random effects are generated as $u_i^s \sim N(0, 0.5)$. Finally, the simulated interval bounds are deterministically constructed as $y_{it}^L = y_{it}^c - y_{it}^r/2$ and $y_{it}^U = y_{it}^c + y_{it}^r/2$.

In the nonlinear scenarios (S1, S2), the true benchmark intercepts are adjusted by the theoretical mean of the generating function to account for the zero-mean identifiability constraint imposed on the B-spline basis. In the linear scenarios (S3, S4), the true functions are set to $f^c(z) = f^r(z) = 0$ to assess the model's ability to shrink the spline components towards zero. We simulate 100 replications for sample sizes $n \in \{50, 150\}$ with $T = 12$.

In each replication, the Gibbs sampler runs for 5000 iterations, discarding the first 2500 as burn-in.

To evaluate predictive performance, we employ a cross-sectional hold-out validation strategy: in each replication, subjects are randomly partitioned into a training set (75%) and a test set (25%). The model is estimated using the full time series of the training subjects. For the hold-out subjects, out-of-sample predictions are strictly generated using the marginal mean structure, setting the unobserved random effects to their prior expected value of zero.

For the B-spline component, we set the basis dimension to $K = 7$. Unlike unpenalized regression splines, the P-splines framework is inherently robust to the choice of K , as curve smoothness is adaptively regularized by the Bayesian penalty parameter τ_s^2 rather than the knot count [31,32]. To empirically validate this, we tested larger knot dimensions ($K = 10$ and $K = 20$) in the nonlinear regimes (S1, S2). As summarized in Table 2, while added parameter complexity slightly increases the RASE, the out-of-sample RMSE remains highly stable. This confirms that the Bayesian second-order random walk prior effectively prevents overfitting, validating $K = 7$ as an optimal specification that balances estimation accuracy with computational efficiency.

Table 1. Comprehensive simulation design and parameter configurations.

Panel A: Common Data Generating Process and Parameters				
Data Generating Process: $y_{it}^s = \beta_1^s + X_{it2}\beta_2^s + X_{it3}\beta_3^s + f^s(Z_{it}) + u_i^s + \epsilon_{it}^s$ ($s \in \{c, r\}$)				
Category	Parameter/Variable		Configuration	
Dimensions	Sample sizes (n, T)		$n \in \{50, 150\}, \quad T = 12$	
Covariates	Linear (X_{it2}, X_{it3})		$U(0, 2)$	
	Smoothing (Z_{it})		$U(0, 1)$	
Fixed Effects	Center model (β^c)		$(8, 7, 6)^\top$	
	Range model (β^r)		$(4, 6, 5)^\top$	
Random Effects	Subject-specific (u_i^s)		$N(0, 0.5)$	
Panel B: Scenario-Specific Configurations				
Scenario	Linearity	Noise Level	Non-parametric Function $f^s(Z)$	Error Dist. ϵ_{it}^s
S1	Nonlinear	Low	$2 \sin(\pi Z_{it})$	$N(0, \sigma_{\epsilon, s}^2 = 0.1)$
S2	Nonlinear	High	$2 \sin(\pi Z_{it})$	$N(0, \sigma_{\epsilon, s}^2 = 0.5)$
S3	Linear	Low	0	$N(0, \sigma_{\epsilon, s}^2 = 0.1)$
S4	Linear	High	0	$N(0, \sigma_{\epsilon, s}^2 = 0.5)$

Table 2. Sensitivity of predictive performance to the number of knots K ($n = 150$).

Scenario	Knots	RASE _c	RASE _r	RMSE _c	RMSE _r
S1	$K = 7$	0.034	0.034	0.785	0.780
	$K = 10$	0.037	0.042	0.774	0.775
	$K = 20$	0.044	0.046	0.768	0.795
S2	$K = 7$	0.052	0.051	1.007	1.007
	$K = 10$	0.058	0.063	1.000	1.003
	$K = 20$	0.072	0.073	0.997	1.019

4.3. Internal Validation of Estimation Performance

In this subsection, we assess the statistical properties of the BHS-CRM using the simulated datasets. The evaluation focuses on three key aspects: parameter estimation accuracy, MCMC convergence, and the goodness-of-fit for the non-parametric component.

- (1) **Parameter Estimation Accuracy:** Tables 3 and 4 show the Bayesian estimates of the BHS-CRM parameters when $n = 50, 150$ and T is fixed at 12, based on 100 repetitions. The accuracy of the estimation is expressed by BIAS and standard deviation (SD). Comparing configurations S1/S3 (low noise, $\sigma_{\epsilon,s}^2 = 0.1$) with S2/S4 (high noise, $\sigma_{\epsilon,s}^2 = 0.5$), the SDs for S1 and S3 are consistently smaller than those for S2 and S4. This indicates that, as expected, the smaller the variance of the error term, the more precise the parameter estimates. Furthermore, as the sample size increases from $n = 50$ to $n = 150$, the accuracy of Bayesian estimation gradually improves. For example, in configuration S1, the SD of β_1^c decreases from 0.1257 to 0.0802, and the SD for β_1^r decreases from 0.1252 to 0.0666. This trend of increasing precision is consistent across all parameters and configurations. In general, the BIAS and SD of all parameters remain at low levels, indicating that the Bayesian estimation performance is robust and the estimators are consistent.

Table 3. The results of Bayesian estimation (BHS-CRM) when $n = 50$.

Config.	Para.	β_1^c	β_2^c	β_3^c	$\sigma_{\epsilon,c}^2$	$\sigma_{u,c}^2$	β_1^r	β_2^r	β_3^r	$\sigma_{\epsilon,r}^2$	$\sigma_{u,r}^2$
S1	BIAS	0.0138	0.0020	−0.0001	0.0067	0.0043	0.0120	−0.0005	0.0001	0.0079	0.0085
	SD	0.1257	0.0301	0.0307	0.0066	0.0980	0.1252	0.0245	0.0301	0.0072	0.1119
S2	BIAS	−0.0063	0.0153	0.0020	0.0027	−0.0111	0.0015	−0.0011	0.0099	−0.0019	0.0042
	SD	0.1605	0.0598	0.0554	0.0314	0.1224	0.1344	0.0633	0.0572	0.0386	0.1102
S3	BIAS	−0.0098	0.0043	0.0006	0.0078	0.0109	−0.0079	0.0036	−0.0024	0.0076	0.0095
	SD	0.1443	0.0264	0.0249	0.0067	0.1048	0.1238	0.0247	0.0283	0.0075	0.1173
S4	BIAS	−0.0153	0.0086	0.0007	0.0034	−0.0005	−0.0004	−0.0005	−0.0040	0.0024	−0.0185
	SD	0.1335	0.0490	0.0615	0.0308	0.1215	0.1549	0.0602	0.0646	0.0344	0.1163

Table 4. The results of Bayesian estimation (BHS-CRM) when $n = 150$.

Config.	Para.	β_1^c	β_2^c	β_3^c	$\sigma_{\epsilon,c}^2$	$\sigma_{u,c}^2$	β_1^r	β_2^r	β_3^r	$\sigma_{\epsilon,r}^2$	$\sigma_{u,r}^2$
S1	BIAS	0.0082	0.0035	0.0008	0.0027	−0.0078	−0.0024	0.0004	0.0000	0.0027	−0.0052
	SD	0.0802	0.0168	0.0152	0.0041	0.0628	0.0666	0.0149	0.0154	0.0045	0.0634
S2	BIAS	0.0084	0.0007	−0.0015	0.0021	−0.0115	0.0056	−0.0008	0.0014	0.0008	0.0067
	SD	0.0931	0.0359	0.0357	0.0199	0.0701	0.0835	0.0350	0.0371	0.0196	0.0706
S3	BIAS	0.0114	0.0020	−0.0022	0.0020	−0.0111	−0.0078	−0.0001	0.0022	0.0020	−0.0098
	SD	0.0717	0.0160	0.0173	0.0039	0.0682	0.0720	0.0162	0.0144	0.0038	0.0600
S4	BIAS	−0.0109	−0.0001	−0.0014	0.0007	0.0005	−0.0115	−0.0043	0.0022	0.0010	−0.0065
	SD	0.0833	0.0377	0.0325	0.0182	0.0737	0.0850	0.0333	0.0383	0.0217	0.0662

- (2) **Convergence Diagnostics:** To investigate the convergence of the proposed algorithm, we visualize the trace plots of 5000 iteration updates for the final configuration (S4, $n = 150$), and the results are shown in Figures 2 and 3. It can be seen that the algorithm converges quickly; the samples exhibit stable, horizontal bands with no discernible trends. The trace plots for both the center and range model parameters demonstrate good mixing and stability after the burn-in period. Note that the smoothing parameter τ_s^2 exhibits characteristic spikes, which is expected behavior for P-splines with random walk priors as it adaptively explores the roughness penalty space. To further rigorously validate convergence across all scenarios, we simulated three parallel chains with highly dispersed initial values for the $n = 150$ configurations. As reported in Table 5, the Gelman-Rubin statistics ($\hat{R}$) for the variance components ($\sigma_{\epsilon,c}^2, \sigma_{u,c}^2$), the smoothing parameter (τ_c^2), and all fixed-effect coefficients ($\beta_1^c, \beta_2^c, \beta_3^c$) are

strictly below the convergence threshold of 1.1 recommended by Gelman et al. [17]. This confirms robust convergence independent of initialization. While the intercept (β_1^c) naturally exhibits slightly higher autocorrelation due to its inherent functional dependence on the centered P-spline basis constraint, its $\hat{R}$ statistic adequately satisfies the standard inferential requirements even in low-noise configurations. Overall, the primary structural and variance parameters mix well across all settings.

Table 5. Multi-chain MCMC Convergence Diagnostics ($\hat{R}$) across Configurations ($n = 150$).

Config.	$\sigma_{\epsilon,c}^2$	$\sigma_{u,c}^2$	τ_c^2	β_1^c	β_2^c	β_3^c
S1	1.0003	1.0010	1.0469	1.0993	0.9999	1.0000
S2	1.0002	1.0003	1.0147	1.0029	1.0000	1.0003
S3	1.0003	1.0003	1.0020	1.0198	1.0000	1.0001
S4	1.0002	1.0003	1.0302	1.0031	1.0000	1.0001

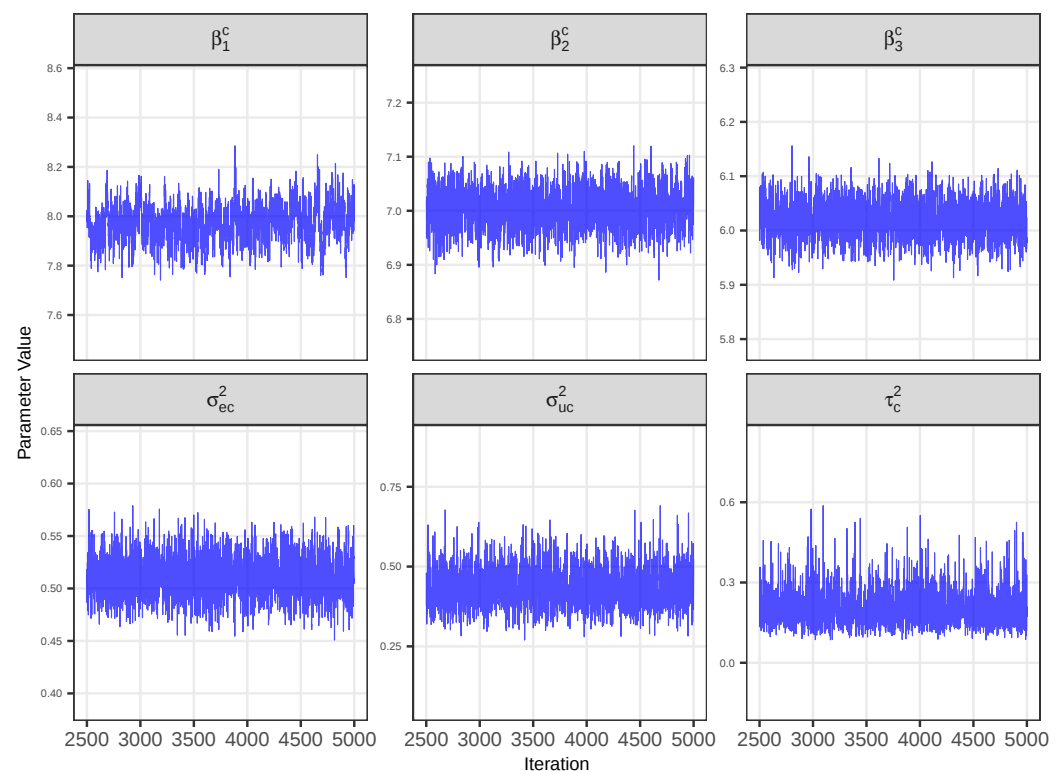


Figure 2. Center Model Parameters (y^c). The trace plot shows stable convergence after the burn-in period.

- (3) **Non-parametric Fit:** To visually assess the fit of the non-parametric component $f^s(z)$, Figures 4 and 5 show the results for $n = 50$ and $n = 150$, respectively. The figures plot the true centered function $f^s(z)$ (blue solid line) against the average estimated function (red dashed line) from the 100 simulations. The grey shaded area represents the 95% pointwise credible band. The results confirm the robustness of our BHS-CRM. In the nonlinear scenarios (S1, S2), the mean estimate closely tracks the true sinusoidal curve. This is quantitatively supported by the low $RASE_c$ values from our simulation (e.g., 0.057 for S1, $n = 50$, as shown later in Table 6). In the linear scenarios (S3, S4), the estimated function is correctly centered at zero, with $RASE_c$ values also approaching zero (e.g., 0.034 for S3, $n = 50$). This demonstrates that our semiparametric model does not overfit and correctly identifies the simpler linear relationship when the nonlinear component is absent. Comparing Figure 4 with

Figure 5, the 95% credible band is visibly narrower for $n = 150$, demonstrating that the precision of the non-parametric estimation improves with a larger sample size.

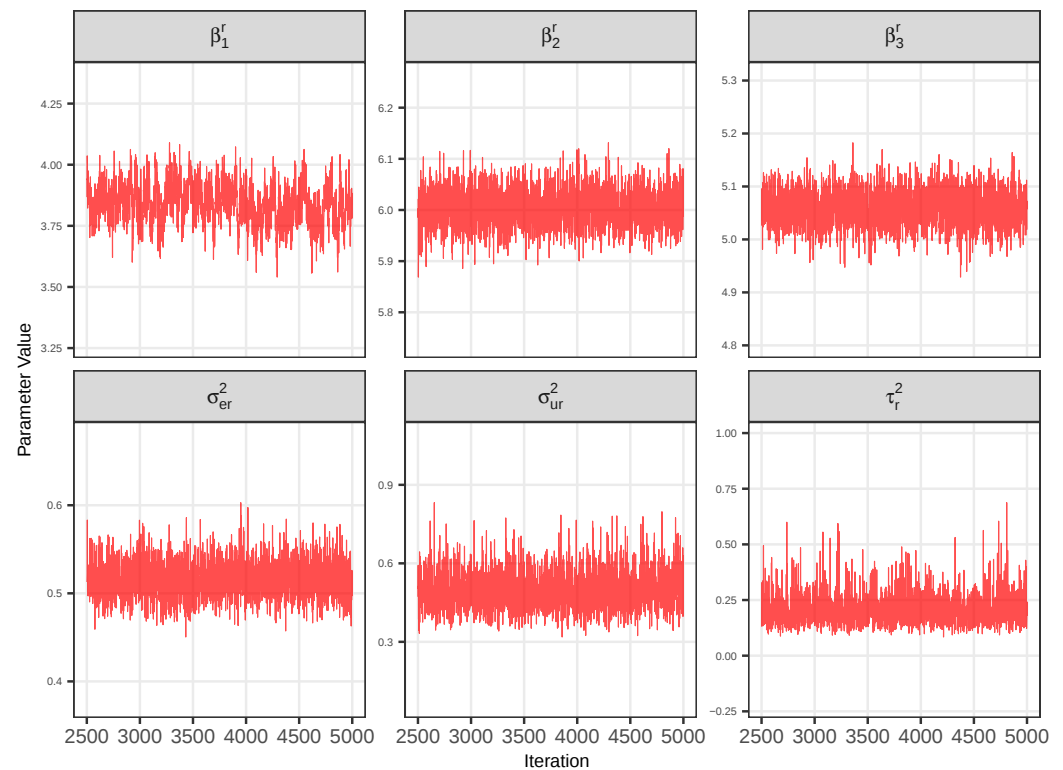


Figure 3. Range Model Parameters (y^r). The trace plot shows stable convergence after the burn-in period.

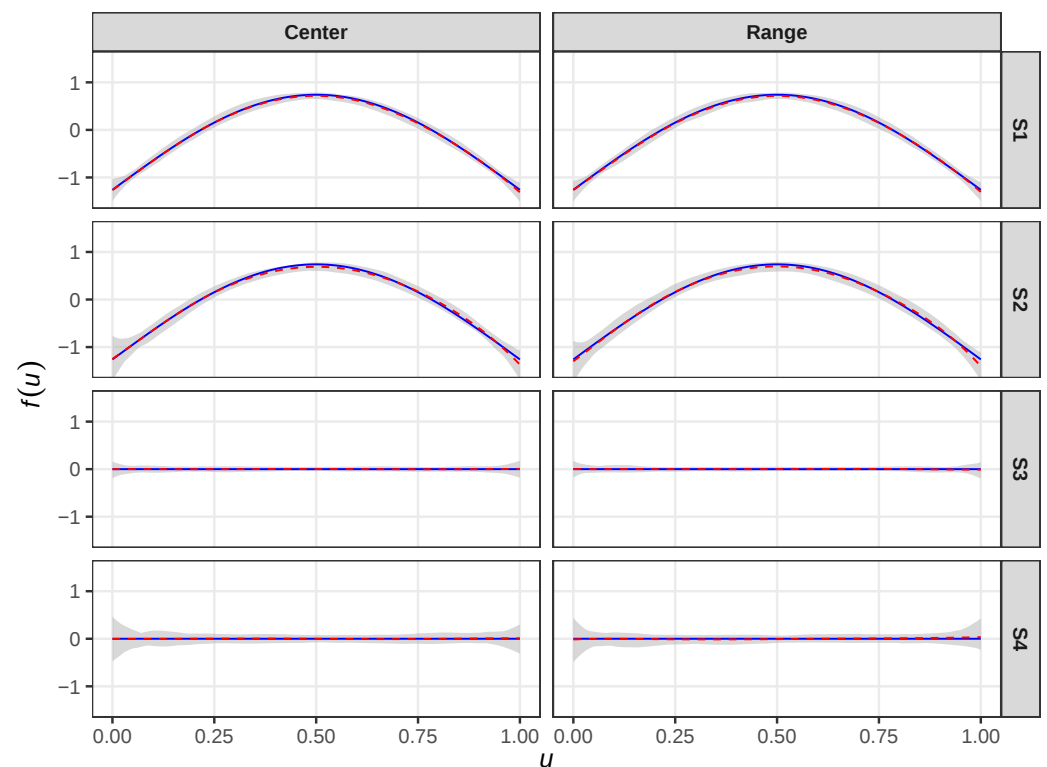


Figure 4. Average fitted non-parametric function for $n = 50$. The blue solid line is the true function $f(z)$, the red dashed line is the mean estimate, and the grey band is the 95% credible bands from 100 simulations.

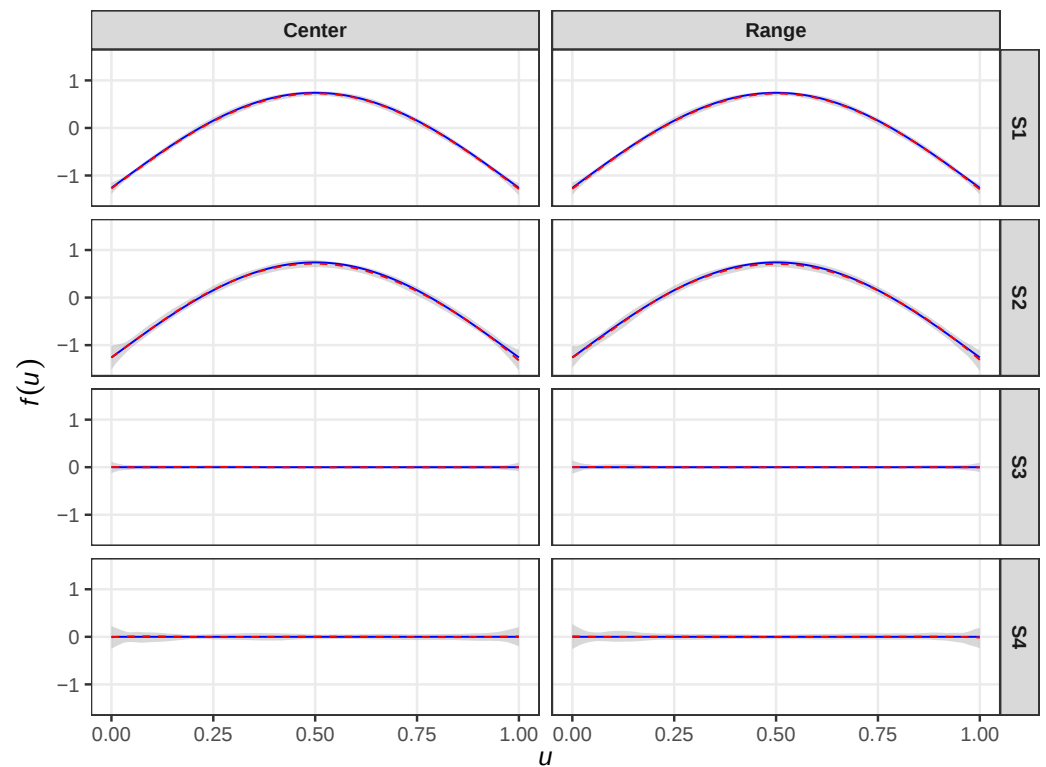


Figure 5. Average fitted non-parametric function for $n = 150$. The blue solid line is the true function $f(z)$, the red dashed line is the mean estimate, and the grey band is the 95% credible bands from 100 simulations.

Table 6. Predictive performance comparison for $n = 50$.

Scen.	Model	RASE _c	RASE _r	RI	R _c ²	R _r ²	RMSE _c	RMSE _r	RMSE _L	RMSE _U	RMSE _H
S1	BHS-CRM	0.057 (0.019)	0.056 (0.020)	0.912 (0.013)	0.979 (0.008)	0.970 (0.010)	0.768 (0.138)	0.776 (0.129)	0.851 (0.146)	0.870 (0.152)	1.620 (0.251)
	PLME-CRM	0.070 (0.023)	0.071 (0.027)	0.907 (0.012)	0.976 (0.007)	0.969 (0.010)	0.811 (0.128)	0.793 (0.140)	0.898 (0.150)	0.907 (0.149)	1.702 (0.236)
	LME-CRM	0.625 (0.000)	0.625 (0.000)	0.887 (0.010)	0.965 (0.008)	0.953 (0.011)	0.998 (0.112)	0.982 (0.114)	1.102 (0.128)	1.123 (0.124)	2.089 (0.204)
S2	BHS-CRM	0.083 (0.027)	0.088 (0.026)	0.886 (0.011)	0.963 (0.010)	0.952 (0.011)	1.027 (0.120)	1.007 (0.104)	1.150 (0.146)	1.136 (0.125)	2.149 (0.226)
	PLME-CRM	0.088 (0.033)	0.086 (0.027)	0.886 (0.011)	0.964 (0.009)	0.948 (0.013)	1.016 (0.120)	1.034 (0.121)	1.147 (0.140)	1.132 (0.135)	2.136 (0.231)
	LME-CRM	0.625 (0.000)	0.625 (0.000)	0.870 (0.010)	0.952 (0.010)	0.931 (0.014)	1.176 (0.111)	1.194 (0.104)	1.322 (0.131)	1.314 (0.123)	2.468 (0.210)
S3	BHS-CRM	0.034 (0.011)	0.033 (0.009)	0.904 (0.014)	0.978 (0.008)	0.972 (0.010)	0.772 (0.130)	0.755 (0.131)	0.860 (0.144)	0.858 (0.149)	1.615 (0.241)
	PLME-CRM	0.033 (0.027)	0.038 (0.025)	0.901 (0.015)	0.978 (0.008)	0.967 (0.012)	0.782 (0.148)	0.803 (0.147)	0.881 (0.155)	0.877 (0.174)	1.647 (0.274)
	LME-CRM	- -	- -	0.904 (0.014)	0.979 (0.007)	0.970 (0.010)	0.759 (0.131)	0.776 (0.127)	0.851 (0.148)	0.852 (0.156)	1.596 (0.243)
S4	BHS-CRM	0.069 (0.026)	0.070 (0.025)	0.878 (0.011)	0.965 (0.008)	0.950 (0.013)	1.003 (0.107)	1.015 (0.122)	1.125 (0.127)	1.122 (0.123)	2.104 (0.196)
	PLME-CRM	0.045 (0.032)	0.047 (0.033)	0.876 (0.012)	0.962 (0.009)	0.950 (0.012)	1.025 (0.122)	1.019 (0.123)	1.146 (0.132)	1.144 (0.136)	2.149 (0.230)
	LME-CRM	- -	- -	0.878 (0.011)	0.964 (0.009)	0.952 (0.011)	1.006 (0.114)	1.002 (0.119)	1.120 (0.118)	1.128 (0.130)	2.109 (0.214)

4.4. Comparative Analysis of Predictive Performance

To rigorously evaluate the predictive advantages of the proposed framework, we benchmark the BHS-CRM against two frequentist alternatives: the LME-CRM and the PLME-CRM. The comparative results are presented in Table 6 ($n = 50$) and Table 7 ($n = 150$). These tables provide a comprehensive evaluation using the metrics defined in Section 4.1. In all performance tables, values in parentheses represent standard deviations (SD). The analysis clearly highlights the robustness and predictive advantages of our model across different scenarios:

- (1) Performance in Nonlinear Scenarios (S1 and S2): We first examine the scenarios where the underlying relationship is nonlinear. The LME-CRM, which assumes a linear trend for $f(Z)$, performs worst across all key metrics. For instance, its $RASE_c$ is 0.625, reflecting a fundamental inability to capture the $2 \sin(\pi z)$ curve. This failure in fitting the non-parametric component leads to significantly poorer performance in all subsequent metrics, including the final interval prediction (RI and $RMSE_H$), when compared to the spline-based models (BHS-CRM and PLME-CRM). This result demonstrates that the semiparametric component of our model is essential for accurately modeling data with nonlinear relationships.
- (2) Robustness in Linear Scenarios (S3 and S4): In addition to the nonlinear settings, it is crucial to assess the model's performance in linear scenarios where the parsimonious LME-CRM represents the true data-generating process. As detailed in Tables 6 and 7, the proposed BHS-CRM not only matches but, in several metrics, marginally outperforms the benchmark LME model. This highly competitive predictive performance indicates that the Bayesian regularization effectively shrinks the spline coefficients toward linearity, preventing overfitting even when the true relationship is simple.

Table 7. Predictive performance comparison for $n = 150$.

Scen.	Model	$RASE_c$	$RASE_r$	RI	R_c^2	R_r^2	$RMSE_c$	$RMSE_r$	$RMSE_L$	$RMSE_U$	$RMSE_H$
S1	BHS-CRM	0.034 (0.011)	0.036 (0.012)	0.911 (0.007)	0.978 (0.005)	0.972 (0.006)	0.790 (0.079)	0.767 (0.077)	0.881 (0.094)	0.875 (0.089)	1.651 (0.146)
	PLME-CRM	0.046 (0.012)	0.046 (0.013)	0.911 (0.007)	0.979 (0.004)	0.970 (0.006)	0.776 (0.082)	0.780 (0.079)	0.868 (0.091)	0.870 (0.079)	1.634 (0.149)
	LME-CRM	0.625 (0.000)	0.625 (0.000)	0.890 (0.006)	0.966 (0.004)	0.954 (0.006)	0.983 (0.065)	0.985 (0.067)	1.100 (0.074)	1.099 (0.070)	2.058 (0.121)
S2	BHS-CRM	0.053 (0.014)	0.052 (0.016)	0.889 (0.006)	0.966 (0.004)	0.953 (0.007)	0.998 (0.059)	0.996 (0.065)	1.113 (0.074)	1.117 (0.058)	2.091 (0.110)
	PLME-CRM	0.055 (0.017)	0.057 (0.019)	0.887 (0.005)	0.965 (0.004)	0.951 (0.006)	1.010 (0.062)	1.015 (0.068)	1.127 (0.073)	1.134 (0.071)	2.117 (0.116)
	LME-CRM	0.625 (0.000)	0.625 (0.000)	0.871 (0.005)	0.953 (0.004)	0.934 (0.007)	1.175 (0.059)	1.180 (0.063)	1.309 (0.066)	1.321 (0.068)	2.463 (0.110)
S3	BHS-CRM	0.020 (0.007)	0.021 (0.007)	0.905 (0.008)	0.979 (0.004)	0.971 (0.006)	0.770 (0.079)	0.777 (0.079)	0.863 (0.096)	0.862 (0.083)	1.615 (0.145)
	PLME-CRM	0.020 (0.015)	0.023 (0.015)	0.902 (0.007)	0.978 (0.003)	0.970 (0.005)	0.788 (0.071)	0.779 (0.073)	0.877 (0.090)	0.881 (0.078)	1.653 (0.134)
	LME-CRM	- -	- -	0.903 (0.007)	0.978 (0.003)	0.971 (0.005)	0.782 (0.072)	0.771 (0.071)	0.869 (0.090)	0.873 (0.076)	1.638 (0.134)
S4	BHS-CRM	0.042 (0.012)	0.044 (0.014)	0.879 (0.006)	0.965 (0.005)	0.952 (0.006)	1.002 (0.062)	1.003 (0.064)	1.123 (0.072)	1.117 (0.070)	2.101 (0.114)
	PLME-CRM	0.030 (0.019)	0.026 (0.019)	0.877 (0.007)	0.964 (0.005)	0.952 (0.007)	1.008 (0.068)	1.005 (0.069)	1.128 (0.080)	1.126 (0.073)	2.114 (0.128)
	LME-CRM	- -	- -	0.878 (0.006)	0.965 (0.004)	0.952 (0.006)	1.000 (0.065)	0.998 (0.066)	1.118 (0.076)	1.117 (0.067)	2.097 (0.120)

- (3) **Comparison with Frequentist Semiparametric Approach:** Finally, we compare the two semiparametric models: our BHS-CRM and the frequentist PLME-CRM. While both models outperform the LME-CRM in nonlinear settings, the BHS-CRM demonstrates a clear competitive edge, particularly in the high-noise regime (S2). Theoretical justification lies in the second-order random walk prior, which functions as a robust stochastic penalty mechanism. Unlike frequentist estimators that rely on deterministic point estimates for the smoothing parameter, the robustness of BHS-CRM stems from integrating over the posterior distribution of the smoothing parameter, offering enhanced robustness in complex data environments.

5. Empirical Analysis

5.1. Data Source and Variable Description

To evaluate the effectiveness of the proposed BHS-CRM in modeling complex financial dynamics, we apply the method to analyze the carbon price volatility in China's Carbon ETS. The dataset comprises daily transaction records from four major pilot markets: Hubei, Guangdong, Shenzhen, and Shanghai. The sample period spans from May 2014 to April 2024, covering the primary developmental phases of China's carbon market.

Let $y_{it} = [y_{it}^L, y_{it}^U]$ denote the interval-valued carbon price for the i -th pilot city at time t , where y_{it}^L and y_{it}^U represent the lowest and highest transaction prices of the day, respectively. Consistent with the definitions provided in the Methodology section, the dependent variables are decomposed into the interval center (y_{it}^c) representing the price level, and the interval range (y_{it}^r) representing the price volatility.

In selecting the explanatory variables (X_{it}), we integrate energy economics theory with market microstructure characteristics. Fundamental determinants include energy prices (Steam Coal, Crude Oil, Natural Gas) and Electricity prices, which capture fuel substitution effects and generation costs. The HS300 Index and Exchange Rate are employed to proxy macroeconomic sentiment. Crucially, given the stylized facts of high-frequency financial data, we explicitly incorporate the first-order lag terms ($y_{i,t-1}^c$ and $y_{i,t-1}^r$) into the covariate vector. This specification is designed to capture the strong market inertia and path dependence inherent in carbon trading, thereby isolating the marginal contributions of fundamental drivers from short-term trading noise.

Temperature (Temp) is selected as the non-parametric covariate (Z_{it}) for the P-spline component. This choice is motivated by the complex, often asymmetric nonlinear relationship between climate conditions and energy demand, which linear models typically fail to capture. To ensure the stability of the MCMC sampling process and facilitate parameter comparison, all continuous variables (both dependent and independent) are standardized to have zero mean and unit variance before estimation.

Specifically, the pooled sample yields a balanced panel of 4276 trading day observations, distributed equally across Guangdong, Hubei, Shanghai, and Shenzhen (1069 observations per market). The carbon price records and other financial/meteorological covariates were compiled from official local carbon exchanges and standard financial databases. To synchronize the daily weather variables with the financial data, we matched the climatic records strictly to the trading calendar, explicitly excluding non-trading days such as weekends and public holidays. Missing price records, which occasionally occur due to zero trading volume, were handled using the last-observation-carried-forward (LOCF) method to maintain temporal continuity. The daily intervals were constructed using the realized intraday highest and lowest transaction prices. Zero-range observations, where the highest and lowest prices are identical, represent valid illiquid trading days and were retained without modification.

Furthermore, considering the data-intensive nature of hierarchical Bayesian inference, we adopt a 90–10% data splitting strategy. The first 90% of the temporal observations serve as the training set to ensure the robust convergence of random effects and spline parameters, while the subsequent 10% are reserved for strict out-of-sample validation. The detailed variable definitions are summarized in Table 8.

Table 8. Description of Variables.

Variable	Symbol	Description
Dependent Variables		
Carbon Price Level	y_{it}^c	Center of daily highest and lowest prices
Price Volatility	y_{it}^r	Range (Full-width) of daily prices
Independent Variables (X_{it})		
Lagged Price	$y_{i,t-1}^s$	One-day lagged terms to capture market inertia
Coal Price	Coal	Daily closing price of steam coal futures
Oil Price	Oil	Daily closing price of crude oil futures
Macro Economy	HS300	CSI 300 Index reflecting market sentiment
Natural Gas	Gas	Daily price of natural gas
Electricity	Elec	Industrial electricity price index
Exchange Rate	Rate	USD/CNY exchange rate
Smoothing Covariate (Z_{it})		
Temperature	Temp	Daily average temperature

Note: The superscript $s \in \{c, r\}$ denotes the center and range models, respectively.

5.2. Bayesian Estimation Results and Mechanism Analysis

Table 9 presents the posterior estimates of the structural parameters for the BHS-CRM, derived from 10,000 MCMC iterations with a 20% burn-in period. To ensure the posterior inference is strictly data-dominated, weakly informative priors are adopted (e.g., $IG(0.01, 0.01)$ for variance components, $IG(1, 0.005)$ for the smoothing parameter, and normal priors with a variance of 10 for fixed effects). An empirical sensitivity analysis perturbing these hyperparameters by an order of magnitude confirms that the posterior estimates are highly robust. Unlike traditional frequentist point estimates, the Bayesian framework provides full posterior distributions. To facilitate direct comparison of effect magnitudes, all coefficients reported are based on standardized variables.

Table 9. Bayesian Posterior Estimates for BHS-CRM Parameters.

Variable	Center Model (y^c)			Range Model (y^r)		
	Mean	SD	95% CrI	Mean	SD	95% CrI
Lagged Y ($y_{i,t-1}^s$)	0.9765 *	0.0032	[0.9702, 0.9829]	0.7309 *	0.0112	[0.7090, 0.7525]
Coal Price (Coal)	0.0184 *	0.0091	[0.0006, 0.0363]	−0.0133	0.0126	[−0.0379, 0.0116]
Oil Price (Oil)	0.0012	0.0043	[−0.0074, 0.0097]	−0.0063	0.0120	[−0.0298, 0.0175]
Macro Index (HS300)	−0.0009	0.0030	[−0.0067, 0.0050]	0.0165	0.0154	[−0.0137, 0.0464]
Natural Gas (Gas)	−0.0002	0.0124	[−0.0241, 0.0239]	−0.0203	0.0265	[−0.0729, 0.0315]
Electricity (Elec)	0.0021	0.0133	[−0.0238, 0.0280]	−0.0091	0.0263	[−0.0618, 0.0421]
Exchange Rate (Rate)	0.0007	0.0038	[−0.0068, 0.0083]	−0.0076	0.0101	[−0.0271, 0.0119]
$\sigma_{\epsilon,s}^2$ (Error Var.)	0.0197	0.0005	[0.0189, 0.0207]	0.3373	0.0078	[0.3226, 0.3529]
$\sigma_{u,s}^2$ (RE Var.)	0.0102	0.0206	[0.0018, 0.0435]	0.0441	0.0665	[0.0079, 0.1821]

Note: * indicates that zero is excluded from the 95% credible interval.

Regarding the fundamental drivers, the results reveal the distinct roles of market inertia and energy costs. As anticipated, the autoregressive term dominates price formation.

In the center model, the lag coefficient is 0.977, indicating a high degree of persistence in carbon price levels. This near-unit root behavior reflects strong path dependence and market inertia, which is a widely documented stylized fact in high-frequency financial time series, as current prices heavily internalize historical information. Similarly, the range model exhibits a persistence coefficient of 0.731, confirming the stylized fact of volatility clustering, albeit with lower memory compared to the price level.

Crucially, distinct from pure time-series models, the BHS-CRM successfully isolates the marginal contribution of fundamental factors. Notably, even after controlling for the strong autoregressive dynamics, the Steam Coal Price (Coal) exhibits a significant positive impact on the carbon price level (Mean = 0.0184). The 95% credible interval for this coefficient ([0.0006, 0.0363]) strictly excludes zero, validating the hypothesis that fluctuations in coal prices significantly spill over into the carbon market. In terms of magnitude, since the variables are standardized, this implies that a one-standard-deviation increase in coal prices is associated with approximately a 0.018 standard-deviation increase in carbon prices. While the coefficient appears numerically small due to the dominance of the lag term, its statistical significance underscores that coal remains the robust fundamental anchor for China's carbon pricing mechanism, as shown in Figure 6.

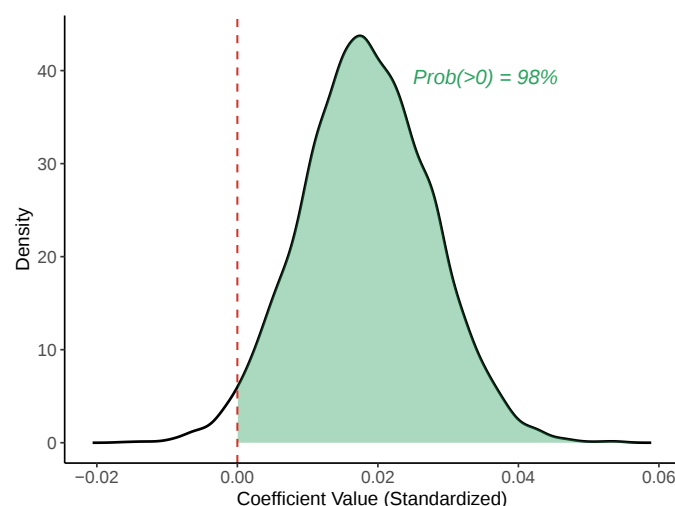


Figure 6. Posterior Density of Coal Price Coefficient. The vertical dashed line represents zero, and the green area indicates the probability mass of a positive effect.

The near-unit-root autoregressive coefficient mathematically explains the majority of the total variance. This is a common characteristic of high-frequency financial time series. However, the inclusion of fundamental and climatic covariates provides crucial structural insights. To quantify their incremental contribution, we estimated a baseline model with only the autoregressive term and random effects. The comparison reveals that the full BHS-CRM further reduces the in-sample residual variance by 0.90% and decreases the out-of-sample RMSE by 0.22% relative to the baseline. Although these absolute predictive gains are modest due to the dominant path dependence, the covariates successfully capture essential economic signals and climate-driven risk premiums hidden in the daily price innovations. This justifies their indispensable role for structural inference.

Turning to the climatic factor, Figure 7 visualizes the estimated nonlinear relationship between temperature and carbon price dynamics. A distinguishing feature of the proposed BHS-CRM is its ability to decouple the impact on the price level from the impact on market volatility. The vertical axis in both panels represents the marginal contribution of temperature to the respective standardized response variables. It is important to note that to satisfy the statistical identifiability constraint, the B-spline basis functions are centered to

have a zero mean over the sample. Consequently, the global average price level is absorbed by the model's intercept, and the solid blue curves depicted here strictly represent the relative deviation from this baseline attributed to temperature variations. The shaded region surrounding the curve denotes the 95% pointwise credible interval.

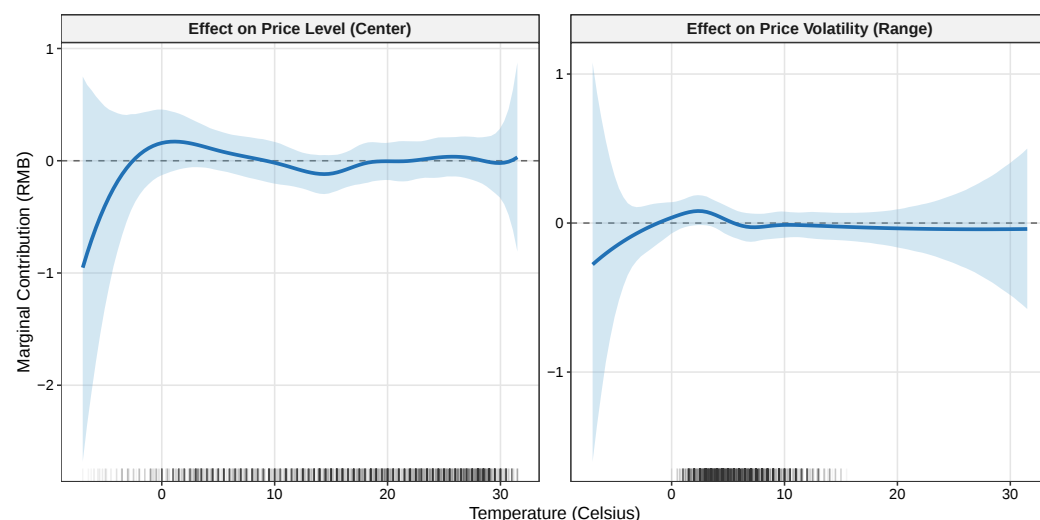


Figure 7. Nonlinear Effects of Temperature on Carbon Price Level and Volatility. The left panel displays the marginal effect on the price level (y^c), while the right panel shows the effect on price volatility (y^r). The solid blue lines indicate the posterior mean estimates, and the shaded areas denote the 95% credible intervals.

The left panel illustrates the marginal effect of temperature on the carbon price level. The estimated trajectory reveals a distinct asymmetric pattern driven by varying market mechanisms. In the low-temperature regime where temperatures fall below 0 degrees Celsius, the posterior mean exhibits a noticeable negative deviation. This finding suggests that severe cold primarily induces a seasonal slowdown in industrial production and outdoor construction activities, which outweighs the potential increase in heating demand. Consequently, these factors exert downward pressure on carbon prices relative to the baseline. Conversely, at the upper end of the temperature distribution exceeding 25 degrees Celsius, the model exhibits substantial uncertainty as evidenced by the significantly widening credible bands. To explicitly illustrate this data dependency, a rug plot along the x-axis maps the empirical density of observed temperatures (-7.0 to 31.5 °C). The widening of the credible boundaries closely corresponds to the extreme scarcity of observations in these tail regions, underscoring the robustness of the Bayesian P-splines: rather than imposing a spurious deterministic trend, the model conservatively estimates broader uncertainty where data is sparse. This highly asymmetric trajectory highlights the necessity of the semiparametric approach; a conventional linear specification would erroneously average out these divergent temperature effects, masking the true climate-price dynamics.

The right panel provides complementary insights by illustrating the response of price volatility to temperature fluctuations. The estimated function reveals that market risk dynamics differ significantly from price level trends. Specifically, the volatility effect becomes negative during extreme cold spells, implying that the market is characterized by both lower prices and reduced trading divergence. However, a distinct rise in volatility is observed as temperatures approach the interval between 0 and 5 degrees Celsius. This phenomenon suggests the existence of a volatility premium near the freezing point where market uncertainty increases regarding critical heating demand thresholds. Similar to the price level analysis, the volatility effect under high-temperature conditions remains neutral with expanding credible bands due to the limited data density in that range.

Finally, we examine the structural heterogeneity across pilot markets using the random effects component. As shown in Table 9, the posterior mean of the random effects variance $\sigma_{u,c}^2$ in the center model is 0.0102. While simple descriptive statistics show large price gaps between cities, the hierarchical model reveals a more nuanced picture.

Figure 8 visualizes the city-specific baseline deviations after controlling for market inertia and fundamental drivers. While the point estimates suggest relative positioning—where Shanghai displays a positive baseline deviation and Hubei a negative one consistent with their respective economic development levels—the 95% credible intervals for all cities substantially overlap with zero. This implies that residual baseline differences among the pilot markets lack statistical significance at the 95% level. The apparent structural heterogeneity in carbon prices is largely absorbed by the strong autoregressive term (y_{t-1}); thus, persistent price gaps are captured dynamically through path dependence rather than static intercepts. Nevertheless, retaining the hierarchical structure remains essential to properly account for unobserved city-specific correlations, preventing the underestimation of standard errors for other coefficients. While the distinct liquidity regimes in Figure 9 suggest heterogeneous transmission mechanisms, the common-slope specification was chosen for model parsimony. Introducing random coefficients across only four pilot markets alongside P-splines could over-parameterize the model and destabilize MCMC convergence, justifying the use of random intercepts to capture baseline structural differences.

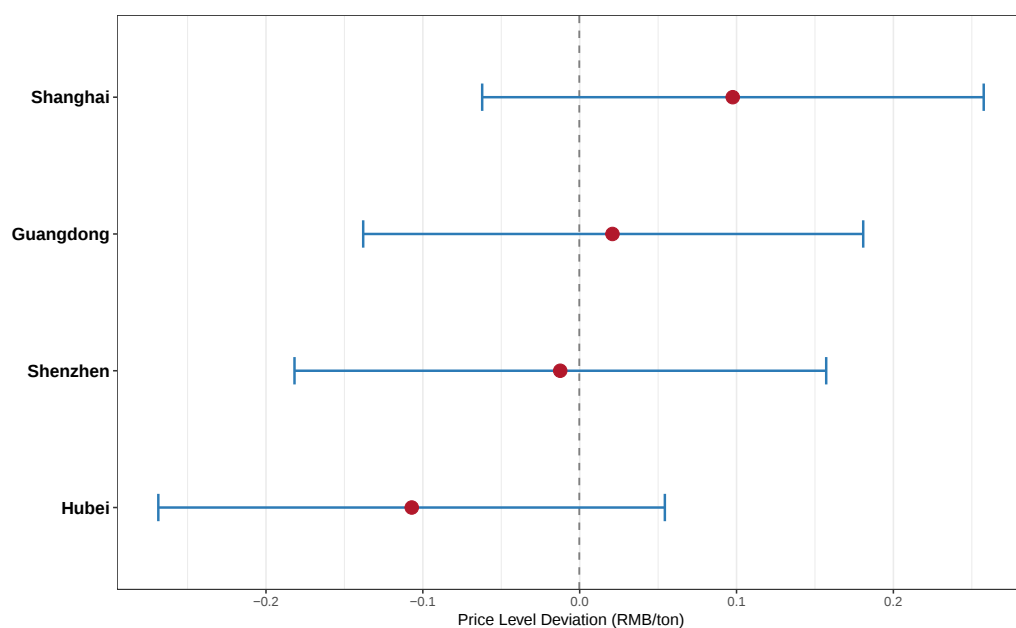


Figure 8. Heterogeneity of Baseline Carbon Prices (Random Effects). The red dots represent the posterior means of city-specific intercepts, and error bars indicate 95% credible intervals.

5.3. Out-of-Sample Predictive Performance Comparison

We strictly evaluate the one-step-ahead out-of-sample predictive performance of BHS-CRM against two competitive benchmarks: the LME-CRM and the PLME-CRM. The evaluation metrics include the RMSE series and the RI score as defined in Section 4.1. Specifically, in this temporal-split procedure, the prediction for each out-of-sample day t is generated utilizing the true observed lagged values ($y_{i,t-1}^s$). Furthermore, since the pilot markets in the test set are identical to those in the training set, we directly apply the posterior means of the subject-specific random effects ($\hat{u}_i^s$) estimated from the training phase.

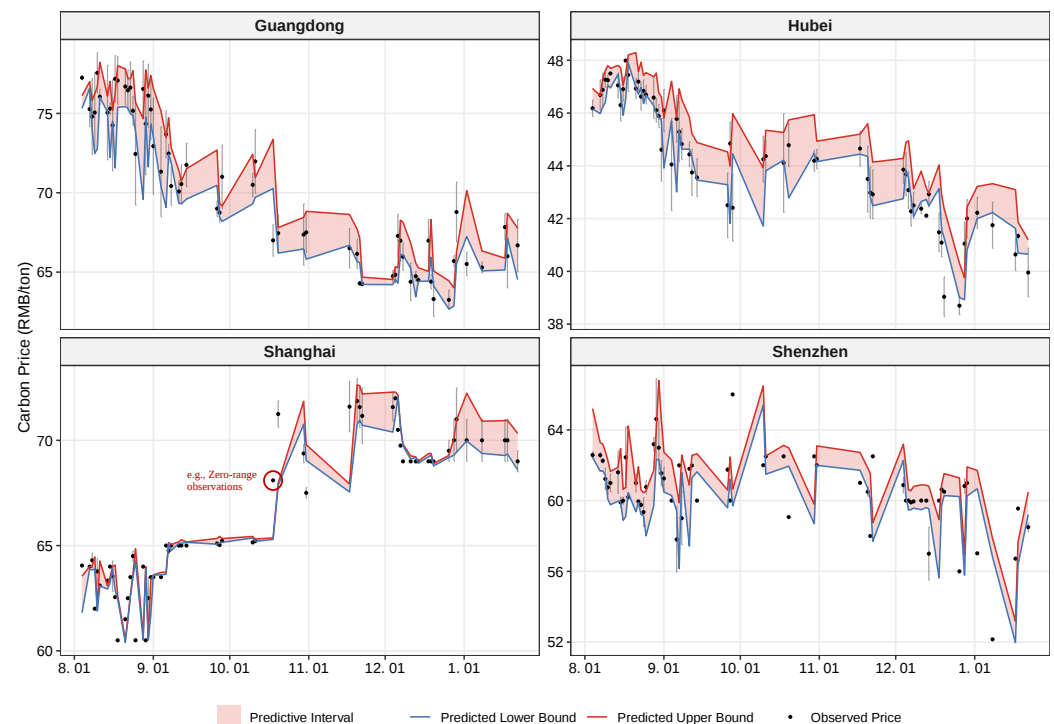


Figure 9. Out-of-sample Interval Predictions for Representative Pilot Markets. The red-shaded regions denote the predicted interval bounds, while vertical bars indicate realized daily price ranges.

In addition, we introduce the Average Width (AW) to assess the informativeness of the constructed intervals, noting that narrower intervals imply higher precision. The AW is defined as:

$$AW = \frac{1}{N_{\text{test}}} \sum_{i,t \in \text{Test}} (\hat{y}_{it}^U - \hat{y}_{it}^L).$$

Note that while the model estimation relied on standardized data to ensure MCMC convergence, all predictive metrics reported below are calculated based on the back-transformed values in the original scale (RMB/ton) to preserve practical interpretability.

Table 10 reports the comprehensive performance metrics. In terms of point forecasting accuracy, BHS-CRM achieves the lowest error in the primary price level prediction ($RMSE_c = 1.5867$), outperforming both LME (1.5928) and PLME (1.5927). Although the linear autoregressive component captures the majority of the variance across all models, BHS-CRM successfully extracts additional predictive power from the nonlinear climatic effects and the hierarchical structure. By effectively filtering out noise through the Bayesian penalty mechanism, it provides a more precise anchor for daily transaction prices compared to the frequentist benchmarks.

Table 10. Comparison of Out-of-Sample Prediction Performance.

Metric	BHS-CRM	LME-CRM	PLME-CRM
$RMSE_c$	1.5867 **	1.5928	1.5927
$RMSE_r$	1.2234	1.2207	1.2207
$RMSE_L$	1.7194 **	1.7230	1.7230
$RMSE_U$	1.6814 **	1.6883	1.6882
$RMSE_H$	3.2309 **	3.2414	3.2413
RI	0.2114 **	0.2059	0.2060
AW	1.3162	1.3106	1.3106

Note: ** indicates statistical significance at the 5% level based on the Diebold–Mariano test.

Regarding volatility prediction (y^r), the linear benchmark LME shows a marginal advantage. This suggests that the volatility magnitude in carbon markets is largely driven by linear autoregressive features, which LME captures efficiently. However, BHS-CRM remains highly competitive with a negligible difference. More importantly, BHS-CRM demonstrates enhanced performance in predicting the interval boundaries (RMSE_L and RMSE_U), resulting in the lowest combined boundary error ($\text{RMSE}_H = 3.2309$).

Crucially, BHS-CRM achieves the highest Relative Intersection score ($\text{RI} = 0.2114$). This metric reflects the degree of geometric overlap between the predicted and realized intervals. While the Average Width (AW) is comparable across models (≈ 1.31), the higher RI score of BHS-CRM indicates that its intervals are structurally more aligned with the true price dynamics. This implies that BHS-CRM strikes a better balance between precision and shape accuracy, effectively capturing the realized price range without relying on unnecessarily wide bounds.

To visually inspect the model's adaptability, Figure 9 presents interval forecasts for four representative pilot markets, including Guangdong, Hubei, Shanghai, and Shenzhen, over the last 60 trading days. The red-shaded regions denote the predicted interval bounds generated by BHS-CRM. Because the response variable is intrinsically interval-valued, these bounds are deterministically reconstructed using the posterior mean estimates of the center and range components, i.e., $[\hat{y}_{it}^c - \hat{y}_{it}^r/2, \hat{y}_{it}^c + \hat{y}_{it}^r/2]$, while vertical bars represent realized daily ranges. Visually, the markets exhibit distinct liquidity profiles. While Guangdong displays continuous volatility, Shanghai is characterized by frequent zero-range observations. Economically, this phenomenon reflects the specific market microstructure of the Shanghai pilot during certain periods. Driven by severe liquidity constraints and low trading volumes, the market frequently recorded only a single clearing price per day, causing the realized daily high and low prices to perfectly converge.

Despite these structural disparities, the BHS-CRM demonstrates robust generalization. In active markets such as Guangdong, the predictive intervals tightly track price dynamics to capture time-varying volatility with precision. Conversely, in low-liquidity regimes like Shanghai, the model avoids overfitting to static noise. The intervals appropriately encompass realized prices without collapsing, confirming that the proposed hierarchical framework effectively leverages information pooling to anchor the baseline while accurately quantifying risks, regardless of the market's liquidity state.

6. Conclusions and Discussion

This paper introduces a BHS-CRM designed to handle the complexities of nonlinearity and heterogeneity in interval-valued panel data. By incorporating Bayesian P-splines within a hierarchical framework, the model estimates central tendencies and volatility dynamics in parallel while effectively managing parameter uncertainty. Simulation results demonstrate the model's advantages over traditional approaches: it improves upon LME-CRM in capturing nonlinear patterns and exhibits greater robustness against noise and overfitting compared to PLME-CRM. These findings validate the BHS-CRM as a reliable alternative for analyzing complex panel data structures where standard linear assumptions fail.

In the empirical analysis of China's Carbon ETS, the model provides nuanced insights into price formation mechanisms. We identify coal prices as a fundamental cost anchor, even after controlling for strong market inertia, and reveal a distinct asymmetric nonlinear relationship between temperature and carbon prices. Furthermore, out-of-sample evaluations confirm that the BHS-CRM achieves a balanced trade-off between interval coverage and width, offering competitive predictive bounds. Consequently, this model serves as an effective approach for risk quantification in energy finance, with potential applications extending to other fields requiring precise uncertainty modeling.

Despite its meaningful contributions, this paper acknowledges certain limitations that point to avenues for future research. First, the Bayesian MCMC inference, while rigorous, is computationally more intensive than standard frequentist approaches, which may pose challenges for large-scale high-frequency applications. Second, while the current BHS-CRM framework relies on an independence assumption to maintain computational efficiency, future research could relax this constraint by developing a fully multivariate extension. For instance, incorporating a joint covariance structure between the center and range components, or explicitly modeling spatial spillovers across regional markets (e.g., via a spatial autoregressive structure), would further enhance the model's capacity to capture complex financial co-movements. Finally, to better capture the heterogeneous transmission mechanisms and distinct liquidity regimes observed across different markets, expanding the model to include random coefficients represents an important avenue for future research.

Author Contributions: Conceptualization, D.X., M.Q. and R.T.; methodology, D.X., M.Q. and R.T.; software, M.Q.; validation, D.X. and M.Q.; formal analysis, D.X. and M.Q.; investigation, M.Q.; resources, D.X. and M.Q.; data curation, M.Q.; writing—original draft preparation, M.Q. and D.X.; writing—review and editing, D.X. and M.Q.; supervision, D.X. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the National Social Science Fund of China (Grant No. 23BTJ069).

Data Availability Statement: The data presented in this study are available on request from the corresponding author.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Moore, R.E. *Interval Analysis*; Prentice-Hall: Englewood Cliffs, NJ, USA, 1966.
2. Billard, L.; Diday, E. *Symbolic Data Analysis: Conceptual Statistics and Data Mining*; John Wiley & Sons: Chichester, UK, 2012.
3. Lima Neto, E.A.; De Carvalho, F.D.A. Centre and range method for fitting a linear regression model to symbolic interval data. *Comput. Stat. Data Anal.* **2008**, *52*, 1500–1515. [[CrossRef](#)]
4. Hao, P.; Guo, J. Constrained center and range joint model for interval-valued symbolic data regression. *Comput. Stat. Data Anal.* **2017**, *116*, 106–138. [[CrossRef](#)]
5. Guan, L.; Li, M. Interval-valued linear regression model with an asymmetric Laplace distribution. *J. Korean Stat. Soc.* **2025**, *54*, 161–193. [[CrossRef](#)]
6. Lim, C. Interval-valued data regression using nonparametric additive models. *J. Korean Stat. Soc.* **2016**, *45*, 358–370. [[CrossRef](#)]
7. Kong, L.; Song, X.; Wang, X. Nonparametric regression for interval-valued data based on local linear smoothing approach. *Neurocomputing* **2022**, *501*, 834–843. [[CrossRef](#)]
8. Sun, L.; Zhu, L.; Li, W.; Chen, Z. Interval-valued functional clustering based on the Wasserstein distance with application to stock data. *Inf. Sci.* **2022**, *606*, 910–926. [[CrossRef](#)]
9. Tang, J.; Mi, C.; Fu, C.; Yao, Q. Novel solution framework for inverse problem considering interval uncertainty. *Int. J. Numer. Methods Eng.* **2022**, *123*, 1654–1672. [[CrossRef](#)]
10. Tang, J.; Cao, L.; Mi, C.; Fu, C.; Liu, Q. Interval assessments of identified parameters for uncertain structures. *Eng. Comput.* **2022**, *38*, 2905–2917. [[CrossRef](#)]
11. Tang, J.; Fu, C.; Mi, C.; Liu, H. An interval sequential linear programming for nonlinear robust optimization problems. *Appl. Math. Model.* **2022**, *107*, 256–274. [[CrossRef](#)]
12. Baltagi, B.H. *Econometric Analysis of Panel Data*, 4th ed.; John Wiley & Sons: Chichester, UK, 2008.
13. Ji, A.B.; Li, Q.Q.; Zhang, J.J. Panel interval-valued data nonlinear regression models and applications. *Comput. Econ.* **2024**, *64*, 2413–2435. [[CrossRef](#)]
14. Engle, R.F.; Rangel, J.G. The spline-GARCH model for low-frequency volatility and its global macroeconomic causes. *Rev. Financ. Stud.* **2008**, *21*, 1187–1222. [[CrossRef](#)]
15. Li, Q.; Racine, J.S. *Nonparametric Econometrics: Theory and Practice*; Princeton University Press: Princeton, NJ, USA, 2007.

16. Liu, Y.; Zou, J.; Zhao, S.; Yang, Q. Model averaging estimation for varying-coefficient single-index models. *J. Syst. Sci. Complex.* **2022**, *35*, 264–282. [[CrossRef](#)]
17. Gelman, A.; Carlin, J.B.; Stern, H.S.; Dunson, D.B.; Vehtari, A.; Rubin, D.B. *Bayesian Data Analysis*, 3rd ed.; CRC Press: Boca Raton, FL, USA, 2013.
18. Zhang, D.; Wu, L.; Ye, K.; Chen, M. Bayesian quantile semiparametric mixed-effects double regression models. *Stat. Theory Relat. Fields* **2021**, *5*, 303–315. [[CrossRef](#)]
19. Xu, M.; Qin, Z. A bivariate Bayesian method for interval-valued regression models. *Knowl.-Based Syst.* **2022**, *235*, 107396. [[CrossRef](#)]
20. Xu, M.; Qin, Z. A Bayesian parametrized method for interval-valued regression models. *Stat. Comput.* **2023**, *33*, 67. [[CrossRef](#)]
21. Ferede, M.M.; Dagne, G.A.; Mwalili, S.M.; Bilchut, W.H.; Engida, H.A.; Karanja, S.M. Flexible Bayesian semiparametric mixed-effects model for skewed longitudinal data. *BMC Med. Res. Methodol.* **2024**, *24*, 56. [[CrossRef](#)]
22. Langat, K.A.; Mwalili, S.M.; Kazembe, L.N. Mixed effects and semi-parametric Bayesian integration models for measurement error correction in the context of fertilizer application levels: A simulation study. *Commun. Math. Biol. Neurosci.* **2024**, *2024*, 105. [[CrossRef](#)]
23. Kowal, D.R.; Wu, B. Monte Carlo inference for semiparametric Bayesian regression. *J. Am. Stat. Assoc.* **2025**, *120*, 1063–1076. [[CrossRef](#)]
24. Cui, E.; Lu, X.; Zhou, J.; Ma, H. A semiparametric Bayesian method for instrumental variable analysis with partly interval-censored time-to-event outcome. *arXiv* **2025**, arXiv:2501.14837.
25. Ruppert, D.; Wand, M.P.; Carroll, R.J. *Semiparametric Regression*; Cambridge University Press: Cambridge, UK, 2003.
26. Hastie, T.J.; Tibshirani, R.J. *Generalized Additive Models*; Chapman and Hall/CRC: Boca Raton, FL, USA, 1990.
27. Su, L.; Ullah, A. Profile likelihood estimation of partially linear panel data models with fixed effects. *Econ. Lett.* **2006**, *92*, 75–81. [[CrossRef](#)]
28. Wood, S.N. *Generalized Additive Models: An Introduction with R*, 2nd ed.; Chapman and Hall/CRC: Boca Raton, FL, USA, 2017.
29. Tian, R.; Xia, M.; Xu, D. Profile quasi-maximum likelihood estimation for semiparametric varying-coefficient spatial autoregressive panel models with fixed effects. *Stat. Pap.* **2024**, *65*, 5109–5143. [[CrossRef](#)]
30. Xu, D.; Zhang, Z.; Wu, L. Bayesian analysis of joint mean and covariance models for longitudinal data. *J. Appl. Stat.* **2014**, *41*, 2504–2514. [[CrossRef](#)]
31. Eilers, P.H.; Marx, B.D. Flexible smoothing with B-splines and penalties. *Stat. Sci.* **1996**, *11*, 89–121. [[CrossRef](#)]
32. Lang, S.; Brezger, A. Bayesian P-splines. *J. Comput. Graph. Stat.* **2004**, *13*, 183–212. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.