

Bayesian system identification for structures considering spatial and temporal correlation

Koune, Ioannis; Rózsás, Árpád; Slobbe, Arthur; Cicirello, Alice

DOI

[10.1017/dce.2023.18](https://doi.org/10.1017/dce.2023.18)

Publication date

2023

Document Version

Final published version

Published in

Data-Centric Engineering

Citation (APA)

Koune, I., Rózsás, Á., Slobbe, A., & Cicirello, A. (2023). Bayesian system identification for structures considering spatial and temporal correlation. *Data-Centric Engineering*, 4(5), Article e22.
<https://doi.org/10.1017/dce.2023.18>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

RESEARCH ARTICLE

Bayesian system identification for structures considering spatial and temporal correlation

Ioannis Koune^{1,2} , Árpád Rózsás², Arthur Slobbe² and Alice Cicirello¹ 

¹Faculty of Civil Engineering and Geosciences, Delft University of Technology, Delft, Netherlands

²Building, Infrastructure and Maritime Unit, TNO, Delft, Netherlands

Corresponding author: Ioannis Koune; Email: i.c.koune@tudelft.nl

Received: 04 May 2023; **Revised:** 18 August 2023; **Accepted:** 18 September 2023

Keywords: Bayesian system identification; correlation; model prediction uncertainty; model calibration; structural health monitoring

Abstract

The decreasing cost and improved sensor and monitoring system technology (e.g., fiber optics and strain gauges) have led to more measurements in close proximity to each other. When using such spatially dense measurement data in Bayesian system identification strategies, the correlation in the model prediction error can become significant. The widely adopted assumption of uncorrelated Gaussian error may lead to inaccurate parameter estimation and overconfident predictions, which may lead to suboptimal decisions. This article addresses the challenges of performing Bayesian system identification for structures when large datasets are used, considering both spatial and temporal dependencies in the model uncertainty. We present an approach to efficiently evaluate the log-likelihood function, and we utilize nested sampling to compute the evidence for Bayesian model selection. The approach is first demonstrated on a synthetic case and then applied to a (measured) real-world steel bridge. The results show that the assumption of dependence in the model prediction uncertainties is decisively supported by the data. The proposed developments enable the use of large datasets and accounting for the dependency when performing Bayesian system identification, even when a relatively large number of uncertain parameters is inferred.

Impact Statement

In Bayesian system identification for structures, simplistic probabilistic models are typically used to describe the discrepancies between measurement and model predictions, which are often defined as independent and identically distributed Gaussian random variables. This assumption can be unrealistic for real-world problems, potentially resulting in underestimation of the uncertainties and overconfident predictions. We demonstrate that in a real-world case study of a twin-girder steel bridge, the inclusion of correlation is decisively favored by the data. In the proposed approach, both the functional form of the probabilistic model and the posterior distribution over the uncertain parameters of the probabilistic model are inferred from the data. A novel efficient log-likelihood evaluation method is proposed to reduce the computational cost of the inference.

 This research article was awarded an Open Materials badge for transparent practices. See the Data Availability Statement for details.

© The Author(s), 2023. Published by Cambridge University Press. This is an Open Access article, distributed under the terms of the Creative Commons Attribution licence (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted re-use, distribution and reproduction, provided the original article is properly cited.

1. Introduction

1.1. Motivation

Structural health monitoring (SHM) methods based on probabilistic approaches have seen significant development in recent years (Farrar and Worden, 2012) and have been applied for system identification and damage detection for various types of structures including bridges (Behmanesh and Moaveni, 2014), rail (Lam et al., 2014), offshore oil and gas installations (Brownjohn, 2007), offshore wind farms (Rogers, 2018) and other civil engineering structures (Chen, 2018). The Bayesian system identification framework established by Beck and Katafygiotis (1998) uses measurements of structural responses obtained from sensors in combination with computational physics models to infer uncertain parameters, calibrate models, identify structural damage, and provide insight into the structural behavior (Huang et al., 2019). In Bayesian statistics the problem is cast as a parameter estimation and model selection problem, often referred to as system identification in the SHM literature (Katafygiotis et al., 1998). Specifically, previous knowledge about the system parameters to be inferred is represented by statistical distributions and combined with measurements to infer the posterior parameter distribution. A key advantage of this approach is that it provides a rigorous framework for combining prior knowledge and data with a probabilistic description of the uncertainties to obtain a posterior distribution over nondirectly observed parameters of interest (the so-called latent variables) using directly observed responses. For example, the rotational stiffness of a support can be estimated based on measured deflections (Ching et al., 2006; Lam et al., 2018).

In parallel with the probabilistic methods for SHM, sensor and monitoring technologies have seen significant progress in recent years. These technologies can provide higher accuracy and improved measurement capabilities, for example, by utilizing fiber optic strain sensors (Ye et al., 2014; Barrias et al., 2016). Fiber optic strain sensors provide measurements with high spatial and temporal resolution as large numbers of sensors with high sampling rates are used in the same structure. System identification is carried out under the assumption that there is sufficient information in the measurements so that the data can overrule the prior assumption on the latent variables. Therefore, utilizing the additional information contained in these measurements can potentially improve the accuracy of our predictions, reduce the uncertainty on the inferred system parameters, and lead to improved physical models that can more accurately capture the structural behavior. However, when using measurements from dense sensor layouts, such as fiber optic strain sensors, the discrepancies between model prediction and observations are expected to be dependent. This dependence has to be considered in the system identification to avoid inaccurate parameter estimation and overconfidence in the model predictions.

1.2. Problem statement

Current approaches in Bayesian inference for structures largely neglect the dependencies in the model prediction error. Instead, it is typically assumed that the prediction error is Gaussian white noise, that is, uncorrelated with zero mean (Lye et al., 2021). When using closely spaced measurements and model predictions, for example, in the case of time series with high sampling rates or spatial data from densely spaced sensors, dependencies may be present in the model prediction errors (Simoen et al., 1998). The strength of the correlation typically depends on the proximity of the measurements in time and the spacing of sensors on the structure. A fictitious example of a simply supported beam where the error between measurement and model prediction for two sensors is modeled as a bivariate Normal distribution, explicitly accounting for the spatial correlation of three different sets of measurements, is shown in Figure 1 for illustration purposes. Disregarding the spatial and temporal measurements correlation, by enforcing the assumption of independence can lead to large errors in the posterior distribution of the inferred parameters, as correlation has been shown to have an impact on the information content of measurements (Papadimitriou and Lombaert, 2012), the maximum likelihood and maximum a posteriori estimates of the parameters of interest, and the posterior uncertainty (Simoen et al., 2013).

To consider correlations in Bayesian system identification for structures poses a number of challenges for the modeler. An appropriate functional form of the prediction error correlation is not known a priori,

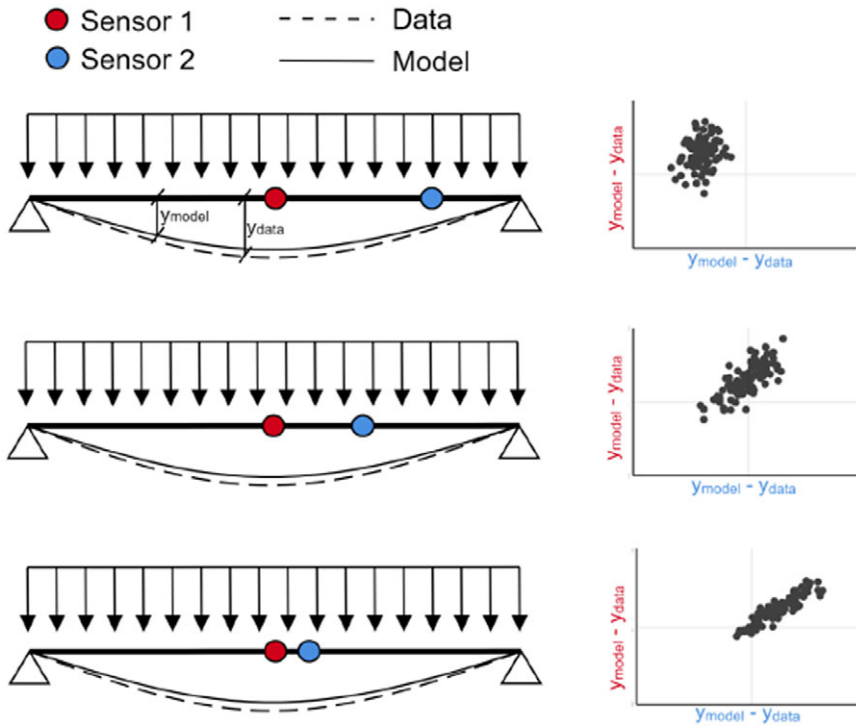


Figure 1. Illustration of the impact of correlation in the model prediction error for the fictitious case of a simply supported beam with two sensors.

and due to the prevalence of the independence assumption, there is limited information available on how to model it. Additionally, it is not known to what degree the correlation is problem-specific. Hence, to perform Bayesian system identification on real-world structures when spatial and/or temporal dependence are present, we identified the following open issues:

1. Appropriate models for the spatial and temporal correlations must be included in the probabilistic model that describes the uncertainties.
2. Bayesian inference must be performed in a computationally efficient manner when large datasets and combined spatial and temporal dependencies are considered.

1.3. Approach

The approach proposed in this article to address the issues mentioned above can be summarized as follows. First, a mathematical model of the data-generating process is formulated. This model is composed of a physical model describing the response of the structure, and a probabilistic model describing the measurement and model prediction error including the spatial and temporal correlation. Both the measurement and model prediction error are taken as normally distributed, and the strength of the correlation is assumed to be dependent on the distance between measurements (in time and/or in space). This dependence is modeled by a set of kernel functions. A pool of candidate models is defined, with each model considering a different kernel function to describe the correlation in the physical model prediction error. Bayesian inference is performed to obtain the posterior distribution of physical and probabilistic model parameters based on the data. The posterior probability and Bayes factor are calculated for each candidate model, making it possible to evaluate how strongly a given model is supported relative to the other candidate models based on the data. The proposed approach is illustrated in Figure 2. More details on the individual building blocks are given in Section 3.

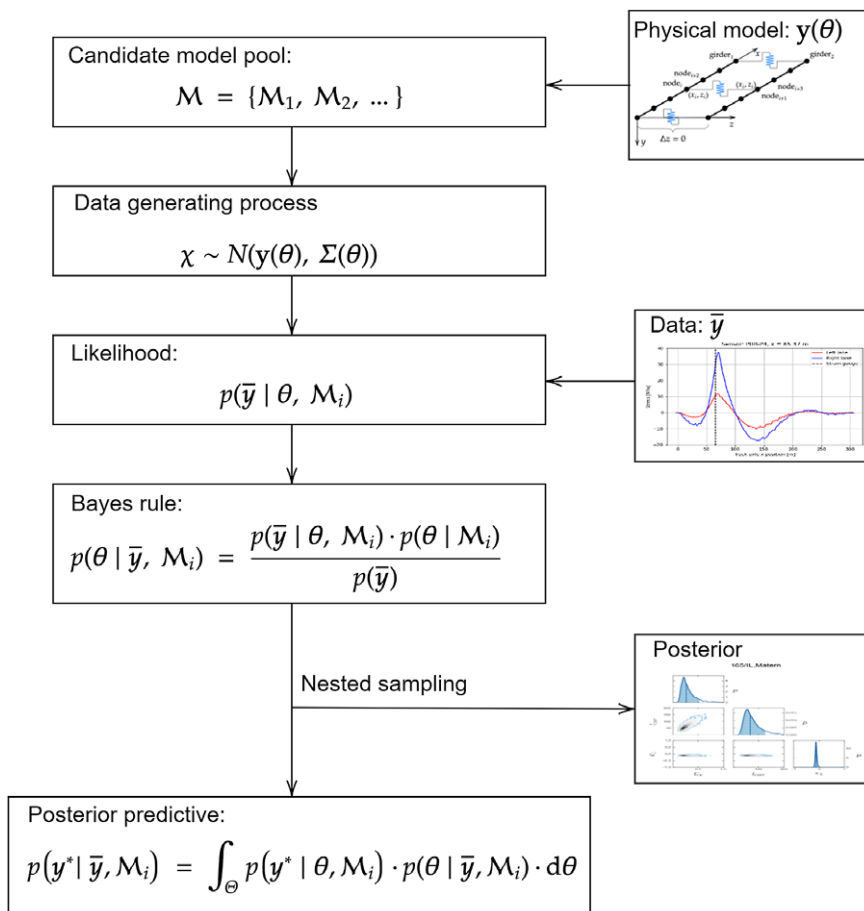


Figure 2. Overview of the Bayesian inference approach used in this work.

Second, a strategy is presented for performing system identification for relatively large datasets ($N > 10^2$ for temporal dependencies and $N > 10^3$ for combined spatial and temporal dependencies) by efficiently evaluating the log-likelihood and the evidence. We propose a procedure for exact and efficient log-likelihood calculation by (i) assuming separability of the spatial and temporal correlation (Genton, 2007); (ii) exploiting the Markov property of the Exponential kernel (Marcotte and Allard, 2018); and (iii) using the nested sampling strategy (Skilling, 2006) to reduce the computational cost of estimating the evidence under each model. The accuracy of the proposed approach is initially investigated in a case study using synthetic data, and subsequently, the feasibility of the approach for its use in real-world cases is demonstrated through a twin-girder steel road bridge case study. In the real-world use case, stress influence lines obtained from controlled load tests are used to estimate the posterior distribution of a set of uncertain, unobservable parameters. The accuracy and uncertainty of the posterior predictive stress distributions obtained from each candidate model are compared, to determine the benefit of using a larger dataset and considering dependencies.

2. Previous work

In the Bayesian system identification literature, it is typically assumed that the prediction error is Gaussian white noise, that is, uncorrelated with zero mean (Mthembu et al., 2011; Chiachio et al., 2015; Pasquier and Smith, 2015; Astroza et al., 2017). In some studies, for example, Goller and Schueller (2011) and

Ebrahimian et al. (2018) the variance of the model prediction error is included in the vector of inferred parameters, however, dependencies are not considered. In other works, such as Simoen et al. (2015), Pasquier and Marcotte (2020), and Vereecken et al. (2022) the parameters that define the uncertainty (with or without considering dependencies) are estimated using a subset of the available data. This approach, however, results in the use of data for inferring nuisance parameters and may not be practical when limited data is available. Examples of inference of the uncertainty parameters can be found in applications outside of structural engineering, for example, in geostatistics (Diggle and Ribeiro, 2002).

To the best of the authors knowledge, Simoen et al. (2013) is the only work concerning model prediction error correlation in Bayesian system identification for structures and investigates the impact of considering dependencies in model prediction error in Bayesian system identification. The aforementioned study presents an approach with many similarities to the proposed one. Bayesian inference and model selection are applied to a pool of candidate models to infer the distribution of uncertain parameters and to determine the strength of the evidence in favor of each model. The physical and probabilistic parameters are inferred for a simple linear regression example as well as a reinforced concrete beam example using modal data. In both cases, the datasets are composed of synthetic observations polluted with correlated noise. Furthermore, the posterior distributions are assumed to be Gaussian, allowing for a computationally efficient asymptotic approximation to be utilized to obtain the posterior and evidence.

We focus on the feasibility of the approach in a practical application with real-world data, where the ground truth of the correlation structure and parameters are not known and the posterior and evidence are not approximated analytically. We instead utilize nested sampling to estimate the evidence, ensuring the applicability of the approach in cases where the Gaussian assumption for the posterior is not valid. Additionally, we address the issue of efficiently calculating the log-likelihood for large datasets with combined spatial and temporal dependencies under the assumption of separable space and time covariance.

3. Methods and tools

3.1. Continuous Bayes theorem

The Bayes theorem of conditional probability for continuous random variables can be written as (Gelman et al., 2013):

$$p(\boldsymbol{\theta}|\bar{\mathbf{y}}, \mathcal{M}) = \frac{p(\bar{\mathbf{y}}|\boldsymbol{\theta}, \mathcal{M}) \cdot p(\boldsymbol{\theta}|\mathcal{M})}{\int_{\Theta} p(\bar{\mathbf{y}}|\boldsymbol{\theta}, \mathcal{M}) \cdot p(\boldsymbol{\theta}|\mathcal{M}) \cdot d\boldsymbol{\theta}}, \quad (1)$$

where $\boldsymbol{\theta}$ is a vector of uncertain parameters; $\bar{\mathbf{y}}$ a vector of observations; $\mathcal{M}$ denotes the model; $p(\boldsymbol{\theta}|\bar{\mathbf{y}}, \mathcal{M})$ is the posterior distribution; $p(\bar{\mathbf{y}}|\boldsymbol{\theta}, \mathcal{M})$ is the likelihood; and $p(\boldsymbol{\theta}|\mathcal{M})$ is the prior.

It can be seen that $p(\boldsymbol{\theta}|\bar{\mathbf{y}}, \mathcal{M})$ describes the posterior distribution of the model parameter set $\boldsymbol{\theta}$ conditional on measurements $\bar{\mathbf{y}}$ under model $\mathcal{M}$. The likelihood term gives the probability of observing $\bar{\mathbf{y}}$ given parameters $\boldsymbol{\theta}$. Finally, the denominator on the right-hand side is known as the evidence, or marginal likelihood, and gives the likelihood of obtaining the measurements conditional on the model $\mathcal{M}$. Obtaining the evidence is necessary for performing Bayesian model selection. In most practical applications this integral is high-dimensional (see e.g., Lye et al., 2021) and computationally intractable. Furthermore, the conventional Markov Chain Monte Carlo (MCMC) methods (Metropolis et al., 1953; Hastings, 1970)—typically used in Bayesian inference—are primarily geared toward estimating the posterior, and do not compute the evidence. The nested sampling method, implemented in the *Dynesty* Python package (Speagle, 2019) is utilized to overcome this limitation. In this approach, the posterior is separated into nested slices of increasing likelihood. Weighted samples are generated from each slice and subsequently recombined to yield the posterior and evidence. Nested Sampling can deal effectively with moderate to high-dimensional problems (e.g., in problems with up to 100 parameters) and multi-modal posteriors. The

reader is referred to Skilling (2006) and Speagle (2019) for detailed information on the nested sampling method.

3.2. Bayesian model selection

The posterior distribution of the parameters for a given set of data is conditional on the model $\mathcal{M}$. Often, multiple models can be defined a priori to describe the observed behavior. To select the most plausible model, a pool of models is defined and inference is performed conditional on each model $\mathcal{M}_i$. The Bayes rule can then be applied to select the most likely model based on the evidence. In Hoeting et al. (1999), the following equation is provided for performing Bayesian model selection:

$$p(\mathcal{M}_i|\bar{\mathbf{y}}) = \frac{p(\bar{\mathbf{y}}|\mathcal{M}_i) \cdot p(\mathcal{M}_i)}{\sum_{i=1}^K p(\bar{\mathbf{y}}|\mathcal{M}_i) \cdot p(\mathcal{M}_i)}, \tag{2}$$

where $p(\mathcal{M}_i|\bar{\mathbf{y}})$ is the posterior probability of model i ; $p(\bar{\mathbf{y}}|\mathcal{M}_i)$ is the evidence under model $\mathcal{M}_i$; and $p(\mathcal{M}_i)$ is the prior probability of model i .

Given a pool of models, selecting the model that best fits the data can straightforwardly be achieved by selecting the model that minimizes a particular error metric between measurements and model outputs. However, simply choosing the model that best fits the data could potentially lead to overfitting: more complicated models would tend to fit the data best, making them the most likely in this approach even if the added complexity provides a negligible benefit. An advantage of Bayesian model selection is that it automatically enforces model parsimony, also known as Occam’s razor as discussed in MacKay (2003) and Beck and Yuen (2004), penalizing overly complex models. It should be emphasized that a high posterior model probability does not necessarily indicate that a particular model provides a good fit with the data, since the model probabilities are conditioned on the pool of candidate models. Therefore, a high posterior model probability can only be interpreted as a particular model being more likely, relative to the other models that are considered. To aid the interpretation of the results, the relative plausibility of two models $\mathcal{M}_1$ and $\mathcal{M}_2$ belonging to a class of models can be expressed in terms of the Bayes factor:

$$R = \frac{p(\mathcal{M}_1|\bar{\mathbf{y}})}{p(\mathcal{M}_2|\bar{\mathbf{y}})} \cdot \frac{p(\mathcal{M}_2)}{p(\mathcal{M}_1)}. \tag{3}$$

An advantage of using the Bayes factor over the posterior model probabilities for model selection is that it can be readily interpreted to indicate the support of one model over another, and thus offers a practical means of comparing different models. The interpretation of Jeffreys (2003) is used in this work, given in Table 1.

3.3. Posterior predictive distribution

Bayesian system identification can be used to obtain point estimates and posterior distributions of uncertain parameters using physical models and measurement data. However, directly using the point estimates of the inferred parameters to make predictions would result in underestimation of the uncertainty and overly

Table 1. Interpretation of the Bayes factor from Jeffreys (2003)

R	Strength of evidence
$<10^0$	Negative
10^0 to $10^{1/2}$	Barely worth mentioning
$10^{1/2}$ to 10^1	Substantial
10^1 to $10^{3/2}$	Strong
$10^{3/2}$ to 10^2	Very strong
$>10^2$	Decisive

confident predictions. This is due to the fact that using point estimates to make predictions disregards the uncertainty in the inferred parameters resulting from lack of data. In contrast, the posterior predictive is a distribution of possible future observations conditioned on past observations taking into account the combined uncertainty from all sources (e.g., modeling and measurement error and parameter uncertainty). The posterior predictive can be obtained as Gelman et al. (2013):

$$p(\mathbf{y}^*|\bar{\mathbf{y}}) = \int_{\Theta} p(\mathbf{y}^*|\boldsymbol{\theta}) \cdot p(\boldsymbol{\theta}|\bar{\mathbf{y}}) \cdot d\boldsymbol{\theta}, \quad (4)$$

where $\mathbf{y}^*$ is a vector of possible future observations.

3.4. Data-generating process

In order to perform system identification, the likelihood function is formulated based on the combination of a probabilistic model and a deterministic physical model. This coupled probabilistic-physical model is used to represent the process that is assumed to have generated the measurements, referred to as the *data-generating process*. Details on the deterministic physical model are provided in Section 5.2. The probabilistic model is used to represent the uncertainties that are inherent when using a model to describe a physical system. The following sources of uncertainty are considered:

- Measurement uncertainty
- Physical model uncertainty

Measurement uncertainty refers to the error between the measured response quantities and the true system response, caused by the combined influence of sensor errors and environmental noise (Kennedy and O'Hagan, 2001). Modeling uncertainty can contain several components and refers to the error between reality and the models used to represent it. These errors arise, for example, due to simplifications in the physical model and numerical approximations.

In this article, we consider data-generating processes based on a multiplicative and additive model prediction error, which will be explained in the following subsections. In the following expressions, Greek letters are used to represent random variables, while bold lower and upper-case letters denote vectors and matrices respectively.

3.4.1. Multiplicative model

The data-generating processes described by equation (5) are obtained by considering the discrepancies between the deterministic model output and the real system response, a process referred to as *stochastic embedding* in Beck (2010). In this model of the data-generating process, a multiplicative prediction error is considered:

$$\boldsymbol{\chi}(\boldsymbol{\theta}) = \mathbf{Y}(\boldsymbol{\theta}_s)\boldsymbol{\eta}_m(\boldsymbol{\theta}_c) + \boldsymbol{\varepsilon}(\boldsymbol{\theta}_c), \quad (5)$$

where $\boldsymbol{\chi}$ is a vector of predictions obtained from the coupled physical-probabilistic model of the data generating process; $\mathbf{Y}$ is a diagonal matrix of physical model predictions obtained as $\mathbf{Y} = \text{diag}(\mathbf{y})$, with $\mathbf{y}$ denoting the corresponding vector of predictions; $\boldsymbol{\eta}_m$ is a vector of multiplicative physical model error factors; $\boldsymbol{\varepsilon}$ is a vector of measurement error random variables; $\boldsymbol{\theta}_s$ is a vector of physical model parameters to be estimated; $\boldsymbol{\theta}_c$ is a vector of probabilistic model parameters to be estimated; and $\boldsymbol{\theta} = \{\boldsymbol{\theta}_s, \boldsymbol{\theta}_c\}$ is the set of combined physical and probabilistic model parameters to be estimated.

In this model formulation, the error in the physical model prediction is assumed to scale with the magnitude of the model output. This assumption is prevalent in the structural reliability literature (Cervenka et al., 2018; Sykorka et al., 2018). The physical model predictions are multiplied by a factor $\boldsymbol{\eta}_m$, expressing the discrepancy between model prediction and reality. A correlated Multivariate Normal distribution with a mean of 1.0 and covariance matrix $\boldsymbol{\Sigma}_\eta$ is assumed for $\boldsymbol{\eta}_m$ as shown in equation (6):

$$\boldsymbol{\eta}_m(\boldsymbol{\theta}_c) \sim \mathcal{N}(1.0, \boldsymbol{\Sigma}_\eta(\boldsymbol{\theta}_c)). \quad (6)$$

The assumption of a Gaussian distribution for $\boldsymbol{\eta}_m$ is made primarily for simplicity and computational convenience. The impact of this assumption is deemed to be outside the scope of this work and is not further examined. The measurement error is taken as independent, identically distributed (i.i.d.) Gaussian random variables, distributed as $\boldsymbol{\varepsilon} \sim \mathcal{N}(0, \sigma_\varepsilon)$. The assumption of Gaussian white noise for the measurement error is prevalent in the literature and is commonly used in Bayesian system identification for structures (see Section 2), stemming from the fact that measurement noise can be considered as a sum of a large number of independent random variables. Modeling the measurement error as i.i.d. realizations from a Normal distribution is therefore justified by the central limit theorem. Utilizing the affine transformation property of the Multivariate Normal distribution we obtain the following model for the data-generating process:

$$\boldsymbol{\chi}_m \sim \mathcal{N}(\mathbf{y}(\boldsymbol{\theta}_s), \mathbf{Y}\boldsymbol{\Sigma}_\eta\mathbf{Y}^T + \sigma_\varepsilon^2\mathbf{I}), \quad (7)$$

with $\mathbf{I}$ being the identity matrix. The residuals between measurements and model predictions are considered as a random field, with the position of each observation defined by a spatial coordinate (representing the location of a sensor) and a temporal coordinate, denoted as x_i and t_i , respectively. The position of an observation $\mathbf{y}$ is described by a two-dimensional vector $\mathbf{x}_i = (x_i, t_i)$, and the random field is represented as a (not necessarily regular) grid of points, as shown in Figure 3 with n denoting the total number of sensors and m denoting the number of observations per sensor over time.

The correlation in the model prediction error between two points $\mathbf{x}_i = (x_i, t_i)$ and $\mathbf{x}_j = (x_j, t_j)$ is obtained as the product of the spatial and temporal correlation, described in terms of the respective kernel functions:

$$\rho_{i,j} = k_x(x_i, x_j; \boldsymbol{\theta}_c) \cdot k_t(t_i, t_j; \boldsymbol{\theta}_c), \quad (8)$$

where $k_x(x_i, x_j; \boldsymbol{\theta}_c)$ and $k_t(t_i, t_j; \boldsymbol{\theta}_c)$ are parametrized by the set of parameters of the probabilistic model $\boldsymbol{\theta}_c$. The standard deviation of the model prediction error at a point i is obtained as $\sigma_i = C_v \cdot y_i$, where C_v denotes the coefficient of variation (COV) of the model prediction error. Calculating the covariance for every pair of points yields a symmetric positive semi-definite covariance matrix $\boldsymbol{\Sigma}_\eta$ that describes the covariance of the physical model prediction error for every point in the random field:

$$\boldsymbol{\Sigma}_p = \mathbf{Y}\boldsymbol{\Sigma}_\eta\mathbf{Y}^T = \begin{bmatrix} \sigma_1^2 & \dots & \sigma_1 \cdot \sigma_N \cdot \rho_{1,N} \\ \vdots & \ddots & \vdots \\ \sigma_N \cdot \sigma_1 \cdot \rho_{N,1} & \dots & \sigma_N^2 \end{bmatrix}, \quad (9)$$

with $N = n \cdot m$.

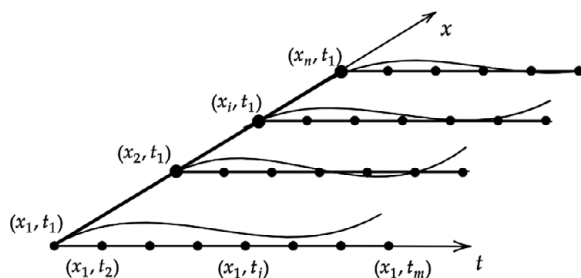


Figure 3. Illustration of space and time coordinate system. Influence lines along the time axis t are obtained for each sensor position x .

3.4.2. Additive model

A coupled probabilistic-physical model of the data-generating process based on a correlated, additive model prediction error is also considered (equation (10)). In this case, the model prediction error is described by an additive term, modeled as a multivariate normal distribution with zero mean and covariance $\Sigma_\eta(\theta_c)$. Similarly to the multiplicative model, the measurement error is represented by a vector of i.i.d. normal random variables, with a mean of zero and standard deviation σ_ε .

$$\chi_a(\theta) = y(\theta_s) + \eta_a(\theta_c) + \varepsilon(\theta_c), \quad (10)$$

where η_a denotes the vector of additive model prediction error.

$$\eta_a(\theta_c) \sim \mathcal{N}(0.0, \Sigma_\eta(\theta_c) + \sigma_\varepsilon^2 \mathbf{I}). \quad (11)$$

4. Efficient log-likelihood evaluation

The assumption of a multiplicative or additive physical model uncertainty factor described by a Gaussian distribution leads to a multi-variate Gaussian likelihood description. For a given covariance matrix Σ , and omitting the dependence on the parameter vector θ from the right-hand side of the equation for brevity, the multivariate normal log-likelihood function can be expressed as:

$$\mathcal{L}(\theta) = -\frac{1}{2} \cdot [\log |\Sigma| + (\bar{y} - y)^T \Sigma^{-1} (\bar{y} - y) + N \cdot \log(2 \cdot \pi)]. \quad (12)$$

The evaluation of the Multivariate Gaussian log-likelihood (equation (12)) requires calculating the determinant ($|\cdot|$) and inverse ($(\cdot)^{-1}$) of the covariance matrix Σ . These operations typically have $O(N^3)$ time complexity and $O(N^2)$ memory requirements for factorizing and storing the covariance matrix respectively, making the direct evaluation of the log-likelihood infeasible for more than a few thousand points.

To address this issue, we present an approach for efficient log-likelihood evaluation under the *multiplicative* model uncertainty with additive Gaussian noise described in Section 3.4.1. For the case of *additive* model uncertainty, described in Section 3.4.2 we utilize an existing approach from the literature. A comparison of the average wall clock time required for log-likelihood evaluation as a function of the size of a 2-dimensional grid of measurements against a naive implementation using the full covariance matrix, for both the additive and multiplicative cases, can be found in Koune (2021). A Python implementation of both methods is available at <https://github.com/TNO/tripty>.

4.1. Efficient log-likelihood evaluation for combined spatial and temporal correlation and multiplicative model prediction uncertainty

To reduce the computational complexity for evaluating the log-likelihood under the *multiplicative* model uncertainty, we propose an approach that utilizes the tridiagonal inverse form of the correlation matrix that can be obtained from the Exponential kernel, as well as the Kronecker structure of the separable space and time covariance matrix.

In the following, it is assumed that the correlation is exponential in time. No assumptions are made regarding the structure of the correlation in space or the number of spatial dimensions. The i, j th element of the temporal covariance matrix Σ_t is obtained as:

$$\Sigma_t^{ij} = C_{v,i} \cdot C_{v,j} \cdot \exp\left(-\frac{\|t_i - t_j\|}{l_{\text{corr}}}\right), \quad (13)$$

where l_{corr} is the correlation length and C_v is the coefficient of variation described in Section 3.4.1. It is shown by Pasquier and Marcotte (2020) that the inverse of the covariance matrix for this kernel function has a symmetric tridiagonal form:

$$\Sigma_t^{-1} = \begin{bmatrix} d_1 & c_1 & & & \\ c_1 & d_2 & c_2 & & \\ & c_2 & d_3 & c_3 & \\ & & \ddots & \ddots & \ddots \end{bmatrix}. \quad (14)$$

Following Cheong (2016), the diagonal vectors of diagonal and off-diagonal terms in equation (14) can be obtained analytically, eliminating the need to form the full correlation and covariance matrices which is often computationally intensive due to the amount of memory and operations required. For a given vector of observations with coordinates $\mathbf{t} = \{t_1, t_2, \dots, t_m\}$ denoting $\Delta t_i = |t_i - t_{i-1}|$ for $i \in [m]$ yields equation (15).

$$a_i = e^{-\lambda \Delta t_i}, \quad (15)$$

where λ is the inverse of the correlation length l_{corr} and a_i is the correlation between points i and $i-1$. The diagonal and off-diagonal elements of the inverse correlation matrix Σ_t^{-1} can then be obtained analytically, eliminating the need for direct inversion of Σ_t and reducing computational complexity and memory requirements:

$$d_1 = \frac{1}{C_{v,1}^2} \cdot \frac{1}{1-a_2^2}, \quad (16)$$

$$d_m = \frac{1}{C_{v,m}^2} \cdot \frac{1}{1-a_m^2}, \quad (17)$$

$$d_{ii} = \frac{1}{C_{v,i}^2} \cdot \left(\frac{1}{1-a_i^2} + \frac{1}{1-a_{i+1}^2} - 1 \right), \quad (18)$$

$$c_{ii-1} = -\frac{1}{C_{v,i} \cdot C_{v,i+1}} \cdot \frac{a_i}{1-a_i^2}. \quad (19)$$

Furthermore, we define the combined space and time covariance which can be obtained as the Kronecker product of the temporal correlation matrix and the spatial correlation matrix, $\Sigma_\eta = \Sigma_t \otimes \Sigma_x$. Using the properties of the Kronecker product, it can be shown that the resulting inverse matrix Σ_η^{-1} has a symmetric block tridiagonal form:

$$\Sigma_\eta^{-1} = \begin{bmatrix} \mathbf{D}_1 & \mathbf{C}_1 & & & \\ \mathbf{C}_1 & \mathbf{D}_2 & \mathbf{C}_2 & & \\ & \mathbf{C}_2 & \mathbf{D}_3 & \mathbf{C}_3 & \\ & & \ddots & \ddots & \ddots \end{bmatrix}. \quad (20)$$

We consider the covariance matrix for the data-generating process defined in equation (7). Expressing the physical model uncertainty covariance matrix Σ_p in terms of the combined space and time covariance Σ_η yields:

$$\Sigma_p = \mathbf{Y} \Sigma_\eta \mathbf{Y}^T. \quad (21)$$

Then Σ_p^{-1} will also be block tridiagonal. However, including additive noise such that $\Sigma_p = \mathbf{Y} \Sigma_\eta \mathbf{Y}^T + \sigma_\epsilon^2 \mathbf{I}$ leads to a dense inverse matrix. To efficiently evaluate the likelihood, we aim to calculate the terms $(\bar{\mathbf{y}} - \mathbf{y})^T \Sigma_p^{-1} (\bar{\mathbf{y}} - \mathbf{y})$ and $|\Sigma_p|$ in equation (12) without explicitly forming the corresponding matrices or directly inverting the covariance matrix, while taking advantage of the properties described previously to reduce the complexity. Algebraic manipulation of the product $(\bar{\mathbf{y}} - \mathbf{y})^T \Sigma_p^{-1} (\bar{\mathbf{y}} - \mathbf{y})$ is performed in order to obtain an expression that can be evaluated efficiently by taking advantage of the Kronecker structure and block symmetric tridiagonal inverse of the covariance matrix. We apply the Woodbury matrix identity given below:

$$(A^{-1} + BC^{-1}B^T)^{-1} = A - AB(C + B^T AB)^{-1}(AB)^T. \quad (22)$$

Substituting $A \rightarrow \Sigma_\varepsilon^{-1}$, $B \rightarrow Y$, and $C^{-1} \rightarrow \Sigma_\eta$, yields:

$$\Sigma_p^{-1} = \Sigma_\varepsilon^{-1} - (\Sigma_\varepsilon^{-1} Y)(\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y)^{-1} (\Sigma_\varepsilon^{-1} Y)^T. \quad (23)$$

Applying the left and right vector multiplication by y , the second term in the r.h.s. of equation (12) becomes:

$$y^T \Sigma_p^{-1} y = y^T \Sigma_\varepsilon^{-1} y - y^T (\Sigma_\varepsilon^{-1} Y)(\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y)^{-1} (\Sigma_\varepsilon^{-1} Y)^T y. \quad (24)$$

In the previous expression, the term $y^T \Sigma_\varepsilon^{-1} y$ can be efficiently evaluated as the product of vectors and diagonal matrices. Similarly, the term $y^T (\Sigma_\varepsilon^{-1} Y)$ can be directly computed and yields a vector. We consider the following term from the r.h.s. of equation (24):

$$(\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y)^{-1} (\Sigma_\varepsilon^{-1} Y)^T y = X. \quad (25)$$

We note that the term $\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y$ is the sum of a symmetric block tridiagonal matrix Σ_η^{-1} and the diagonal matrix $Y^T \Sigma_\varepsilon^{-1} Y$. We can therefore take advantage of efficient algorithms for Cholesky factorization of symmetric block tridiagonal matrices and for solving linear systems using this factorization to compute X . Furthermore, the Cholesky factors obtained previously are also used to reduce the computational cost of evaluating the determinant $|\Sigma| = |\Sigma_\varepsilon + Y \Sigma_\eta Y^T|$. Applying the determinant lemma for Σ_p yields equation (26):

$$|\Sigma_\varepsilon + Y \Sigma_\eta Y^T| = |\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y| \cdot |\Sigma_\eta| \cdot |\Sigma_\varepsilon|. \quad (26)$$

The determinant $|\Sigma_\eta|$ can be calculated efficiently by utilizing the properties of the Kronecker product, given that $|\Sigma_\eta| = |\Sigma_t \otimes \Sigma_x|$. Furthermore, Σ_ε is a diagonal matrix meaning that the determinant can be trivially obtained. Finally, we have previously calculated the Cholesky factorization of the term $\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y$. Using the fact that the determinant of a block triangular matrix is the product of the determinants of its diagonal blocks and the properties of the determinant, the first expression in the r.h.s. of equation (26) can be computed with:

$$|\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y| = |LL^T| = |L| \cdot |L^T| = |L|^2, \quad (27)$$

where the matrix L is the lower triangular Cholesky factor of $\Sigma_\eta^{-1} + Y^T \Sigma_\varepsilon^{-1} Y$. Since each block L_{ii} is also triangular, the evaluation of the determinant has been reduced to evaluation of the determinant of each triangular block L_{ii} , which is equal to the product of its diagonal elements.

Using the above calculation procedure, an efficient solution can also be obtained for the case of only temporal correlation, where Σ_t^{-1} has the symmetric tridiagonal form given in equation (14). The term $\Sigma_t^{-1} + Y^T \Sigma_\varepsilon^{-1} Y$ will be the sum of a symmetric tridiagonal and a diagonal matrix. From a computational viewpoint, this property is advantageous as it allows for a solution to the system of equations equation (25) with $O(N)$ operations using the Thomas algorithm (Quarteroni et al., 2007). Alternatively, for improved efficiency and numerical stability, a Cholesky decomposition can be applied to solve the linear system and calculate the determinants of the symmetric tridiagonal terms in equation (26).

4.2. Efficient log-likelihood evaluation for combined spatial and temporal correlation and additive model prediction uncertainty

To reduce the computational complexity of the log-likelihood evaluation in the case of additive model prediction uncertainty and combined spatial and temporal correlation, we use an approach that utilizes the properties of the Kronecker product and the eigendecomposition of the separable covariance matrix. For a detailed description of this approach, the reader is referred to Stegle et al. (2011).

5. Description of the IJssel bridge case study

5.1. Description of the structure

The IJssel bridge is a twin-girder steel plate road bridge that carries traffic over the river IJssel in the direction of Westervoort. It consists of an approach bridge and a main bridge, of which the latter is of interest in this case. The main bridge has a total length of 295 m and five spans with lengths of 45, 50, 105, 50, and 45 m. In total, the bridge has 12 supports. An elevation view of the structure is shown in Figure 4. The supports at pillar H are hinges, while the rest are roller bearings in the longitudinal directions. The roller bearings at pillars G and K can resist uplift forces. The deck structure of the bridge is composed of two steel girders with variable height, ranging from 2.4 to 5.3 m, and cross-beams with a spacing of approximately 1.8 m. The main girders and cross beams support the steel deck plate. The deck plate has a thickness of 10 or 12 mm and 160×8 mm longitudinal bulb stiffeners. The cross beams are placed with a center-to-center distance of 1.75 to 1.80 m and are composed of a 500×10 mm web with a 250×12 mm welded flange. The cross beams are tapered in the parts that extend beyond the main girders and the beam height is reduced to 200 mm at the beam ends. The two main girders are coupled with K-braces located below every second or third cross beam, with a distance of 5.4 m on average.

5.2. Description of the physical model

A two-dimensional twin-girder finite element (FE) model based on Euler-Bernoulli beam elements is used to model the IJssel bridge, shown in Figure 5. Each element has four degrees of freedom (DOFs): two translations and two rotations. The variable geometrical properties of the steel girders along the longitudinal axis are taken into account by varying the structural properties of the individual beam elements, where each element has a prismatic cross-section. In addition to the main girder, half the width of the deck plate and the corresponding longitudinal stiffeners are also considered in the calculation of the structural properties for each cross-section along the x -axis. The maximum beam element length can be specified in order to approximate the variable geometry of the main girders along the length of the bridge to the required precision. A maximum element length of 2.0 m was used for all simulations, resulting in a

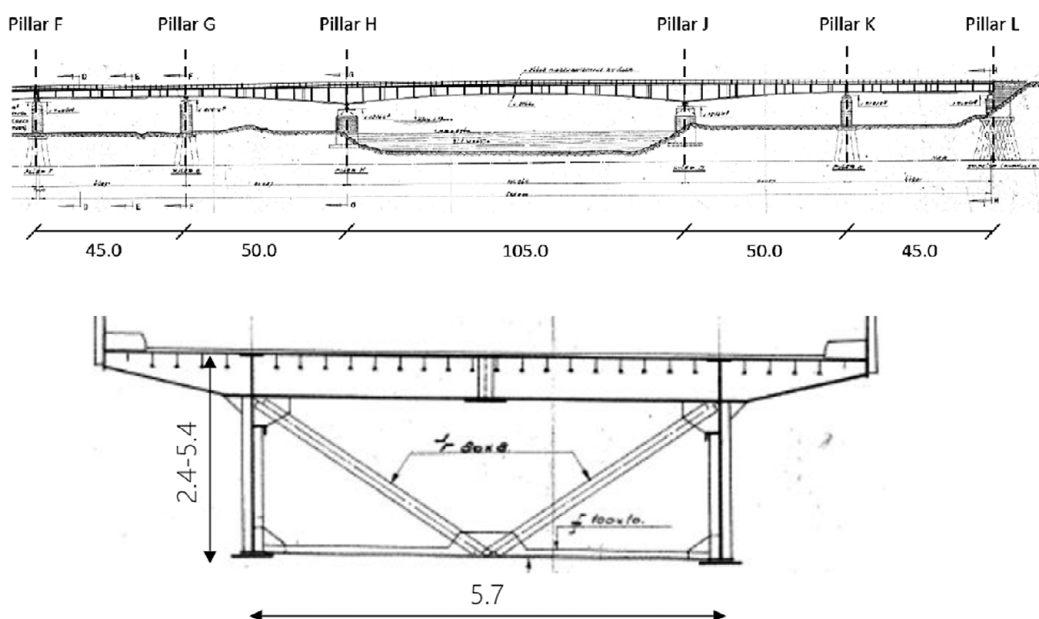


Figure 4. Elevation view of the IJssel bridge (top), and typical cross-section including K-brace (bottom), with lengths shown in meters (from Rijkswaterstaat).

model with $n_{\text{node}} = 193$ nodes and $n_{\text{DOF}} = 386$ DOFs. The coupling between the main girders due to combined stiffness of the deck, crossbeams, and K-braces is simulated by vertical translational springs, placed approximately at the positions of the K-braces that connect the two main girders of the IJssel bridge. Six pinned supports are specified for each girder at locations corresponding to pillars F through L. Independent rotational springs are defined at the supports to simulate the friction at the support bearings and to account for the possibility of partial locking.

During the measurement campaign, the bridge was closed for traffic and only loaded by a heavy weighted truck at the left and right lanes. To account for the position of the truck along the transverse direction (z -axis)—which is not included in the FE model—each load is multiplied by a value of the Lateral Load Function (LLF), as illustrated in Figure 5 (right). The LLF is taken as a linear function, with slope and intercept coefficients such that: (i) a point load applied directly on a girder does not affect the other girder; and (ii) a load applied at the center of the bridge deck is equally distributed between the left and right girders.

5.3. Measurements

The data used in this study is obtained from a measurement campaign performed under controlled load tests. A total of 34 strain measurement sensors were placed on the top and bottom flanges of the steel girders, the cross beams, the longitudinal bulb stiffeners, and the concrete approach bridge, to measure the response of the structure to traffic loads. In this article, we use only a subset of sensors that are placed on the center of the bottom flange of the right main girder, since they measure the global response of the bridge. These sensors are denoted with the prefix “H” (Figure 6). The exact position of each considered sensor along the length of the bridge and the sensor label are provided in Table 2. It should be noted that the authors were given access to the experimental data and details regarding the sensor network, data acquisition system, and experimental procedure of the measurement campaign, however, the location of the sensors was not chosen specifically for the purposes of the present study. Additional information regarding the measurement campaign can be found in Appendix B.

Time series of the strain ε at a sampling rate of 50 Hz are obtained from each sensor and postprocessed to yield the corresponding influence lines. These strain influence lines are converted to stress influence lines using Hooke's law $\sigma = E \cdot \varepsilon$, where σ denotes the stress and E denotes Young's modulus. The latter is taken as $E = 210$ GPa, as specified in the IJsselbridge design. Each sensor yields two influence lines, one for each lane that was loaded during the measurement campaign. Linear interpolation is performed to

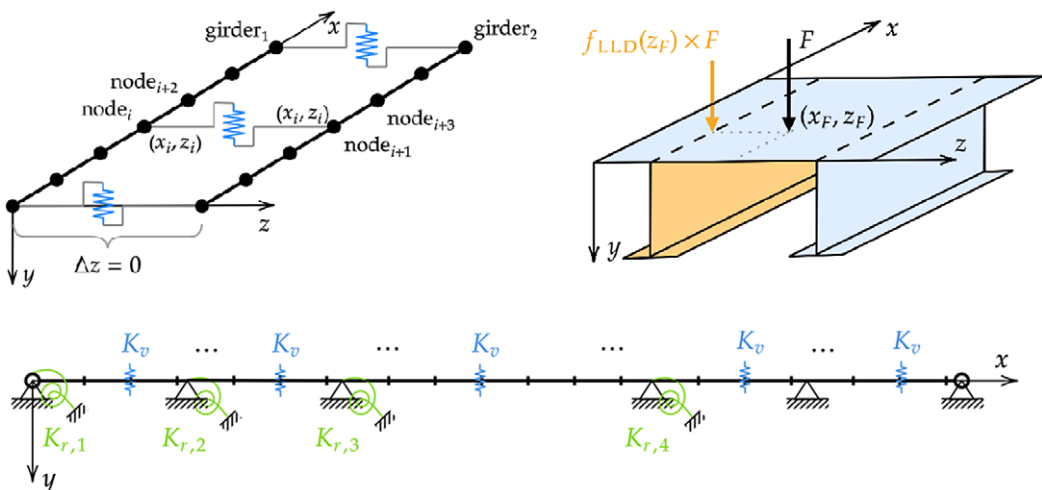


Figure 5. Illustration of the IJsselbridge FE model (left), lateral load function (right), and parametrization of the FE model (bottom).

obtain the stresses at the locations along the length of the bridge corresponding to the locations of the nodes of the FE model. The processed influence lines are plotted in Figure 7.

It is noted that significant discrepancies were observed between the measurements and FE model predictions for a number of sensors during a preliminary comparison of model predictions using a model fitted to the measurements with conventional numerical optimization. It was determined, after verifying the validity of the measurements, that this can be attributed to the simplicity of the model, which was not able to capture the structural behavior at a number of sensor locations. The following list is a summary of the physical simplifications and assumptions that could potentially contribute to the observed discrepancies:

- Only limited number of structural elements are explicitly modeled, with elements such as stiffeners, K-braces, the steel deck, and cross-beams being only considered implicitly (e.g., by modifying the cross-sectional properties of the elements representing the main girders), or omitted entirely.
- The 3D distribution of forces within the elements is neglected, loads are lumped to the closest node and supports are considered as points, potentially overestimating the stresses.
- Although likely negligible, the stiffness of the deck and cross-beams between the two girders is only considered implicitly as lumped stiffness in the vertical springs.
- Variations in the geometry and cross-section properties of the K-braces along the length of the bridge are not taken into account.
- Shear lag in the deck is not modeled.

The uncertainty in the physical model prediction resulting from the aforementioned simplifications and misspecifications is accounted for in the Bayesian inference by considering the additive and multiplicative error models, introduced in Section 3.4.

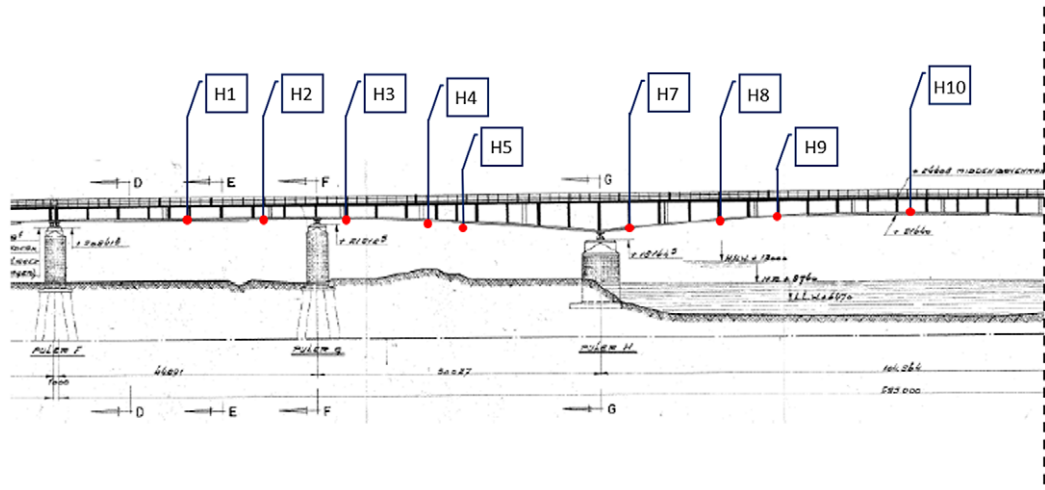


Figure 6. Approximate location of sensors on the right girder. The prefix “H” is used to denote the sensors on the main structure of the IJsselbridge. Adapted from a Rijkswaterstaat internal report.

Table 2. Names, labels, and positions of strain gauges placed on the IJsselbridge main girder

Sensor	H1	H2	H3	H4	H5	H7	H8	H9	H10
Position (m)	20.42	34.82	47.70	61.97	68.60	96.80	113.35	123.90	147.50

Note. The positions are measured from pillar F (see Figure 4).

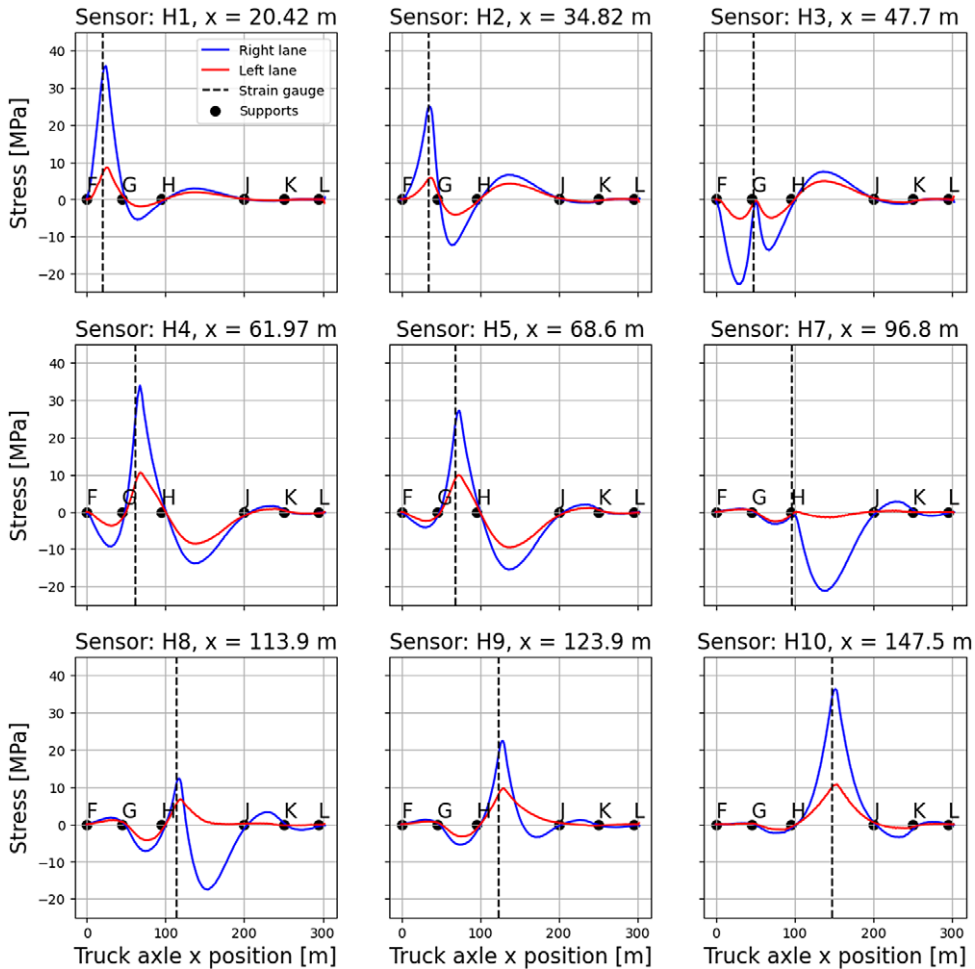


Figure 7. Stress influence lines obtained from the measurement campaign. The blue and red lines correspond to the measured response for different truck positions in the transverse direction of the bridge.

5.4. Correlation functions

Correlation functions also referred to as kernels or kernel functions in the literature and throughout this work, are positive definite functions of two Euclidean vectors $k(\mathbf{x}, \mathbf{x}'; \boldsymbol{\theta}_c)$ (Duvenaud, 2014) that describe the correlation between points $\mathbf{x}$ and $\mathbf{x}'$. Table 3 provides a summary of the kernel functions used to model the correlation in the model prediction error in this article. These kernel functions were chosen due to their wide adoption in statistical applications, ease of implementation, and small number of parameters. Additionally, these kernel functions were empirically found to result in more accurate posterior and posterior predictive distributions among a group of candidate kernel functions in Koune (2021).

5.5. Physical model parameters

The set of physical model parameters to be inferred, $\boldsymbol{\theta}_s$, is shown in Figure 5. The choice of the uncertain physical model parameters is based primarily on sensitivity analyses, the damage mechanisms expected to affect the behavior of the structure and consultation with steel bridge experts. Independent rotational springs are defined at supports F, G, H, and J, with the corresponding rotational spring stiffnesses denoted as $K_{r,1}$ through $K_{r,4}$, to simulate the friction at the support bearings and to account for the possibility of partial locking. An analysis of the sensitivity of the stress response to the stiffness of the support bearings

Table 3. List of correlation functions and corresponding parameters

Kernel	Shorthand	$k(\mathbf{x}, \mathbf{x}')$	Parameters
Independent	IID	$\begin{cases} 1 & \text{if } \mathbf{x} = \mathbf{x}' \\ 0 & \text{if } \mathbf{x} \neq \mathbf{x}' \end{cases}$	—
Radial basis	RBF	$\exp\left(-\frac{\ \mathbf{x}, \mathbf{x}'\ ^2}{2l_{\text{corr}}^2}\right)$	l_{corr}
Exponential	EXP	$\exp\left(-\frac{\ \mathbf{x}, \mathbf{x}'\ }{l_{\text{corr}}}\right)$	l_{corr}

Table 4. Description and uniform prior distribution bounds of physical model parameters

Parameter	Unit	Description	Lower bound	Upper bound
$\log_{10}(K_{r,1-4})$	kNm/rad	Rotational spring stiffness at the supports F to J (see Figure 5)	4.0	10.0
$\log_{10}(K_v)$	kN/m	Stiffness of the vertical translational springs, representing the K-braces	0.0	8.0

indicated that the stiffness at supports K and L has negligible influence on the stress at the sensor locations. Therefore, the bearings at supports K and L are assumed to be hinges with no rotational stiffness. Furthermore, the riveted connections present in the K-braces result in high uncertainty on their stiffness, while additionally this stiffness was found to have a significant effect on the FE model response. The stiffness parameter of the vertical springs, K_v , is assumed to be equal for all vertical springs along the length of the bridge. Both the vertical and rotational spring stiffness values span several orders of magnitude. We therefore represent these parameters in terms of their base-10 logarithms. In addition to the more natural and convenient representation, the parametrization by the base-10 logarithm yields a reduction of the relative size of the prior to the posterior, and therefore a reduction of the number of samples required for convergence when using nested sampling.

Table 4 summarizes the prior distributions that are used for $\log_{10}(K_r)$ and $\log_{10}(K_v)$. The prior distributions are determined using a combination of engineering judgment and sensitivity analysis. Additional details of the sensitivity analysis for the rotational and vertical spring stiffness parameters can be found in Appendix A. It should be noted that the impact of the physical model parameterization on the Bayesian model selection is not considered in this study. The feasibility of inferring parameters of the probabilistic model, and selecting the most likely model from a pool of candidate models, when a large number of physical model parameters is present is left as a potential topic for future work.

5.6. Probabilistic model parameters

Table 5 provides an overview of the set of probabilistic model parameters to be inferred, θ_c , and their prior distributions. Details of the probabilistic model formulations for each of the cases investigated are provided in Sections 6 and 7. The support (domain) of the priors must be defined such that it is possible to capture the structure of the correlations in the residuals between model predictions and measurements. It should be noted that choosing the priors for the parameters of the probabilistic model is not a simple task as no information on the correlation structure is available. Furthermore, a poor choice of prior can significantly impact the inference and prediction, leading to wide credible intervals (Fuglstad et al., 2018). Uniform priors are chosen with supports that are wide enough to capture a range of correlations that are expected to be present in the measurements.

5.7. Notation

The coupled probabilistic-physical models considered in the case studies presented in Sections 6 and 7 are distinguished by the model prediction error which can be multiplicative equation (5) or additive equation (10), the kernel function used in the probabilistic model and the number of measurements used in the inference per influence line. A shorthand notation will be adopted to refer to the models for the remainder of this article. Models will be referred to by the kernel that describes the temporal correlation (as shown in Table 3) and the type of model prediction error considered, that is, multiplicative or additive, denoted by the suffixes “M” and “A,” respectively. As an example, following this notation, a model with multiplicative error where complete independence between the errors is considered will be written as IID-M.

6. Exploratory analyses on the IJsselbridge case study using synthetic measurements

6.1. Description

Initially, a synthetic example is studied to explore the impact of the size of the dataset on the feasibility of inferring the functional form the probabilistic model and the posterior distribution of the uncertain parameters. A pool of candidate probabilistic models with different correlation functions is formed, as described in Table 6. For each probabilistic model, a number of synthetic datasets with varying size (i.e., varying spatial and temporal discretization) are generated by evaluating the response of the physical model and contaminating this response with random samples drawn from the probabilistic model for a prescribed set of ground truth values of the parameters. Bayesian inference is then performed for all models for each dataset, using the nested sampling method described in Section 3.1 to estimate the posterior distribution and evidence. The resulting point estimates of the parameters of the probabilistic model are compared to the ground truth, under the assumption that the probabilistic model used to generate each dataset is known a priori. Additionally, the most likely model corresponding to each dataset is determined by comparing the estimated evidence under each probabilistic model for different dataset sizes. The aim of this synthetic case study is twofold:

- To investigate the impact of the size of the dataset, and the functional form of the probabilistic model, on the accuracy of maximum a posteriori (MAP) estimates of the parameters of the probabilistic model.
- To gain insight into the size of the dataset required to correctly identify the probabilistic model from a pool of candidate models.

The example is structured as follows: The response of the physical model is evaluated for a set of ground truth parameters for increasing numbers of sensors per span, and considering an equal number of measurements in time. Denoting the number of sensors per span with N_x and the number of points per influence line by N_t (with each sensor yielding one influence line for each of the two traffic lanes), the resulting physical model response is a rectilinear grid with $N_x = N_t$. In this manner, the physical model response is evaluated for $N_x = \{1, 2, \dots, 10\}$ sensors per span, with each sensor yielding two influence

Table 5. Description and uniform prior distribution bounds of probabilistic model parameters

Parameter	Unit	Description	Lower bound	Upper bound
C_v	(—)	COV of the multiplicative model prediction error	0.0	1.0
σ_{model}	MPa	Standard deviation of the additive model prediction error	0.0	5.0
σ_{meas}	MPa	Standard deviation of the additive measurement error	0.0	1.0
$l_{\text{corr},t}$	m	Temporal correlation lengthscale	0.0	300.0
$l_{\text{corr},x}$	m	Spatial correlation lengthscale	0.0	300.0

Table 6. Overview of models used in the case with synthetic measurements

Model	Shorthand	Temporal correlation	Spatial correlation	θ_c
$\mathcal{M}_1$	IID-M	Independent	Independent	$C_v, \sigma_{\text{meas}}$
$\mathcal{M}_2$	RBF-M	Radial Basis	Exponential	$C_v, \sigma_{\text{meas}}, l_{\text{corr},x}, l_{\text{corr},t}$
$\mathcal{M}_3$	EXP-M	Exponential	Exponential	$C_v, \sigma_{\text{meas}}, l_{\text{corr},x}, l_{\text{corr},t}$
$\mathcal{M}_4$	IID-A	Independent	Independent	σ_{model}
$\mathcal{M}_5$	RBF-A	Radial basis	Exponential	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},x}, l_{\text{corr},t}$
$\mathcal{M}_6$	EXP-A	Exponential	Exponential	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},x}, l_{\text{corr},t}$

Note. See Table 3 for the meaning of the abbreviations and the details of the correlation function.

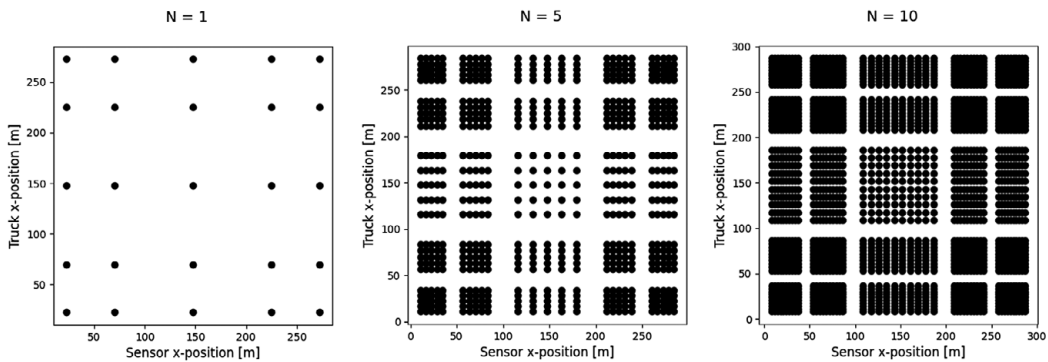


Figure 8. Rectilinear grid of sensor and measurement positions considered in the synthetic case study for $N_x = N_t = \{1, 5, 10\}$.

lines, each with $N_t = \{1, 2, \dots, 10\}$ measurements respectively. The sensors are placed such that they are equally spaced within each span, with the distance between sensors taken as $L_{\text{span}}/(N_x + 1)$ (Figure 8), where L_{span} denotes the length of the span. The physical model response is then contaminated with noise samples drawn from the probabilistic models summarized in Table 6. To reduce the impact of the random sampling of the noise on the synthetic case study results, 50 samples are generated from each probabilistic model and for each grid size, and the resulting MAP estimates and the evidence for each model are averaged over the samples. The number of simulations was chosen as a compromise between having a large enough sample size to minimize the effects of the random sampling in the synthetic dataset, and the computational cost of performing multiple Bayesian inference realizations for a large number of models over a range of grid sizes.

6.2. Results

Bayesian inference is performed for the models $\mathcal{M}_1$ to $\mathcal{M}_6$ listed in Table 6. At each grid size, we examine if the correlation parameters are correctly inferred, when performing inference with the true probabilistic model used to generate the measurements, by computing the relative error between the known ground truth for each probabilistic model and the mean MAP point estimate obtained from the Bayesian inference. The scatter of the MAP estimates for each ground truth model, expressed in terms of the COV is also computed in order to quantify the impact of the random sampling as a function of the grid size. The results are shown in Figure 9. A grid size of 10×10 (corresponding to $N = 2$ sensors per span) is sufficient to obtain an accurate point estimate of the multiplicative model prediction uncertainty COV parameter C_v , the standard deviation of the additive model prediction uncertainty σ_{model} and the standard deviation of the measurement uncertainty σ_{meas} (i.e., the relative error is below 0.10). For models with

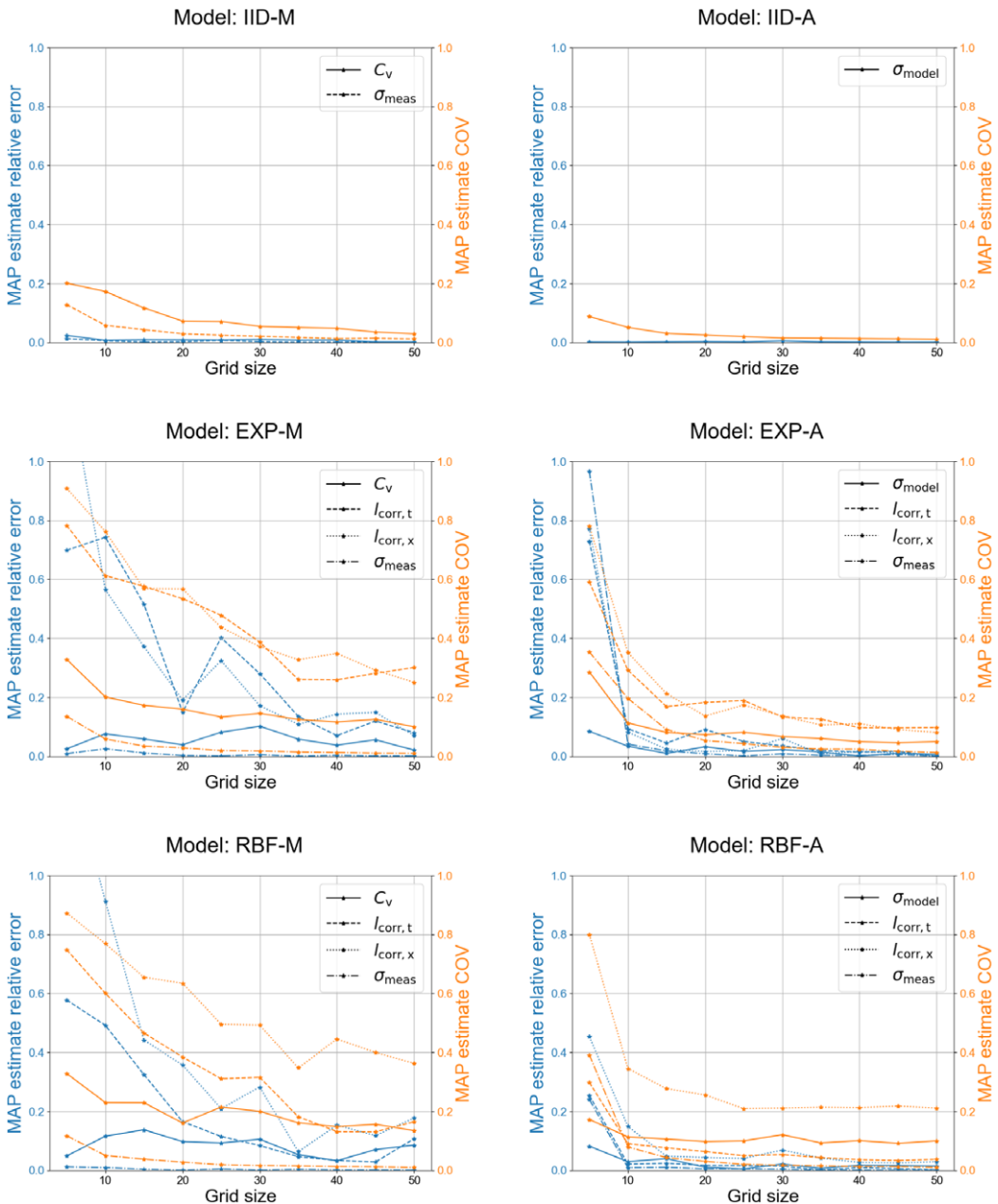


Figure 9. Relative error of the mean MAP estimates of probabilistic model parameters compared to the ground truth, and COV of MAP estimates as a function of grid size.

multiplicative model prediction uncertainty, both the relative error and the COV of the MAP estimate of the spatial and temporal correlation length parameters are heavily dependent on the grid size. This is not the case for models with additive prediction uncertainty, where the relative error is not significantly affected when increasing the grid resolution beyond 10×10 . This could potentially be explained by considering that the multiplicative error structure introduces a dependency of the probabilistic model on the physical model. We speculate that this additional complexity may hinder the estimation of the posterior distribution of the parameters of multiplicative models, resulting in slower convergence (with

respect to the grid size) of the point estimates of the probabilistic model parameters to the ground truth values. It is noted that the accuracy of the obtained point estimates for the correlation length parameters is also expected to depend on their size relative to the distances between points on the measurement grid. However, this dependence is not further examined in this article.

By calculating the log-evidence for each of the models $\mathcal{M}_1$ to $\mathcal{M}_6$ we investigate if the correct probabilistic model is identified from the pool of candidate models for each dataset. The log of the average evidence—based on 50 generated datasets—obtained for each probabilistic model and for each ground truth model and grid size is shown in Table 7. The correct ground truth model is identified with a single sensor per span in the case of additive model prediction uncertainty. For multiplicative models, the ground truth is correctly identified for two to three sensors per span.

Additionally, the posterior probability of the ground truth model p_{gt} , and the identification accuracy (i.e., the percentage of realizations for which the ground truth model obtains the highest posterior probability) for increasing refinement of the grid of sensors are also shown in Table 7. The posterior probabilities are obtained using equation (2), with the evidence for each model taken as the average across the independent realizations. It is observed that multiplicative models generally require finer discretization and larger datasets, compared to additive models, to achieve similar levels of accuracy. For the additive models, the ground truth model is recovered with perfect accuracy for three or more sensors per span, with the exception of the IID-A case. The comparatively low posterior probability of the ground truth model observed for the IID-A case can be attributed to the fact that all of the considered models are able to describe uncorrelated additive Gaussian noise.

6.3. Conclusions

Based on the results presented previously it can be seen that, although the estimated evidence for each model is sensitive to the size of the dataset, the correct probabilistic model used to generate the data can be identified in all cases, for as few as three sensors per span. Furthermore, it is evident that the correlation structure and the size of the dataset impact the accuracy of the posterior distributions and point estimates of the parameters describing the uncertainty. For additive probabilistic models with correlation, the MAP estimates of the parameters converge faster (i.e., for a smaller number of measurements) to the known ground truth values, compared to the multiplicative probabilistic models.

7. Analyses on the IJsselbridge case study using real-world measurements

7.1. Analysis considering a single-sensor

An initial analysis is performed considering data from the H4 sensor (see Figure 7), with the aim of assessing the feasibility of performing system identification while considering dependencies in the model prediction uncertainties, and to examine the benefit of the additional data compared to a reference case where only four hand-picked measurements from the largest peaks and troughs of each influence line are considered. A second analysis with multiple sensors is performed to demonstrate the feasibility of considering large datasets under combined spatial and temporal dependencies and to determine the efficiency of the block Cholesky log-likelihood evaluation presented in Section 4.1.

Table 8 provides an overview of the considered models in the real-world case study, including their labels, correlation structure, dataset size, and parameters. Compared to the synthetic case study, two “reference” analyses are added, referred to as “REF,” in which we adopt a small dataset with 4×2 points. These analyses represent a typical application of Bayesian system identification for structures, where only a limited number of manually selected measurements are included in the dataset. In these reference analyses, we assume that the discrepancies between measurement and model prediction are fully independent. The four largest (in absolute value) peaks per influence line are chosen to form this dataset, such that they maximize the amount of information regarding the parameters of interest while being spaced far enough apart to be considered independent. The assumption of independence is based mainly on engineering judgment, which is often the case in applications where a limited amount of measurements makes it infeasible to assess their independence.

Table 7. Log of the mean evidence, posterior probability of the ground truth model, and identification accuracy per model as a function of the number of sensors per span for different ground truth models, averaged over 50 randomly generated datasets

N	1	2	3	4	5	6	7	8	9	10
<u>IID-M</u>	−66.50	− 267.45	− 587.77	− 1050.10	− 1635.36	− 2357.95	− 3198.63	− 4210.59	− 5329.55	− 6539.73
RBF-M	− 66.08	−273.26	−594.69	−1057.69	−1644.28	−2365.88	−3207.05	−4220.28	−5340.00	−6550.03
EXP-M	−66.44	−272.45	−594.31	−1057.76	−1644.18	−2365.74	−3206.78	−4219.59	−5340.46	−6549.96
IID-A	−74.11	−292.36	−634.86	−1156.65	−1758.60	−2602.01	−3562.04	−4604.75	−5903.94	−7294.54
RBF-A	−76.82	−294.67	−638.15	−1160.86	−1764.36	−2609.82	−3564.69	−4608.00	−5913.43	−7302.32
EXP-A	−76.87	−294.77	−637.77	−1160.42	−1764.56	−2609.67	−3564.56	−4607.45	−5913.02	−7302.19
p_{gt}	0.28	0.99	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00
Acc.	0.62	0.96	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00
IID-M	−72.97	−262.56	−569.84	−1024.87	−1624.73	−2283.02	−3133.34	−4035.82	−5174.47	−6323.96
<u>RBF-M</u>	− 71.06	−260.07	− 547.66	− 968.03	− 1536.89	− 2159.71	− 2947.94	− 3854.52	− 4875.21	− 6022.68
EXP-M	−71.61	− 259.56	−549.72	−970.29	−1542.18	−2162.22	−2956.65	−3864.16	−4886.69	−6027.02
IID-A	−80.98	−270.99	−594.33	−1091.26	−1701.98	−2344.76	−3233.32	−4073.89	−5341.58	−6585.27
RBF-A	−82.65	−275.41	−598.35	−1080.67	−1673.37	−2303.21	−3182.04	−4018.69	−5187.90	−6314.14
EXP-A	−82.50	−275.12	−598.38	−1082.46	−1674.78	−2307.78	−3185.81	−4021.77	−5201.11	−6327.58
p_{gt}	0.58	0.36	0.89	0.91	0.99	0.92	1.00	1.00	1.00	0.99
Acc.	0.44	0.72	0.86	0.92	0.88	0.96	0.96	0.98	1.00	1.00
IID-M	−61.64	−266.86	−590.31	−1033.72	−1619.24	−2317.31	−3158.95	−4162.92	−5233.14	−6482.46
RBF-M	− 59.92	−263.45	−579.58	−999.59	−1550.39	−2227.75	−3030.65	−3944.78	−5015.33	−6114.11
<u>EXP-M</u>	−60.67	− 262.29	− 575.65	− 996.73	− 1541.55	− 2216.66	− 3018.81	− 3918.58	− 4975.22	− 6090.35
IID-A	−67.82	−296.88	−614.16	−1092.10	−1689.28	−2426.96	−3270.71	−4385.12	−5494.79	−6846.72
RBF-A	−70.18	−292.18	−613.42	−1089.07	−1684.85	−2377.04	−3216.52	−4245.98	−5317.11	−6551.88
EXP-A	−70.26	−291.23	−613.26	−1088.96	−1684.45	−2378.49	−3214.03	−4250.08	−5321.82	−6551.58
p_{gt}	0.29	0.76	0.98	0.95	1.00	1.00	1.00	1.00	1.00	1.00
Acc.	0.30	0.82	0.98	0.90	0.98	0.98	1.00	1.00	1.00	1.00
IID-M	−110.91	−457.77	−1032.15	−1822.18	−2882.95	−4150.56	−5678.20	−7375.49	−9373.51	−11566.59
RBF-M	−110.69	−457.87	−1033.20	−1822.45	−2885.07	−4150.93	−5678.80	−7376.49	−9373.58	−11566.89

Continued

Table 7. Continued

N	1	2	3	4	5	6	7	8	9	10
EXP-M	−110.84	−457.80	−1033.07	−1822.57	−2885.07	−4150.86	−5678.61	−7376.35	−9373.13	−11566.89
<u>IID-A</u>	−108.57	−454.82	−1029.84	−1819.42	−2881.06	−4147.81	−5674.38	−7371.99	−9369.32	−11563.00
RBF-A	−110.11	−457.30	−1031.99	−1822.60	−2885.24	−4151.45	−5678.68	−7376.64	−9371.79	−11566.81
EXP-A	−109.91	−457.29	−1031.70	−1822.44	−2885.15	−4151.30	−5678.59	−7376.45	−9371.84	−11566.58
p_{gt}	0.56	0.76	0.69	0.80	0.82	0.82	0.93	0.93	0.82	0.89
Acc.	0.80	0.92	0.92	0.94	0.96	0.94	1.00	0.96	1.00	0.98
IID-M	−116.12	−449.47	−1004.21	−1773.42	−2680.87	−3635.65	−5189.07	−6853.51	−8621.59	−10818.35
RBF-M	−116.25	−449.20	−1003.54	−1727.56	−2590.33	−3540.85	−5144.94	−6612.07	−8247.39	−9938.49
EXP-M	−116.22	−449.45	−1003.71	−1743.65	−2621.39	−3569.22	−5165.94	−6711.81	−8354.34	−10122.47
IID-A	−113.77	−446.42	−1000.72	−1770.73	−2681.02	−3635.19	−5184.55	−6849.24	−8619.79	−10818.07
<u>RBF-A</u>	−102.52	−327.30	−693.87	−1171.40	−1758.91	−2431.99	−3295.85	−4260.14	−5318.40	−6524.88
EXP-A	−104.22	−336.12	−709.04	−1197.58	−1795.79	−2490.29	−3375.42	−4337.38	−5441.78	−6677.48
p_{gt}	0.85	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00
Acc.	0.76	0.98	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00
IID-M	−114.34	−442.17	−996.64	−1777.73	−2726.68	−3959.78	−5384.25	−7211.06	−9076.12	−11267.86
RBF-M	−114.40	−442.10	−996.48	−1775.52	−2707.97	−3911.51	−5211.81	−7000.00	−8721.13	−10778.94
EXP-M	−114.48	−441.95	−996.60	−1775.57	−2715.95	−3918.94	−5241.51	−7056.89	−8805.41	−10875.05
IID-A	−111.73	−439.29	−993.23	−1774.09	−2723.19	−3956.71	−5381.91	−7207.27	−9076.46	−11280.75
RBF-A	−110.85	−393.23	−837.71	−1432.90	−2120.02	−3013.09	−4039.86	−5175.60	−6460.53	−7814.57
<u>EXP-A</u>	−110.56	−389.56	−814.55	−1383.96	−2068.01	−2906.41	−3836.56	−4944.19	−6174.49	−7484.61
p_{gt}	0.47	0.98	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00
Acc.	0.54	0.92	1.00	1.00	1.00	1.00	1.00	1.00	1.00	1.00

Note. Underscores denote the ground truth model used to generate the data and bold type denotes the highest evidence.

Table 8. Overview of models used in the case with real-world measurements

Model	Shorthand	Temporal correlation	Dataset size ^a	θ_c
$\mathcal{M}_1$	IID-M	Independent	193×2	$C_v, \sigma_{\text{meas}}$
$\mathcal{M}_2$	RBF-M	Radial basis	193×2	$C_v, \sigma_{\text{meas}}, l_{\text{corr},t}$
$\mathcal{M}_3$	EXP-M	Exponential	193×2	$C_v, \sigma_{\text{meas}}, l_{\text{corr},t}$
$\mathcal{M}_4$	REF-M	Independent	4×2	$C_v, \sigma_{\text{meas}}$
$\mathcal{M}_5$	IID-A	Independent	193×2	σ_{model}
$\mathcal{M}_6$	RBF-A	Radial basis	193×2	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},t}$
$\mathcal{M}_7$	EXP-A	Exponential	193×2	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},t}$
$\mathcal{M}_8$	REF-A	Independent	4×2	σ_{model}

Note. See Table 3 for the meaning of the abbreviations and the details of the correlation function.

^aThe factor 2 in the dataset size is included to indicate that each sensor yields two influence lines, one for each controlled loading test as discussed in Section 5.3.

Inference is performed using the nested sampling technique, yielding an estimate of the posterior distribution and evidence for each of the models $\mathcal{M}_1 - \mathcal{M}_8$. The IID-M, REF-M, IID-A, and REF-A models consider complete independence in the model prediction uncertainty, while models RBF-M, EXP-M, RBF-A, and EXP-A assume dependencies modeled by exponential and radial basis kernel functions. The uniform prior distributions assumed for the physical model parameters and uncertainty parameters are listed in Tables 4 and 5, respectively. It is noted that in models IID-M and IID-A, complete independence is assumed for the model prediction uncertainty, despite the dense spacing of the measurements. It is therefore expected that the uncertainty in the posterior distributions will be significantly underestimated if dependencies are present in the model prediction error. The posterior distributions obtained by this model are included for the purpose of comparing the inferred means of the parameters with other models that assume dependence, as well as illustrating the effect of increasing the number of points under the independence assumption on the posterior distributions.

Figures 10 and 11 show the posterior distribution of highest density (HD) credible intervals (CIs) for models with multiplicative model prediction uncertainty and additive model prediction uncertainty. Comparing the reference models $\mathcal{M}_4$ and $\mathcal{M}_8$ with the other models, we observe that for the additive probabilistic model, the posterior CIs of the physical model parameters are wider, indicating that the additional information contained in the full dataset of 193×2 data points can result in reduced uncertainty in the posterior distributions of the parameters of interest. This can also be observed for parameters $\log_{10}(K_{r,1})$ and $\log_{10}(K_{r,4})$ in the case of multiplicative model prediction uncertainty. Furthermore, for the RBF-M and EXP-M models higher point estimates are obtained for C_v compared to both the IID-M and REF-M models, indicating that including correlation parameters in the vector of uncertain parameters to be inferred can affect the inference of other parameters of interest, or that it requires an increase in the size of the dataset to obtain accurate estimates and low uncertainty in the posterior distribution of parameters of the probabilistic model.

The number of parameters considered in the Bayesian inference, and the probabilistic model of the data-generating process used to obtain the likelihood function can have a significant impact on the number of likelihood function evaluations required to achieve convergence when performing Bayesian inference using nested sampling. The total number of function evaluations (NFE) per model is shown in the second column of Table 9. A range of approximately $29 \cdot 10^3$ to $109 \cdot 10^3$ likelihood evaluations are needed to achieve convergence depending on the model. The models where correlation is considered generally require a larger number of likelihood evaluations compared to models where the model prediction errors are taken as independent. Although no clear conclusions arise in this case regarding the impact of the number of parameters, correlation structure, and dataset size on the NFE, it is found that the reference cases, and the cases where independence is assumed, generally require considerably fewer likelihood function evaluations compared to models where correlation is considered.

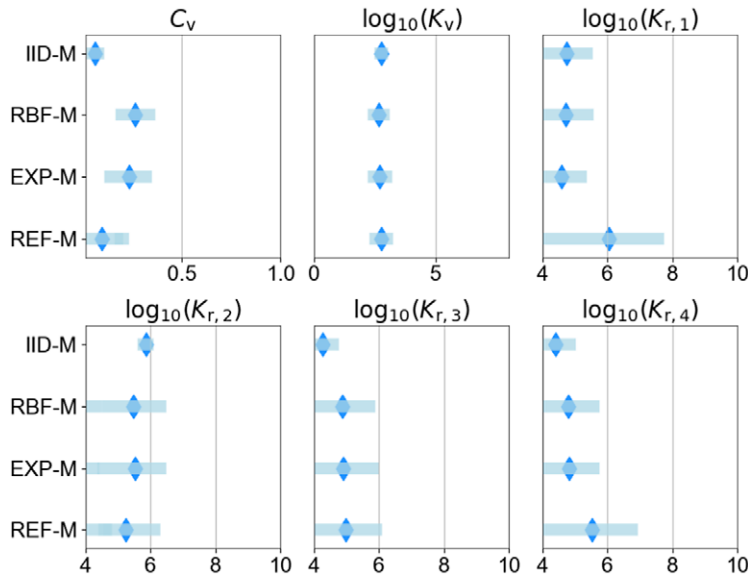


Figure 10. Comparison of posterior mean and 90% HD CIs for models with multiplicative uncertainty structure.

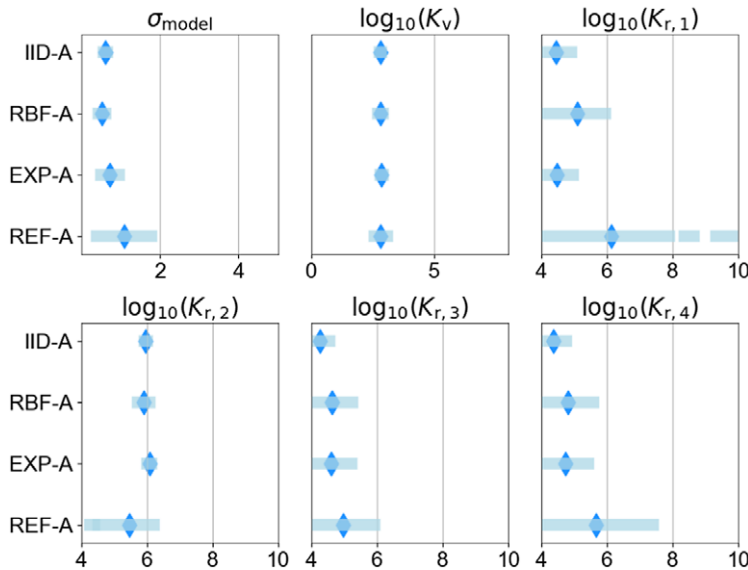


Figure 11. Comparison of posterior mean and 90% HD CIs for models with additive uncertainty structure.

Nested sampling also yields a noisy estimate of the evidence, which can be used to determine the most likely model $\mathcal{M}_i$ within the candidate pool of models. Based on these estimated evidences, the posterior probability of each model in $\mathcal{M}$ is calculated using equation (2). All models in $\mathcal{M}$ are considered a priori equally likely. Table 9 provides the log-evidence, posterior probability, and interpretation of the Bayes factor per model. It is noted that the reference models $\mathcal{M}_4$ and $\mathcal{M}_8$ are not included in the model selection due to the different dataset used with these models, and therefore no evidence, posterior probabilities, or Bayes factors are obtained. It is possible to observe that the assumption of complete independence in the

Table 9. NFE required for convergence rounded to the nearest thousandth, log-evidence, posterior probability, and Bayes factors per model

Model	Shorthand	NFE ($\cdot 1000$)	$\log(\mathcal{Z})$	$p(\mathcal{M})$	Evidence against
$\mathcal{M}_1$	IID-M	64	-358.28	0.00	Decisive
$\mathcal{M}_2$	RBF-M	93	-58.06	0.00	Decisive
$\mathcal{M}_3$	EXP-M	76	-126.95	0.00	Decisive
$\mathcal{M}_4$	REF-M	51	—	—	—
$\mathcal{M}_5$	IID-A	41	-381.15	0.00	Decisive
$\mathcal{M}_6$	RBF-A	109	322.22	0.00	Decisive
$\mathcal{M}_7$	EXP-A	92	349.55	1.00	Barely worth mentioning
$\mathcal{M}_8$	REF-A	29	—	—	—

Note. All models included in the model selection are assumed a priori equally likely.

model prediction uncertainty is not supported by the evidence, which is reflected in the low values of the log evidence for the IID-M and IID-A models. The additive model with exponentially correlated model prediction uncertainty (EXP-A) is the most likely model, with practically 100% posterior probability. This indicates that the data decisively supports the inclusion of correlation. It should be emphasized that the posterior probability only reflects the likelihood of a model compared to other models in the candidate pool, and cannot be used to draw conclusions on the validity of a model in general.

As mentioned in Section 5.6, the choice of prior distributions for the probabilistic model parameters can have a significant impact on the posterior and posterior predictive distributions. It was observed that the models with exponential correlation for the case of a single sensor are particularly sensitive to prior distribution of the correlation length parameter, due to the joint unidentifiability of the marginal variance and the correlation length parameters (Fuglstad et al., 2018). The corresponding joint posterior distributions for the EXP-M and EXP-A models are shown in Figure 12. It can be seen that specifying a large upper bound on the prior of the correlation length will result in wide credible intervals in the posteriors of the corresponding σ_{model} and C_v parameters. This effect is less pronounced in the case of multiple sensors presented in Section 7.2 and furthermore does not affect the RBF kernel in either the single or multiple sensor case.

7.2. Analysis considering multiple sensors

In order to evaluate the feasibility of the approach under combined spatial and temporal dependencies, Bayesian inference is performed considering influence lines from multiple sensors. For the reference models REF-M and REF-A, four measurements are selected per influence line, with each sensor yielding two influence lines. Specifically, one peak stress measurement per span is selected from the four spans with the largest absolute peak stresses. For the remaining models, 193 measurements per influence line are considered from locations along each influence line corresponding to the longitudinal positions of nodes in the FE model. The dataset includes influence lines for truck trial runs on the left and right lanes for sensors H1, H2, H4, H5, H9, and H10. The remaining sensor data exhibits structural behavior that cannot be captured by the twin-girder FE model presented in Section 5.2 and is therefore discarded.

Table 10 provides an overview of the considered models in the analyses using data from multiple sensors. In Section 7.1, with data coming from one sensor we only dealt with temporal correlation. Here, both spatial and temporal correlation is present in the model prediction error. For all models in $\mathcal{M}$ the temporal correlation is described by an exponential kernel function. The spatial correlation is only considered in the models $\mathcal{M}_2$, $\mathcal{M}_3$, $\mathcal{M}_6$, and $\mathcal{M}_7$, also using an exponential kernel function. In the other models, spatial correlation is not taken into account, with each strain gauge along the length of the bridge considered fully independent. By comparing the posterior model probabilities for models with and without spatial correlation, we can evaluate the validity of the spatial correlation assumption, and determine which description of the model prediction uncertainties is best supported by the data.

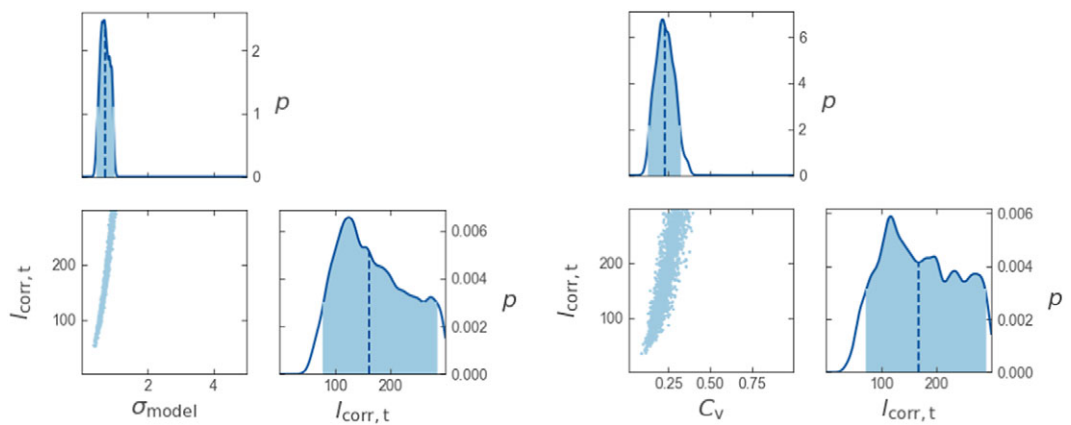


Figure 12. Joint unidentifiability of the correlation length and model prediction uncertainty parameters for the exponential kernel considering additive (left) and multiplicative (right) model prediction uncertainty.

Table 10. Overview of models used in the case with real-world measurements using multiple sensors

Model	Shorthand	Temporal correlation	Spatial correlation	Dataset size ^a	θ_c
$\mathcal{M}_1$	IID-M	Independent	Independent	193×12	$C_v, \sigma_{\text{meas}}$
$\mathcal{M}_2$	RBF-M	Radial basis	Exponential	193×12	$C_v, \sigma_{\text{meas}}, l_{\text{corr},t}, l_{\text{corr},x}$
$\mathcal{M}_3$	EXP-M	Exponential	Exponential	193×12	$C_v, \sigma_{\text{meas}}, l_{\text{corr},t}, l_{\text{corr},x}$
$\mathcal{M}_4$	REF-M	Independent	Independent	4×12	$C_v, \sigma_{\text{meas}}$
$\mathcal{M}_5$	IID-A	Independent	Independent	193×12	σ_{model}
$\mathcal{M}_6$	RBF-A	Radial basis	Exponential	193×12	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},t}, l_{\text{corr},x}$
$\mathcal{M}_7$	EXP-A	Exponential	Exponential	193×12	$\sigma_{\text{model}}, \sigma_{\text{meas}}, l_{\text{corr},t}, l_{\text{corr},x}$
$\mathcal{M}_8$	REF-A	Independent	Independent	4×12	σ_{model}

Note. See Table 3 for the meaning of the abbreviations and the details of the correlation function.
^aThe factor 12 in the dataset size is included to indicate that each of the six sensors yields two influence lines, one for each controlled loading test as discussed in Section 5.3.

The assumption of exponentially correlated model prediction uncertainty in space provides a computational advantage, compared to other kernel functions. Due to poor scaling of the computational complexity of the multivariate Gaussian log-likelihood with the size of the dataset, the use of a conventional approach to evaluate the log-likelihood (e.g., by factorizing the full covariance matrix) would result in significant computational cost even for datasets with a few thousand observations when correlation is taken into account. To alleviate this issue, the efficient log-likelihood evaluation approaches described in Section 4 are utilized for models RBF-M, EXP-M, RBF-A, and EXP-A. For models IID-M and IID-A, efficient log-likelihood evaluation is trivially obtained due to the diagonal structure of the covariance matrix.

The posterior distribution credible intervals (CIs) for the multiplicative and additive uncertainty models are shown in Figures 13 and 14, respectively. For the models with additive model prediction uncertainty, it is observed that the reference case REF-A generally results in wider CIs for the inferred parameters, which is expected given the smaller dataset used in this case. Conversely, IID-A typically yields narrower CIs for the physical model parameters. Under a multiplicative model prediction uncertainty, the reference model REF-M yields the lowest point estimate for the COV parameter C_v . Given that only a few hand-selected peaks are considered for the reference models, it is likely that the model prediction uncertainty is underestimated for the REF-M model due to the omission of measurements across the entire length of the bridge, and particularly at locations near the supports where the

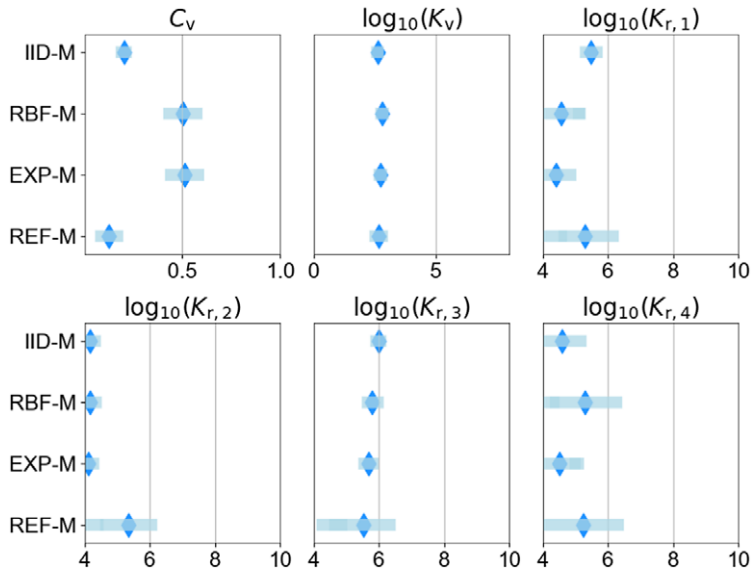


Figure 13. Comparison of posterior mean and 90% highest density credible intervals for models with multiplicative uncertainty structure.

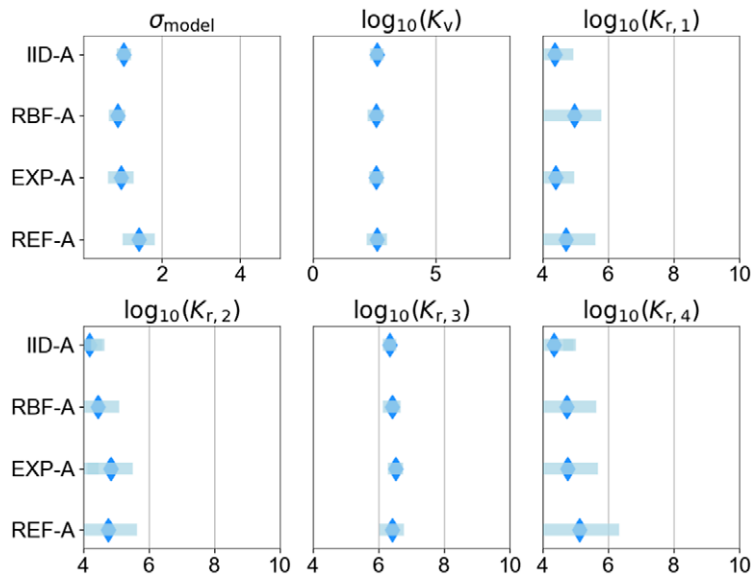


Figure 14. Comparison of posterior mean and 90% highest density credible intervals for models with additive uncertainty structure.

predicted stress is close to zero. This conclusion is supported by observing that IID-M, while having the same probabilistic model as REF-M and a larger dataset, yields a higher point estimate for C_v . Despite the lower C_v , the REF-M model results in wider CIs for the physical model parameters.

The estimated log evidence for each coupled probabilistic-physical model is provided in Table 11. The highest log-evidences are obtained for the models with correlation and additive model prediction uncertainty, RBF-A and EXP-A, while models with multiplicative model prediction uncertainty result in lower evidence, indicating that the mechanisms that contribute to the error between measurements and

Table 11. NFE required for convergence rounded to the nearest thousandth, log-evidence, posterior probability, and Bayes factors per model

Model	Shorthand	NFE ($\cdot 1000$)	$\log(\mathcal{Z})$	$p(\mathcal{M})$	Evidence against
$\mathcal{M}_1$	IID-M	78	−2312.29	0.00	Decisive
$\mathcal{M}_2$	RBF-M	193	−94.12	0.00	Decisive
$\mathcal{M}_3$	EXP-M	140	−440.28	0.00	Decisive
$\mathcal{M}_4$	REF-M	64	—	—	—
$\mathcal{M}_5$	IID-A	48	−3400.96	0.00	Decisive
$\mathcal{M}_6$	RBF-A	204	1693.24	1.00	Barely worth mentioning
$\mathcal{M}_7$	EXP-A	196	1058.73	0.00	Decisive
$\mathcal{M}_8$	REF-A	30	—	—	—

Note. All models included in the model selection are assumed a priori equally likely.

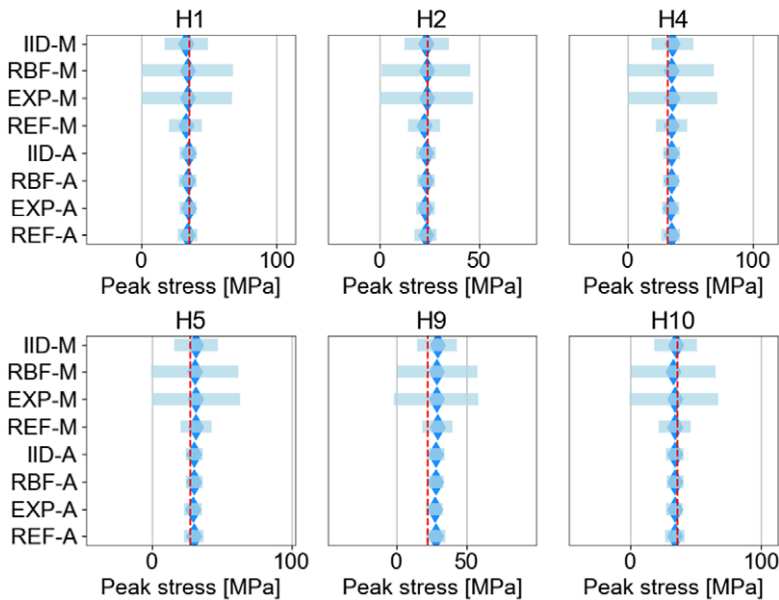


Figure 15. Comparison of median and 90% credible intervals of the posterior predictive stress distribution per model and sensor at the location of peak stress for a truck on the right lane. The red dashed lines denote the measurements.

physical model predictions are closer to being additive in nature for this specific case study. It can also be seen that models where correlation is not considered result in significantly lower evidence. The impact of the probabilistic model and number of parameters on the number of likelihood function evaluations required for convergence is also more pronounced in this case. As shown in the second column of Table 11, the EXP and RBF models require approximately two to four times more function evaluations as the corresponding IID models. This highlights the necessity of utilizing an efficient log-likelihood evaluation method in order to reduce the computational cost to a level that makes this approach feasible in practice. It is also emphasized that the applicability of the proposed approach is limited to cases where a computationally cheap physical model or an efficient surrogate model is available.

The degree to which the size of the dataset and the assumptions on the probabilistic model affect the quality of the prediction are reflected in the posterior predictive distribution. For each model, 2000 samples are drawn from the posterior predictive distribution of the stress influence line. The median and

90% CIs at the location of peak stress for each influence line are plotted in Figure 15, considering a truck load on the right lane. Significant variations in the widths of the posterior predictive CIs are observed, with the multiplicative models yielding wider CIs compared to the additive models. This may be attributed to the sensitivity of the posterior predictive distributions of the multiplicative models to the uncertainty of the COV parameter C_v . Small uncertainties on the posterior of C_v can lead to wide predictive credible intervals at locations where the model response has a large magnitude, such as the peak stress locations shown in Figure 15. Additionally, it is likely that the multiplicative model may not be an appropriate description of the measured data for the considered case, leading to wide posterior distributions of C_v , and consequently wide posterior predictive distributions. This provides additional motivation for performing the model selection procedure described in this study. Furthermore, considering the full dataset and complete independence is shown to result in lower uncertainty in the posterior predictive compared to models where dependence is taken into account for multiplicative models. This is not the case for additive models where no significant differences in the posterior predictive CIs are observed. These results, combined with the calculated evidence and posterior probabilities per model shown in Table 11 indicate that the assumption of complete independence in the model prediction error can result in overconfident posterior and posterior predictive distributions for certain models.

8. Conclusions

An approach is proposed to carry out Bayesian system identification of structures when spatial and temporal correlations are present in the model prediction error, and large datasets (in the order of 10^2 to 10^4 measurements) are available. To address the issue of the computationally expensive likelihood function, which becomes a bottleneck for large datasets, an approach based on the properties of the exponential kernel function is proposed. It is applicable to both additive and multiplicative model prediction error. Moreover, nested sampling is utilized to compute the evidence under each model and apply Bayesian model selection. The nested sampling method has been shown to be effective for high-dimensional and multi-modal posteriors, indicating that the approach presented in this work is applicable to problems with large numbers of uncertain parameters. Investigating the impact of the physical model parameters on Bayesian inference and model selection could be an interesting avenue for future research.

We conducted a case study using both synthetic and real-world measurements to evaluate the feasibility of performing Bayesian system identification for structures with large datasets. Our proposed approach considers spatial and temporal correlation in the model prediction error. The synthetic example aims to investigate the impact of the dataset size and the assumed structure of the correlation on the inference of the parameters of the probabilistic model, as well as the feasibility of inferring the true probabilistic model from a pool of candidate models, particularly when the number of measurements available is limited. In this example, Bayesian inference and model selection are performed for a set of coupled probabilistic-physical models while refining the spatial and temporal resolution of a dataset. The size of the dataset ranges from 25 to 2500 measurements. It is demonstrated that the most likely probabilistic model (which in this case is the a priori known model used to generate the data) can be identified from a pool of candidate models using datasets with as few as 25 measurements. The size of the dataset is found to have an impact on both the identification of the correct probabilistic model and the accuracy of the point estimates obtained for the uncertain probabilistic model parameters, with larger datasets leading to higher accuracy in both tasks. The structure (e.g., multiplicative or additive) and kernel function considered in the probabilistic model are also found to have a significant influence on the accuracy of the obtained point estimates, particularly for the correlation length parameters. Whereas a relative error below 10% is achieved in the MAP estimates of the probabilistic model parameters of additive models for grids with as few as 10×10 measurements, grid sizes of 50×50 or larger are required to obtain similar accuracy with multiplicative models.

Through the IJsselbridge case study, it is demonstrated how Bayesian system identification can be feasibly performed when spatial and temporal dependence might be present. Datasets with up to approximately 2300 measurements are used to infer the uncertain parameters of the coupled physical-

probabilistic models, and Bayesian model selection is applied to determine the most probable model within a pool of candidate models. In addition to updating the physical model and quantifying the uncertainty in the physical model parameters, this approach makes it possible to infer the correlation structure and obtain an improved description of the measurement and model prediction errors. Through the study of the posterior predictive credible intervals and the corresponding model posterior probabilities, both the single and multiple sensor use cases highlight the importance of considering dependencies in the probabilistic model formulation, and demonstrate the benefits of performing Bayesian model selection to determine the posterior probability of different probabilistic models. It is found that the use of a few selected peaks under the i.i.d. assumption for the model prediction error in Bayesian inference can lead to insufficient data and inability to infer all of the parameters of interest, therefore limiting the number and type of parameters that can be identified. Using the full dataset under the assumption of dependence allows for additional parameters to be inferred. In both real-world examples, the estimated Bayes factors and posterior probabilities overwhelmingly favor additive error models where correlation is considered, with the EXP-A and RBF-A models obtaining practically 100% posterior probability for the single and multiple sensor cases respectively.

The proposed approach is based on a multiplicative or additive physical model error. These errors are assumed to follow a Gaussian distribution which determines the likelihood of our model. However, this assumption may not always be realistic in real-world applications. We are currently working on making our model more flexible to better accommodate real-world scenarios.

Acknowledgments. We are grateful to the Dutch Ministry of Public Works and Transport (Rijkswaterstaat) for providing the measurement data used in this study, as well as technical information regarding the IJssel bridge.

Author contribution. Conceptualization: I.K., A.R., A.S., A.C.; Data curation: I.K.; Formal analysis: I.K.; Methodology: I.K., A.R., A.S., A.C.; Supervision: A.R., A.S., A.C.; Writing—original draft: I.K.; Writing—review and editing: A.R., A.S., A.C. All authors approved the final submitted draft.

Competing interest. The authors declare none.

Data availability statement. The code used for the synthetic use case is available at https://github.com/JanKoune/bayesian_si_with_correlation. Restrictions apply to the availability of the real-world measurements, which were used with permission from Rijkswaterstaat.

Funding statement. This publication is part of the project LiveQuay: Live Insights for Bridges and Quay walls (with project number NWA.1431.20.002) of the research programme NWA UrbiQuay which is (partly) financed by the Dutch Research Council (NWO). Part of this research was conducted during the TNO ERP-SI/DT: TNO Early Research Program—Structural Integrity/Digital Twin, use case Steel Bridge. A.C. gratefully acknowledges the financial support provided by the Alexander von Humboldt Foundation Research Fellowship for Experienced Researchers supporting part of this research.

Ethical standard. The research meets all ethical guidelines, including adherence to the legal requirements of the study country.

References

- Astroza R, Ebrahimi H, Li Y and Conte JP (2017) Bayesian nonlinear structural FE model and seismic input identification for damage assessment of civil structures. *Mechanical Systems and Signal Processing* 93, 661–687.
- Barrias A, Casas JR and Villalba S (2016) A review of distributed optical fiber sensors for civil engineering applications. *Sensors* 16(5), 748.
- Beck JL (2010) Bayesian system identification based on probability logic. *Struct. Control and Health Monitoring* 17, 825–847.
- Beck J and Katafygiotis L (1998) Updating models and their uncertainties. I: Bayesian statistical framework. *Journal of Engineering Mechanics* 124, 455–461.
- Beck J and Yuen K (2004) Model selection using response measurements: Bayesian probabilistic approach. *Journal of Engineering Mechanics* 130(2), 192–203.
- Behmanesh I and Moaveni B (2014) Probabilistic identification of simulated damage on the Dowling hall footbridge through Bayesian finite element model updating. *Structural Control and Health Monitoring* 22(3), 463–483.
- Brownjohn J (2007) Structural health monitoring of civil infrastructure. *Philosophical Transactions of the Royal Society A* 365, 589–622.
- Cervenka V, Cervenka J and Kadlec L (2018) Model uncertainties in numerical simulations of reinforced concrete structures. *Structural Concrete* 19(6), 2004–2016.

- Chen H-P** (2018) *Structural Health Monitoring of Large Civil Engineering Structures*. Hoboken, NJ: Wiley-Blackwell.
- Cheong S** (2016) *Parameter Estimation for the Spatial Ornstein-Uhlenbeck Process with Missing Observations*. PhD Thesis, University of Wisconsin Milwaukee UWM.
- Chiachío J, Chiachío M, Saxena A, Sankararaman S, Rus G and Goebel K** (2015) Bayesian model selection and parameter estimation for fatigue damage progression models in composites. *International Journal of Fatigue* 70, 361–373.
- Ching J, Muto M and Beck JL** (2006) Structural model updating and health monitoring with incomplete modal data using Gibbs sampler. *Computer-Aided Civil and Infrastructure Engineering* 21(4), 242–257.
- Diggle PJ and Ribeiro PJ** (2002) Bayesian inference in Gaussian model-based geostatistics. *Geographical and Environmental Modelling* 6(2), 129–146.
- Duvenaud DK** (2014) *Automatic Model Construction with Gaussian Processes*. PhD Thesis, University of Cambridge.
- Ebrahimian H, Astroza R, Conte JP and Papadimitriou C** (2018) Bayesian optimal estimation for output-only nonlinear system and damage identification of civil structures. *Structural Control and Health Monitoring* 25(4), e2128.
- Farrar C and Worden K** (2012) *Structural Health Monitoring: A Machine Learning Perspective*. Hoboken, NJ: Wiley.
- Fuglstad G-A, Simpson D, Lindgren F and Rue H** (2018) Constructing priors that penalize the complexity of Gaussian random fields. *Journal of the American Statistical Association* 114(525), 445–452.
- Gelman A, Carlin JB, Stern HS, Dunson DB, Vehtari A and Rubin DB** (2013) *Bayesian Data Analysis*, 3rd Edn. London: Chapman and Hall/CRC.
- Genton MG** (2007) Separable approximations of space-time covariance matrices. *Environmetrics* 18, 681–695.
- Goller B and Schueller GI** (2011) Investigation of model uncertainties in Bayesian structural model updating. *Journal of Sound and Vibration* 330, 6122–6136.
- Hastings WK** (1970) Monte Carlo sampling methods using Markov chains and their applications. *Biometrika* 57(1), 97–109.
- Hoeting JA, Madigan D, Raftery AE and Volinsky CT** (1999) Bayesian model averaging: A tutorial. *Statistical Science* 14(4), 382–401.
- Huang Y, Shao C, Wu B, Beck JL and Li H** (2019) State of the art review on Bayesian inference in structural system identification and damage assessment. *Advances in Structural Engineering* 22(6), 1329–1351.
- Jeffreys H** (2003) *Theory of Probability*, 3rd Edn. Oxford Classic Texts in the Physical Sciences. Oxford: Clarendon Press.
- Katafygiotis LS, Papadimitriou C and Lam H-F** (1998) A probabilistic approach to structural model updating. *Soil Dynamics and Earthquake Engineering* 17, 495–507.
- Kennedy MC and O'Hagan A** (2001) Bayesian calibration of computer models. *Journal of the Royal Statistical Society: Series B* 63(3), 425–464.
- Koune I** (2021) *Bayesian System Identification for Structures Considering Spatial and Temporal Dependencies*. MSc Thesis, Delft University of Technology.
- Lam HF, Hu Q and Wong MT** (2014) The Bayesian methodology for the detection of railway ballast damage under a concrete sleeper. *Engineering Structures* 81, 289–301.
- Lam H-F, Yang J-H and Au S-K** (2018) Markov chain Monte Carlo-based Bayesian method for structural model updating and damage detection. *Structural Control and Health Monitoring* 25(4), e2140.
- Lye A, Ciciello A and Patelli E** (2021) Sampling methods for solving Bayesian model updating problems: A tutorial. *Mechanical Systems and Signal Processing* 159, 107760.
- MacKay DJC** (2003) *Information Theory, Inference and Learning Algorithms*. Cambridge: Cambridge University Press.
- Marcotte D and Allard D** (2018) Gibbs sampling on large lattice with GMRF. *Computers & Geosciences* 111, 190–199.
- Metropolis N, Rosenbluth AW, Rosenbluth MN, Teller AH and Teller E** (1953) Equation of state calculations by fast computing machines. *The Journal of Chemical Physics* 21(6), 1087–1092.
- Mthembu L, Marwala T, Friswell MI and Adhikari S** (2011) Model selection in finite element model updating using the Bayesian evidence statistic. *Mechanical Systems and Signal Processing* 25, 2399–2412.
- Papadimitriou C and Lombaert G** (2012) The effect of prediction error correlation on optimal sensor placement in structural dynamics. *Mechanical Systems and Signal Processing* 28, 105–127.
- Pasquier P and Marcotte D** (2020) Robust identification of volumetric heat capacity and analysis of thermal response tests by Bayesian inference with correlated residuals. *Applied Energy* 261, 114394.
- Pasquier R and Smith IFC** (2015) Robust system identification and model predictions in the presence of systematic uncertainty. *Advanced Engineering Informatics* 29(4), 1096–1109.
- Quarteroni A, Sacco R and Saleri F** (2007) *Numerical Mathematics*. Berlin: Springer.
- Rogers TJ** (2018) *Towards Bayesian System Identification: With Application to SHM of Offshore Structures*. PhD Thesis, University of Sheffield.
- Simoen E, Papadimitriou C, De Roeck G and Lombaert G** (1998) *Influence of the Prediction Error Correlation Model on Bayesian FE Model Updating Results*. Life-Cycle and Sustainability of Civil Infrastructure Systems.
- Simoen E, Papadimitriou C and Lombaert G** (2013) On prediction error correlation in Bayesian model updating. *Journal of Sound and Vibration* 332(18), 4136–4152.
- Simoen E, Roeck GD and Lombaert G** (2015) Dealing with uncertainty in model updating for damage assessment: A review. *Mechanical Systems and Signal Processing* 56–57, 123–149.
- Skilling J** (2006) Nested sampling for general Bayesian computation. *Bayesian Analysis* 1(4), 833–860.

- Speagle JS** (2019) *Dynesty: A dynamic nested sampling package for estimating bayesian posteriors and evidences*. *arXiv: 1904.02180v1*.
- Stegle O, Lippert C, Mooij JM, Lawrence ND and Borgwardt KM** (2011) Efficient inference in matrix-variate Gaussian models with iid observation noise. *Advances in Neural Information Processing Systems 24*, 630–638.
- Sykorka M, Krejsa J, Ilcch J, Prieto M and Tanner P** (2018) Uncertainty in shear resistance models of reinforced concrete beams according to fib MC2010. *Structural Concrete 19*, 284–295.
- Vereecken E, Slobbe A, Rózsás Á, Botte W, Lombaert G and Caspele R** (2022) Efficient Bayesian model selection and calibration using field data for a reinforced concrete slab bridge. *Structure and Infrastructure Engineering*, DOI: [10.1080/15732479.2022.2131847](https://doi.org/10.1080/15732479.2022.2131847), 1–19.
- Ye XW, Su YH and Han JP** (2014) Structural health monitoring of civil infrastructure using optical fiber sensing technology: A comprehensive review. *The Scientific World Journal 2014*, 1–11.

Appendix A. Sensitivity analysis

For $\log_{10}(K_r)$ we assume that any stiffness value from zero (practically pinned support) to infinity (fixed support) is equally likely. In practice, the stiffness must be finite and therefore a high value is specified as an upper limit instead. This value is determined by calculating the peak stress predicted at each sensor location as a function of $\log_{10}(K_r)$. The upper bound of the support for the prior distribution is chosen as the point where any further increase has a negligible effect on the calculated influence line. The prior of $\log_{10}(K_v)$ is determined in a similar manner. All values between zero (practically uncoupled main girders) and infinity (fully coupled main girders) are considered equally likely, and the difference in peak stress for a truck load applied to the left and right lane is calculated. The analyses described previously are performed for the sensors H1, H5, and H10, which are the closest to each of the first three midspans (see Figure 6). The results for $\log_{10}(K_{r,4})$ and $\log_{10}(K_v)$ are shown in Figure A1.

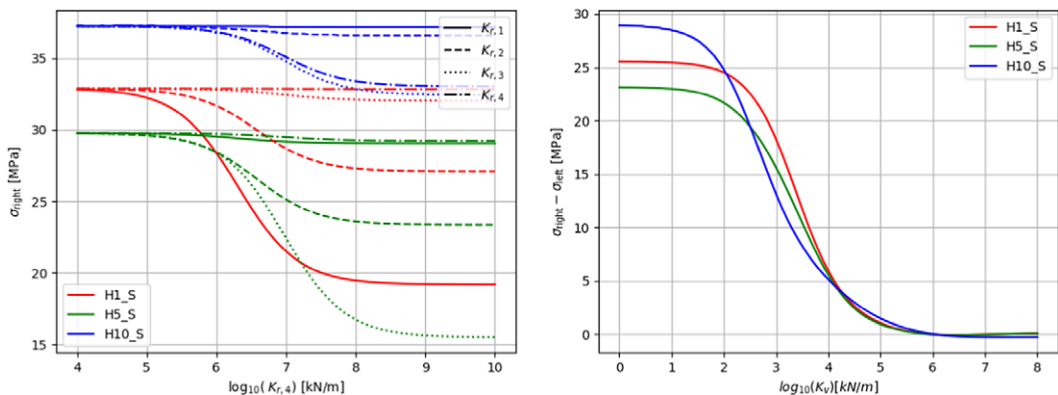


Figure A1. Peak stress response at selected sensors as a function of $\log_{10}(K_r)$ (left) and $\log_{10}(K_v)$ (right).

Appendix B. Measurements

Measurements are obtained by 34 sensors connected by a total of 9 fiber optic lines to an interrogator sampling at a frequency of 50.0 Hz. A total of six tests were performed with trucks driving over the left or right lane at a constant speed, with the truck transverse position roughly corresponding to that of the right or left girder depending on the test. Both the transverse position and speed were manually controlled. Different load tests with various truck speeds were performed. A summary of these tests is provided in Table B1. In this article, we only consider the low-speed tests (i.e., 20 km/h) to measure the structural response under approximately static loading conditions. The axle distance and load per axle of the truck used to perform the controlled loading tests are provided in Table B2.

The truck center of mass is calculated by assuming that the front axle take 12% of the total load, with the remaining axles taking 22% of the load. The center of mass is calculated as:

$$x_{CM} = \frac{\sum w_i \cdot x_i}{\sum w_i}. \quad (\text{B.1})$$

During processing of the measurement data, it was found that the truck speed deviated from the assumed 20 km/h and this deviation should be accounted for in the processing. To implement the correction it was assumed that the influence line peak for each sensor occurs when the truck center of mass coincides with the sensor longitudinal position. The time difference Δt between the peaks of

Table B1. Controlled loading test parameters

Time start (CET)	Lane	Speed (km/h)
21:56:55	Right	20
22:05:55	Left	20
22:21:30	Left	80
22:29:12	Left	80
22:41:25	Right	80
22:49:15	Right	80

Table B2. Properties of truck used in controlled load tests

Axle no.	Axle distance (m)	Load per axle (kN)
1	2.06	59.35
2	1.83	108.82
3	1.82	108.82
4	1.82	108.82
5	—	108.82

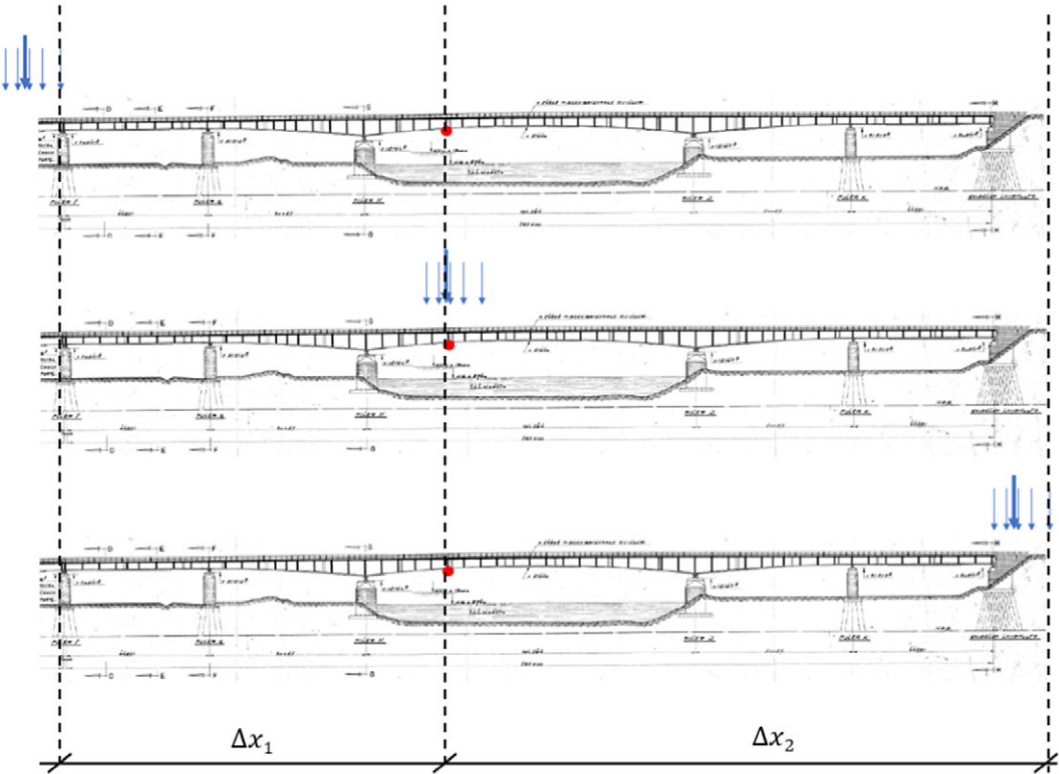


Figure B1. Load position at influence line start (top), peak (middle), and end (bottom).

sensors H1 and H10 was measured. The distance Δx between the two sensor positions was then divided by Δt to obtain the truck velocity for the left and right lanes equal to $v_l = 21.18$ km/h and $v_r = 21.66$ km/h respectively. The influence lines are obtained by applying a time window to the strain time series. The window start and end times correspond to the first truck axle entering the bridge and the last truck axle leaving the bridge respectively, as shown in [Figure B1](#). The time corresponding to the start and end position can be determined using the known distances Δx_1 and Δx_2 and the truck speed calculated previously. A -0.1 s shift was applied to the right lane measurements to minimize the discrepancies between the measured and predicted stress influence lines.