

SANDIA REPORT

SAND2018-1889

Unlimited Release

Printed February 24, 2018

Application of Bayesian Model Selection for Metal Yield Models Using ALEGRA and Dakota

Teresa N. Portone, John H. J. Niederhaus and Jason J. Sanchez, Laura P. Swiler

Prepared by

Sandia National Laboratories

Albuquerque, New Mexico 87185 and Livermore, California 94550

Sandia National Laboratories is a multimission laboratory managed and operated by National Technology and Engineering Solutions of Sandia, LLC., a wholly owned subsidiary of Honeywell International, Inc., for the U.S. Department of Energys National Nuclear Security Administration under contract DE-NA0003525.

Approved for public release; further dissemination unlimited.



Sandia National Laboratories

Issued by Sandia National Laboratories, operated for the United States Department of Energy by National Technology and Engineering Solutions of Sandia, LLC.

NOTICE: This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government, nor any agency thereof, nor any of their employees, nor any of their contractors, subcontractors, or their employees, make any warranty, express or implied, or assume any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represent that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise, does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government, any agency thereof, or any of their contractors or subcontractors. The views and opinions expressed herein do not necessarily state or reflect those of the United States Government, any agency thereof, or any of their contractors.

Printed in the United States of America. This report has been reproduced directly from the best available copy.

Available to DOE and DOE contractors from
U.S. Department of Energy
Office of Scientific and Technical Information
P.O. Box 62
Oak Ridge, TN 37831

Telephone: (865) 576-8401
Facsimile: (865) 576-5728
E-Mail: reports@adonis.osti.gov
Online ordering: <http://www.osti.gov/bridge>

Available to the public from
U.S. Department of Commerce
National Technical Information Service
5285 Port Royal Rd
Springfield, VA 22161

Telephone: (800) 553-6847
Facsimile: (703) 605-6900
E-Mail: orders@ntis.fedworld.gov
Online ordering: <http://www.ntis.gov/help/ordermethods.asp?loc=7-4-0#online>



Application of Bayesian Model Selection for Metal Yield Models Using ALEGRA and Dakota

Teresa N. Portone
The University of Texas at Austin
Institute for Computational Engineering and Sciences
portone@ices.utexas.edu

John H. J. Niederhaus¹, Jason J. Sanchez² and Laura P. Swiler³
Sandia National Laboratories
P.O. Box 5800, MS-1323
Albuquerque, NM 87185-1323

¹ Multiphysics Applications (1446)

² Computational Shock and Multiphysics (1443)

³ Optimization and Uncertainty Quantification (1441)

Abstract

This report introduces the concepts of Bayesian model selection, which provides a systematic means of calibrating and selecting an optimal model to represent a phenomenon. This has many potential applications, including for comparing constitutive models. The ideas described herein are applied to a model selection problem between different yield models for hardened steel under extreme loading conditions.

Acknowledgment

This work was jointly supported by the U.S. Army Research Laboratory and the Verification and Validation program within the Advanced Simulation and Computing (ASC) program at Sandia.

The bulk of the study was performed using the ALEGRA [13] and Dakota [1] libraries. Features not available in ALEGRA and Dakota, as well as postprocessing, were performed using SciPy [9] and NumPy [22], and figures were produced using Matplotlib [6].

Contents

1	Introduction	9
2	Bayesian Uncertainty Quantification and Calibration	11
3	Bayesian Model Selection	15
3.1	Approximating model evidence	16
3.1.1	Monte Carlo integration	17
3.1.2	Laplace approximation	17
3.1.3	KDE approximation	17
4	Model selection workflow	19
4.1	Determining parameter uncertainty representations	19
4.2	Surrogate construction	20
4.3	Sensitivity analysis	21
4.4	Bayesian calibration	21
4.5	Model selection	22
5	Flow stress model selection for high-velocity impact of steel	25
5.1	Johnson-Cook strength model	25
5.2	Zerilli-Armstrong strength model	26
5.3	Steinberg-Guinan-Lund strength model	26
5.4	Calibration data experimental configuration	27
5.4.1	Observational uncertainty model	27
5.5	Parameter uncertainty models	28

5.5.1	JC model parameters	28
5.5.2	ZA model parameters	29
5.5.3	Steinberg-Guinan-Lund model parameters	29
5.6	Numerical configuration	30
6	Results	33
6.1	Johnson-Cook Bayesian Calibration Results	33
6.2	Zerilli-Armstrong Bayesian Calibration Results	34
6.3	Steinberg-Guinan-Lund Bayesian Calibration Results	34
6.4	Model selection results	35
6.5	Conclusion	36
 Appendix		
A	Variance-based global sensitivity analysis	53
A.1	Sensitivity analysis for dimension reduction	54
A.2	Sensitivity analysis for experimental design	54
	 References	 54

List of Figures

2.1	This is an example of Bayes's rule where the model is the identity ($d = \theta + \epsilon$). We see that the posterior distribution falls in the region where the likelihood and the prior distributions overlap, more heavily weighted toward the likelihood distribution. This is because the likelihood distribution was more tightly peaked around its mean. Something more precise could be said about the data than about the prior, so it had a stronger effect on the posterior distribution.	13
4.1	A flowchart of the Bayesian model selection process for models M_1 , M_2 and M_3 .	23
5.1	A quadratic was fit to the data and used to estimate the data uncertainty....	28
5.2	The cell size $1e-4$ was selected because the numerical error was small compared to the measurement error, but the model still ran quickly enough to perform the hundreds of evaluations required to construct a surrogate model for the UQ framework.	31
5.3	Representative ALEGRA simulation results for 1250-m/s impact using the JC model for HzB steel.	32
6.1	Prior (uniform) and posterior probabilities for each model. Each color indicates a different method of approximating the model evidence, used to compute the posterior model probabilities. The different methods are described in Chapter 3.	35
6.2	Sobol indices for JC.	38
6.3	Prior vs. posterior distributions for JC.	39
6.4	Posterior MCMC chains for JC.	40
6.5	Joint marginal samples of the posterior distribution for JC model parameters.	41
6.6	Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.	42
6.7	Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of JC model parameters.	42

6.8	Sobol indices for ZA.	43
6.9	Prior vs. posterior distributions for ZA.	44
6.10	Joint marginal samples of the posterior distribution for ZA model parameters.	45
6.11	Posterior MCMC chains for ZA.	46
6.12	Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.	47
6.13	Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of ZA model parameters.	47
6.14	Sobol indices for SGL.	48
6.15	Prior vs. posterior distributions for SGL.	49
6.16	Joint marginal samples of the posterior distribution for SGL model parame- ters.	50
6.17	Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.	51
6.18	Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of SGL model parameters.	51
6.19	Posterior MCMC chains for SGL.	52

Chapter 1

Introduction

This report documents the use of ALEGRA [13] and Dakota [1] to perform a Bayesian model selection of the best model to reproduce the depth of penetration (d.o.p.) as a function of time for a hardened steel plate being impacted by a tungsten rod. It is common in realistic engineering problems that there are multiple constitutive models for a single phenomenon, none of which is highly reliable in all situations. Bayesian model selection provides a means of measuring which model is best to represent that phenomenon, given the uncertainties present in the modeling process, such as measurement error and parameter uncertainty.

The report will proceed as follows: in Chapter 2, the Bayesian uncertainty quantification framework is quickly summarized; in Chapter 3, Bayesian model selection and other model selection criteria are outlined; in Chapter 4, the workflow for performing the Bayesian model selection amongst the material models is described; in Chapter 5, the material models under consideration for the study described in this report are introduced; and finally, results of the model selection are reported in 6.

Chapter 2

Bayesian Uncertainty Quantification and Calibration

Uncertainties arise in the modeling process for many different reasons. Two examples are measurement uncertainty and parameter uncertainty. Measurement uncertainty is widely acknowledged and has historically been reported in the form of confidence intervals. Model parameter uncertainty arises because parameters, with the exception of fundamental physical constants such as the speed of light, cannot be determined exactly. If a parameter is a physical constant that can be measured directly, it will still have measurement uncertainty; if it is unobservable, it must be inferred indirectly from uncertain data. If it appears in a closure model, there may not be any “true” value the parameter should attain, because the model is only an approximate representation of reality. In that case, there is uncertainty about which value of the parameter is “best,” in some sense. Uncertainty quantification accounts for these uncertainties by modeling them with probability distributions and propagates those uncertainties to the solution of interest, either a solution of a computational model or a solution of an inverse problem.

Consider a model $\mathcal{M}$ with parameters $\boldsymbol{\theta}$. The model parameters are uncertain, but there may be some available prior information. The uncertain parameters can be modeled using probability densities. For instance, if a parameter is known to fall in a specific range of values, it can be modeled as a uniform random variable falling in that interval, or if it can be measured and there is a sample mean and standard deviation, it can be modeled as a normal distribution. These probability densities describing the prior information available about the model parameters are called prior probability densities. See [7] for considerations in choosing a prior and the implications of using maximum-entropy priors. If the parameters are statistically independent, the joint prior distribution over all the model parameters is defined as

$$p(\boldsymbol{\theta}) \equiv \prod_{i=1}^{N_{\theta}} p(\theta_i).$$

Now consider data corresponding to an observable quantity predicted by the model. The data is also uncertain due to experimental and measurement variations, which are generally

assumed to be independent and normally distributed with zero mean and variance σ^2 :

$$d_i = s_i + \epsilon_i, \quad \epsilon_i \sim \mathcal{N}(0, \sigma^2), \quad i = 1, \dots, N_d,$$

where s_i represents the “true” signal that would have been observed if there were no observational error. Assuming the model $\mathcal{M}$ is a correct description of the underlying physics, the signal is equivalent to

$$s_i \equiv \mathcal{O}[\mathcal{M}(\boldsymbol{\theta}^*)]_i,$$

where $\mathcal{O}$ represents the process of mapping the model output to the quantity that is being observed, such as the time-trace of penetration of a penetrator, and $\boldsymbol{\theta}^*$ represents the “true” values of the model parameters.

With this substitution, the likelihood that the observed data d_i arose from the model, given a specific set of parameters $\boldsymbol{\theta}$, is related to the discrepancy between the model prediction and the data, since

$$\begin{aligned} d_i &= \mathcal{O}[\mathcal{M}(\boldsymbol{\theta})]_i + \epsilon_i \\ &\Downarrow \\ d_i - \mathcal{O}[\mathcal{M}(\boldsymbol{\theta})]_i &= \epsilon_i \sim \mathcal{N}(0, \sigma^2). \end{aligned}$$

Substituting this discrepancy into the measurement error model defines what is called the likelihood density:

$$\mathcal{L}(\boldsymbol{\theta}) \equiv p(\mathbf{d}|\boldsymbol{\theta}) \equiv \frac{1}{(2\pi\sigma)^{N_d/2}} \exp\left(-\frac{1}{2\sigma^2} \|\mathbf{d} - \mathcal{O}[\mathcal{M}(\boldsymbol{\theta})]\|_2^2\right).$$

Finally, Bayes’s theorem provides us with the probability of the model parameters $\boldsymbol{\theta}$ given observational data $\mathbf{d}$ and prior distribution $p(\boldsymbol{\theta})$, called the posterior distribution:

$$p(\boldsymbol{\theta}|\mathbf{d}) = \frac{p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{\int p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta}}.$$

A simple example of an application of Bayes’s theorem is shown in Figure 2.1.

Assuming the model is correct, the maximum *a posteriori* (MAP) point, $\hat{\boldsymbol{\theta}}$, which is the point that maximizes the posterior distribution, should be the true model parameter values, $\boldsymbol{\theta}^*$. With increasing amounts of data, the posterior distribution becomes tightly peaked around $\boldsymbol{\theta}^*$, and a dirac delta in the limit of infinite data.

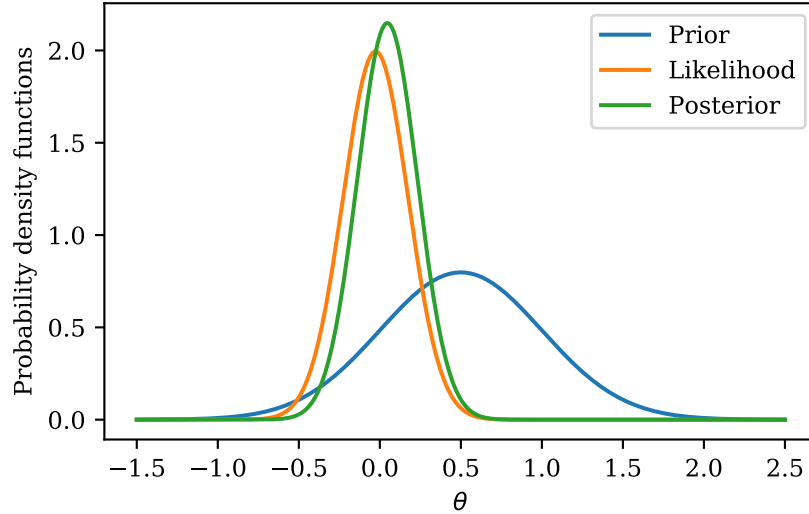


Figure 2.1. This is an example of Bayes’s rule where the model is the identity ($d = \theta + \epsilon$). We see that the posterior distribution falls in the region where the likelihood and the prior distributions overlap, more heavily weighted toward the likelihood distribution. This is because the likelihood distribution was more tightly peaked around its mean. Something more precise could be said about the data than about the prior, so it had a stronger effect on the posterior distribution.

An analytical expression for the posterior distribution is not typically available because of the nonlinear nature of models for complex physical processes. Because of this, it is common to generate sample draws from the posterior distribution using the Markov Chain Monte Carlo (MCMC) method and estimate its density using the samples. The details of this will not be described further here. The important thing to note is that each sample from the posterior distribution requires one or more model evaluations, which can be prohibitive if the relevant models are computationally expensive. Dakota will be used to generate these samples and for other analysis tools that will be needed for the model selection. For more comprehensive introductions regarding UQ and MCMC, see [17, 3].

Chapter 3

Bayesian Model Selection

In the previous discussion, it was assumed that the model $\mathcal{M}$ was a correct description of the underlying physics of the problem, and the only uncertainties arose from incomplete knowledge of the model parameters. A much more common situation, especially during the process of model development, is that there is more than one model that might explain the physical phenomenon.

This situation arises often in material modeling. For instance, as will be described in Chapter 5, many models and model variations have been proposed to represent the von Mises flow stress, such as the Johnson-Cook (JC), Zerilli-Armstrong (ZA), and Steinberg-Guinan-Lund (SGL) constitutive models. Different models perform better in different situations and in predicting certain quantities. Certainly none of the models is a perfect representation of reality. This report focused on determining which of the three models could best reproduce depth-of-penetration data as a function of time.

We will denote such a collection of models $\{\mathcal{M}_i\}_{i=1}^{N_m}$, where N_m is the number of models under consideration. Each model has its own set of parameters, which we will denote $\{\theta_i\}_{i=1}^{N_m}$, and each model can have a different number of parameters, as is the case for JC, ZA, and SGL. These parameters might be, for instance, the coefficient and exponent in the work-hardening term, or initial yield strength. We will denote the number of model parameters for model $\mathcal{M}_i$ as N_{θ}^i . In this situation, one can imagine there is a “true” model that produced some observation data,

$$d_i = \mathcal{O}[m(\theta)]_i + \epsilon_i, \quad i = 1, \dots, N_d,$$

and one can ask “what is the probability that $m = \mathcal{M}_i$, given the observed data $\mathbf{d}$?” In the case of this report, the data $\mathbf{d}$ is time-series observations of the nose of a long tungsten-alloy rods impacting a hardened steel plate at high velocity, described in Chapter 5 and reported in [2]. The observations of model output $\mathcal{O}[\mathcal{M}_i(\theta_i)]$ are the location of a Lagrangian material tracer at the interface of the tungsten-alloy rod and the steel plate, taken at different times during an ALEGRA simulation, using one of the constitutive models and one set of values for its parameters.

This collection of models can be thought of as a discrete probability space, where each model has some probability of being the true model, and their probabilities sum to one.

Bayes's rule with discrete probabilities for the models can be written

$$P(m = \mathcal{M}_i | \mathbf{d}) = \frac{p(\mathbf{d} | m = \mathcal{M}_i) P(m = \mathcal{M}_i)}{\sum_{j=1}^{N_m} p(\mathbf{d} | m = \mathcal{M}_j) P(m = \mathcal{M}_j)}, \quad (3.1)$$

where p represents probability densities of continuous random variables and P represents probabilities of discrete random variables. For notational simplicity, the probabilities will be denoted in terms of $\mathcal{M}_i$ instead of $m = \mathcal{M}_i$, for the remainder of the report.

The likelihood of the model $\mathcal{M}_i$ is determined by averaging the likelihood of its model parameters over all possible values of the model parameters, according to their prior distributions:

$$p(\mathbf{d} | \mathcal{M}_i) = \int p(\mathbf{d} | \boldsymbol{\theta}_i, \mathcal{M}_i) p(\boldsymbol{\theta}_i) d\boldsymbol{\theta}_i. \quad (3.2)$$

This value is often called the evidence, or the marginal likelihood of the model. The likelihood in the integral, $p(\mathbf{d} | \boldsymbol{\theta}_i, \mathcal{M}_i)$, is the likelihood the data $\mathbf{d}$ would be observed, assuming $\mathcal{M}_i$ was the true model m and assuming a specific set of values $\boldsymbol{\theta}_i$. Mathematically, it is

$$p(\mathbf{d} | \boldsymbol{\theta}_i, \mathcal{M}_i) = (2\pi\sigma^2)^{-N_{\theta}^i/2} \exp \left(-\frac{1}{2\sigma^2} \|\mathbf{d} - \mathcal{O}[\mathcal{M}_i(\boldsymbol{\theta}_i)]\|_2^2 \right).$$

Typically all models are assigned equal prior probabilities; that is, $P(\mathcal{M}_i) = N_m^{-1}$. Assuming equal prior probabilities for each model, the posterior probability for each model is simplified to

$$P(\mathcal{M}_i | \mathbf{d}) = \frac{p(\mathbf{d} | \mathcal{M}_i)}{\sum_{j=1}^{N_m} p(\mathbf{d} | \mathcal{M}_j)}. \quad (3.3)$$

We see that the key component to Bayesian model selection will lie in computing, or accurately approximating, the model evidence. Three methods are detailed herein, and all were computed for this study. Much of this discussion is a summary of the ideas described in detail in [23].

3.1 Approximating model evidence

There are several methods to approximate the model evidence. The methods that were used in this study are described below. Multiple approximation methods were used because some are relatively cheap compute but are less reliable, while others are more expensive but more reliable.

3.1.1 Monte Carlo integration

The most brute-force, generally applicable method uses Monte Carlo integration:

$$p(\mathbf{d}|\mathcal{M}_i) = \int p(\mathbf{d}|\boldsymbol{\theta}_i, \mathcal{M}_i)p(\boldsymbol{\theta}_i)d\boldsymbol{\theta} \approx \frac{1}{N_s} \sum_{j=1}^{N_s} p(\mathbf{d}|\boldsymbol{\theta}_i^{(j)}, \mathcal{M}_i), \quad \boldsymbol{\theta}^{(j)} \sim p(\boldsymbol{\theta}).$$

The process amounts to randomly sampling from the prior distribution of the parameters for model $\mathcal{M}_i$, evaluating the likelihood, and averaging the results. The error converges as $N_s^{-1/2}$, which is not ideal, but the cost doesn't balloon as the parameter dimension increases. This approximation can take many more samples than one might expect, because if only a small region of the parameters' prior probability space maps to high likelihood regions, many of the samples will evaluate to approximately zero. See, e.g. [24], for background on Monte Carlo integration.

3.1.2 Laplace approximation

A popular alternate approach employs a Laplace approximation of the integral, which approximates the posterior distribution with a Gaussian centered at the maximum *a posteriori* (MAP) point, $\hat{\boldsymbol{\theta}}$, where the inverse covariance of the Gaussian is approximated with the Hessian of the negative log posterior distribution, denoted by H [11]. Then the Laplace approximation is

$$\int p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta} \approx p(\mathbf{d}|\hat{\boldsymbol{\theta}})p(\hat{\boldsymbol{\theta}})(2\pi)^{N_{\theta}/2} \left| \det(H(\hat{\boldsymbol{\theta}})) \right|^{-1/2}.$$

This approximation is relatively accurate as long as the posterior of the parameters is well approximated by a Gaussian. If the posterior distribution is multimodal or significantly non-Gaussian, the approximation may be suspect. The method requires no sampling. The bad news here is that it requires a Hessian of the log posterior, or else an approximation of the Hessian. This requires first and second derivatives of the model with respect to its parameters, which are often quite difficult to obtain. Even so, a finite-difference approximation might not be as costly as running an expensive Monte Carlo simulation. For the study, derivatives were taken with respect to a surrogate model of the original model, for which analytical derivatives were available.

3.1.3 KDE approximation

The approach advocated by Wasserman in [23] is to compute a Kernel Density Estimate (KDE) of the posterior distribution of a model's parameters, given samples of the posterior

distribution (details on this subject are available in [16]). This can be used to approximate the evidence by noting that

$$p(d|\boldsymbol{\theta}) = \frac{p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{\int p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta}} \implies \int p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta} = \frac{p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{p(d|\boldsymbol{\theta})}.$$

The expression on the right-hand side is true for any $\boldsymbol{\theta}$, including $\hat{\boldsymbol{\theta}}$. Denoting the KDE approximation of the posterior as $\hat{p}(\boldsymbol{\theta}|\mathbf{d})$, the evidence can thus be approximated as

$$\int p(\mathbf{d}|\boldsymbol{\theta})p(\boldsymbol{\theta})d\boldsymbol{\theta} = \frac{p(\mathbf{d}|\hat{\boldsymbol{\theta}})p(\hat{\boldsymbol{\theta}})}{p(d|\hat{\boldsymbol{\theta}})}.$$

This approach works generally because it does not assume the posterior is Gaussian or uni-modal, although computing a KDE can become prohibitive if the number of model parameters becomes too large. Since this feature is only available through Dakota for independent variables, the KDE functionality in SciPy [9] was used, which implements the algorithm in [16].

Chapter 4

Model selection workflow

The workflow for model selection will fall into the following broad categories: determining prior distributions for model parameters (either using nominal values or deterministic calibration), surrogate construction, sensitivity analysis, and model evidence approximation. A flowchart of this process, once the subset of model parameters to be calibrated for each model has been chosen, is provided in Figure 4.1. In the following sections, the purpose and methodology behind each category will be discussed.

4.1 Determining parameter uncertainty representations

Model parameter uncertainty representations are used as prior densities for the Bayesian calibration of each model. The determination of these prior densities is an important part of the model selection process and can affect the solution of the model selection problem. This is because the model evidence is essentially an average over likelihood values, weighted by the prior density. Care should be taken to begin with as accurate a description of the parameter uncertainty as possible. For instance, if reasonable ranges a parameter should take are available, a uniform distribution over that range could be used. If a parameter is known to be positive, a log-Normal distribution can be used. If no prior information can be brought to bear, be sure to begin with a broad prior density, to signify great uncertainty in the parameter values.

For this study, all uncertain parameters were modeled as multivariate Gaussians, with means either set to nominal values or to deterministically calibrated values, and with standard deviations assumed to be 5%. This was chosen for simplicity and had no basis in prior information about the parameters' uncertainty. In general one might wish to use a larger standard deviation, such as 10%, when it is not clear what values the model parameters should take.

If experimental uncertainty in Hopkinson bar, tensile, etc. test data is accounted for when determining the model parameter values in the first place using Bayesian calibration, then the experimental uncertainty will be propagated to the model parameters. The posterior distributions of the model parameters after Bayesian calibration against these canonical tests could then be used as the prior distributions for a model selection study.

4.2 Surrogate construction

Bayesian calibration, sensitivity analysis, and MC integration require hundreds of thousands, even millions of model evaluations. For computationally-intensive models such as those used for this study, using the model for these purposes is infeasible. The solution was to use surrogate models, constructed in Dakota using strategically-chosen samples of the yield model parameters. In the context of this work, surrogate models are interpolative fits to response data, given variations in parameter values. They are popular for performing UQ studies on complex systems where the underlying computational model is expensive to run, because they provide a means of producing the thousands of samples that are required for such studies at a fraction of the computational cost. The sampling technique employed for the study is Latin Hypercube Sampling (LHS), which requires fewer samples to approximate statistics than purely random sampling.

LHS is a semi-random exploration of parameter space. To produce N samples, each dimension of parameter space is partitioned into N equal-probability regions. One sample is drawn for each dimension according to its distribution in that partition, and the samples from each dimension are randomly paired with each other. If the samples are projected onto one dimension, there will always be only one sample per partition. For further details on how Latin Hypercube samples are produced, there is a very nice explanation in Section 2 of [20].

Latin Hypercube sampling can be advantageous for the purposes of surrogate construction it requires fewer samples than standard MC to obtain accurate statistics. The key point that the number of samples does not increase with the number of parameters. It is important to note that because of the equal-probability partitioning of the parameter space in each dimension, it is nontrivial to produce more samples. In order to maintain this property, the existing partitions must again be split equally. In the case of using Dakota, this means that the number of sample points would need to be doubled in order to keep the partitions equal. A heuristic for selecting the number of LHS samples is to have approximately 20 samples per model parameter. For the study, slightly more than 20 samples per parameter were used for each model.

The two types of surrogates used for this study were Gaussian Process (GP) models and Polynomial Chaos Expansion (PCE) models. In broad strokes, both are constructed by fitting to samples of model output. However, the way they are constructed differs. For more information on the different types of surrogates and how they are constructed, see [12] and [15] for GP models and [25] and [4] for PCEs.

Unfortunately, analytical results regarding approximation error for surrogates are largely unavailable. This has to do with the fact that the performance of the surrogate depends on the scenario in which it is being used. For instance, GPs and PCE expansions do not perform well when there is a discontinuity in the model output space. However, it is difficult to know *a priori* when such challenges will arise. Because of this, it is important to vet the surrogate before using it as a substitute for the original model. Given the minimal effort to

generate different types of surrogates, once model evaluations have been computed, it is a good idea to build two or three different types and compare the result for consistency between the methods. See [21] for discussion of ways to improve and ensure surrogate accuracy in practice.

4.3 Sensitivity analysis

The number of model parameters were relatively small for this study, none of them having more than 10 parameters. However, for more complicated models, it would be useful to perform a global sensitivity analysis to determine which model parameters significantly affect the model output and which do not. The method of sensitivity analysis used in this paper is a variance-based sensitivity analysis, in which Sobol indices are computed. Sobol indices express, roughly, the fraction of variation in model outputs that can be attributed to variations in each model parameter.

An example of the result of such a sensitivity analysis, displaying the Sobol indices for the parameters of the Zerilli-Armstrong yield model, as well as their effect on depth of penetration (d.o.p.) as a function of time, are shown in Figure 6.8. The result of this sensitivity analysis shows that for all three models, the parameters that are related to initial yield strength significantly affect the depth of penetration, which agrees with intuition. For more complicated problems where this intuition is not available, variance-based sensitivity analysis can help shed light on which parameters to calibrate and treat as uncertain for the model selection process.

The computation of Sobol indices requires thousands of sample response evaluations (in this case, the depth of penetration at each observation time), so PCE and GP surrogates were used for the sensitivity analysis. The samples used for the sensitivity analysis cannot be reused for the model selection study, if some of the parameters are deemed insignificant and are not going to be calibrated. With this in mind, the LHS samples generated for the sensitivity analysis used a coarser mesh for the ALEGRA solves than was used for those used for the Bayesian calibration in order to minimize computational cost. The cell size was chosen so that ALEGRA runs were fast while maintaining the same qualitative behavior of the d.o.p. compared to the converged mesh. See Appendix A for more details regarding the definition and computation of Sobol indices.

4.4 Bayesian calibration

For the model selection study, LHS samples were generated on the target mesh size and were used to construct a GP surrogate, as mentioned in Section 4.2. The surrogate was constructed and used for the Bayesian calibration using Dakota. The calibration was performed using Markov Chain Monte Carlo (MCMC), as mentioned in Chapter 2, using Dakota. A pre-run

optimization for the MAP point was performed to start the chain in a high-probability region of the posterior density. This MAP point was extracted and used to compute the Laplace approximation of the model evidence, discussed in 3.1.2. MCMC chains of length $1e6$ were produced for the study.

The result of the MCMC-based Bayesian calibration is a set of samples of the model parameters, distributed according to the posterior distribution. These samples were used to generate a KDE approximation of the posterior distribution, used for the KDE-based approximation of the model evidence, described in 3.1.3. Note that the KDE approximation is only as good as the quality of the samples used in the approximation. A good discussion of diagnostic checks and analyzing chain convergence is provided in [5].

4.5 Model selection

As discussed in Chapter 3, the most important part of performing a Bayesian model selection is approximating the model evidence as defined in (3.2) and appearing in the model posterior probability in (3.1). Two of the three methods used in this study for approximating the model evidence are generated using the MCMC chain produced during Bayesian calibration. The third approach, using MC integration, was performed by sampling from the specified prior distributions and evaluating a GP surrogate model of the DOP. These evaluations were then used to evaluate the likelihood density for each sample. In total, $1e6$ samples were used for each MC integration. A running approximation of the integral was reported every $1e5$ samples to check for statistical convergence.

For each method of approximating the evidence, a posterior probability for each model was computed according to (3.3), which assumed uniform prior model probabilities. By construction, the posterior probabilities sum to 1 for each of the model evidence approximations. In this case, the posterior distributions for the model parameters were unimodal and approximately Gaussian, so the different methods of approximation assigned highest probability to the same model (see Figure 6.1). However, if there is significant discrepancy between the approximated evidences for a model, it is possible that the posterior distributions are not as “nice.” In that case, one should check for multimodality or a significantly non-Gaussian shaped posterior distribution.

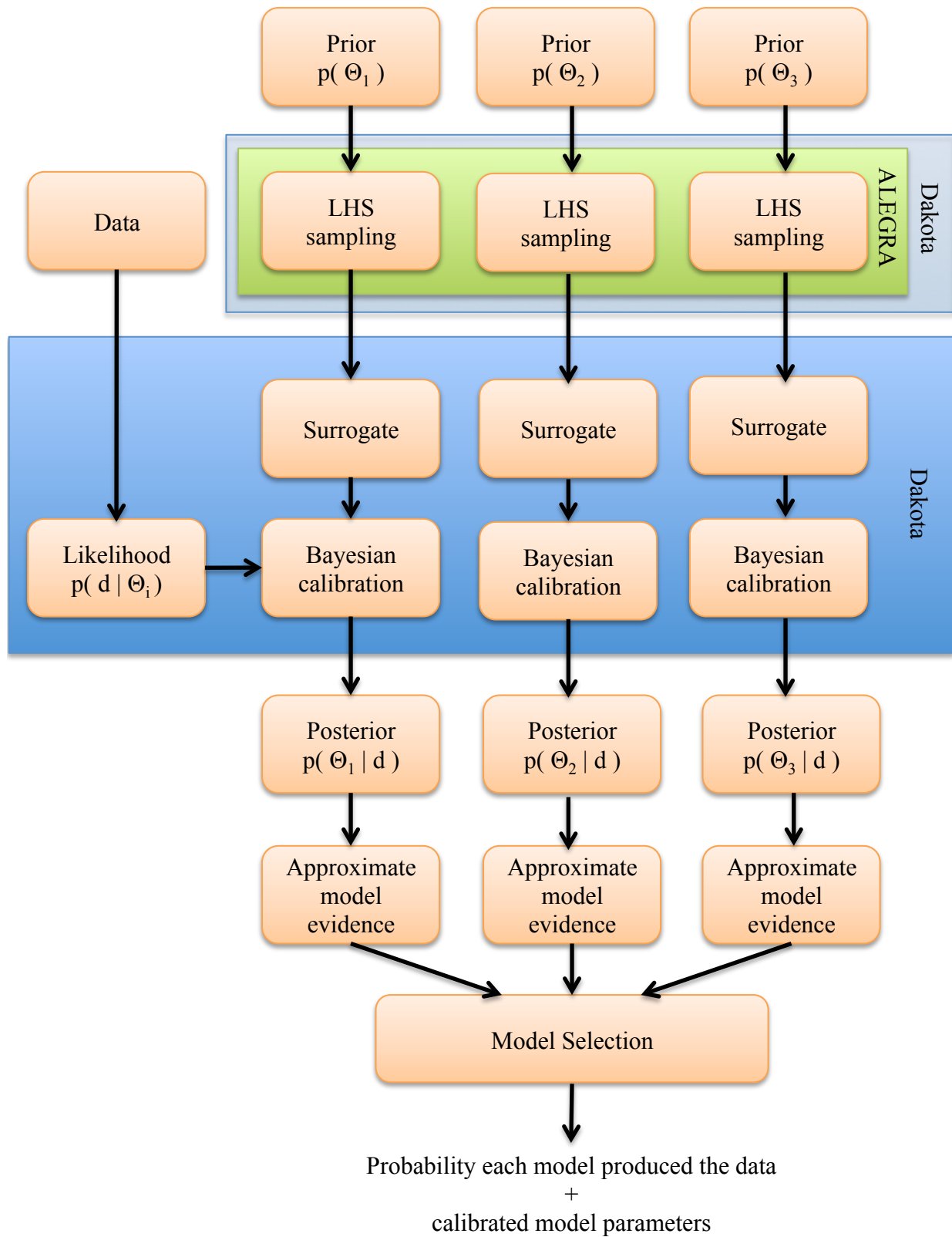


Figure 4.1. A flowchart of the Bayesian model selection process for models M_1 , M_2 and M_3 .

Chapter 5

Flow stress model selection for high-velocity impact of steel

For this study we applied these ideas to a simple model selection problem between material models, evaluating their ability to reproduce depth of penetration as a function of time. We will be considering three popular strength models used for metals subject to large strains, high strain rates, and high temperatures: the Johnson-Cook, Zerilli-Armstrong, and Steinberg-Guinan-Lund constitutive models [8, 26, 19]. Each model is briefly described below.

5.1 Johnson-Cook strength model

The Johnson-Cook strength model [8] is an empirical model, which defines the von Mises flow stress as

$$\sigma = [A + B\epsilon^n][1 + C \ln \dot{\epsilon}^*][1 - T^{*m}], \quad (5.1)$$

where the meaning of the state variables are defined in Table 5.1. The first term represents the stress as a function of strain when $\dot{\epsilon}^* = 1$ and $T^* = 0$. The second and third terms represent the effects of strain rate and temperature, respectively.

ϵ	equivalent plastic strain
$\dot{\epsilon}^* = \dot{\epsilon}/\dot{\epsilon}_0$	dimensionless plastic strain rate, $\dot{\epsilon}_0 = 1.0s^{-1}$
$T^* = \frac{T - T_{\text{room}}}{T_m - T_{\text{room}}}$	homologous temperature, T_m melting point

Table 5.1: The terms in the Johnson-Cook strength model.

5.2 Zerilli-Armstrong strength model

The Zerilli-Armstrong strength model is a constitutive model based on theoretical results related to a physical model of the crystal structure of the material (body-centered or face-centered crystals) [26]. It models the flow stress as

$$\sigma = \Delta\sigma'_G + kl^{-\frac{1}{2}} + (c_1 + c_2\epsilon^{\frac{1}{2}}) \exp(-c_3T + c_4T \ln \dot{\epsilon}) + c_5\epsilon^N,$$

where again ϵ is the equivalent plastic strain. The first order terms $\Delta\sigma'_G$ and $kl^{-\frac{1}{2}}$ are often lumped together. The combined parameter will be denoted c_0 . The parameter c_2 is related to the face-centered case, while c_1, c_5 are related to the body-centered case. Thus, if $c_2 \neq 0$, then $c_1, c_5 = 0$ and vice-versa. Substituting in c_0 for the first-order terms will give us our working definition of the ZA model:

$$\sigma = c_0 + (c_1 + c_2\epsilon^{\frac{1}{2}}) \exp(-c_3T + c_4T \ln \dot{\epsilon}) + c_5\epsilon^N. \quad (5.2)$$

5.3 Steinberg-Guinan-Lund strength model

The Steinberg-Guinan-Lund strength model is a semi-empirical constitutive model that accounts for the effects of pressure and temperature on the yield strength Y and the shear modulus G [19]. Note that $Y = \sigma$ from previous discussion; the notation is different for historical reasons. The version of the model that was considered neglects dependence on strain rate based on experimental observations that rate-dependent effects are negligible at strain rates above $\dot{\epsilon} \geq 10^5 s^{-1}$. The (rate-independent) SGL model in [19] is defined as

$$\begin{aligned} Y &= Y_0 f(\epsilon) \left[1 + \left(\frac{Y'_p}{Y_0} \right) \frac{P}{\eta^{1/3}} - \left(\frac{G'_T}{G_0} \right) (T - T_{room}) \right], \\ Y_{max} &\geq Y_0 f(\epsilon), \\ f(\epsilon) &= [1 + \beta\epsilon]^n, \\ G &= G_0 \left[1 + \left(\frac{G'_p}{G_0} \right) \frac{P}{\eta^{1/3}} - \left(\frac{G'_T}{G_0} \right) (T - T_{room}) \right], \end{aligned}$$

where ϵ is the equivalent plastic strain, η is compression (initial specific volume v_0 divided by specific volume v), β and n are work-hardening parameters, Y_{max} is the largest value for Y in the literature (this limit will not be reached in most cases). The variables subscripted with 0 are their values at the reference state ($T = 300$ K, $P = 0$, $\epsilon = 0$). The primed parameters are derivatives with respect to their subscripts, evaluated at the reference state. The work-hardening function $f(\epsilon)$ is entirely empirical, but SGL includes a pressure dependence that is not accounted for in the ZA or JC models.

The ALEGRA implementation of SGL assumes that

$$\frac{Y'_p}{Y_0} \approx \frac{G'_p}{G_0},$$

as was mentioned in [19]. Then the final form of SGL used for the study is

$$\begin{aligned} Y &= Y_0 f(\epsilon) \frac{G}{G_0}, \\ Y_{max} &\geq Y_0 f(\epsilon), \\ f(\epsilon) &= [1 + \beta \epsilon]^n, \\ G &= G_0 \left[1 + \frac{AP}{\eta^{1/3}} - B(T - T_{room}) \right]. \end{aligned} \tag{5.3}$$

5.4 Calibration data experimental configuration

The experimental data that will be used to calibrate the models is presented in [2]. For each observation, a long tungsten-alloy rod penetrated and stagnated in an effectively semi-infinite plate of hardened steel. The target material was HzB steel, which is a hardened steel found in some ballistic protection applications. A time trace of the penetration was collected by repeatedly shooting rods at the same initial velocity and observing the depth of penetration at successively later times. The time history was collected for two different penetration velocities. See Figure 5.1 for a time-trace at one of the impact velocities. The study only used the lower-velocity data in [2], with an impact velocity of 1250 m/s.

5.4.1 Observational uncertainty model

The observational standard deviation for the data is reported as .01 cm in [2], but there is a distinct, nonphysical “kink” in the observational data, which calls into question whether there were some uncertainties not accounted for in that estimate. The issue with misreported measurement error is that the relative effect of the likelihood density on the posterior distribution increases with decreasing measurement error. This can cause an overfitting to the data that could lead to nonphysical results.

To attempt to address this missing uncertainty, a polynomial regressive fit was performed, and a standard deviation of approximately .06 cm was estimated using the data’s deviation from that experimental fit. See the difference in the uncertainties for these two approaches in Figure 5.1. The study used the estimated standard deviation from the polynomial fit. The y -intercept was not constrained to the origin, so the standard deviation may even still be smaller than it should be.

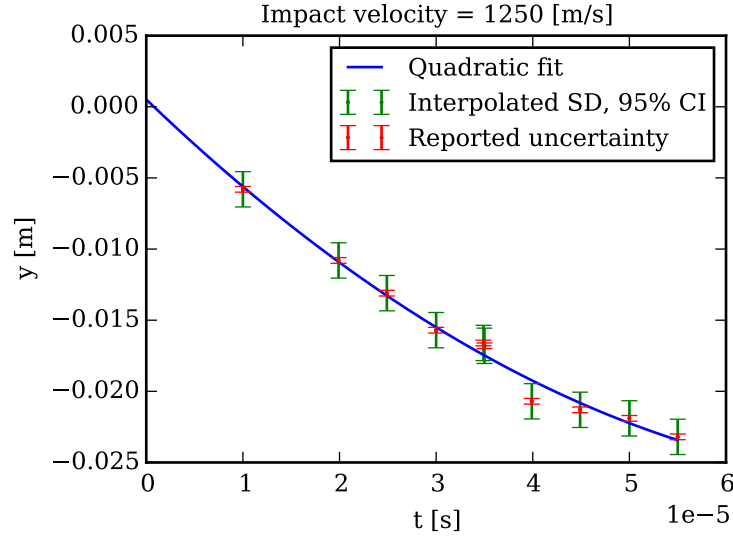


Figure 5.1. A quadratic was fit to the data and used to estimate the data uncertainty.

5.5 Parameter uncertainty models

As mentioned in Section 4.1, the parameters for all models will be assumed distributed according to normal distributions with means at the values listed in Tables 5.5.1, 5.5.2, and 5.5.3, and standard deviations of 5%. The mean values were based on information in ALEGRA material model libraries for HzB or comparable steels. In the case when nominal values were not available for the model parameters, the prior means were determined using a global optimization scheme, starting with parameterizations for the comparable steels. The global optimization was run using the ALEGRA/Dakota embedded interface.

The choice to use normal distributions to represent parameter uncertainty and the choice to optimize for initial parameter values was for the sake of simplicity and because there was little prior information regarding the uncertainty in the model parameters. It is always preferable to use any reasonable arguments that can be made regarding bounds on the model parameter values or prior information about parameter uncertainties and values gleaned from previous fits to experimental data, such as Hopkinson bar, tensile, etc. tests. Doing so helps to improve the fidelity of the results.

5.5.1 JC model parameters

The authors of [2] compared their experiments to a numerical simulation using the JC model and parameters for high-hard steel, which is similar to the HzB steel they used in their

experiments. Those parameters and those for the tungsten alloy rod are presented in Table 5.5.1 and were used as initial values in the calibration study performed for the JC model.

	A [GPa]	B [GPa]	n	C	m
Tungsten alloy	1.51	0.177	0.12	0.016	1.00
High-hard steel	1.50	0.569	0.22	0.003	1.17

Table 5.5.1: JC model parameters used in [2]. Flow stress $\sigma = [A + B\epsilon^n][1 + C \ln \epsilon^*][1 - T^{*m}]$.

5.5.2 ZA model parameters

The values to use for the ZA were not as obvious, because predetermined constants were not available. The starting constants were decided by choosing the most-similar material that had a ZA parametrization. After performing a sensitivity analysis to determine which parameters were significant, those parameters, c_1 , c_3 , and c_5 , were calibrated deterministically using the ALEGRA/Dakota interface using the ZA model directly (versus using a surrogate model). The best set of parameters out of 100 global optimization steps were used as the means for the prior distributions for the remainder of the study, and the same 5% standard deviation was used for this case. The ZA parameter values for HzB steel that were used for the model selection study are listed in Table 5.5.2. All the parameters were considered uncertain and were be calibrated, except for c_2 , which remained fixed at 0.

c_0 [Pa]	c_1 [Pa]	c_3 [K ⁻¹]	c_4 [K ⁻¹]	c_5 [Pa]	N
5.0e7	1.524e9	1.649e-3	4.5e-5	1.690e9	6.2e-1

Table 5.5.2: ZA model parameters for HzB steel used for the model selection study. Flow stress $\sigma = c_0 + c_1 \exp(-c_3 T + c_4 T \ln \epsilon) + c_5 \epsilon^N$.

5.5.3 Steinberg-Guinan-Lund model parameters

The SGL model required several parameters to be set. Similar to ZA, a parametrization for HzB steel was unavailable, so starting parameter values were determined from a similar material. After a sensitivity analysis, it was determined that the only parameter that significantly affected the depth-of-penetration behavior was the initial yield stress, Y_0 . Since a value for this quantity was provided in the experimental paper, that value was used as the prior mean distribution. The rest of the parameters were left at the values for the similar material.

Y_0 [Pa]	β	n	Y_{max}	G_0 [Pa]	A [Pa $^{-1}$]	B [K $^{-1}$]
1.45e+09	2.0	5.0e-1	2.5e+09	7.18e+10	2.06e-11	3.1501e-04

Table 5.5.3: The initial prior mean values used in the model selection study for the SGL model parameters, as defined in Equation 5.3.

5.6 Numerical configuration

The experimental configuration described in [2] was replicated and run in ALEGRA using cylindrical coordinates and assuming axisymmetry. Dimensions included a rod length of 5 cm, rod radius of 0.2 cm, target plate thickness of 2.9 cm, and target plate radius of 4 cm. The depth of penetration as a function of time was computed by placing a Lagrangian tracer at $r = 0$ at the interface between the rod and the steel target. The mesh was refined until the difference in the time history of the tracer was smaller than the data’s measurement error (see Figure 5.2) and the behavior of the tracer had more or less converged. The cell size was 1e-4, equating to 40 elements spanning the radius of the tungsten-alloy rod. Representative results from the ALEGRA simulations are included here in Figure 5.3, showing snapshots of the density field during penetration for the JC case with parameters from Table 5.5.1.

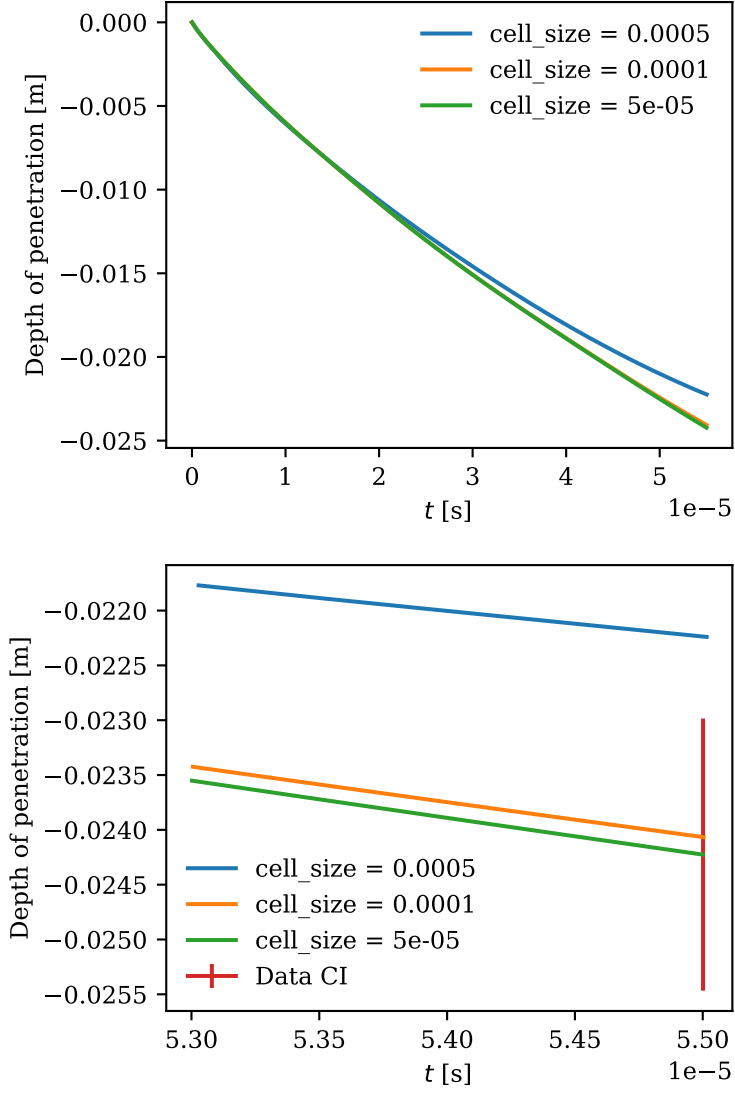


Figure 5.2. The cell size $1e-4$ was selected because the numerical error was small compared to the measurement error, but the model still ran quickly enough to perform the hundreds of evaluations required to construct a surrogate model for the UQ framework.

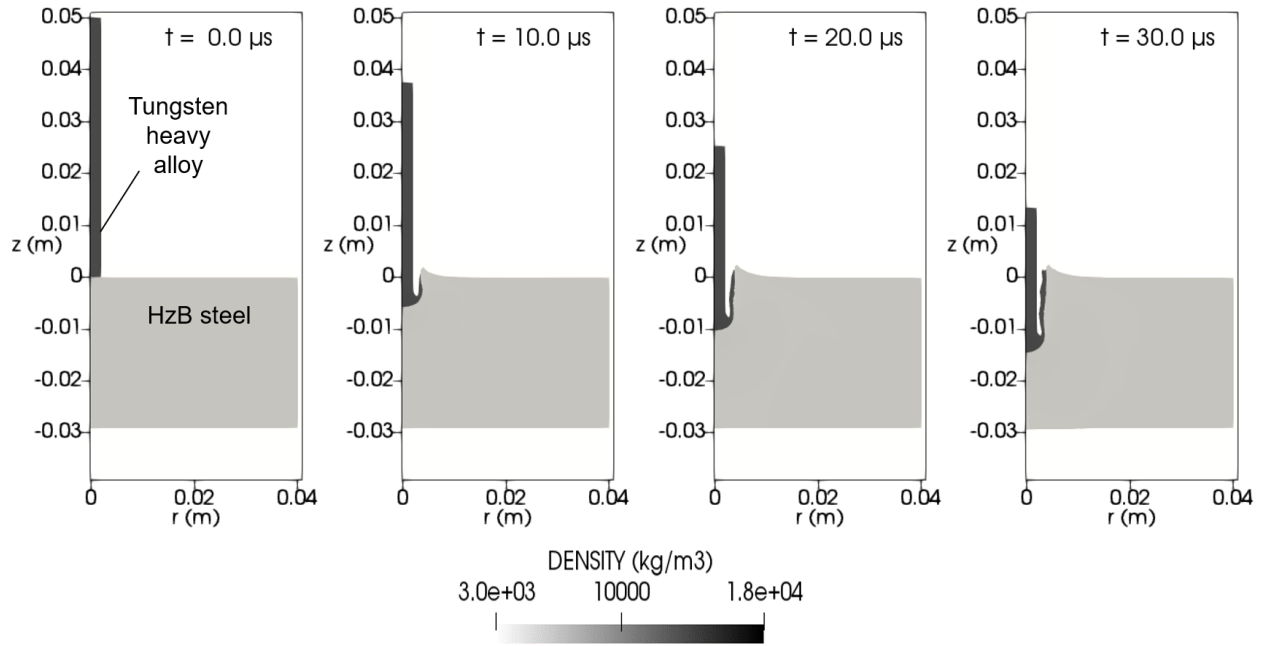


Figure 5.3. Representative ALEGRA simulation results for 1250-m/s impact using the JC model for HzB steel.

Chapter 6

Results

To reiterate, the Bayesian model selection was performed to select between the Johnson-Cook (JC), Zerilli-Armstrong (ZA), and Steinberg-Guinan-Lund (SGL) yield models. The model selection was performed both with all of the model parameters discussed in Chapter 5, as well as with only those determined to be significant after a sensitivity analysis. The results of the Bayesian calibration will be reported for each model individually, then the posterior probabilities will be reported for all three models together. The results for JC will be described in detail, to explain the significance of each figure. The results for ZA and SGL present the same figures and so will be discussed in less detail.

6.1 Johnson-Cook Bayesian Calibration Results

First, 256 LHS samples were generated for the Johnson-Cook model using ALEGRA’s embedded Dakota interface. These samples were used to generate a Gaussian Process (GP) that was then sampled extensively to compute Sobol Indices, which are shown in Figure 6.2. A brief explanation of Sobol indices, how they are computed, and the distinction between “main” and “total effects” is available in Appendix A. We see that JC’s model parameter A , which represents the yield stress, has the most significant impact on the depth of penetration (d.o.p.) at each time observation. Generally this sort of result could be used to motivate limiting calibration to only that parameter. As a proof of concept, all the parameters were calibrated for this study, to show that the parameters that were informed by data corresponded to those deemed significant by the sensitivity analysis.

This is demonstrated in Figure 6.3, which shows the prior and posterior distributions for each of the JC model parameters. The only one that was significantly changed by the calibration is the initial yield strength, A . As was discussed in Chapter 2, the width of a distribution indicates the degree of uncertainty in that quantity’s value. The fact that the other distributions did not change at all is an indication that the uncertainty in the other parameters was not affected at all by comparing against the data. In effect, nothing was learned about the other model parameter values.

The MCMC chains and marginal joint posterior samples for the JC model parameters are presented in Figures 6.4 and 6.5. The joint posterior samples in Figure 6.5 show a slight

dependence between JC's model parameters A and B . This is to be expected, since they appear in the same strain-dependent term of the JC model.

The posterior mean and 95% confidence interval of the depth of penetration as a function of time, plotted against the data with its estimated uncertainty, is shown in Figure 6.6. Finally, the mean posterior depth of penetration is compared to the predicted depth of penetration before any calibration (deterministic or Bayesian), along with the observational data, in Figure 6.7.

6.2 Zerilli-Armstrong Bayesian Calibration Results

The result of the sensitivity analysis for ZA show, in Figure 6.8, that c_1, c_3 and c_5 are significant to the behavior of the depth of penetration as a function of time. c_1 appears in a term related to initial yield strength, so its significance agrees with what was found for the other models. In contrast to the other two models, c_3 and c_5 , which appear on the ZA thermal softening and work-hardening terms respectively, are deemed significant as well. It is interesting that the parameters associated with work hardening were not deemed similarly significant in JC and SGL. One potential explanation for this is that only in ZA does the work-hardening term appear independently, without being multiplied by other dependencies, such as temperature.

Calibration results show C_1 and C_5 to be informed by the data, while C_3 was modified only slightly (see Figure 6.9). The mean of the work-hardening exponent N shifted slightly, but the support of the prior and posterior distributions is essentially the same width, indicating that it was not informed by the data (which is equivalent to saying it was not significant to the d.o.p.). The marginal joint posterior plots of the MCMC samples show that the samples of C_1 and C_5 are correlated, indicating a dependence between C_1 and C_5 (see Figure 6.10).

6.3 Steinberg-Guinan-Lund Bayesian Calibration Results

Strikingly, for the model with the most parameters, only one parameter was significant to the behavior of the d.o.p. Once again, the sensitivity analysis results, in Figure 6.14, show that the significant model parameter is the initial yield strength, Y_0 . The calibration results in Figure 6.15 clearly show that only Y_0 was informed by the data.

6.4 Model selection results

The results of this model selection study are presented in Figure 6.1. All the models were assumed to be equally likely to start. All methods of approximating the model evidence produced the same final result in terms of which model was more likely to have produced the observational data. For this particular study, the ZA model was determined to be most probable, followed by SGL and finally JC. Bayesian model selection tends to select for the most constrained model that can best explain the data. One can argue that ZA has the most theoretical basis, followed by SGL, while JC is entirely empirical. This result is thus in agreement with the idea that Bayesian model selection will tend to select the most “physical” model.

On the other hand, in Figures 6.6, 6.12, and 6.17, we see that the ZA model was the closest to reproducing the ‘kink’ in the data, starting at $4\text{e-}5$ seconds. This could be the reason that the ZA model was considered most probable. It would be interesting to repeat this study, omitting the $4\text{e-}5$ seconds data point, to determine how much the outlier affected the final result. The process of Bayesian calibration minimizes misfit with all available data, so it is important to be conscientious the quality of data used for such calibrations.

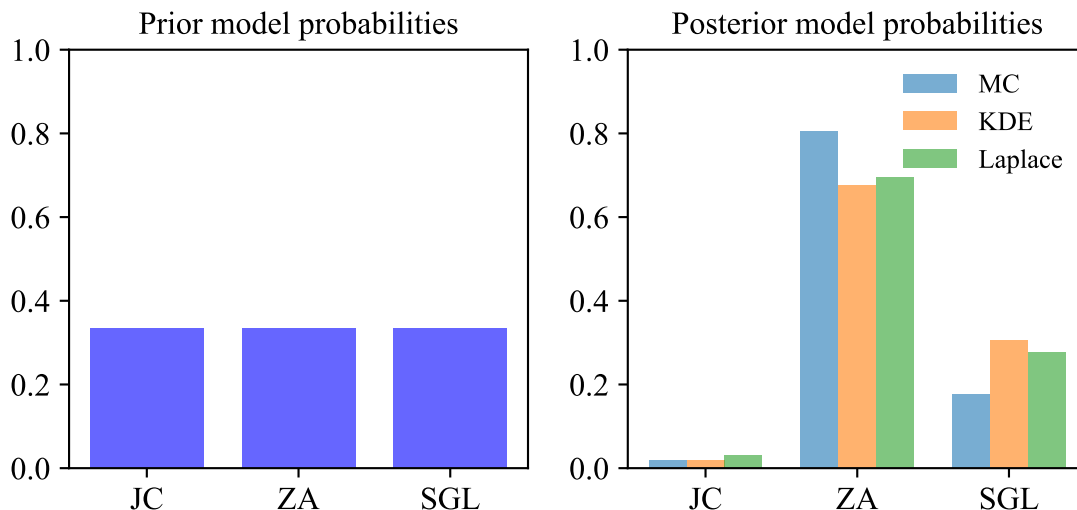


Figure 6.1. Prior (uniform) and posterior probabilities for each model. Each color indicates a different method of approximating the model evidence, used to compute the posterior model probabilities. The different methods are described in Chapter 3.

6.5 Conclusion

The results of this model selection study indicate that ZA may be most appropriate for representing the time-dependent behavior of depth-of-penetration in high-strain-rate, high-temperature impacts such as the experiment in [2]. It should be noted, however, that the prior distributions used for the model parameters were not necessarily accurate descriptions of the underlying parameter uncertainty. Even so, it is certainly worth considering ZA for such applications in the future.

The result of the Bayesian calibrations for this study reduced the initial yield strength in all three models, consistent with the reported value of 1.45 GPa in [2], which indicates that HzB steel is softer than other hardened steels that appear in the ALEGRA material libraries. Parameter values for ZA and SGL were not available for HzB steel in the ALEGRA material libraries prior to this study, so their calibrated values may be useful for future modeling efforts.

Broadly speaking, Bayesian model selection provides an intuitive, systematic way to calibrate and select between models. Calibration of model parameters is done in a systematic, automatic way, as part of the model selection process, once appropriate uncertainty representations are chosen for model parameters and data. Because the uncertainty representations affect the final outcome of the model selection process, care should be taken to have as accurate representations of uncertainty as possible. The results of the model selection provide relative weights for each model that can be used for model averaging, if desired.

Overall, Bayesian model selection provides a framework for calibrating and selecting the best in a collection of models to represent a phenomenon, while accounting for uncertainties in the process. One should note, however, that current model selection paradigms do not evaluate extrapolative ability, besides implementing Occam’s razor. A predictive validation of the chosen model is still advisable, if the model will need to extrapolate outside of the situation in which it was calibrated. These ideas are discussed in detail in [10].

	A [GPa]	B [GPa]	n	C	m
Prior mean values	1.50	0.569	0.22	0.003	1.17
Posterior mean values	1.44	0.565	0.22	0.003	1.17

Table 6.1. Prior vs. posterior mean parameter values for JC.

	c_0 [Pa]	c_1 [Pa]	c_3 [K ⁻¹]	c_4 [K ⁻¹]	c_5 [Pa]	N
Prior mean values	5.00e+7	1.80e+9	1.50e-3	4.50e-5	1.20e+9	6.20e-1
Posterior mean values	4.98e+7	1.49e+9	1.67e-3	4.48e-5	1.68e+9	6.10e-1

Table 6.2. Prior vs. posterior mean parameter values for ZA.

	Y_0 [Pa]	β	n	Y_{max}	G_0 [Pa]	A [Pa ⁻¹]	B [K ⁻¹]
Prior mean values	1.45e+9	2.0	5.0e-1	2.5e+09	7.18e+10	2.06e-11	3.15e-04
Posterior mean values	1.41e+9	2.0	4.99e-1	2.50e+09	7.14e+10	2.05e-11	3.14e-04

Table 6.3. Prior vs. posterior mean parameter values for SGL.

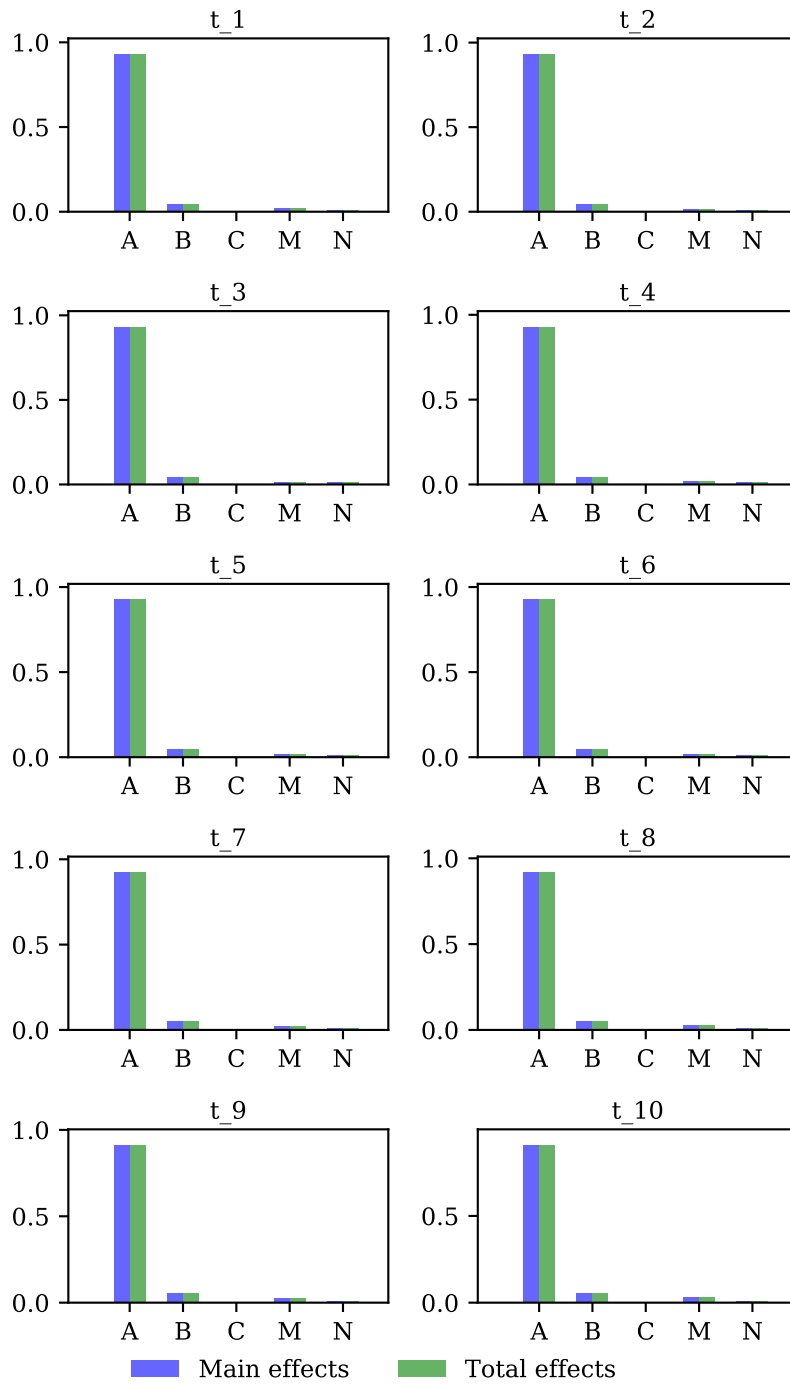


Figure 6.2. Sobol indices for JC.

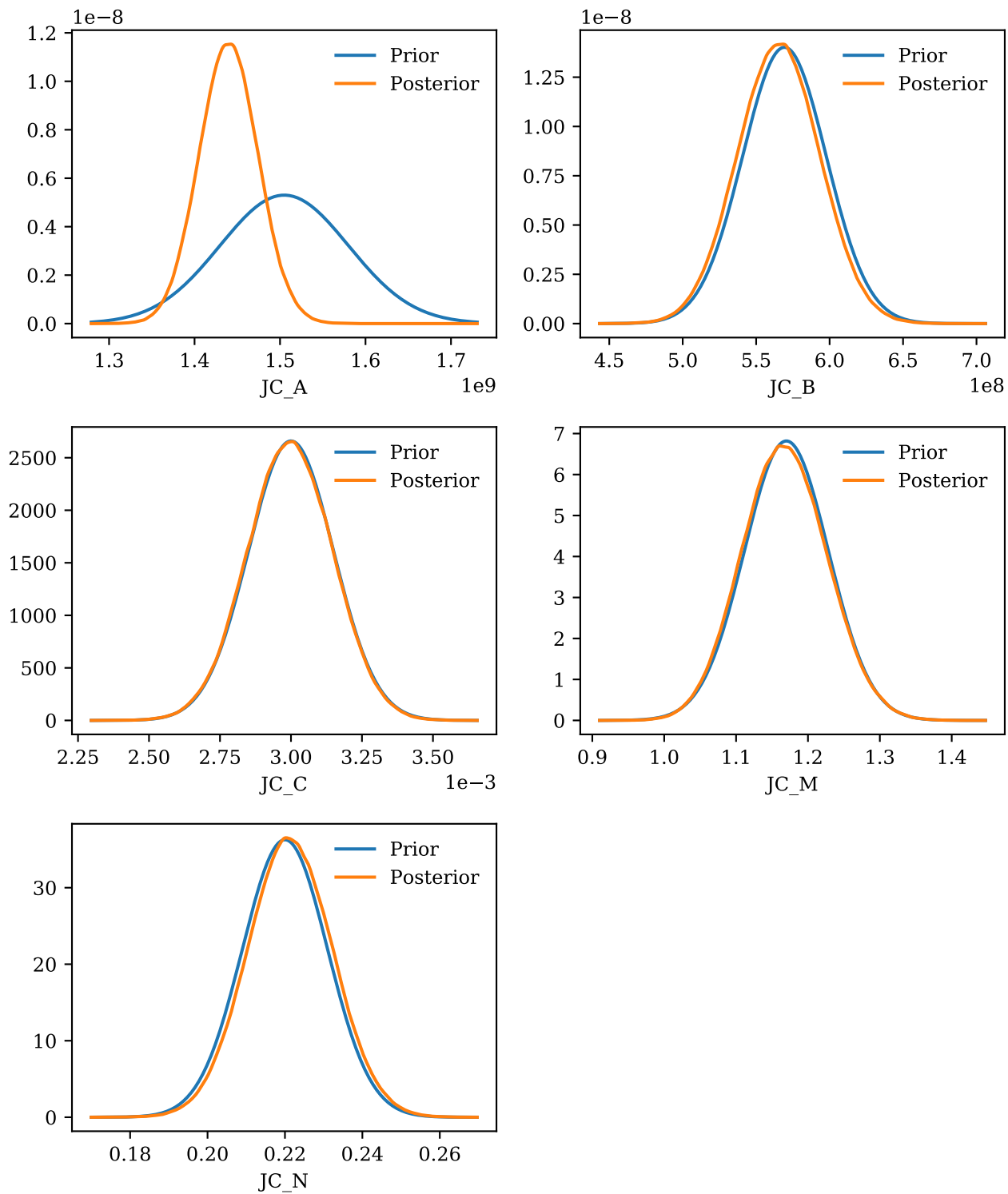


Figure 6.3. Prior vs. posterior distributions for JC.

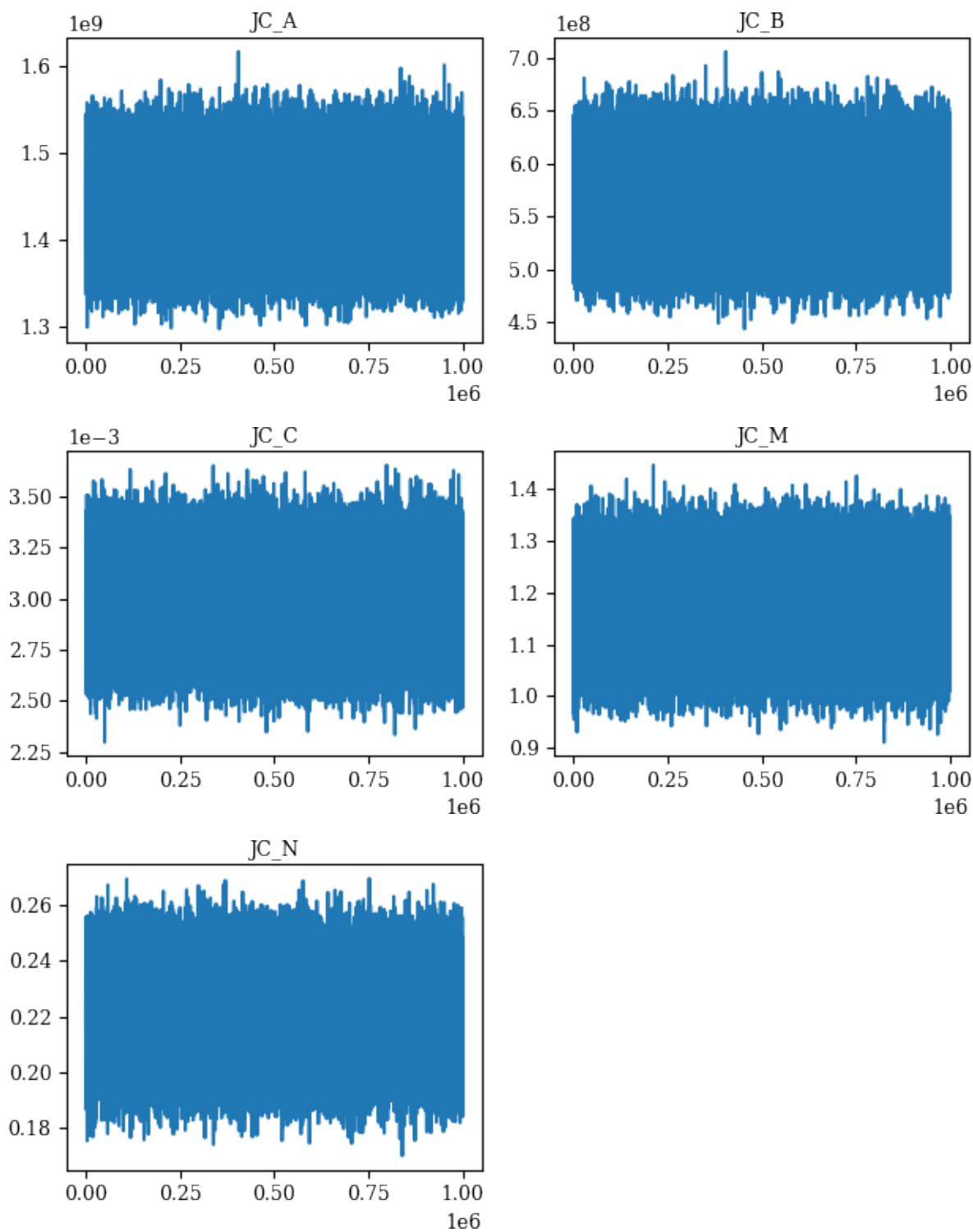


Figure 6.4. Posterior MCMC chains for JC.

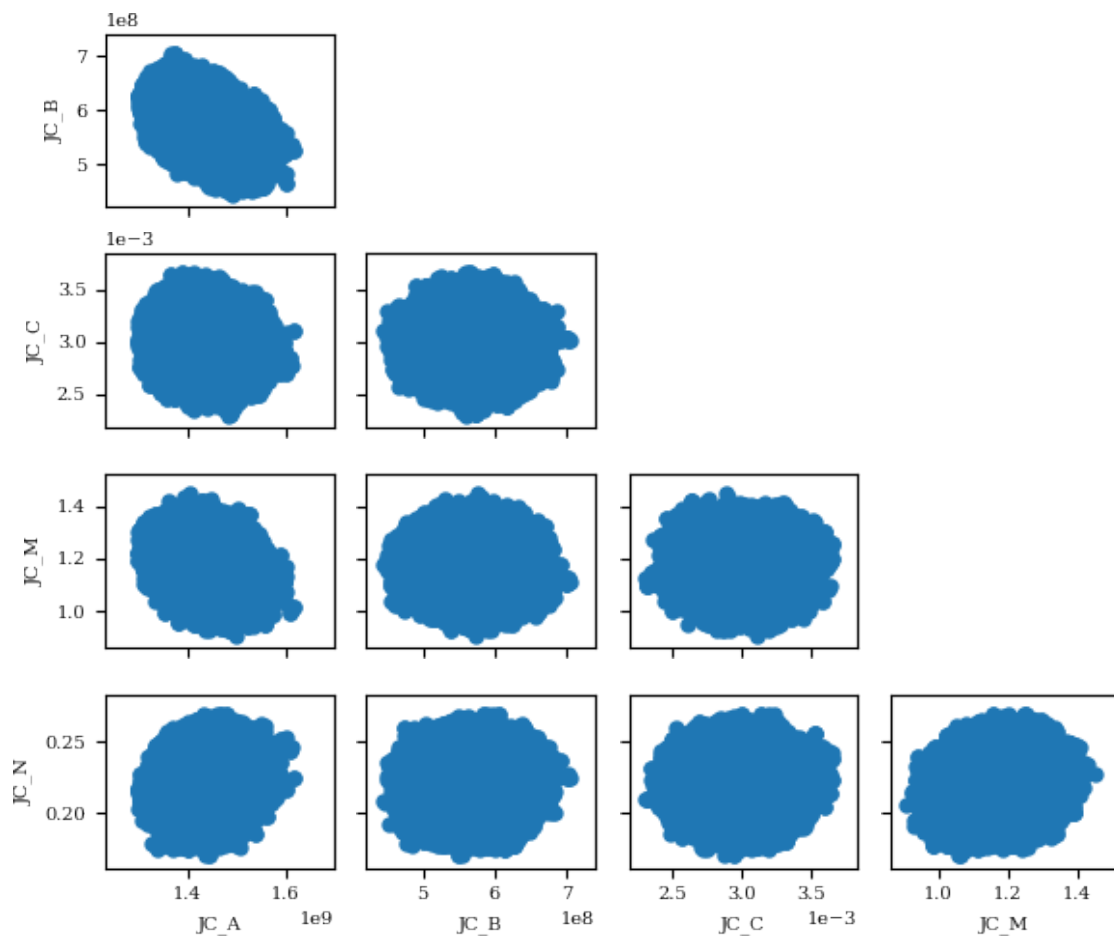


Figure 6.5. Joint marginal samples of the posterior distribution for JC model parameters.

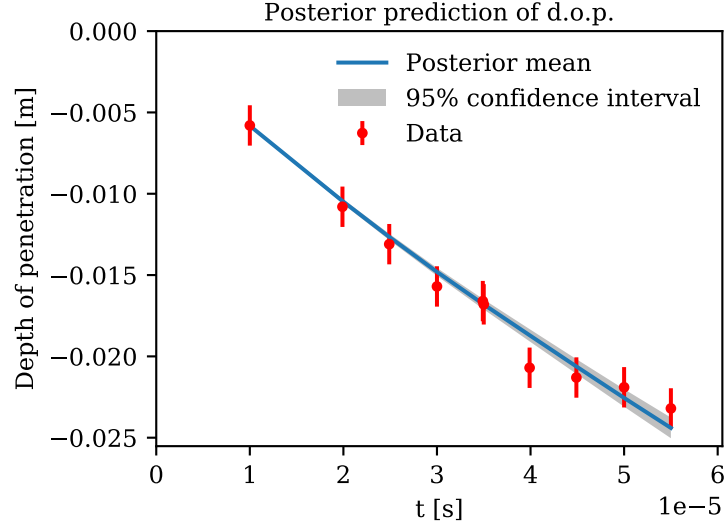


Figure 6.6. Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.

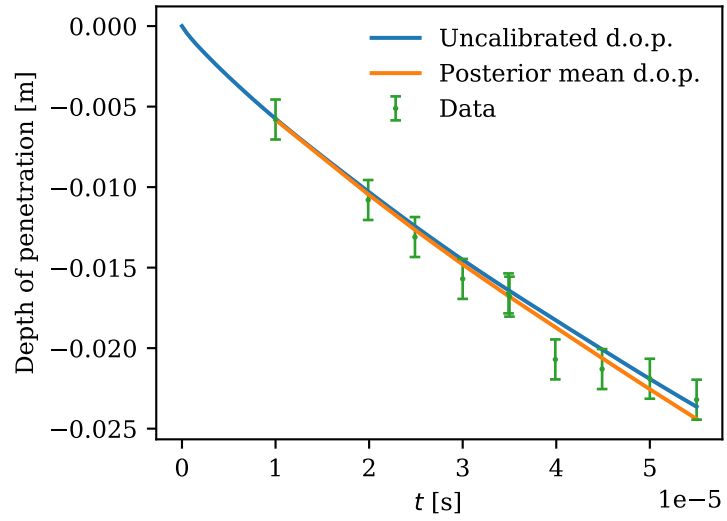


Figure 6.7. Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of JC model parameters.

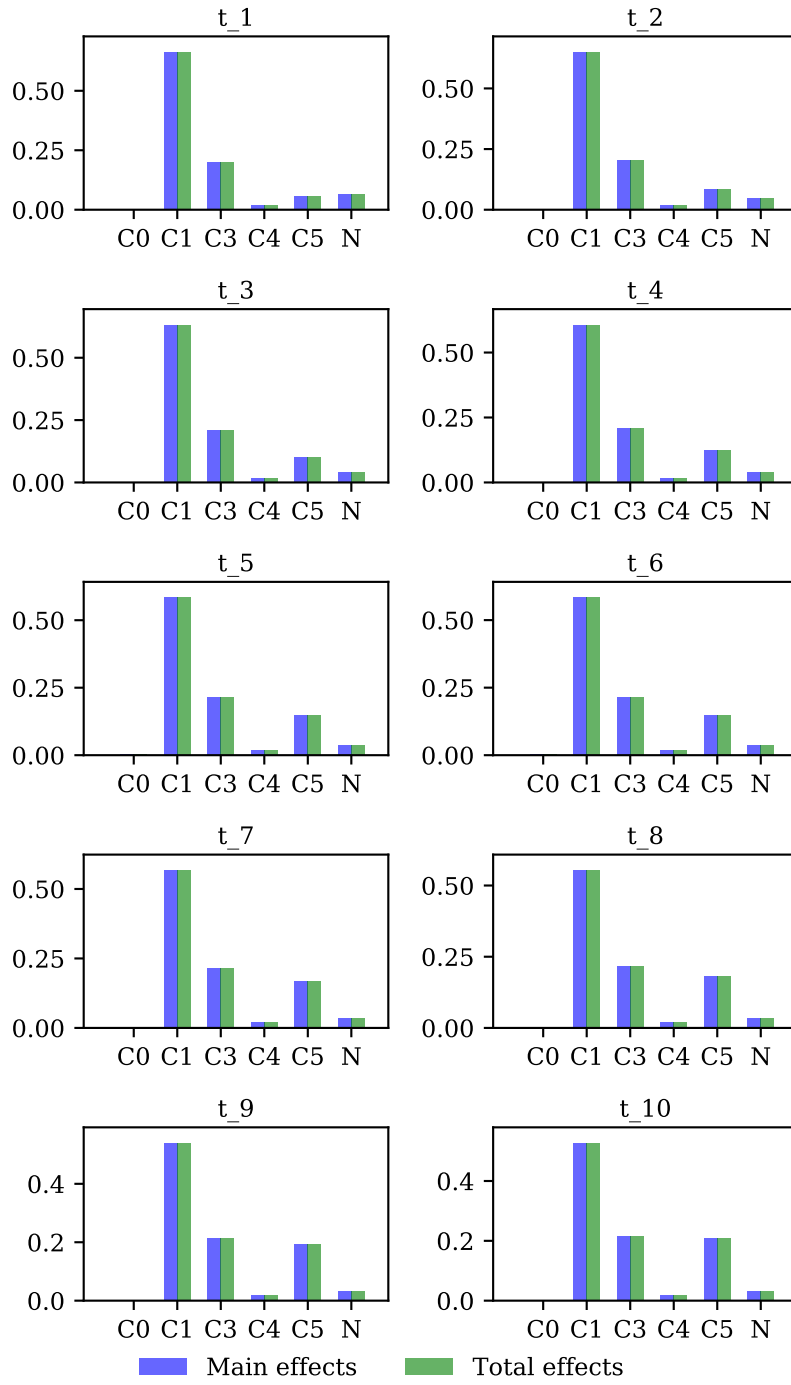


Figure 6.8. Sobol indices for ZA.

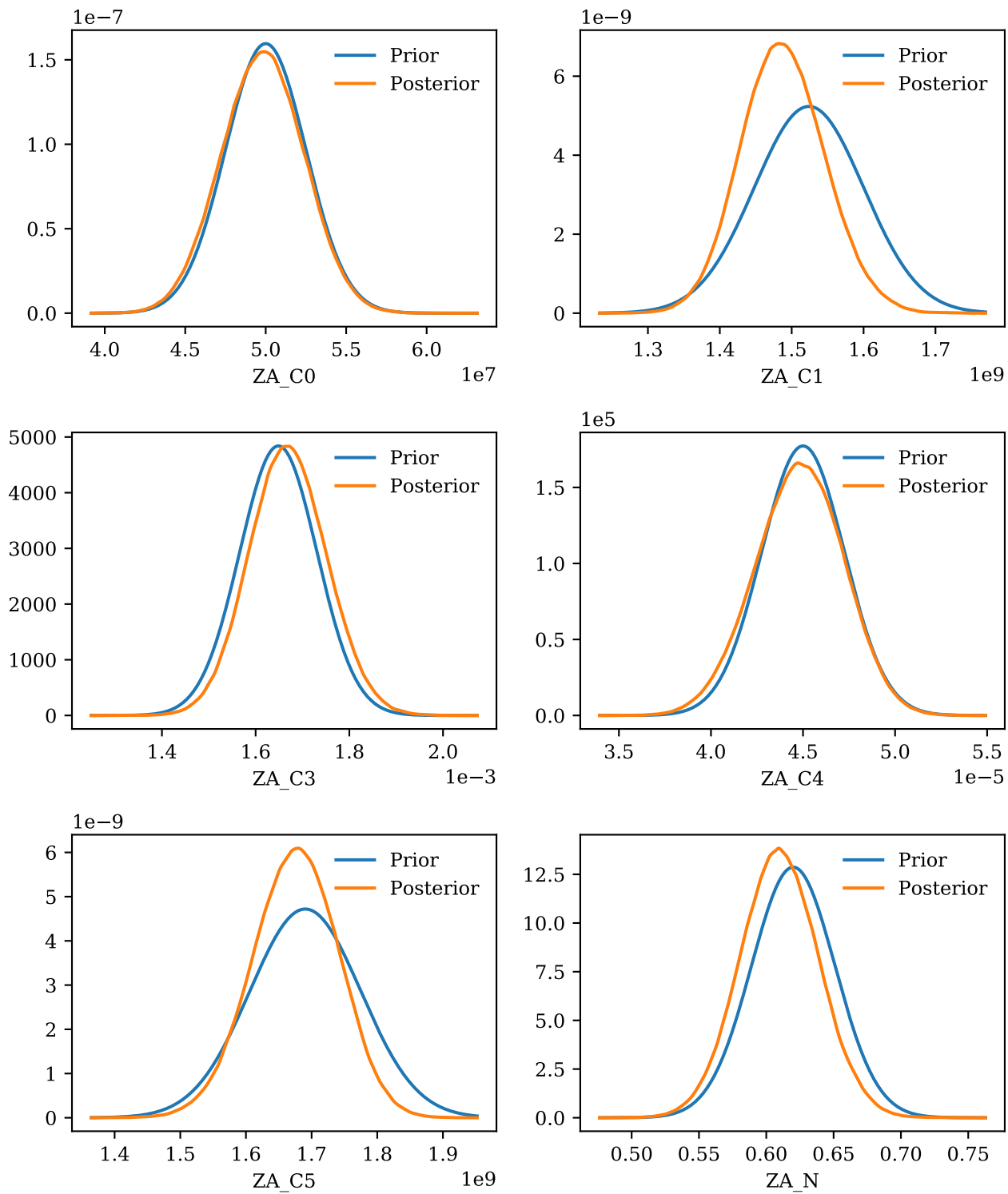


Figure 6.9. Prior vs. posterior distributions for ZA.

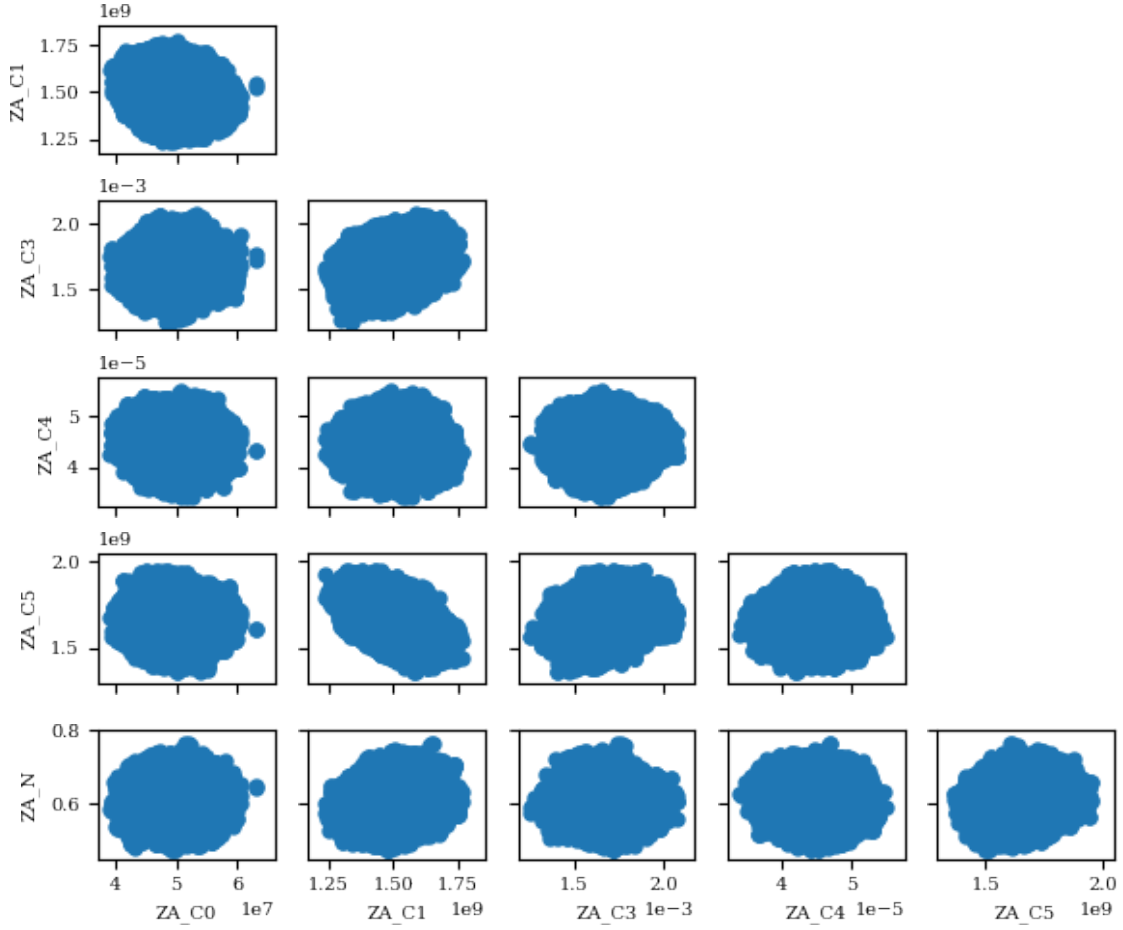


Figure 6.10. Joint marginal samples of the posterior distribution for ZA model parameters.

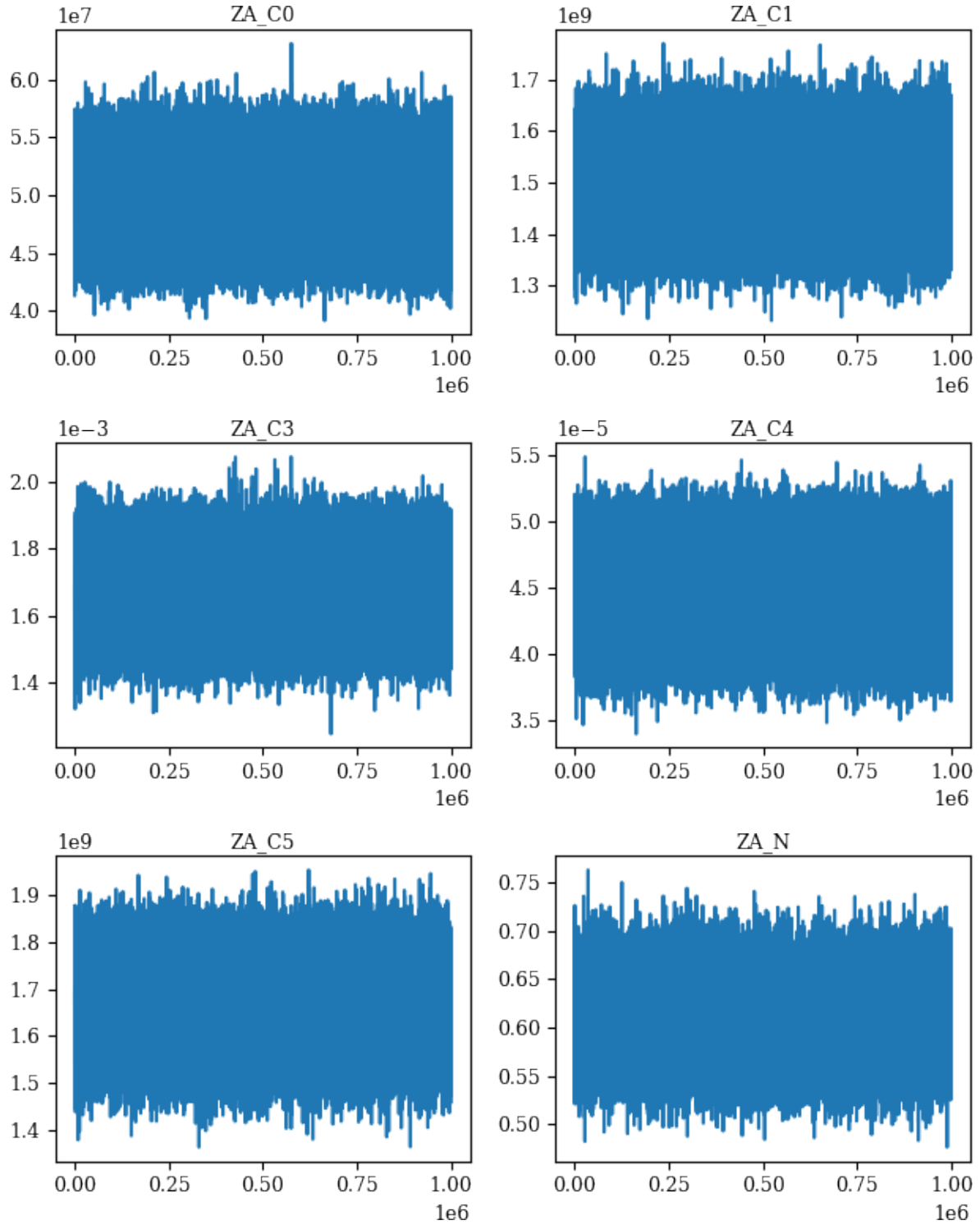


Figure 6.11. Posterior MCMC chains for ZA.

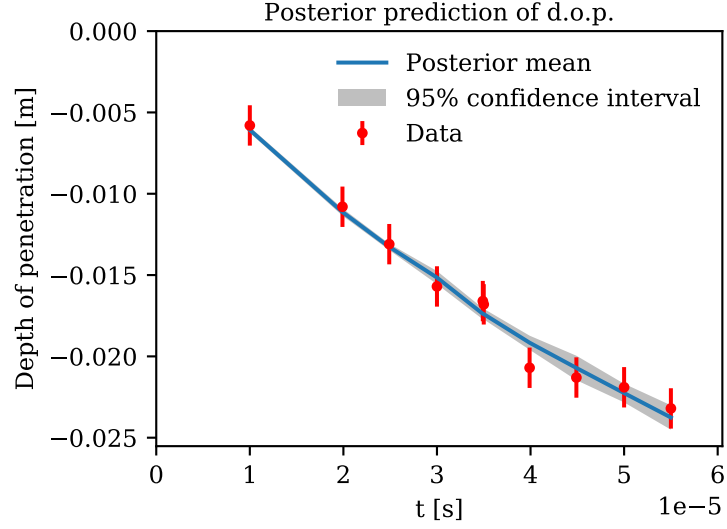


Figure 6.12. Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.

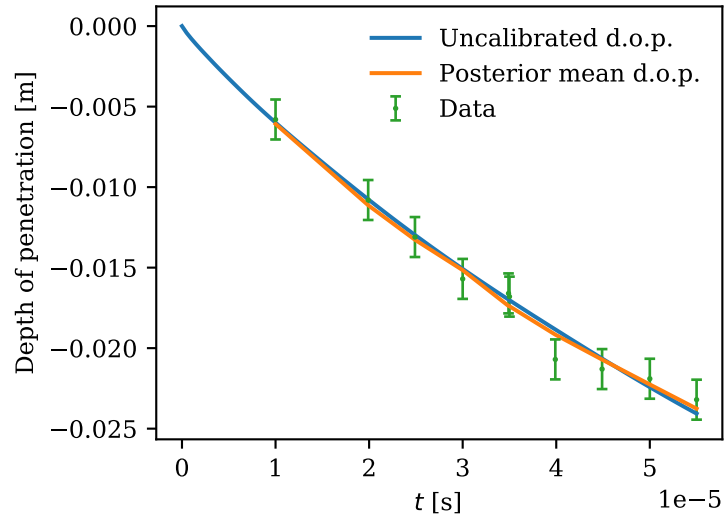


Figure 6.13. Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of ZA model parameters.

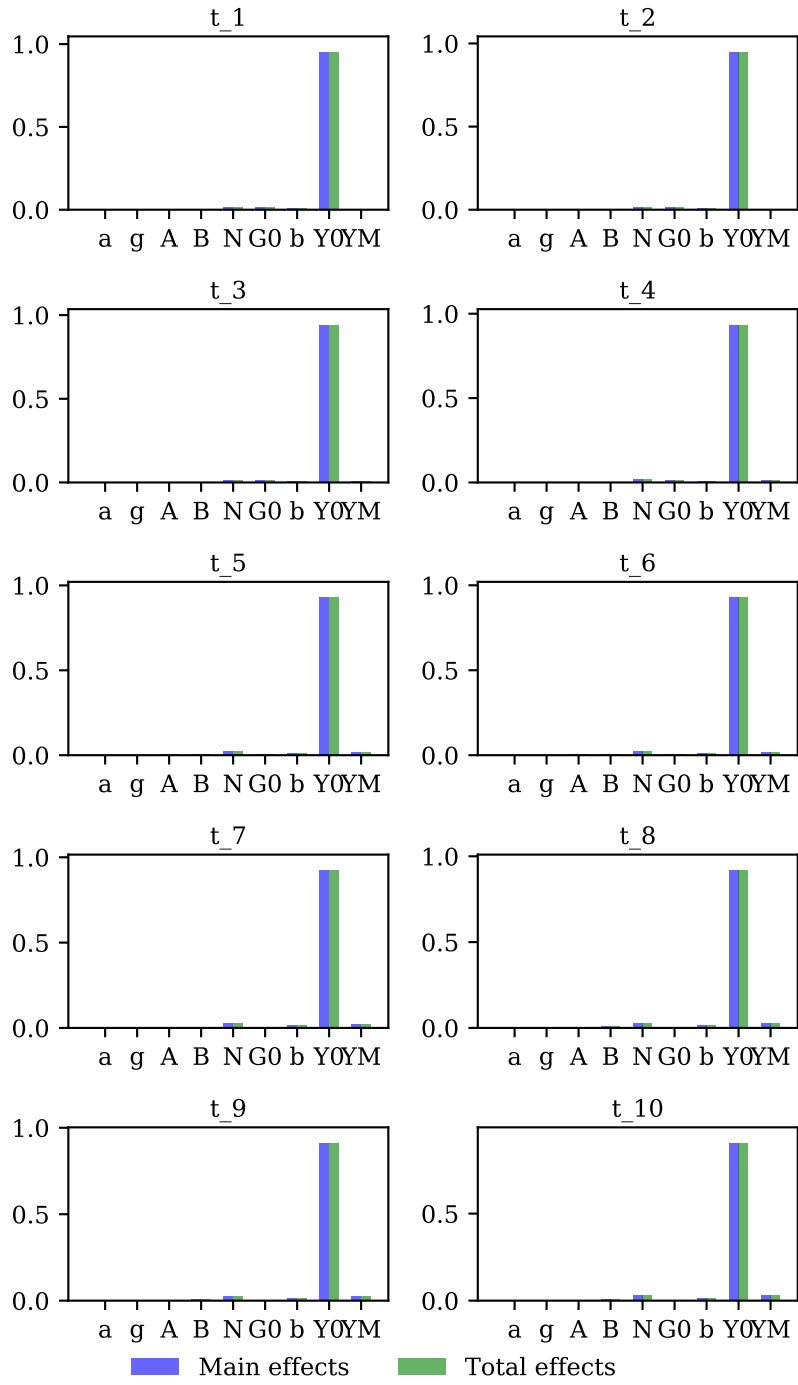


Figure 6.14. Sobol indices for SGL.

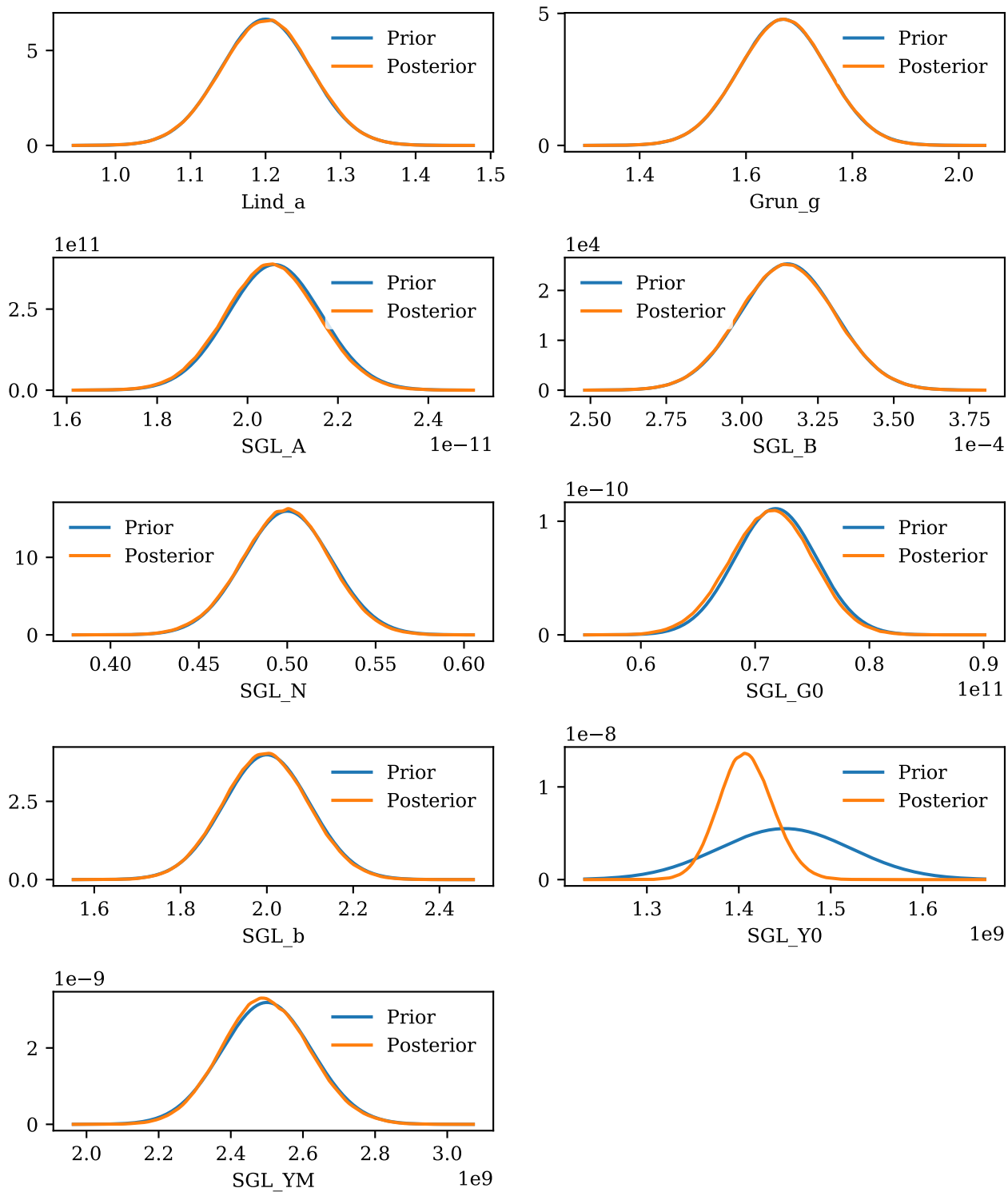


Figure 6.15. Prior vs. posterior distributions for SGL.

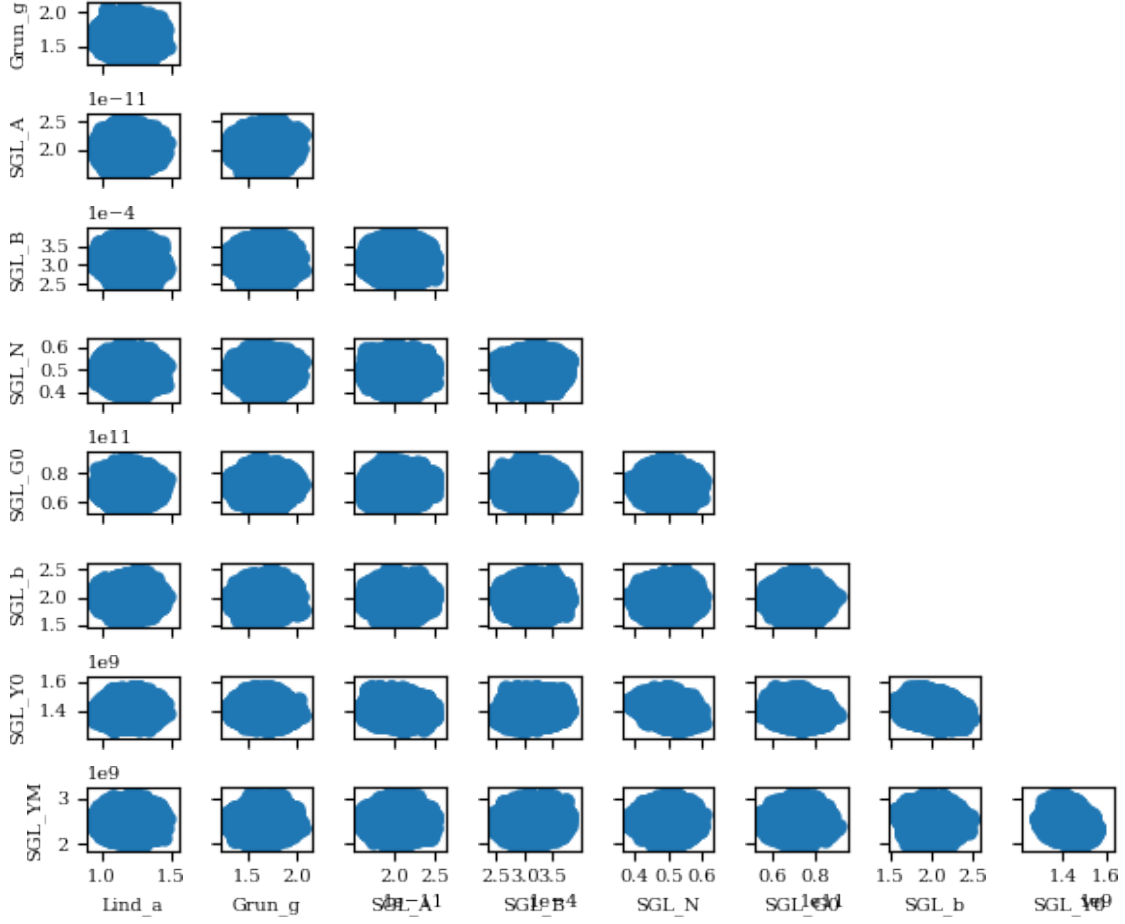


Figure 6.16. Joint marginal samples of the posterior distribution for SGL model parameters.

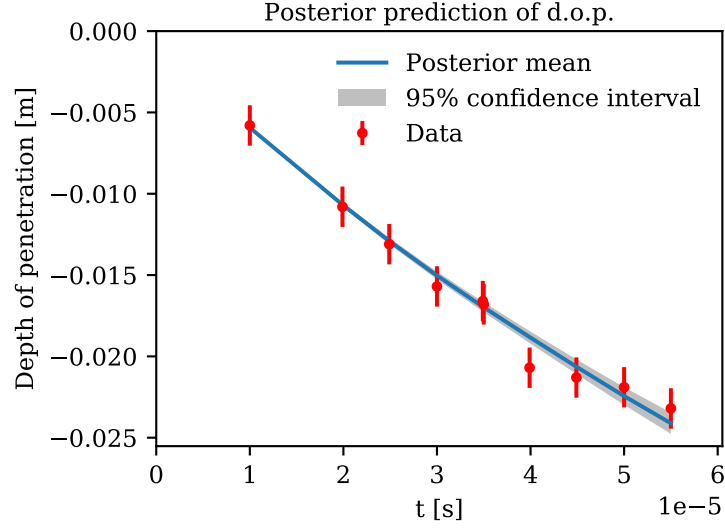


Figure 6.17. Mean and 95% confidence interval for the depth of penetration as a function of time after Bayesian calibration.

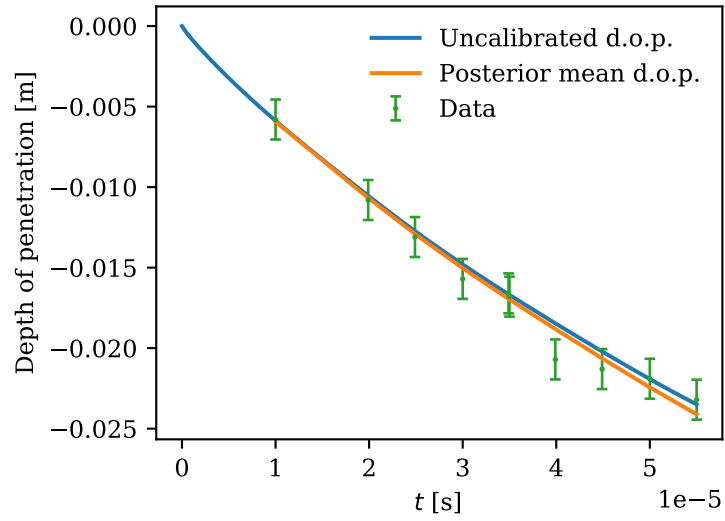


Figure 6.18. Initial d.o.p before any calibration compared to the mean d.o.p. after Bayesian calibration of SGL model parameters.

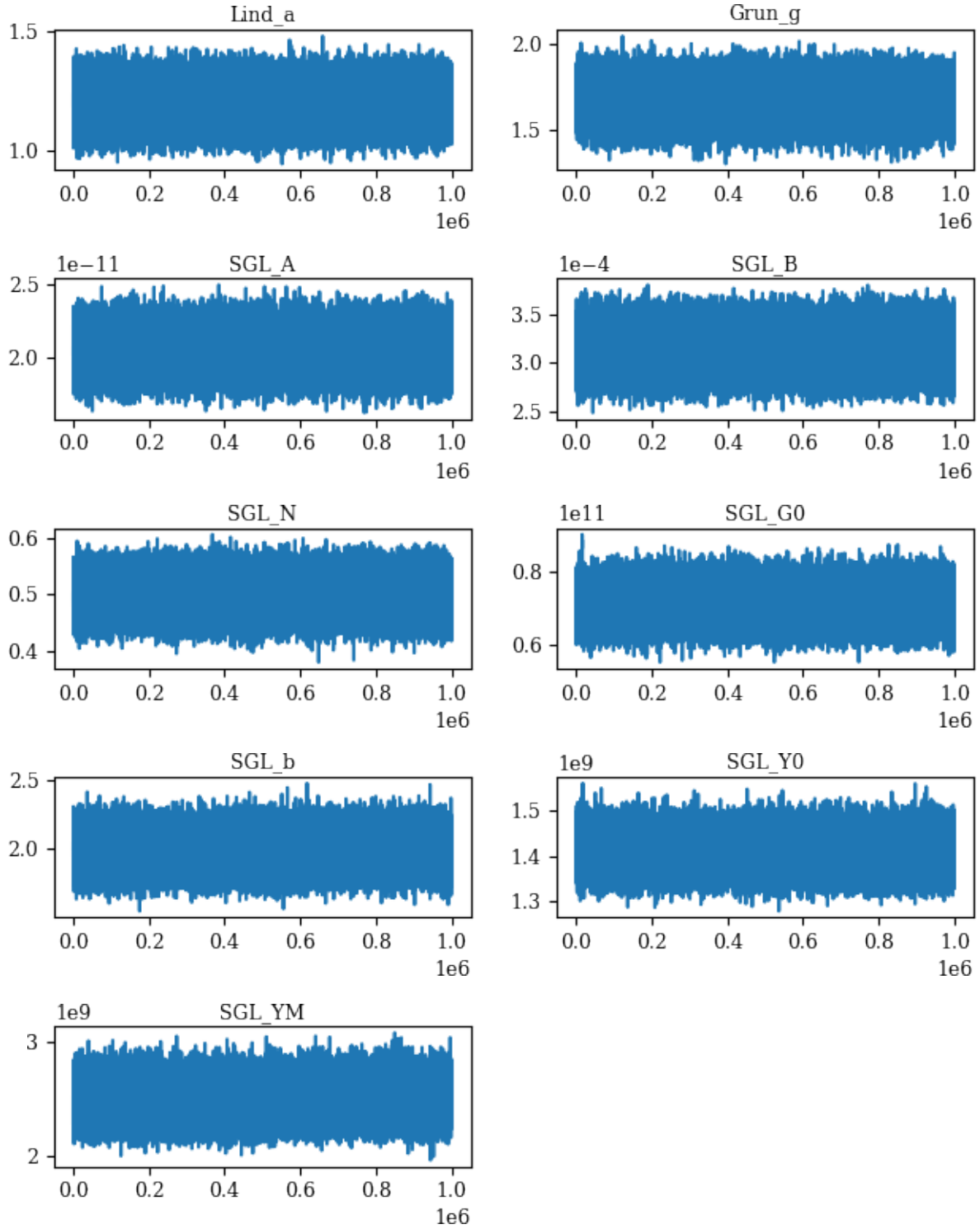


Figure 6.19. Posterior MCMC chains for SGL.

Appendix A

Variance-based global sensitivity analysis

There are many approaches to sensitivity analysis, but the approach taken here is to use global, variance-based sensitivity analysis as first described by Sobol in [18]. An explanation of the technical details that provides some intuition is provided in [14]. The basic concept is this: if we consider our model output $\mathcal{O}(\boldsymbol{\theta})$ as a function of uncertain variables $\boldsymbol{\theta}$, the output will vary with the uncertain parameters. One measure of variation caused by one parameter θ_i is

$$V_{\theta_i} [\mathbb{E}_{\boldsymbol{\theta}_{\sim i}} [\mathcal{O} | \theta_i]].$$

The inner expression represents fixing the parameter θ_i and taking the average over all possible values of the other parameters (i.e. all θ_j where $j \neq i$). Then the outer expression represents taking the variance over all possible values of θ_i .

The so-called “main effects” Sobol index S_i is defined as

$$S_i = \frac{V_{\theta_i} [\mathbb{E}_{\boldsymbol{\theta}_{\sim i}} [\mathcal{O} | \theta_i]]}{V[\mathcal{O}]}$$

and can be seen as a measure of the amount that variations in θ_i alone contribute to the total variance in $\mathcal{O}$.

The other Sobol index that is commonly used in practice is the “total effect” index, defined as

$$T_i = 1 - \frac{V_{\boldsymbol{\theta}_{\sim i}} [\mathbb{E}_{\theta_i} [\mathcal{O} | \boldsymbol{\theta}_{\sim i}]]}{V[\mathcal{O}]} = \frac{V[\mathcal{O}] - V_{\boldsymbol{\theta}_{\sim i}} [\mathbb{E}_{\theta_i} [\mathcal{O} | \boldsymbol{\theta}_{\sim i}]]}{V[\mathcal{O}]}$$

The second equality is perhaps more enlightening as to what this index is communicating. In analogy to the previous discussion,

$$V_{\boldsymbol{\theta}_{\sim i}} [\mathbb{E}_{\theta_i} [\mathcal{O} | \boldsymbol{\theta}_{\sim i}]]$$

is the variance in $\mathcal{O}$ induced by varying all the model parameters except θ_i . Then subtracting that quantity from the total variance leaves behind the contribution to variance from varying θ_i alone as well as varying θ_i along with model parameter with which it is correlated.

Both measures provide a description of the percent of the variation in $\mathcal{O}$ that is induced by each model parameter. The observable $\mathcal{O}$ in this paper is the depth of penetration of a tungsten rod into a steel plate as a function of time. One can compute the Sobol indices as defined above by doing a parameter study such as the LHS sampling described in 4.2 and using surrogates as described also in 4.2 to approximate the variances. This was done for the ZA model, for each depth of penetration at each observation time, as shown in Figure 6.8. The depth of penetration as a function of time was most sensitive to model parameters c_1 and c_3 and c_5 parameters, which are associated with initial yield strength, thermal softening, and work hardening. This information could be useful in many ways, but two specifically will be discussed here.

A.1 Sensitivity analysis for dimension reduction

If one is attempting to calibrate a set of model parameters, it can be useful to know which parameters a quantity of interest (QOI) or observable is sensitive to. If the quantity of interest (e.g. the depth of penetration as a function of time) is only sensitive to two out of five of the model parameters, one can fix the parameters to which it is insensitive at their nominal values and proceed with calibrating only those which were deemed significant.

For the models considered in this paper this isn't crucial, but when one begins considering models with a high level of parametrization, traditional optimization methods and UQ methods would begin to struggle. In such a context, sensitivity analysis can be an incredibly useful tool in identifying on which model parameters to focus calibration efforts.

A.2 Sensitivity analysis for experimental design

On the other hand, one may wish to know what sort of data is required to best inform a parameter. In this scenario, sensitivity analysis can be used to identify which observable quantities are sensitive to the parameter of interest. For instance, as shown in Figure 6.2, it would not be helpful to use depth of penetration as a function of time to determine the exponents in the JC model, because the depth of penetration was not sensitive to their variations. Obviously, physical intuition and expert guidance on which parameters will be informed by which experiments is preferred, but sensitivity analysis can help to fill the gaps when such knowledge is unavailable.

References

- [1] B.M. Adams, L.E. Bauman, W.J. Bohnhoff, K.R. Dalbey, M.S. Ebeida, J.P. Eddy, M.S. Eldred, P.D. Hough, K.T. Hu, J.D. Jakeman, J.A. Stephens, L.P. Swiler, D.M. Vigil, and T.M. Wildey. Dakota, A Multilevel Parallel Object-Oriented Framework for Design Optimization, Parameter Estimation, Uncertainty Quantification, and Sensitivity Analysis: Version 6.0 Users Manual. Technical Report SAND2014-4633, Sandia National Laboratories, July 2014. July 2014. Other versions released December 2009 (Version 5.0), December 2010 (Version 5.1), November 2011 (Version 5.2), February 2013 (Version 5.3), May 2013 (Version 5.3.1), November 2013 (Version 5.4), November 2014 (Version 6.1), May 2015 (Version 6.2), November 2015 (Version 6.3), May 2016 (Version 6.4), November 2016 (Version 6.5), May 2017 (Version 6.6), November 2017 (Version 6.7).
- [2] Charles E Anderson, Volker Hohler, James D Walker, and Alois J Stilp. Time-resolved penetration of long rods into steel targets. *International Journal of Impact Engineering*, 16(1):1–18, 1995.
- [3] Steve Brooks, Andrew Gelman, Galin Jones, and Xiao-Li Meng. *Handbook of Markov Chain Monte Carlo*. CRC press, 2011.
- [4] M. S. Eldred and J. Burkardt. Comparison of non-intrusive polynomial chaos and stochastic collocation methods for uncertainty quantification. In *Proceedings of the 47th AIAA Aerospace Sciences Meeting and Exhibit*, number AIAA-2009-0976, Orlando, FL, January 5–8, 2009.
- [5] Ghassan Hamra, Richard MacLehose, and David Richardson. Markov chain Monte Carlo: an introduction for epidemiologists. *International journal of epidemiology*, 42(2):627–634, 2013.
- [6] John D Hunter. Matplotlib: A 2D graphics environment. *Computing In Science & Engineering*, 9(3):90–95, 2007.
- [7] E. T. Jaynes. Prior probabilities. *IEEE Transactions on Systems Science and Cybernetics*, 4(3):227–241, 1968.
- [8] Gordon R Johnson and William H Cook. A constitutive model and data for metals subjected to large strains, high strain rates and high temperatures. In *Proceedings of the 7th International Symposium on Ballistics*, volume 21, pages 541–547. The Netherlands, 1983.
- [9] Eric Jones, Travis Oliphant, Pearu Peterson, et al. SciPy: Open source scientific tools for Python, 2001–. [Online; accessed December 7, 2017].

- [10] Todd A Oliver, Gabriel Terejanu, Christopher S Simmons, and Robert D Moser. Validating predictions of unobserved quantities. *Computer Methods in Applied Mechanics and Engineering*, 283:1310–1335, 2015.
- [11] Costas Papadimitriou, James L Beck, and Lambros S Katafygiotis. Asymptotic expansions for reliability and moments of uncertain systems. *Journal of Engineering Mechanics*, 123(12):1219–1229, 1997.
- [12] C. E. Rasmussen and C. K. I Williams. *Gaussian process for machine learning*. MIT Press, 2006.
- [13] Allen C Robinson, Thomas A Brunner, Susan Carroll, Richard Drake, Christopher J Garasi, Thomas Gardiner, Thomas Haill, Heath Hanshaw, David Hensinger, Duane Labreche, et al. ALEGRA: An arbitrary Lagrangian-Eulerian multimaterial, multi-physics code. *AIAA Paper*, 1235:2008, 2008.
- [14] Andrea Saltelli, Paola Annoni, Ivano Azzini, Francesca Campolongo, Marco Ratto, and Stefano Tarantola. Variance based sensitivity analysis of model output. design and estimator for the total sensitivity index. *Computer Physics Communications*, 181(2):259–270, 2010.
- [15] T. Santner, B. Williams, and W. Notz. *The design and analysis of computer experiments*. Springer, New York, NY, 2003.
- [16] David W Scott. *Multivariate density estimation: theory, practice, and visualization*. John Wiley & Sons, 2015.
- [17] Ralph C Smith. *Uncertainty quantification: theory, implementation, and applications*, volume 12. SIAM, 2013.
- [18] Ilya M Sobol. Sensitivity estimates for nonlinear mathematical models. *Mathematical Modelling and Computational Experiments*, 1(4):407–414, 1993.
- [19] DJ Steinberg, SG Cochran, and MW Guinan. A constitutive model for metals applicable at high-strain rate. *Journal of Applied Physics*, 51(3):1498–1504, 1980.
- [20] Laura Painton Swiler and Gregory Dane Wyss. A user’s guide to Sandia’s Latin Hypercube Sampling software: LHS UNIX library/standalone version. Technical Report SAND2004-2439, Sandia National Laboratories, July 2004.
- [21] Felipe AC Viana, Christian Gogu, and Raphael T Haftka. Making the most out of surrogate models: tricks of the trade. *Proceedings of the 2010 ASME IDETC and CIE, Montreal, Canada, Paper No. DETC2010-28813*, 2010.
- [22] Stéfan van der Walt, S Chris Colbert, and Gael Varoquaux. The NumPy array: a structure for efficient numerical computation. *Computing in Science & Engineering*, 13(2):22–30, 2011.

- [23] Larry Wasserman. Bayesian model selection and model averaging. *Journal of mathematical psychology*, 44(1):92–107, 2000.
- [24] Stefan Weinzierl. Introduction to Monte Carlo methods. *arXiv preprint hep-ph/0006269*, 2000.
- [25] D. Xiu. *Numerical Methods for Stochastic Computations: A Spectral Method Approach*. Princeton University Press, 2010.
- [26] Frank J Zerilli and Ronald W Armstrong. Dislocation-mechanics-based constitutive relations for material dynamics calculations. *Journal of Applied Physics*, 61(5):1816–1825, 1987.

Distribution

External Distribution:

University of Texas at Austin, Institute for Computational Engineering and Sciences

Attn:

Prof. Robert Moser (electronic copy)

Teresa Portone (electronic copy)

U.S. Army Research Laboratory

Attn:

Robert L. Doney (electronic copy)

Brady Aydelotte (electronic copy)

R. Brian Leavy (electronic copy)

Internal Distribution:

Glen Hansen, 01443 (electronic copy)

John Niederhaus, 01446 (electronic copy)

Jason Sanchez, 01443 (electronic copy)

Laura Swiler, 01441 (electronic copy)

Brian Adams, 01441 (electronic copy)

Kathryn Maupin, 01441 (electronic copy)

Mike Eldred, 01441 (electronic copy)

Jim Stewart, 01440 (electronic copy)

Dan Turner, 01441 (electronic copy)

Technical Library, 09536 (electronic copy)

