

Minerva Access is the Institutional Repository of The University of Melbourne

Author/s:

Monsalve-Bravo, GM;Lawson, BAJ;Drovandi, C;Burrage, K;Brown, KS;Baker, CM;Vollert, SA;Mengersen, K;McDonald-Madden, E;Adams, MP

Title:

Analysis of sloppiness in model simulations: Unveiling parameter uncertainty when mathematical models are fitted to data

Date:

2022-09-23

Citation:

Monsalve-Bravo, G. M., Lawson, B. A. J., Drovandi, C., Burrage, K., Brown, K. S., Baker, C. M., Vollert, S. A., Mengersen, K., McDonald-Madden, E. & Adams, M. P. (2022). Analysis of sloppiness in model simulations: Unveiling parameter uncertainty when mathematical models are fitted to data. Science Advances, 8 (38), <https://doi.org/10.1126/sciadv.abm5952>.

Persistent Link:

<https://hdl.handle.net/11343/327084>

License:

CC BY

MATHEMATICS

Analysis of sloppiness in model simulations: Unveiling parameter uncertainty when mathematical models are fitted to data

Gloria M. Monsalve-Bravo^{1,2,3*}, Brodie A. J. Lawson^{4,5,6,7}, Christopher Drovandi^{4,5,6}, Kevin Burrage^{5,6,7,8}, Kevin S. Brown^{9,10}, Christopher M. Baker^{11,12,13}, Sarah A. Vollert^{4,5,6}, Kerrie Mengersen^{4,5,6}, Eve McDonald-Madden^{1,2†}, Matthew P. Adams^{3,4,5,6†}

Copyright © 2022
The Authors, some
rights reserved;
exclusive licensee
American Association
for the Advancement
of Science. No claim to
original U.S. Government
Works. Distributed
under a Creative
Commons Attribution
License 4.0 (CC BY).

This work introduces a comprehensive approach to assess the sensitivity of model outputs to changes in parameter values, constrained by the combination of prior beliefs and data. This approach identifies stiff parameter combinations strongly affecting the quality of the model-data fit while simultaneously revealing which of these key parameter combinations are informed primarily by the data or are also substantively influenced by the priors. We focus on the very common context in complex systems where the amount and quality of data are low compared to the number of model parameters to be collectively estimated, and showcase the benefits of this technique for applications in biochemistry, ecology, and cardiac electrophysiology. We also show how stiff parameter combinations, once identified, uncover controlling mechanisms underlying the system being modeled and inform which of the model parameters need to be prioritized in future experiments for improved parameter inference from collective model-data fitting.

INTRODUCTION

A single biological cell is itself a complex system, as is an organism made up of such cells, as is an ecosystem of those organisms interacting with one another. Despite the diversity of systems composing our world, many of these share similar structural and functional features that can be unraveled through computer simulation (1–3). Consequently, modeling and simulation have become increasingly important to understand and predict the underlying behavior of systems across different scales (3–5), including molecules (6), cells (7, 8), engineered processes (9), and astrophysical phenomena (10). Continuous advances in model descriptions of reality together with the model fit to experimental data have improved the fidelity of computer experiments and made them much more predictive (1, 2). However, the cost of this fidelity is an increase in the number of model parameters (4), and a greater risk that these parameters cannot be uniquely identified (3, 11–13). For statistical models of familiar form, one may be able to formally determine how and to what extent parameters can possibly be identified. Lewbel (14) provides

many such examples. When it comes to complex models defined, for example, in terms of the solution of a set of differential equations, however, a more practical approach will often be required for parameter estimation (3, 11). Expectedly, a substantial amount of uncertainty in parameter values often remains after even a very successful fit of the model to data (15–17).

Sensitivity analysis and uncertainty quantification comprise a whole field dedicated to learning about how model behavior is controlled by their parameters (18–20). These techniques can be used to assess the sensitivity of the model-data fit to changes in parameter values either in a local sense, around a single point (i.e., the set of best-fit parameter values), or in a global sense, across all plausible parameter values consistent with the available data (11, 15, 21). An alternative approach is Bayesian inference (22, 23), an increasingly used modeling technique that accounts for collective parameter uncertainty constrained by the combination of both data and prior beliefs (5, 7, 8, 17, 24, 25). However, regardless of the approach taken to characterize the effects of changes in parameter values on model outputs, critical model parameters are often considered as individuals in terms of their impact on the model behavior (19). Sensitivity analysis typically considers the derivative of model outputs with respect to the parameters (15, 18, 21), while a Bayesian posterior is analyzed predominantly in terms of its marginal distributions (7, 17). When combinations of model parameters are considered, it is largely in terms of crude numerical scores (19–21). Unfortunately, model parameters that are not very constrained by the data are often assumed not to have a strong influence on model predictions, although it is the case of many systems that certain combinations of seemingly unconstrained model parameters are more narrowly constrained by the data than any of the individual model parameters (11–13, 16).

Model parameters can act together or against each other, and often must be understood in terms of their combinations (13, 15). Parameter combinations that significantly influence model predictions, called stiff eigenparameters, essentially act as emergent “control knobs” for the model: Predictions are possible without precise

¹School of Earth and Environmental Sciences, The University of Queensland, St Lucia, QLD 4072, Australia. ²Centre for Biodiversity and Conservation Science, The University of Queensland, St Lucia, QLD 4072, Australia. ³School of Chemical Engineering, The University of Queensland, St Lucia, QLD 4072, Australia. ⁴Centre for Data Science, Queensland University of Technology, Brisbane, QLD 4001, Australia. ⁵ARC Centre of Excellence for Mathematical and Statistical Frontiers, Queensland University of Technology, Brisbane, QLD 4001, Australia. ⁶School of Mathematical Sciences, Queensland University of Technology, Brisbane, QLD 4001, Australia. ⁷ARC Centre of Excellence for Plant Success in Nature and Agriculture, Queensland University of Technology, Brisbane, QLD 4001, Australia. ⁸Department of Computer Science, University of Oxford, Oxford OX1 3QD, UK. ⁹Department of Pharmaceutical Sciences, Oregon State University, Corvallis, OR 97331, USA. ¹⁰Department of Chemical, Biological, & Environmental Engineering, Oregon State University, Corvallis, OR 97331, USA. ¹¹School of Mathematics and Statistics, The University of Melbourne, Parkville, VIC 3010, Australia. ¹²Melbourne Centre for Data Science, The University of Melbourne, Parkville, VIC 3010, Australia. ¹³Centre of Excellence for Biosecurity Risk Analysis, The University of Melbourne, Parkville, VIC 3010, Australia.

*Corresponding author. Email: g.monsalvebravo@uq.edu.au

†These authors contributed equally to this work.

knowledge of individual parameter values as it is these stiff eigenparameters that are tightly constrained by the data (16, 26). Conversely, the model-data fit may also be relatively insensitive to some other parameter combinations, called sloppy eigenparameters (12, 13), which hence are poorly constrained by the data (16, 27). Recently, efforts have been made to unravel these connections among parameters through the expanding literature on model sloppiness (12, 28–31). Methods to analyze model sloppiness seek to expose the sensitivities of the model-data fit to changes in sets of parameter values by characterizing the topography of the surface describing how the model-data fit depends on the model parameters in the vicinity of the best-fit parameter values (15, 16). However, thus far, such methods have primarily focused on the field of systems biology where there is little prior knowledge of parameter values (11, 26, 27), and so, the sensitivities of the model-data fit to changes in parameter values remain to be considered in the context where prior information is also available (e.g., from experts or previous studies) to inform parameter values (11, 32–34).

In this work, we propose a comprehensive approach to characterize local and global sensitivities of the model-data fit to changes in parameter values. This is achieved by bringing a Bayesian inference perspective (22, 23) to the analysis of sloppiness that consequently leads to the robust identification of the stiff eigenparameters. In this way, analysis of sloppiness gains the ability to incorporate prior information and to look beyond the curvature at a single point (i.e., the set of best-fit parameter values) in an uncertainty-informed way. Meanwhile, Bayesian inference gains a tool to identify well-constrained combinations of parameters that can be otherwise hidden when considering the uncertainty in individual model parameters, critical when the number of parameters to be estimated is large.

As part of our comprehensive approach, we extend the usage of two well-established Bayesian approaches to dimensionality reduction (35–38) to define the sensitivity matrix that underlies the analysis of model sloppiness, suitably calculated using the posterior samples generated by Bayesian inference. The first definition uses the covariance of the posterior samples to inform parameter space curvature in a nonlocalized manner (12, 39), with ties to classical principal component analysis (PCA) (36). This approach has appeared in works analyzing model sloppiness but only in the context of uninformative priors (12, 15, 40). Considering it here in the Bayesian context with informative priors, we identify the need for the second approach that uses the dimension reduction idea from Cui *et al.* (35) to conveniently separate the effect of any prior information from that of the data. Using this novel adaptation of Bayesian techniques for dimensionality reduction to analyze model sloppiness, we illustrate how to identify the combinations of parameters driving model behavior in applications beyond systems biology and in a manner that acknowledges separately the available information [e.g., via expert knowledge (33)].

We focus our attention on the fit of three deterministic models to data with closed-form likelihood functions, although models that do not satisfy this criterion may also be analyzed using some of the methods presented here (further details in Discussion). Thus, we first highlight the advantages of our approach using the well-known Michaelis-Menten model of enzyme kinetics (41). We then apply it to a well-studied ecosystem network from Australia (a relatively data-poor system) (42) and a model for the action potential (AP) of heart cells (characterized by complex dynamical behavior) (43). In these latter two applications, different aspects of the interaction

between model and data are revealed by the analysis of sloppiness that are otherwise hidden by the individual techniques we bring together here. Last, we illustrate how stiff eigenparameters, once identified, can be used to design future experiments to improve parameter inference from collective model-data fittings and identify controlling mechanisms underlying the systems being modeled.

RESULTS

Our comprehensive analysis of sloppiness identifies the sensitivities of the model-data fit to changes in parameter values either in the region local to a point of interest in parameter space (standard approach; see Materials and Methods) or in the global sense, across all plausible parameter values consistent with available information (Bayesian approach; see Materials and Methods). Here, our results illustrate the benefits of using both standard and Bayesian approaches together to identify critical parameter combinations (stiff eigenparameters) that readily acknowledge the source of information (i.e., prior and/or data). To do so, we first analyze sloppiness in a biochemical model with three parameters (motivating example), known to suffer from poor parameter identifiability even when an excellent amount and quality of data are used to estimate model parameters (27, 44, 45). Then, we analyze sloppiness in an ecological four-species dynamic model with 20 parameters (case study 1), representing a typical dilemma in ecology of having too many parameters to be practically estimated well using noisy time-series data (2, 17, 44). Last, we analyze sloppiness in a cardiac electrophysiology model with nine parameters (case study 2), representing complex systems with strong nonlinear dynamics (7, 8).

Motivating example: The Michaelis-Menten kinetics Critical parameter combinations are readily identified by the analysis of sloppiness

The ubiquitous Michaelis-Menten model of biochemistry (41) is a perfect example to demonstrate the benefits of both understanding parameter dependence through the lens of model sloppiness and bringing a Bayesian approach to the topic (step i; see Materials and Methods). This model describes the dependence of an enzyme-catalyzed reaction rate v on substrate concentration $[S]$ as (46)

$$v = \frac{k_{\text{cat}} [E_T] [S]}{K_M + [S]} = \frac{k_{\text{cat}} [E_T]}{1 + K_M/[S]} \quad (1)$$

where parameters k_{cat} and $[E_T]$ together dictate the maximum rate of reaction (v_{max}), while K_M controls the substrate concentrations at which saturation effects become significant (45).

From the right-hand side of Eq. 1, it is already clear that there are two rate-limiting regimes, one in which the reaction rate simplifies to zero-order kinetics with respect to substrate at high $[S]$, and the other one in which the reaction rate simplifies to first-order kinetics at low $[S]$ (41, 46). To illustrate our methods, we thus consider two noisy synthetic datasets (step ii; see Materials and Methods) representing these two well-known rate-limiting regimes: The first dataset (A) consists of five measurements obtained beyond the saturation point, while the second dataset (B) consists of five measurements obtained before saturation has any apparent impact on the model behavior (Fig. 1). Both datasets fail to describe the full behavior represented by Eq. 1 and thus suitably highlight the well-known parameter identifiability issues in this model (27, 45).

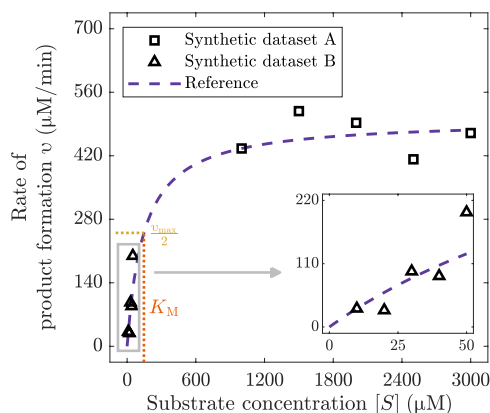


Fig. 1. Synthetic data generated using a measurement error of $\varepsilon = 25\%$ versus noiseless model prediction based on Eq. 1 using reference parameter values $k_{\text{cat}} = 100 \text{ min}^{-1}$, $[E_T] = 5 \text{ } \mu\text{M}$, and $K_M = 146.7 \text{ } \mu\text{M}$. Dataset A is obtained at a relatively high $[S]$, and dataset B is obtained at a relatively low $[S]$. Neither dataset can inform the full rate of reaction v across the range of $[S]$ (see also Figures' Supplementary Legends).

In dataset A, measurements only inform the reaction rate at saturation, $v \approx k_{\text{cat}}[E_T]$, and so, nothing can be learned about parameter K_M . While this tendency could also be identified by traditional sensitivity analysis (18) or by inspecting the posterior variance for this parameter obtained from Bayesian inference (45, 47), approaches for model sloppiness go a step further. By identifying key directions in the space of the log parameters, as encoded by the eigenvectors and eigenvalues of a sensitivity matrix, model sloppiness identifies that dataset A only informs the product of the remaining two parameters in Eq. 1, $k_{\text{cat}}[E_T]$. Regardless of whether a traditional definition (matrices **H** or **L**; see Materials and Methods) or any of the Bayesian definitions (matrices **P** and **G**; see Materials and Methods) of the sensitivity matrix is taken, a single eigenvalue dominates, with parameter combination denoted $\hat{\theta}_1 = k_{\text{cat}}[E_T]$ being the corresponding eigenparameter (Table 1, scenario 1). This is not, however, visible in the parameter marginals when Bayesian inference is used to fit the model to data, even in this simple problem (fig. S1).

Analogously, model sloppiness successfully identifies the parameter combination governing the rate of reaction in the nonsaturating regime (Fig. 1, dataset B). Given that this dataset is taken at low substrate concentration ($[S] \ll K_M$), Eq. 1 reduces to a linear dependence $v \approx (k_{\text{cat}}[E_T]/K_M)[S]$, and coefficient $k_{\text{cat}}[E_T]/K_M$ is the dominant eigenparameter (Table 1, scenario 3), which uncovers the nature of the poor parameter identifiability in this model. However, in this scenario and in the second scenario for dataset A (Table 1), we choose informative priors that cause the Bayesian approaches to model sloppiness (matrices **P** and **G**) to lead to different dominant eigenparameters. We explore the information provided by these approaches that take into account both prior and data to inform model parameters in the following section.

A Bayesian perspective reveals whether stiff parameter combinations are informed by the data or are influenced by the prior

Often, values for model parameters are meaningfully constrained by known feasible ranges or by expert information (32–34), which can potentially change both the most plausible set of values for the parameters and the nature of the new information provided by the

data. To demonstrate how the Bayesian approach to analyzing model sloppiness addresses this, we consider different scenarios where the reaction rate data (Fig. 1) are now coupled with prior information, and thus highlight how the stiff eigenparameters obtained using our two definitions of the sensitivity matrix (matrices **P** and **G**) together reveal whether parameter values are informed by the data or are influenced by the prior. We first fit Eq. 1 to dataset A (steps iii and iv; see Materials and Methods), considering a multivariate log-normal distribution for the model parameters that sets the value of one parameter (K_M) far away from its reference value (fig. S2). As a result, the posterior correctly concentrates around the reference parameter values used to generate the data (Fig. 2A, first and second panels), except for the poorly specified parameter (K_M) for which the prior renders it unable to (Fig. 2A, third panel). Here, prior and posterior distributions for parameter K_M are approximately equivalent (overlapping), thus reflecting that the data collected at saturation are uninformative to this parameter value. However, by examining the curvature of the posterior via its inverse covariance matrix **P** (steps v and vi; see Materials and Methods), this parameter emerges as the stiffest eigenparameter (Table 1, scenario 2). Thus, as prior and posterior distributions for parameter K_M are overlapped (Fig. 2A, third panel), this method reveals that the information already contained in the prior is dominating that provided by the data.

To learn the data informativity on model parameters while simultaneously acknowledging any prior information, we use the likelihood-informed subspace (LIS) method. This approach works by transforming the effects of the prior on the curvature of parameter space (35, 48), leaving only the effects of the data via the likelihood (further details in Materials and Methods). By doing so, the LIS method produces a sensitivity matrix (**G**) that identifies the region in parameter space where the informativity of the data prevails over that of the prior information (35, 48). For example, by imposing an informative prior for parameter K_M in this scenario, the method (matrix **G**) recognizes that no additional information is gained about this parameter from dataset A through the model-data fitting process, and so, it returns the same dominant eigenparameter $\hat{\theta}_1 = k_{\text{cat}}[E_T]$ (Fig. 2A, fourth panel) as the methods (matrix **H** or **L**) that ignore the prior altogether (Table 1, scenario 2). A natural question is then what does the LIS method provide that is not already given by a standard analysis of sloppiness? The key benefit is that if prior information does change the most plausible (prior-informed) region of parameter space, and the model behaves differently in this region, the LIS method will identify the directions in parameter space where the data are most “informative” relative to the prior, as we discuss next.

In scenario 3 (Table 1), we fit Eq. 1 to dataset B (steps iii and iv; see Materials and Methods) considering a combination of uniform and log-normal prior distributions that strongly specify values of parameters $[E_T]$ and K_M well and badly (fig. S3), respectively. Given that dataset B only constrains the value of combination of parameters $k_{\text{cat}}[E_T]/K_M$ (Fig. 2B, fourth panel), the extreme values of the parameters selected by unconstrained maximum likelihood estimation (MLE) highlight the importance of specifying plausible ranges for parameters via a Bayesian prior (Fig. 2B, MLE in the second and third panels). As for Bayesian inference, the posterior distribution simply fixes the value of the parameter k_{cat} (Fig. 2B, first panel) to a value that constrains well eigenparameter $k_{\text{cat}}[E_T]/K_M$ (Fig. 2B, fourth panel). Similar to scenario 2, model sloppiness, as implied by the posterior covariance method (matrix **P**), selects one of the parameters

Table 1. Comparison of the stiffest eigenparameter $\hat{\theta}_1$ (associated with the largest eigenvalue λ_1) for three different chosen parameter priors to fit the Michaelis-Menten model (Eq. 1) to data (Fig. 1). Each $\hat{\theta}_1$ is identified via Eq. 8 after obtaining eigenvalues (fig. S4) and eigenvectors of sensitivity matrices **H (or **L**), **P**, and **G**. Sensitivity matrices return different stiffest eigenparameters $\hat{\theta}_1$ with change of the prior distributions and dataset used to fit the model.**

Synthetic data	Scenario	Prior distribution	Stiffest eigenparameter, $\hat{\theta}_1$		
			H (or L)	P	G
Dataset A	1	Uniform		$k_{\text{cat}}[E_T]$	$k_{\text{cat}}[E_T]$
	2	Multivariate log-normal	$k_{\text{cat}}[E_T]$	K_M	
Dataset B	3	Uniform and log-normal	$k_{\text{cat}}[E_T]/K_M$	K_M	k_{cat}

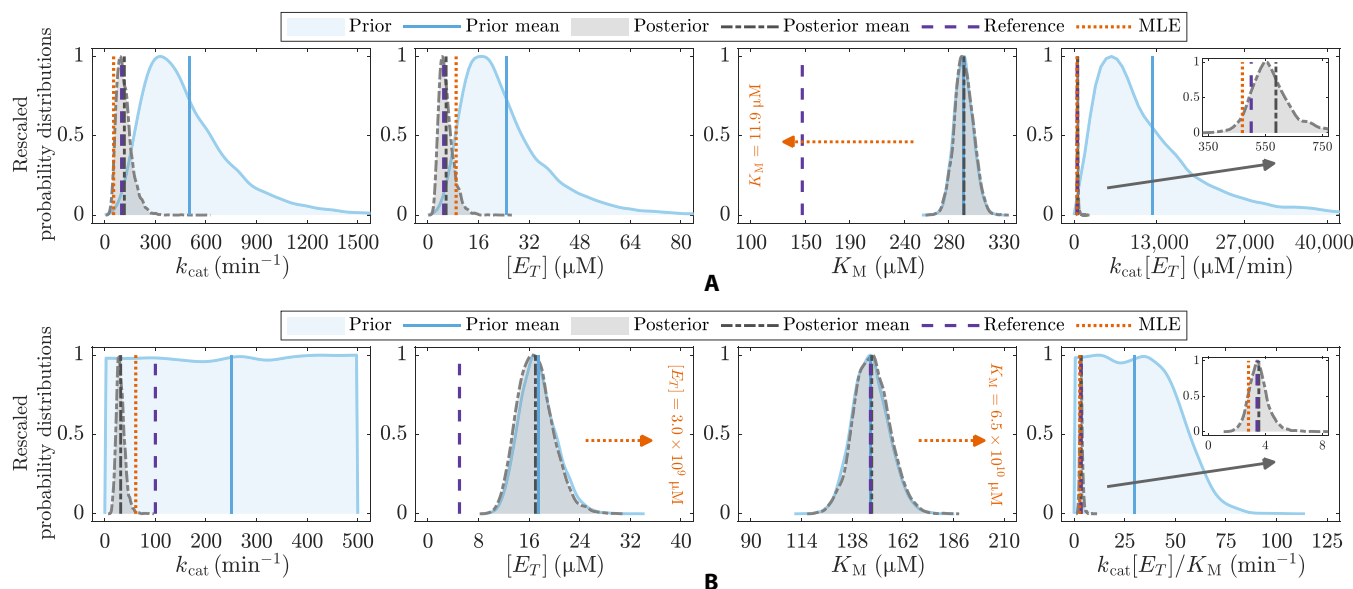


Fig. 2. Prior and posterior distributions for the model parameters together with the stiffest eigenparameters from scenarios 2 and 3 (Table 1) compared to their associated sets of best-fit values (MLE) and reference values (see also insets). (A) Scenario 2 parameters k_{cat} and $[E_T]$ together with stiffest eigenparameters $\hat{\theta}_1 = K_M$ (matrix **P) and $\hat{\theta}_1 = k_{\text{cat}}[E_T]$ (matrices **H** or **L** and **G**). (B) Scenario 3 parameter $[E_T]$ together with stiffest eigenparameters $\hat{\theta}_1 = k_{\text{cat}}$ (matrix **G**), $\hat{\theta}_1 = K_M$ (matrix **P**), and $\hat{\theta}_1 = k_{\text{cat}}[E_T]/K_M$ (matrices **H** or **L**). Parameter combinations $k_{\text{cat}}[E_T]$ and $k_{\text{cat}}[E_T]/K_M$ are well constrained by the data in scenarios 2 and 3, respectively. Parameter K_M is well constrained by the prior (posterior and prior overlapping) in both scenarios, and parameter k_{cat} is well constrained by the data relative to the prior in scenario 3. The best-fit values for parameter K_M lie far away from its reference values in both scenarios. As parameter combination $k_{\text{cat}}[E_T]/K_M$ is well constrained by the data in (B), the posterior distribution for parameter k_{cat} is left-shifted from the reference to compensate for parameter $[E_T]$ that is right-shifted from the reference in fig. S3A (see also Figures' Supplementary Legends).**

strongly specified by the prior, K_M (Fig. 2B, third panel), as the stiffest eigenparameter (Table 1). In this scenario, the LIS method (matrix **G**) instead identifies that dataset B acts only to fix the value of parameter k_{cat} and selects it as the dominant eigenparameter. That is, in contrast to the standard analysis of sloppiness only considering the likelihood surface, the LIS method uncovers new information provided by the data when there is prior parameter knowledge. Thus, the Bayesian methods together clarify whether the model parameters (or eigenparameters) are informed by the data or are significantly influenced by the prior beliefs.

Case study 1: Ecosystem network

A global perspective to analyzing sloppiness reveals true informativity of the data

Unlike the simple motivating example considering two well-known rate-limiting regimes that readily unveiled the controlling

eigenparameters (Fig. 2, fourth panels), with much larger models, inferring the parameter combinations that are more or less sensitive to the model-data fit can be difficult from a simple model inspection. To illustrate this, as a more complex case study from ecology, we use a well-known four-species ecosystem network model (42) that includes two threat species (foxes and rabbits), one threatened species (native mammals), and a basal species (pasture), as depicted in Fig. 3A. This ecosystem model consists of four discrete-time equations (based on ordinary differential equations) and eight constitutive equations (table S1) whose 20-parameter point estimates (table S2) were inferred from several studies at two semi-arid locations in Australia (42). Here, we thus seek to illustrate key benefits of the Bayesian analysis of sloppiness for data-poor systems, characterized by low quality and amount of observed data because of practical limitations (1, 2, 12, 17). To do so, we first fit the ecosystem network model (table S1) to noisy synthetic time-series data using

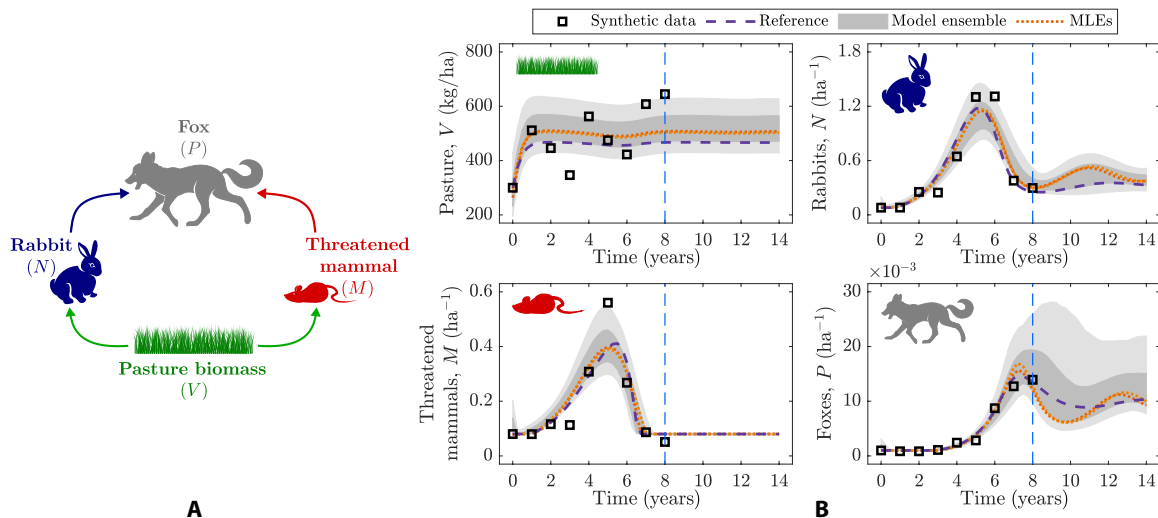


Fig. 3. Ecosystem interaction network, consisting of pasture (V), rabbits (N), foxes (P), and threatened species (M). (A) Ecosystem network in which arrows indicate the direction of the energy transfer associated with the species interaction. (B) Synthetic time-series data for ecological abundance with measurement error of $\epsilon = 25\%$ together with noiseless model prediction using reference parameter values (table S2), model predictions using two sets of best-fit parameter values (MLEs), and model ensemble predictions using all plausible parameter values (fig. S5). The ecosystem network model fits the synthetic time-series data, with the model ensemble propagating parameter uncertainty into species abundance predictions (see also Figures' Supplementary Legends).

both MLE and Bayesian inference (steps i to iv; see Materials and Methods), considering a multivariate log-normal prior distribution for the model parameters (fig. S5).

After fitting the model to data, model predictions (Fig. 3B) based on a model ensemble (shaded regions), considering all plausible parameter values (fig. S5), enclose both the simulated noisy data ($\square$ symbols) and true ecosystem dynamic behavior (dashed profiles). They also enclose predictions based on two sets of best-fit parameter values (dotted profiles) obtained from starting the MLE algorithm at two different initial parameter values (step iii; see Materials and Methods). Furthermore, parameter marginals (fig. S5) enclose these two separate point estimates and also show that most of the model parameters are poorly constrained by the data.

In addition to quantifying parameter uncertainty, a global perspective to the problem of fitting models to data can benefit the inference of critical parameter combinations that control the quality of the model-data fit. For example, while local changes in the topography of the surface described by the likelihood function in the vicinity of the two sets of best-fit parameter values (fig. S5) mislead inference of stiff eigenparameters through the standard analysis of sloppiness (cf. $\hat{\theta}_i, i = 1, 2, 3$ in Table 2 from matrices $\mathbf{H}$ or $\mathbf{L}$, evaluated at the different sets of best-fit values θ_1^* and θ_2^*), the Bayesian methods (matrices $\mathbf{P}$ and $\mathbf{G}$) fully characterize the structure of this surface by considering all plausible parameter values (steps v and vi; see Materials and Methods), informed by the combination of both data and prior beliefs. In this way, differences between dominant eigenparameters from Bayesian sensitivity matrices $\mathbf{P}$ and $\mathbf{G}$ (Table 2) also demonstrate that the prior is influencing the most plausible region of parameter space, which thus implies that the surface described by the posterior distribution (Eq. 9) and the likelihood function (Eq. 15) are different locally and globally.

Analysis of sloppiness brings new insights to Bayesian parameter inference

Combining model sloppiness together with Bayesian inference reveals critical parameter combinations that can be otherwise lost

when only considering the uncertainty in individual model parameters through Bayesian inference. After the fit of the ecological model to data, for example, parameter marginals (fig. S5) illustrate that only a few of the model parameters (R , V_0 , D_{III} , and ϵ) are well constrained by the data, which suggests that these parameters have a strong influence on the quality of model-data fit. Instead, the Bayesian analysis of sloppiness (matrices $\mathbf{P}$ and $\mathbf{G}$) identifies that it is combinations of parameters c_N , a_N , c_M , a_M , c_P , and a_P that are the most constrained by the available data (Fig. 4). The prior distribution appears to be weakly informing the three stiffest eigenparameters $\hat{\theta}_1$, $\hat{\theta}_2$, and $\hat{\theta}_3$ (Table 2) since the first eigenparameter $\hat{\theta}_1$ from the posterior covariance method (matrix $\mathbf{P}$) also corresponds to the third eigenparameter $\hat{\theta}_3$ from the LIS method (matrix $\mathbf{G}$), while the quotient ($\hat{\theta}_3/\hat{\theta}_2$) and product ($\hat{\theta}_2\hat{\theta}_3$) of the second and third eigenparameters $\hat{\theta}_2$ and $\hat{\theta}_3$ from the posterior covariance method (matrix $\mathbf{P}$) approximate the first and second eigenparameters $\hat{\theta}_1$ and $\hat{\theta}_2$ from the LIS method (matrix $\mathbf{G}$), respectively.

For this system, the identified stiff eigenparameters (Fig. 4) do not appear together in single terms within the model (table S1). However, parameter ratios c_X/a_X (or a_X/c_X) with $X = N, M, P$ arise as part of the dominant eigenparameters (Table 2) as they appear in separate terms with opposite sign in this model (table S1). As a result, there is a compensation effect between values of parameters a_X and c_X that has two key implications for the model predictions. First, the model-data fit is highly informative for characterizing growth dynamics (r_N , r_M , r_P) of rabbits (a_N/c_N), threatened mammals (a_M/c_M), and foxes (a_P/c_P), which is likely to significantly affect animal species abundances (N , M , and P). Second, analysis of sloppiness reveals that by measuring either the maximum rate of decrease (a_M) or increase (c_M) of the threatened mammal density (also applies for rabbits and foxes), collective model-data fit will inform values of the other parameter to a similar extent, as we discuss in the next section.

Bayesian analysis of sloppiness readily informs future experimental design

Bayesian analysis of sloppiness unveils hidden parameter interdependencies that can help design future experiments for improved

Table 2. Comparison of the stiff eigenparameters $\hat{\theta}_1$, $\hat{\theta}_2$, and $\hat{\theta}_3$ (associated with the largest eigenvalues λ_1 , λ_2 , and λ_3) considering a multivariate log-normal prior distribution for the parameters to fit the ecosystem model (table S1) to data (Fig. 3B). Each $\hat{\theta}_1$, $\hat{\theta}_2$, and $\hat{\theta}_3$ is identified via Eq. 8 after obtaining eigenvalues (fig. S6) and eigenvectors from matrices **H (or **L**), **P**, and **G**. Stiff eigenparameters from matrix **H** (or **L**) are obtained at two sets of best-fit parameter values θ_1^* and θ_2^* (fig. S5). Matrix **H** (or **L**) returns different stiff eigenparameters when evaluated at two distinct sets of best-fit parameter values, while matrices **P** and **G** return different stiff eigenparameters because the prior influences the model-data fit.**

Eigenparameter $\hat{\theta}_i$	Sensitivity matrices			
	H or L Evaluated at		P	G
	θ_1^*	θ_2^*		
1	$(c_N/a_N)(a_M/c_M)^{0.4}$	$(c_M/a_M)(a_N/c_N)^{0.9}$	$a_p^{0.9}/c_p$	$(c_M/a_M)(a_N/c_N)^{0.9}$
2	$(c_M/a_M)(c_N/a_N)^{0.4}$	$(c_N/a_N)(c_M/a_M)^{0.9}$	c_N/a_N	$(c_N/a_N)(c_M/a_M)^{0.9}$
3	$c_p/(a_p^{0.9} d_N^{0.4})$	$c_p/a_p^{0.9}$	c_M/a_M	$c_p/a_p^{0.9}$

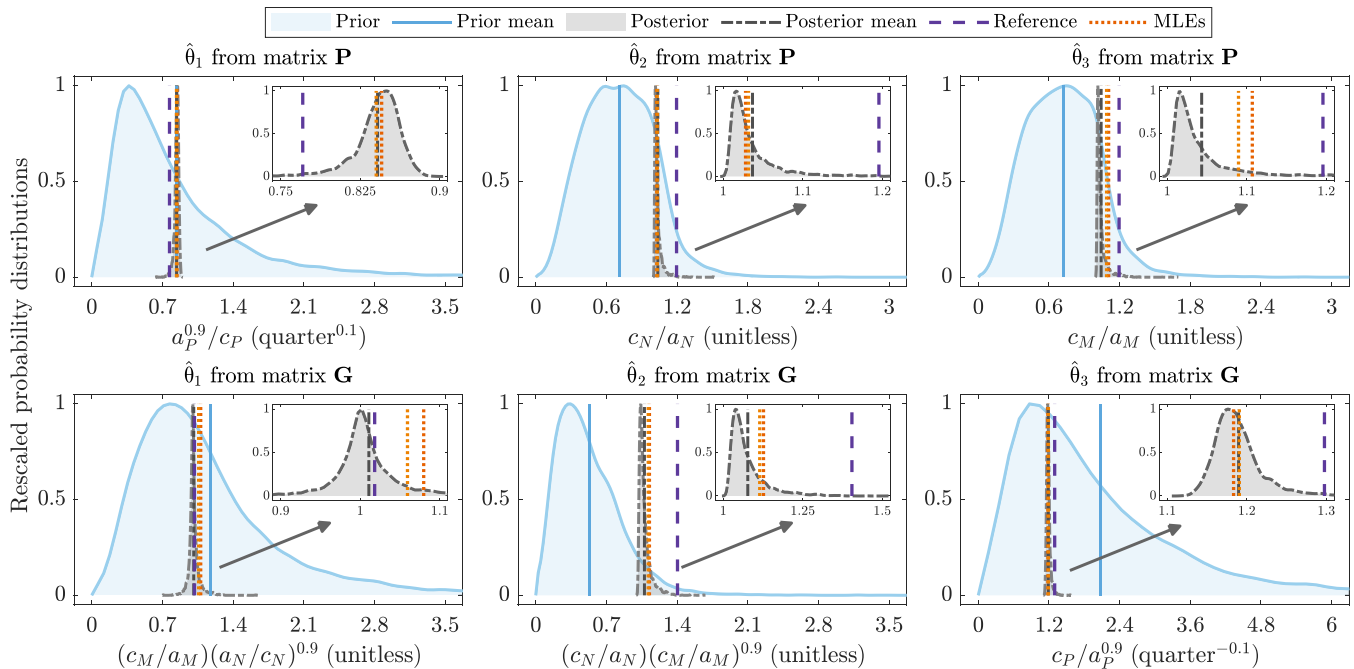


Fig. 4. Prior and posterior distributions for eigenparameters $\hat{\theta}_1$, $\hat{\theta}_2$, and $\hat{\theta}_3$ (Table 2) compared to their reference values and associated sets of best-fit values (MLEs) (see also insets). Posterior distributions for all eigenparameters and their associated MLEs closely match the true values (see also Figures' Supplementary Legends).

parameter inference. For example, given that the posterior covariance method (matrix **P**) reveals that the ratio of parameters $a_p^{0.9}/c_p$ is the stiffest eigenparameter (Table 2), this ratio also indicates that parameters a_p and c_p are approximately linearly related, $a_p^{0.9} \propto c_p$ (Fig. 5A). Here, an analogous tendency is seen for the stiffest eigenparameter from the LIS method (matrix **G**), $(c_M/a_M)(a_N/c_N)^{0.9}$ (Table 2), with parameter c_M and combination of parameters $(1/a_M)(a_N/c_N)^{0.9}$ being instead inversely related, $c_M \propto [(1/a_M)(a_N/c_N)^{0.9}]^{-1}$ (Fig. 5B). In addition, many samples from the posterior distributions are seen to lead to the same value of the log-likelihood function (no apparent color change across posterior distribution samples in Fig. 5), with the two sets of best-fit parameter values ($\times$ symbols) and reference (true) values ($+$ symbols) falling within the corresponding

posterior distribution sample. This tendency indicates that every value for the model parameter (or combination of parameters) on one side of the relation (e.g., $a_p^{0.9}$ and c_M) has a corresponding constrained estimate for the parameter (or combination of parameters) on the other side of the relation [e.g., c_p and $(1/a_M)(a_N/c_N)^{0.9}$]. Similar tendencies are seen for the remaining eigenparameters (fig. S7).

In addition to identifying compensation effects between subsets of parameters (Fig. 5 and fig. S7) that lead to similar model outputs (Fig. 3B), analysis of sloppiness also reveals that prioritizing improvement of the estimates of any of the parameters (or parameter combinations) on one side of the proportionality relationship will immediately improve estimation of parameters (or parameter combinations) on the other side. As an example of this, we considered a

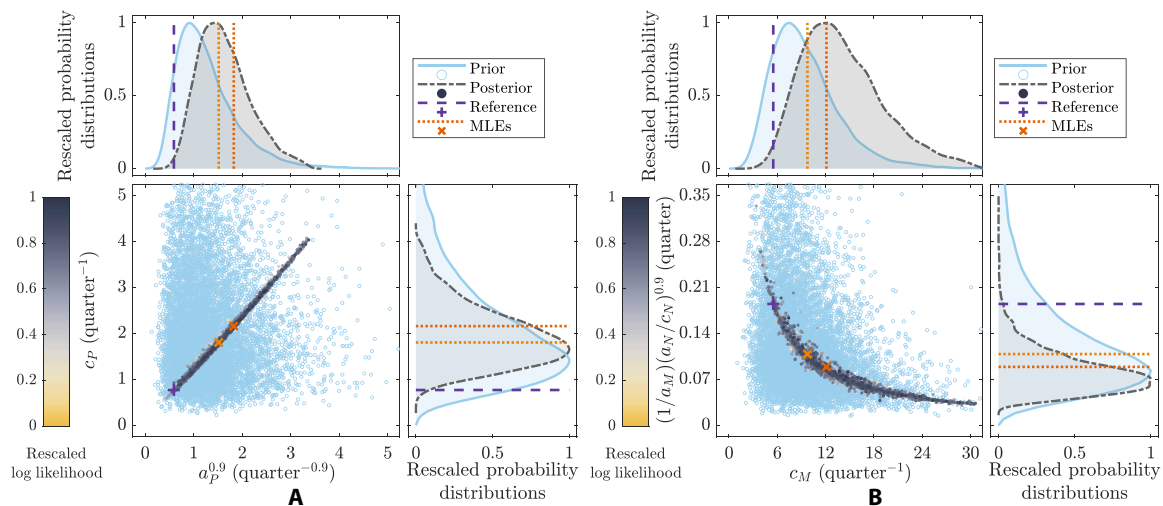


Fig. 5. Bivariate scatterplots of the prior and posterior distributions (shaded regions) for the first stiffest eigenparameter obtained from the Bayesian methods (Table 2) compared to their reference parameter values and sets of best-fit values (MLEs). (A) $\hat{\theta}_1$ and $\hat{\theta}_3$ from matrices **P and **G**, respectively, and (B) $\hat{\theta}_1$ from matrix **G**. Many samples of the posterior distribution yield similar values of the log-likelihood function (see also Figures' Supplementary Legends).**

multivariate log-normal prior distribution (fig. S8), which is very informative for parameters a_N , a_M , and a_P to fit the ecosystem network model to data. These prior conditions act as improved estimates for parameters a_N , a_M , and a_P , obtained from either expert elicitation or parameter-specific experiments [e.g., spotlight counts (42)]. After the model-data fit (fig. S9), parameters on the other side of the relations (c_N , c_M , and c_P) are also found to be more narrowly constrained. The percentage coefficients of variation (CV) for the posterior distributions of parameters c_N , c_M , and c_P range between 7 and 8% when a more informative prior distribution is specified for parameters a_N , a_M , and a_P , which are much lower than those obtained (ranging between 30 and 50%) when a vague multivariate log-normal prior distribution is instead specified (fig. S5). Thus, combining Bayesian inference together with the analysis of sloppiness reveals parameter interdependencies that can be strategically exploited to efficiently improve individual parameter inference using less additional data than might be otherwise expected.

Case study 2: Cardiac electrophysiology

Key controlling mechanisms for complex systems are uncovered by analysis of sloppiness

The previous section considered an ecological model as an example of a system characterized by a moderately large number of parameters and poor access to data. A separate class of systems is that for which data are more readily available, but the dynamics that produce the data manifest in complex sensitivities to their controlling parameters. For these systems, the challenge is often how to summarize these nonlinear dynamics in a meaningful, actionable way, and so, stiff eigenparameters identified by analyzing model sloppiness have a recognizable potential. However, so far, model sloppiness has primarily been considered for models characterized by large numbers of fundamental interactions, such as the Michaelis-Menten kinetics that describe the cell signaling network analyzed in the foundational work of Brown *et al.* (12, 13). Here, we seek to demonstrate the usefulness and purpose of stiff eigenparameters in systems where the constituent dynamics themselves, and not only their interactions, are complex and unwieldy.

As an example of such a system, we consider the Beeler-Reuter (BR) model (43), which describes the AP of a cardiac ventricular myocyte, the pattern of highly regulated ion flow that creates the depolarization and subsequent repolarization governing the heart-beat. This cardiac cell model consists of eight nonlinear ordinary differential equations, six constitutive equations (table S3), and nine parameters (table S4). Although an older model, the BR model captures many of the most important electrophysiological features of the ventricular AP (49), and interest remains regarding its sensitivity to changes in its parameters (8, 25). Cardiac AP models are critical for mechanistically understanding arrhythmia (50), and the issue of parameter variability is fundamental to understanding the differential effects of antiarrhythmic treatments within a population (51) or the cardiotoxicity of other pharmacological agents (52).

The AP is summarized by the time course of a cell's transmembrane potential in response to stimulation and can be recorded by an electronic measurement device at good temporal resolution and without much noise (e.g., synthetic data in Fig. 6A). The complexity in these models rests with the way multiple ion channels—each with its own set of time-adaptive, nonlinear voltage-gated dynamics—combine additively to determine the total ion flow that produces the AP (table S3). The most commonly varied parameters are the relative levels of expression for these different ion channels (53), and so, rather than describing fundamental quantities such as rates of production or destruction, model parameters in this context describe the extent to which a variety of complex and strongly nonlinear dynamics contribute to the system behavior.

For this system, Bayesian model-data fitting (steps iii and iv; see Materials and Methods) produces an ensemble of plausible values for model parameters that recapture the data extremely well (Fig. 6A). Most of the individual parameters are well constrained by the AP data (as seen from their marginal distributions; fig. S10), although none emerges as substantially more important than all others. Analyzing model sloppiness to consider parameters in terms of their combinations (steps v and vi; see Materials and Methods), however, reveals that the combination of parameters $A_{K_1}^{0.9} A_{K_1}^{0.4}/g_s$ is the primary driver of the AP dynamics. This eigenparameter's corresponding

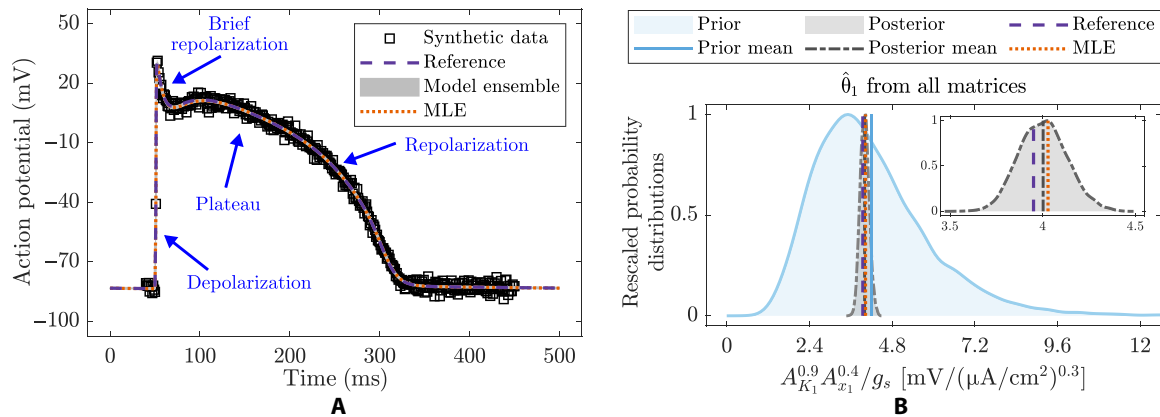


Fig. 6. BR model fit to synthetic time-series data together with the identified stiffest eigenparameter. (A) Synthetic AP data with measurement error of $\sigma = 2$ mV and time resolution of 1 ms (1 kHz) together with noiseless model prediction using reference parameter values (table S4), model predictions using the set of best-fit parameter values (MLE), and model ensemble predictions using all plausible parameter values. (B) Prior and posterior distributions for the stiffest eigenparameter $\hat{\theta}_1 = A_{K_1}^{0.9} A_{x_1}^{0.4} / g_s$ obtained from all matrices compared to their associated set of best-fit values and reference values (see also inset). Both MLE-based and Bayesian inference-based model predictions overlap the true AP dynamics and synthetic data. All methods lead to the same stiffest eigenparameter. The posterior distribution for the stiffest eigenparameter and their associated best-fit value (MLE) closely match the true eigenparameter value (see also Figures' Supplementary Legends).

eigenvalue eclipses the value of the others (fig. S11), and accordingly, its value is extremely well specified by the population of plausible parameter values (Fig. 6B). This eigenparameter and its relative importance are identified by both the standard and Bayesian approaches for model sloppiness, owing to the use of a relatively uninformative prior and the fact that the data are highly informative about the model parameters. Unlike the Michaelis-Menten kinetics or the ecosystem network model examples, here, all approaches for model sloppiness are similarly suitable because of the strong informativeness of the data relative to the prior.

The key eigenparameter has a natural electrophysiological interpretation. Parameters A_{K_1} and g_s describe the relative strengths of the primary outward and inward (i.e., counteracting) currents active during the plateau and repolarization phases that compose the bulk of the AP (Fig. 6A), and so, they appear in the eigenparameter as a ratio. Here, the third parameter A_{x_1} contributes to the eigenparameter to a lesser extent and appears as a product with parameter A_{K_1} , owing to their shared role in describing strengths of the outward potassium currents that drive repolarization. The three currents I_{K_1} , I_{x_1} , and I_s (table S3), associated with these three model parameters (A_{K_1} , A_{x_1} , and g_s , respectively), exhibit nonlinear dynamics (fig. S12). Thus, it is unexpected how well the primary actions of these three currents (I_{K_1} , I_{x_1} , and I_s) can be summarized by a simple product of parameters with exponents ($A_{K_1}^{0.9} A_{x_1}^{0.4} g_s^{-1}$), whose value strongly dictates whether the model output reproduces the data (Fig. 7).

Analysis of model sloppiness naturally uncovers this result, by revealing the precise way in which the three currents I_{K_1} , I_{x_1} , and I_s (table S3) act together and thus highlighting the importance of their balance by assigning a much higher eigenvalue to their eigenparameter than any other. Without considering the curvature of the log parameters, however, this relationship is not easily observed. The precise relationship between A_{K_1} , A_{x_1} , and g_s remains hidden from view in standard Bayesian bivariate analysis (8, 25), and even when directly plotting the values of posterior samples for these three parameters against one another (fig. S13). Such a relationship is also not obvious from the model definition, where none of the three parameters A_{K_1} , A_{x_1} , and g_s appear as products or quotients with one

another, nor do the coefficients of their addition correspond to the exponents found in the governing eigenparameter.

Analysis of sloppiness uncovers knowledge limitations in mathematical models fitted to data

As also observed in the ecological application (Fig. 5), the existence of a strong eigenparameter(s) introduces a pronounced structure to the space of plausible parameter sets (Fig. 8). The nature of the eigenparameter implies a strong linear interdependency between combination of parameters $A_{K_1}^{0.9} A_{x_1}^{0.4}$ and parameter g_s , as seen in the posterior samples found by the Bayesian inference (Fig. 8). Identifying these critical structures introduced by the model-data fitting process is key to understanding the information provided by the data on the model parameters. Cardiac electrophysiology is a particularly important example as parameter identifiability is a well-established issue for AP models (5, 8). Thus, owing to sloppiness in parameter estimation such as that found and quantified here, even perfect AP data (Fig. 6A) can imply multiple different parameter estimations (Fig. 7B and fig. S14), with consequences that then emerge under pathological conditions or in response to drug treatments (54).

As in many other disciplines, in cardiac electrophysiology, it can be difficult to design further experiments and/or to target experiments to learn specific parameter values. To this end, our comprehensive approach to model sloppiness does uncover the deficiencies in the available information through the identification of the critical eigenparameters. For example, once these critical eigenparameters are identified, the model can be used to simulate scenarios considering extreme system conditions (fig. S14) that are theoretically still plausible given current data (Fig. 7B). This concept might be even more applicable where a model's computational runtime limits the feasibility of Bayesian inference. Even when a posterior of plausible parameter value sets cannot be realistically generated, standard analysis of sloppiness can still quickly identify directions in parameter space of poor information. In this way, simulations can be carried out along directions of poor knowledge (e.g., perpendicular to the linear relationship depicted in Fig. 8) to further justify the conclusions of simulation studies against the uncertainty that remains in the parameters after the model-data fit.

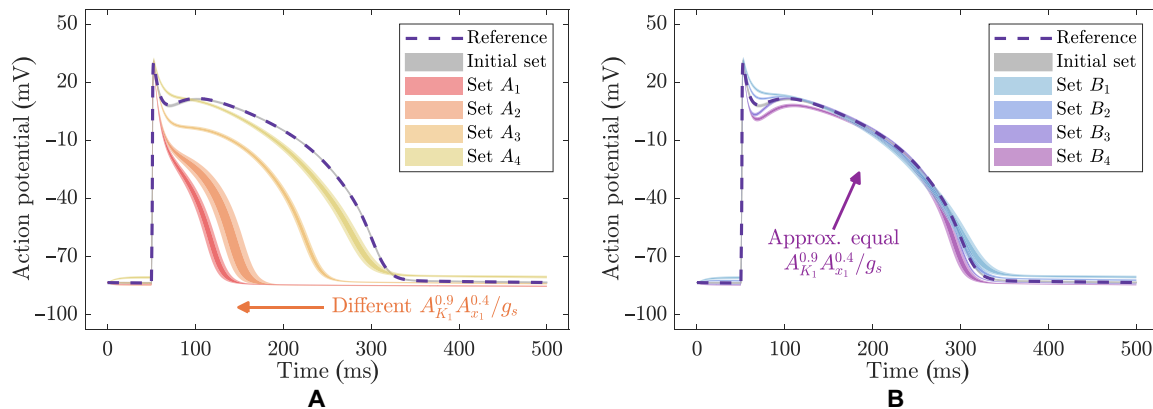


Fig. 7. Effect of varying parameter values that are part of the stiffest eigenparameter $\hat{\theta}_1 = A_{K_1}^{0.9} A_{x_1}^{0.4} / g_s$ on the BR model prediction. (A) Predictions obtained with four sets of parameter values that change the value of $\hat{\theta}_1$ (sets A_1 to A_4 ; fig. S14). (B) Predictions obtained with four sets of parameter values that keep approximately constant the value of $\hat{\theta}_1$ (sets B_1 to B_4 ; fig. S14). Similar predictions are obtained with change in the parameter values when the value of $\hat{\theta}_1$ is kept approximately constant, while in the opposite case predictions strongly differ from those obtained with the initial parameter set (fig. S10) (see also Figures' Supplementary Legends).

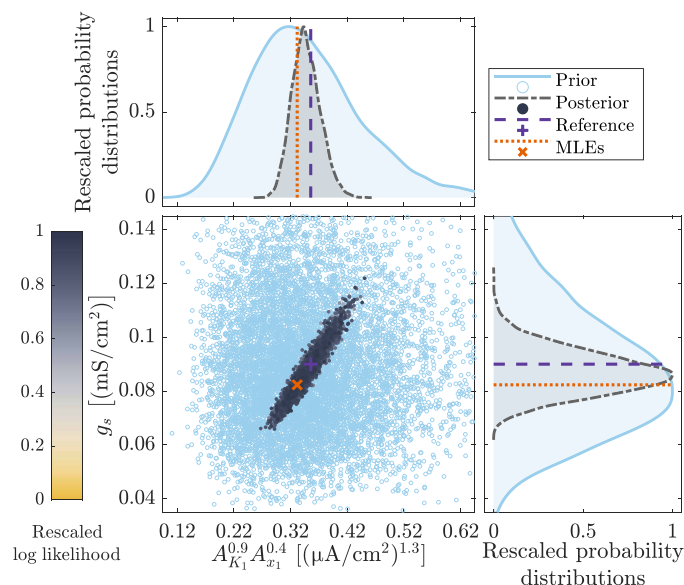


Fig. 8. Bivariate scatterplot for the prior and posterior distributions (shaded regions) for combination of parameters $A_{K_1}^{0.9} A_{x_1}^{0.4}$ and g_s compared to sets of reference parameter values and best-fit values (MLE). Many samples of the posterior distribution yield similar values of the log-likelihood function with $A_{K_1}^{0.9} A_{x_1}^{0.4} \propto g_s$ (see also Figures' Supplementary Legends).

Ever-present knowledge limitations about parameter values in cardiac electrophysiology have motivated studies in which virtual populations consisting of many models with varying parameter values are used to explore how populations as a whole, characterized by variable data, respond to different treatments or conditions (51, 52). This has included a Bayesian methodology for calibrating such populations (7). The Bayesian framework of model sloppiness, which provides a more global sense of parameter space curvature in the plausible region defined by the given data (Fig. 8), could be applied to the “posteriors” of such population-calibration processes, and thus provide a unique way to identify the combinations of parameters that are constrained (or not constrained) by the process of fitting these models to data exhibiting variability.

DISCUSSION

Recognizing the influence of prior information on the quality of model-data fit

The use of informative priors has been shown to help constrain model parameters when mathematical models are fitted to data in many Bayesian inference applications (7, 25, 32, 33, 55). Despite this advantage, the usage of uninformative priors has predominated in the context of analyzing model sloppiness (12, 13, 15). In such a context, vague uniform priors, spanning several orders of magnitude, have been used to prevent potential optimization algorithm failures (11, 12, 16), rather than reflecting their true purpose: accounting for preexisting knowledge about the parameter values (33, 34). Here, we introduced how to account for informative priors when analyzing model sloppiness, with our example results illustrating how this approach identifies the relative effect of informative priors on the quality of the model-data fit. Specifically, the LIS method (matrix **G**) was shown to reveal directions in parameter space where the posterior differs most strongly from the prior (Fig. 2B), while the posterior covariance method (matrix **P**) was shown to reveal directions in parameter space that are strongly informed by the combination of both data and priors (Figs. 2A and 4). In addition, the Bayesian analysis of sloppiness (matrices **P** and **G**) was shown to provide equivalent results to those based on earlier approaches (matrices **H** and **L**) when uninformative (vague) priors are used in the implementation of Bayesian inference (Table 1) and when the data are very informative for the model parameters (Fig. 6 and fig. S11). Consequently, we have demonstrated that the Bayesian approach to analyzing sloppiness complements earlier approaches (12, 13) in that the effects of prior beliefs on the quality of the model-data fit can be segregated when all methods are used together. This then clarifies which of the model parameters (or parameter combinations) are predominantly informed by the data or the prior.

In the motivating Michaelis-Menten kinetics example and the cardiac electrophysiological application, we specifically showed that if stiff eigenparameters obtained from all methods (matrices **H** or **L**, **P**, and **G**) are similar, priors are weakly informative, and so, stiff eigenparameters are largely constrained by the data (Table 1, scenario 1, and fig. S11). We also showed in the motivating example that if stiff eigenparameters obtained from the LIS method (matrix **G**)

are similar to those obtained from the standard method (matrices **H** or **L**) but different from those obtained from the posterior covariance method (matrix **P**), then critical parameter combinations associated with the posterior covariance method (matrix **P**) are significantly influenced by the priors (Table 1, scenario 2). Last, we showed in the same motivating example that if stiff eigenparameters obtained from the standard method (matrices **H** or **L**) differ from those obtained from the Bayesian methods (matrices **P** and **G**), then priors may also be influencing the quality of the model-data fit. Under such conditions, stiff eigenparameters obtained from the LIS method (matrix **G**) are informed by the data relative to the prior, and those from the posterior covariance (matrix **P**) are mostly constrained by the prior (Table 1, scenario 3) (a topographical interpretation of these findings is also provided in the next section). In this way, we also demonstrated that our methods are well suited not only for applications where there is little prior knowledge about the parameter values (12, 13, 15, 27) but also for those where prior beliefs can be confidently included as part of the model-data fitting process (5, 7, 32, 33, 55).

In our implementation of Bayesian inference, we specifically considered combinations of vague and informative uniform and/or log-normal prior distributions (figs. S1A to S3A, S5, S8, and S10), as these types of priors are traditionally used in ensemble modeling applications in biochemistry (11, 12, 40, 47), ecology (17, 34, 55), and biology (7, 8, 25). While implementation of Bayesian inference with application-specific prior distributions is beyond the scope of this work, the Bayesian methods (matrices **P** and **G**) to analyzing model sloppiness may not be limited to the types of priors discussed here. We note, however, that for applications using heavy-tailed and/or sparsity-promoting priors in the implementation of Bayesian inference, we anticipate that more and/or better quality data would be required to obtain critical parameter combinations. Under this condition, the posterior covariance method (matrix **P**) is expected to reveal data-informed stiff parameter combinations at least when the data are considerably more informative than the prior, such as in the Michaelis-Menten kinetics example (Table 1, scenario 1). However, for implementation of the LIS method (matrix **G**), careful estimation of the prior covariance matrix for the logarithms of parameters Ω (Eq. 12) would be required for the successful inference of stiff parameter combinations. Here, an interesting direction for future work would be to apply the prior normalization technique recently introduced by Cui *et al.* (56) in the context of Bayesian inverse problems to transform heavy-tailed priors into standard Gaussian distributions, to then implement the LIS method with these prior transformations in the context of analyzing model sloppiness.

Last, we note that the sensitivity matrices **P** and **G** (Eqs. 11 and 12, respectively) are logarithmically based on this work, since standard methods for analyzing model sloppiness have been usually applied with the doubly logarithmic Hessian (Eq. 5) owing to their history in analyzing complex systems describing physical processes (11, 15, 29). This is a transformation not typically used in the classical implementation of PCA (36, 38); however, it is not uncommon to the implementation of the LIS method for Bayesian dimensionality reduction applications (35, 48). While such a conveniently chosen transformation implies the use of single-sign prior distributions (i.e., either positive or negative for each parameter) in the implementation of Bayesian inference, it also conveys three key advantages for the Bayesian analysis of model sloppiness: (i) It reflects the positivity constraints on model parameters (true of the majority of parameters

characterizing physical systems), (ii) it prevents inconsistencies in scaling between parameters (due to different orders of magnitude) from affecting the analysis of sloppiness (11), and (iii) it provides a basis to identify stiff eigenparameters as products and/or quotients (combinations) of bare model parameters with different power indices whose magnitude informs the relative parameter importance in the combination (Eq. 8) (13). If, despite these advantages, prior distributions spanning negative to positive values are required for a given application, nonlogarithmic versions of sensitivity matrices **P** and **G** may be used as part of the analysis of model sloppiness. Under this condition, eigendecomposition on such matrices will instead reveal stiff eigenparameters as linear combinations (summations and subtractions) of bare model parameters premultiplied by different coefficients whose magnitude represents the relative parameter importance in the combination. However, large differences in the magnitude of model parameters may mask the true stiff eigenparameters from these nonlogarithmic-based sensitivity matrices, so great care must be taken to use them for the analysis of sloppiness.

Characterizing the topography of the surface described by the model-data fit

Our work significantly adds to the literature on sensitivity analysis (18, 21), which, in the context of models fitted to data, largely focuses on “locally” investigating the curvature of the surface described by the likelihood function (11, 12, 15, 16, 27), around the best-fit parameter values (MLE). Thus, a key contribution of the Bayesian approach to analyzing sloppiness is that it accounts for changes in the curvature of this surface “globally” upon considering potentially plausible model parameter sets at a finite distance away from the best-fit parameter values (7, 22, 25). In addition, our implementation of Bayesian inference advances upon earlier such implementations for analyzing model sloppiness (12, 13). In these earlier works, a type of Markov chain Monte Carlo (MCMC) algorithm, with uninformative priors and started at the set of best-fit parameter values (MLE), was used to characterize the likelihood surface in the vicinity of the preidentified MLE (11, 15). However, we have shown here that different MLEs can lead to completely different locations on the surface describing the likelihood function (Fig. 5 and fig. S7), which can mislead inference of stiff eigenparameters (Table 2). More so, for systems in which the topography of the model-data fit function is very rugged, local optima can misguide the optimization algorithm (13, 15, 16), for example, as illustrated by Fernández Slezak *et al.* (57) in fitting a model of avascular tumor growth to noisy data using different optimization algorithms. Hence, convergence issues become the bottleneck for the identification of stiff eigenparameters via standard approaches to analyzing sloppiness. Instead, Bayesian inference as implemented here does not rely upon a single set of best-fit parameter values to characterize the surface describing the quality of the model-data fit (see Materials and Methods). Rather, all posterior samples contribute to the description of the surface topology. This stochastic exploration of the posterior distribution can reduce the risk of convergence to a local optimum (24, 58), with the added value that the Bayesian sensitivity matrices also acknowledge any effect of prior beliefs on the most plausible region in parameter space (Figs. 2, 4, and 6B).

In our example results, we specifically illustrated that comparison of stiff eigenparameters obtained from both the standard (matrices **H** or **L**) and Bayesian (matrices **P** or **G**) methods can reveal whether the topography of the surface described by the likelihood

function is globally and locally similar as well as whether such a surface is similar to that described by the posterior distribution. For example, if stiff eigenparameters obtained from all methods (matrices $\mathbf{H}$ or $\mathbf{L}$, $\mathbf{P}$, and $\mathbf{G}$) are similar, the shape of the surface described by the likelihood function and posterior distribution is locally and globally similar (Table 1, scenario 1, and Fig. 6B). Instead, if stiff eigenparameters obtained from the standard methods (matrices $\mathbf{H}$ or $\mathbf{L}$) are similar to those obtained from the LIS method (matrix $\mathbf{G}$) but differ from those obtained from the posterior covariance method (matrix $\mathbf{P}$), the shape of the surface described by the likelihood function is locally and globally similar but different from that described by the posterior distribution (Table 1, scenario 2, and Table 2, with matrix $\mathbf{H}$ or $\mathbf{L}$ evaluated at θ_2^*). Alternatively, if stiff eigenparameters obtained from all methods (matrices $\mathbf{H}$ or $\mathbf{L}$, $\mathbf{P}$, and $\mathbf{G}$) are different, the shape of the surface described by the likelihood function is not only locally and globally different but also different from the surface described by the posterior distribution (Table 1, scenario 3, and Table 2, with matrix $\mathbf{H}$ or $\mathbf{L}$ evaluated at θ_1^*). We note, however, that while differences between the shape of surfaces describing the posterior distribution and likelihood function are associated with the effects of priors on the quality of the model-data fit, identifying whether the likelihood is locally and globally similar is crucial when multiple (but also very different) parameter sets lead to similar values of the likelihood function (Fig. 5). This is a situation that is likely to occur when there are limited data to inform model parameters (15, 59), for which the Bayesian sensitivity matrices are thoroughly informed by the data and the prior.

We also note that analysis of model sloppiness, including our new Bayesian perspective on the topic, describes the topography of the likelihood surface using the eigenvectors of the sensitivity matrix. Analogous to the use of PCA for dimension reduction (36), this can be viewed in a sense as a linearized description of the topography. However, extending beyond this linearized view would require techniques that produce eigenvectors expressed in terms of the original parameters [as opposed to, say, in a reproducing kernel Hilbert space (60)]. Rather, owing to the connections between the different sensitivity matrices and inverse covariance matrices, methods for improved covariance matrix estimation appear to be a promising direction for extending the way model sloppiness describes topography. For example, the graphical LASSO algorithm estimates sparse inverse covariance matrices that enforce conditional independence between some parameters. This might assist in the identification of stiff eigenparameters similar to how it can assist problems such as classification (61).

Improving parameter identifiability by designing experiments based on identified parameter interdependencies

Careful experimental design can improve ambiguous parameter inferences or even biased model predictions (15, 27, 40). In the context of analyzing model sloppiness, much effort has been devoted to studying effects on parameter identifiability by increasing the quality and quantity of the data used to fit the model (29, 30, 59, 62). For example, Apgar *et al.* (30) carefully designed complementary experiments that constrained parameter values in the model of Brown *et al.* (12, 13). To achieve this, they modified some of the model inputs to create different synthetic datasets that were more informative for some of the model parameters than others, but when used together, all model parameters could be constrained within 10% of their true values. However, these computational experiments still required

considerably more data than those typically obtained in practice (15, 27). Instead, we have shown here that the identification of critical parameter interdependencies may more efficiently improve parameter inference when prior knowledge about related (interdependent) model parameters is strategically improved through expert elicitation or new experiments (fig. S8).

We also showed in the cardiac electrophysiological application that if experiments are designed to modify parameter values as well as the values of the stiff eigenparameters, these new experiments are likely to provide new information about the system (Fig. 7A). However, if experiments are designed to change parameter values and instead keep the values of the stiff eigenparameters approximately constant, these new experiments are unlikely to provide new information about the system (Fig. 7B). We note that if the design of parameter-specific experiments is not practically possible (29, 30, 59), the posterior covariance matrix (inverse of matrix $\mathbf{P}$) can still be used to measure the increase in parameter identifiability obtained by increasing the quantity and quality of data. Furthermore, this technique has been recently used in optimal Bayesian experimental design (63).

Identifying critical parameter combinations in stochastic settings

In our example results, we identified critical parameter combinations through the analysis of sloppiness for three different deterministic models (Eq. 1 and tables S1 and S3) fitted to data, in which we also treated the error structure as having been correctly specified by the modeler (Eq. 15). However, implementing such an approach for stochastic models could be a potential area for future research. In stochastic models, randomness often manifests beyond just noise of known structure being added to a deterministic output (64). More so, incorporation of a stochastic component can be used to include effects of random fluctuations into otherwise deterministic models, for example, to represent temporal variations in gene expression in cardiac electrophysiology models (65), similar to the one considered in this work (case study 2). However, such models present a challenge for understanding model sloppiness. Although methods suitable for stochastic models [such as minimum distance estimation (66)] can be used to statistically estimate the values of their parameters, with no closed-form version of the likelihood function available, nor any guarantee of its smoothness, one cannot reasonably evaluate and analyze the Hessian at this point as per the standard approach. Here, the Bayesian perspective on model sloppiness may provide a remedy for these issues. By adopting a likelihood-free method (67, 68), posterior samples may still be generated, at which point the posterior covariance method (matrix $\mathbf{P}$) can be used to identify stiff eigenparameters in the context of both data and prior, as we have demonstrated (Tables 1 and 2).

In contrast to the posterior covariance method (matrix $\mathbf{P}$), the LIS-based approach presented here (matrix $\mathbf{G}$), separating the analysis of model sloppiness from the effects of the prior, does rely upon large numbers of point-wise evaluations of the Hessian matrix (Eq. 12). To rectify this for stochastic models, one may formulate approximations to the matrix $\mathbf{G}$ that avoid calculation of the Hessian by instead attempting to directly remove the effects of the prior from matrix $\mathbf{P}$. For example, the sensitivity matrices formed by subtracting the posterior inverse covariance Σ^{-1} from the prior inverse covariance Ω^{-1} (69) or by premultiplying the posterior covariance Σ by the inverse prior covariance Ω^{-1} (70) have been used in the

context of Bayesian inference to understand the posterior in connection with the prior. Although these alternative approaches solve related eigenproblems, we put forward here the sensitivity matrix $\mathbf{G}$, formed by premultiplication and postmultiplication of the Hessian matrix by the Cholesky factors (Eq. 12) of the prior covariance matrix $\mathbf{\Omega}$. This matrix factorization leads directly to the eigendirections (parameter combinations) where the data are most informative relative to the prior (further discussion in Materials and Methods).

Recognizing knowledge limitations in mathematical models fitted to data

Regardless of how good a mathematical model is, its predictions are only as useful as its known limitations (2, 55, 71). Here, by recognizing knowledge limitations in mathematical models fitted to data, our work also adds to the literature of model development and simulation (2, 3, 9). For example, the identified stiff parameter combinations in the cardiac electrophysiological application uncovered a hidden controlling mechanism of the system (Fig. 6B) that dictated the success or failure of the model output accurately reproducing the experimental data (Fig. 7). In practical applications, identifying this type of model behavior would inform which of the model parameters are important for model reduction (26) or need to be prioritized in future experimental designs (29, 30, 40). Furthermore, the implementation of Bayesian inference to fit the model to data brings the added benefit of quantifying the uncertainty in both parameter values (figs. S1A to S3A, S5, S8, and S10) and model predictions (Figs. 3B and 6A and figs. S1B to S3B and S9). Hence, this work constitutes an example of how this advanced model-data fitting technique can be exploited to reveal the hidden geometry of parameter uncertainty and its effects on model predictions—a topic of growing interest within many fields of science (2–4, 18) that has thus far been hindered due to concerns about system complexity and limited data accessibility (2, 55, 71).

MATERIALS AND METHODS

To assist with the subsequent mathematical description, we first summarize how sloppiness of a model is analyzed in its standard, non-Bayesian, form. Then, we describe how it can be analyzed via a Bayesian framework. Last, we describe the procedure followed in Results to identify the stiff eigenparameters according to both standard and Bayesian approaches to analyzing model sloppiness.

Standard (non-Bayesian) approach to analyzing sloppiness

The standard approach to analyzing sloppiness involves three key steps (12, 13): (i) obtaining the best-fit parameter values θ^* by fitting the model to data, (ii) calculating the sensitivity matrix $\mathbf{S}$ evaluated at the best-fit parameter values θ^* , and (iii) identifying the eigenparameters that are more or less sensitive to the model-data fit through eigendecomposition of the sensitivity matrix $\mathbf{S}$. We detail these three steps for analyzing sloppiness using the standard approach as follows.

Step 1. Obtaining values of the model parameters by fitting the model to data

Let us assume that a deterministic model $y_{\text{model}} = f(\mathbf{x}, \theta)$, with known structure f , a known vector of input conditions $\mathbf{x} \in \mathbb{R}^{N_x}$ of dimension N_x (e.g., representing the spatial and/or temporal location at which the model is evaluated, and/or the external conditions that alter the model output), and parameterized by a vector $\theta \in \mathbb{R}^{N_\theta}$ of

dimension N_θ , has been proposed to explain a dataset y_{obs} that consists of N_{obs} independent observations $y_{\text{obs}} = (y_{\text{obs},1}, y_{\text{obs},2}, \dots, y_{\text{obs},N_{\text{obs}}})$, where $y_{\text{obs},k}$ represents the k th observation in this dataset, $k \in \{1, 2, \dots, N_{\text{obs}}\}$. Here, the problem of uniquely estimating values of parameter set θ given data y_{obs} depends on whether the deterministic model $y_{\text{model}} = f(\mathbf{x}, \theta)$ is structurally and, ultimately, practically identifiable, discussed in detail elsewhere (39, 44). However, regardless of the source of parameter unidentifiabilities, the standard approach to model sloppiness considers identifiability of parameters in the context of their best-fit values (11, 13). Typically, a likelihood-based approach is taken, in which the modeler specifies an error structure that then formalizes this notion of best fit (26, 31). For example, a common choice is to assume that errors are independent and with Gaussian distribution, each having mean zero and a specified standard deviation (SD) that could be observation specific, σ_k (12, 15, 27, 30). Under these conditions, the likelihood takes the form (26, 40)

$$\mathcal{L}(y_{\text{obs}} | \theta) = \prod_{k=1}^{N_{\text{obs}}} \frac{1}{\sqrt{2\pi} \sigma_k} \exp \left[-\frac{1}{2} \left(\frac{y_{\text{obs},k} - y_{\text{model},k}(\mathbf{x}, \theta)}{\sigma_k} \right)^2 \right] \quad (2)$$

where $y_{\text{model},k}(\mathbf{x}, \theta)$ is the model's prediction of an equivalent noiseless observation for y_k given parameters θ and input conditions $\mathbf{x}$. The advantage of this likelihood-based approach is the ability to specify a given error structure that produces a tractable likelihood function, for example, incorporating heteroscedasticity in the data by varying σ_k with each observation (17, 55) as in Eq. 2, or even choosing appropriate error distributions for more specific model-data fitting problems. While appropriate specification of the error structure could potentially depend on domain knowledge, Eq. 2 serves as a broadly applicable choice (8, 11, 17, 29, 62).

The values of the model parameter vector θ that maximize the likelihood function $\mathcal{L}(y_{\text{obs}} | \theta)$ are altogether called the MLE, here denoted as $\theta^* \equiv \theta_{\text{MLE}}^* = \text{argmax}_{\theta} \mathcal{L}(y_{\text{obs}} | \theta)$ (29). We note, however, that a standard least-squares regression may be cast as maximizing a Gaussian likelihood by enforcing homoscedastic errors $\sigma_k = \sigma$ and introducing (12, 15)

$$C(\theta) = -\log \mathcal{L}(y_{\text{obs}} | \theta) =$$

$$\frac{1}{2} N_{\text{obs}} \log(2\pi) + \frac{1}{2} N_{\text{obs}} \log \sigma^2 + \frac{1}{2} \sum_{k=1}^{N_{\text{obs}}} \left(\frac{y_{\text{obs},k} - y_{\text{model},k}(\mathbf{x}, \theta)}{\sigma} \right)^2 \quad (3)$$

where $C(\theta)$ is the cost function. The first two terms in Eq. 3 are independent of the parameter values, so $C(\theta) \propto \sum_{k=1}^{N_{\text{obs}}} (y_{\text{obs},k} - y_{\text{model},k}(\mathbf{x}, \theta))^2 + \text{constant}$. Thus, minimizing the cost function $C(\theta)$ in Eq. 3 to find the best-fit parameter values is equivalent to maximizing the Gaussian likelihood function in Eq. 2 to find the MLE. Furthermore, as a MLE under these conditions, the ordinary least-squares solution is an estimator in the large sample limit, achieving the minimal variance specified by the Cramér-Rao lower bound (72).

Step 2. Calculating the sensitivity matrix

The standard approach to analyzing sloppiness obtains the sensitivity matrix $\mathbf{S}$ by investigating how the cost function $C(\theta)$ in Eq. 3 varies with respect to the parameter vector θ in the vicinity of the MLE $\theta^* = \theta_{\text{MLE}}^*$. To do so, this matrix is obtained by a Taylor expansion of $C(\theta)$ around the best-fit parameter values while differentiating with respect to the logarithm of the parameters ($\log \theta$), which yields (13, 15, 30)

$$C(\boldsymbol{\theta}) \approx C(\boldsymbol{\theta}^*) + \nabla C(\boldsymbol{\theta}^*) \cdot (\log \boldsymbol{\theta} - \log \boldsymbol{\theta}^*) + \frac{1}{2} (\log \boldsymbol{\theta} - \log \boldsymbol{\theta}^*)^\top \mathbf{H} (\log \boldsymbol{\theta} - \log \boldsymbol{\theta}^*) \quad (4)$$

where the gradient $\nabla C(\boldsymbol{\theta}^*)$ of the cost function is zero at the best-fit parameter values by definition so that the sensitivity of the model fit to changes in parameter values is characterized by the Hessian matrix $\mathbf{H}$ defined in Eq. 4, whose elements are given by (12, 13)

$$H_{ij} = - \left. \frac{\partial^2 \log \mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta})}{\partial \log \theta_i \partial \log \theta_j} \right|_{\boldsymbol{\theta}=\boldsymbol{\theta}^*} \quad (5)$$

with i and j both taking integer values ranging from 1 to N_θ . Thus, the Hessian matrix describes the quadratic behavior of the cost function $C(\boldsymbol{\theta})$ infinitesimally close to the point $\boldsymbol{\theta}^*$, and thus, it is considered here as one of the matrices that could be used as the sensitivity matrix $\mathbf{S}$ for analyzing model sloppiness (12, 15, 30). However, since evaluating second-order derivatives can be computationally expensive, the sensitivity matrix can also be approximated by the Levenberg-Marquardt Hessian $\mathbf{L}$ at a much lower computational cost, following (11, 12, 15)

$$L_{ij} = \left. \sum_{k=1}^{N_{\text{obs}}} \frac{\partial r_k}{\partial \log \theta_i} \frac{\partial r_k}{\partial \log \theta_j} \right|_{\boldsymbol{\theta}^*=\boldsymbol{\theta}_{\text{MLE}}^*} \quad (6)$$

where the residual error r_k for the k th observation is calculated via $r_k = [y_{\text{obs},k} - y_{\text{model},k}(\mathbf{x}, \boldsymbol{\theta})]/\sigma_k$, and the first derivatives in Eq. 6 can be evaluated by first-order finite differences or by integrating sensitivity equations for ordinary differential equation models (11, 29). The Levenberg-Marquardt Hessian $\mathbf{L}$ corresponds to the Gauss-Newton approximation of the Hessian $\mathbf{H}$ in Eq. 5, guaranteed to be positive semidefinite (39). Matrix $\mathbf{L}$ is also equal to the observation information matrix evaluated at the MLE, which itself is a sample-based version of the Fisher information matrix (26), whose connections with information theory have been well considered elsewhere (3, 15, 26). The Levenberg-Marquardt Hessian $\mathbf{L}$ thus is a more computationally convenient sensitivity matrix $\mathbf{S}$ for analyzing sloppiness, although as with the Hessian matrix $\mathbf{H}$ it only considers the curvature of the likelihood surface infinitesimally close to the MLE.

Step 3. Identifying the eigenparameters that are more or less sensitive to the model-data fit

To identify the stiff eigenparameters, eigenvalues λ_n and eigenvectors $\mathbf{v}_n$ are obtained via eigendecomposition of the sensitivity matrix $\mathbf{S}$ or via singular value decomposition if numerical stability is an issue (12). Each of the $n = 1, 2, \dots, N_\theta$ eigenvectors $\mathbf{v}_n$ is mutually orthogonal so that eigenparameters can be conveniently expressed as linear combinations of natural logarithms of model parameters, following (13)

$$\alpha_n = \sum_{j=1}^{N_\theta} (v_n)_j \log \theta_j \quad (7)$$

where $(v_n)_j$ is the j th element of the n th eigenvector $\mathbf{v}_n$ of the sensitivity matrix. Thus, each eigenparameter $\hat{\theta}_n$ can be simply represented as the product and/or quotient of bare model parameters raised to an index given by the elements of eigenvector $\mathbf{v}_n$, by rewriting Eq. 7 as (13)

$$\hat{\theta}_n := \exp(\alpha_n) = \prod_{j=1}^{N_\theta} \theta_j^{(v_n)_j} \quad (8)$$

with stiff eigenparameters $\hat{\theta}_n$ associated with the largest eigenvalues λ_n and sloppy (soft) eigenparameters associated with the smallest eigenvalues. The magnitude of each element $(v_n)_j$, $j = 1, \dots, N_\theta$ of the eigenvector $\mathbf{v}_n$ in Eq. 8 therefore indicates the relative contribution of bare parameter θ_j to eigenparameter $\hat{\theta}_n$. If eigenvectors are normalized, each $(v_n)_j$ takes values between -1 and 1 inclusive so that all $\hat{\theta}_n$ are products of bare parameters having exponents with magnitudes that do not exceed unity. Any factors $\theta_j^{(v_n)_j}$ in Eq. 8 having relatively low magnitudes for $(v_n)_j$ [e.g., $|(v_n)_j| \leq 0.2$] contribute little to the eigenparameter's value; thus, these small factors $\theta_j^{(v_n)_j}$ can be practically excluded from the product (12). Hence, each of the $n = 1, 2, \dots, N_\theta$ eigenparameters $\hat{\theta}_n$ may depend strongly on only a few bare parameters that may be importantly related to each other.

Bayesian approach to analyzing sloppiness

In the context of fitting models to data, Bayesian inference provides a coherent statistical framework to estimate probability distributions for model parameters, constrained by the combination of data and prior beliefs (22, 33). Thus, if the model-data fitting problem is recast as a Bayesian inference problem, the final estimates for the probability distribution of parameters $\boldsymbol{\theta}$, based on all of the data $\mathbf{y}_{\text{obs}}$, are called the posterior distribution $\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}})$. To apply Bayesian inference, we require definition of both a likelihood function $\mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta})$ and a prior distribution $p(\boldsymbol{\theta})$. An example of the former of these was defined in Eq. 2 (i.e., Gaussian likelihood), while the latter of these represents the initial beliefs about the parameter values, which are often based on earlier studies, or in their absence, they are based on expert knowledge (33, 73). Once both likelihood function and prior distribution are defined, Bayes' theorem is then used to obtain the posterior distribution, following (23)

$$\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}}) = \frac{\mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta}) p(\boldsymbol{\theta})}{\int_{\boldsymbol{\theta}} \mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta}) p(\boldsymbol{\theta}) d\boldsymbol{\theta}} \quad (9)$$

Here, the denominator is a multidimensional integral over the parameter space, $\boldsymbol{\theta}$, that serves as a normalizing constant but is, however, difficult to calculate directly or often intractable (22, 23, 58). Therefore, several methods that avoid calculation of this constant have been developed to sample from the posterior distribution, including MCMC sampling (74), sequential Monte Carlo (SMC) sampling (58), approximate Bayesian computation (ABC) (67), variational Bayesian inference (75), Laplace approximation (76), and many others. For the purposes of this section, we simply assume that the posterior has been successfully sampled, and thus, we hereafter discuss practical aspects of computing the sensitivity matrix within a Bayesian framework. Thus, analogous to the standard approach to analyzing sloppiness, the Bayesian approach consists of three steps: (i) obtaining an estimate of the posterior distribution $\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}})$ by fitting the model to data, (ii) calculating a Bayesian-based sensitivity matrix $\mathbf{S}$ from the posterior distribution $\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}})$, and (iii) identifying the eigenparameters that are more or less sensitive to the model-data fit through eigendecomposition of the Bayesian-based sensitivity matrix $\mathbf{S}$.

Exact details of the first step above depend on the posterior-computation method chosen, while the third step is the same as the third step of the standard approach. Thus, we focus here on the second step, for which we adapt two Bayesian methods for dimensionality reduction to obtain sensitivity matrices for analyzing model sloppiness. These are described as follows.

Posterior covariance method

The posterior covariance method is based on the application of PCA (36). This technique uses eigendecomposition of a sensitivity matrix (a covariance matrix) to reduce the dimensionality of large datasets, which thus identifies the dataset components that account for the largest amount of variance (38). In our context, the dataset of interest is a Bayesian ensemble of plausible parameter values, which we obtain from the posterior distribution for the parameters. Thus, if PCA is applied on this specific dataset, eigenvectors and eigenvalues of the posterior covariance matrix inform the variability of the model-data fit to changes in parameter values. However, given that we seek to identify the eigenparameters that are well constrained by the available data (i.e., those that have less variability), we instead calculate the sensitivity matrix $\mathbf{S}$ as the PCA Hessian matrix $\mathbf{P}$ that is based on the inverse of the posterior covariance matrix $\mathbf{\Sigma}$ (12)

$$\mathbf{P} = \mathbf{\Sigma}^{-1} \quad (10)$$

where the matrix $\mathbf{\Sigma}$ is calculated in terms of the natural logarithms of model parameters $\log \boldsymbol{\theta}$, with this transformation required in Eq. 8 to characterize stiff/sloppy eigenparameters as products or quotients of the bare model parameters. This is a key advantage of the posterior covariance method over more sophisticated dimensional reduction techniques [e.g., kernel PCA (60) and/or ISOMAP (77)], in which mappings back to original parameter space are not typically sought, and thus, the associated eigenparameters describing the lower dimensional parameter space are not readily interpretable. Hence, eigendecomposition of the PCA Hessian matrix $\mathbf{P}$ identifies which eigenparameters are more or less constrained by the combination of both data and prior beliefs. Specifically, eigenvectors of matrix $\mathbf{P}$ with large eigenvalues indicate stiff eigenparameters, while eigenvectors with small eigenvalues indicate sloppy eigenparameters.

We note that if Monte Carlo methods such as MCMC sampling (74), SMC sampling (58), or ABC (67) are used to approximate the posterior as a set of M equally weighted samples $\{\boldsymbol{\theta}_m\}_{m=1}^M$, the required posterior covariance matrix $\mathbf{\Sigma}$, calculated with respect to the natural logarithms of parameters $\log \boldsymbol{\theta}$, can be estimated using the sample covariance matrix $\hat{\mathbf{\Sigma}}$

$$\mathbf{\Sigma} \approx \hat{\mathbf{\Sigma}} = \frac{1}{M-1} \sum_{m=1}^M (\log \boldsymbol{\theta}_m - \log \bar{\boldsymbol{\theta}})(\log \boldsymbol{\theta}_m - \log \bar{\boldsymbol{\theta}})^T \quad (11)$$

where $\log \bar{\boldsymbol{\theta}} = \frac{1}{M} \sum_{m=1}^M \log \boldsymbol{\theta}_m$ is the estimated posterior mean for the natural logarithm of parameters. If Monte Carlo methods are overly computationally expensive, fast approximate methods such as variational Bayesian inference or Laplace approximation (76) can be used as an alternative to provide a rapid estimate of the posterior covariance matrix. However, these fast approximate methods provide a rapid, albeit possibly biased, estimate of the posterior covariance matrix (76).

LIS method

The LIS method proposed here has its origins in the Bayesian parameter reduction literature, specifically from the work of Cui *et al.* (35) who developed a method for Bayesian inverse problems that identifies the directions in parameter space where the data are most informative relative to the prior. The motivation for Cui *et al.* (35) was to develop an approximate but accelerated MCMC algorithm that samples over a lower-dimensional subspace, called the LIS, to avoid sampling from directions of prior variability that the likelihood

does not inform (48). The LIS is constructed on the idea that the Hessian of the log-likelihood can be compared to the prior covariance to then identify directions in parameter space along which the posterior distribution differs strongly from the prior, i.e., directions that are likelihood informed (78). Thus, we adapt here the approach used by Cui *et al.* (35) to construct the LIS to define a sensitivity matrix in our context.

Our goal is to make the sensitivity matrix dependent primarily on the data and eliminate effects of the prior distribution. To achieve this, we first assume that the covariance matrix $\mathbf{\Omega}$ of the prior distribution for the logarithms of parameters is known and that this matrix can be Cholesky factored to a lower triangular matrix $\mathbf{L}_p$ such that $\mathbf{L}_p \mathbf{L}_p^T = \mathbf{\Omega}$. Then, by following Cui *et al.* (35), we define the prior-preconditioned Hessian matrix $\mathbf{\Psi}(\boldsymbol{\theta})$ as

$$\mathbf{\Psi}(\boldsymbol{\theta}) = \mathbf{L}_p^T \mathbf{H}(\boldsymbol{\theta}) \mathbf{L}_p \quad (12)$$

for parameter vector $\boldsymbol{\theta}$, with elements of $\mathbf{H}(\boldsymbol{\theta})$ given by Eq. 5. We note that Cui *et al.* (35) used a multivariate Gaussian prior to define the prior-preconditioned Hessian matrix $\mathbf{\Psi}(\boldsymbol{\theta})$ in Eq. 12, which is needed in that context to approximate the posterior distribution as the product of a lower-dimensional posterior defined on the LIS and the prior distribution marginalized onto a complementary subspace (48, 78). However, given that our purpose is to identify the directions that are data informed, and not to approximate a posterior distribution, the LIS definition is not restricted to multivariate Gaussian priors in our application. Thus, we obtain an expression for the LIS, used here to define the LIS-based sensitivity matrix $\mathbf{G}$, by integrating over the prior-preconditioned Hessian matrix with respect to the posterior (35), which yields

$$\mathbf{G} = \int_{\boldsymbol{\theta}} \mathbf{\Psi}(\boldsymbol{\theta}) \pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}}) d\boldsymbol{\theta} \quad (13)$$

Given that Eq. 13 involves an integral over N_{θ} -dimensional space, then if the posterior is approximated by a Monte Carlo method (e.g., MCMC, SMC, or ABC) as a set of M equally weighted samples $\{\boldsymbol{\theta}_m\}_{m=1}^M$, the LIS-based sensitivity matrix $\mathbf{G}$ can instead be estimated as

$$\mathbf{G} \approx \frac{1}{M} \sum_{m=1}^M \mathbf{\Psi}(\boldsymbol{\theta}_m) \quad (14)$$

where each $\mathbf{\Psi}(\boldsymbol{\theta}_m)$ is calculated via Eq. 12 with the Hessian matrix $\mathbf{H}(\boldsymbol{\theta}_m)$ of the negative log-likelihood function evaluated at each posterior sample $\boldsymbol{\theta}_m$ via Eq. 5 or approximated by the Levenberg-Marquardt Hessian $\mathbf{L}(\boldsymbol{\theta}_m)$ via Eq. 6 to reduce computational cost in the calculation of matrix $\mathbf{G}$ (35). Furthermore, we note that these M matrices $\mathbf{H}(\boldsymbol{\theta}_m)$ are all left-multiplied by $\mathbf{L}_p^T$ and right-multiplied by $\mathbf{L}_p$, with the M resulting $\mathbf{\Psi}(\boldsymbol{\theta}_m)$ matrices averaged to obtain the sensitivity matrix $\mathbf{G}$. As a result, eigendecomposition of this prior-informed sensitivity matrix $\mathbf{G}$ can reveal which eigenparameters are strongly informed by the data relative to the prior, i.e., directions in parameter logspace where the posterior differs most strongly from the prior (78).

Demonstrating how to analyze model sloppiness using examples of models fitted to synthetic data

In this section, we describe the six-step procedure used to analyze model sloppiness in the examples discussed in Results. This six-step procedure incorporates both approaches discussed above, i.e., the

local sensitivity analysis around the best-fit parameter values (standard approach) and the global sensitivity analysis considering all plausible parameter values consistent with the available data (Bayesian approach). Each step of the procedure describes specific details of the examples considered in Results.

Step i. Defining the model form

We consider deterministic models of the form $\mathbf{y}_{\text{model}}(\mathbf{x}, \boldsymbol{\theta})$, where $\mathbf{x} \in \mathbb{R}^{N_x}$ is a vector of input conditions, $\boldsymbol{\theta} \in \mathbb{R}^{N_\theta}$ is a vector of model parameters, and $\mathbf{y}_{\text{model}} \in \mathbb{R}^{N_y}$ is a vector of model outputs (see step 1, standard approach). Here, N_x and N_y are the number of model inputs and outputs, respectively.

Step ii. Generating synthetic data to fit the model

We generate measurement data for the motivating example and ecological application by adding heteroscedastic noise with variance proportional to the observation, which follows a truncated normal distribution $y_{\text{obs},j}(x_i) \sim \mathcal{N}(\mu_{y_j}(x_i), \sigma_{y_j}(x_i))$ with mean $\mu_{y_j}(x_i) = y_{\text{model},j}(x_i, \boldsymbol{\theta}_R)$, SD $\sigma_{y_j}(x_i) = \epsilon y_{\text{model},j}(x_i, \boldsymbol{\theta}_R)$, and lower truncation bound of zero on each of the synthetic observations $y_{\text{obs},j}(x_i)$ (17). Here, $\boldsymbol{\theta}_R$ is the vector containing the reference (true) values for the model parameters, ϵ is a user-defined measurement error ranging between 0 and 100%, and noise is added to the j th model output associated with the i th set of input conditions. Alternatively, we generate measurement data for the cardiac electrophysiological application by adding homoscedastic noise, which follows a normal distribution $y_{\text{obs},j}(x_i) \sim \mathcal{N}(\mu_{y_j}(x_i), \sigma_{y_j}(x_i))$ with mean $\mu_{y_j}(x_i) = y_{\text{model},j}(x_i, \boldsymbol{\theta}_R)$ and constant SD $\sigma_{y_j}(x_i) = \sigma$ (8). In each case, measurement error and sampling frequency (number of measurements) are chosen according to typical experimental conditions. As later discussed in detail in step iv, the choice of error structure used for synthetic data generation is also used to define the form of the likelihood function for each case. That is, the error structure is treated as having been correctly specified by the modeler.

Step iii. Defining the vector of unknown model parameters and their prior distributions

We define the vector of unknown model parameters $\boldsymbol{\theta}$ consisting of (i) the model parameters, (ii) model initial conditions (only considered in the ecological model), and (iii) measurement error ϵ or SD σ following the type of noise added to the synthetic data. Then, we specify prior distributions for the parameters $p(\boldsymbol{\theta})$ using either positive uniform or multivariate log-normal probability distributions (8, 25, 34, 47, 55), as follows.

In the Michaelis-Menten kinetic example, three different joint prior distributions $p(\boldsymbol{\theta}) \equiv p(k_{\text{cat}}, [E_T], K_M, \epsilon)$ are used for the three parameters k_{cat} , $[E_T]$, and K_M of the model and measurement error ϵ . The first joint prior consists of a uniform prior for each parameter; the second joint prior consists of multivariate log-normal priors for all parameters, with the prior of parameter K_M being badly specified; and the third joint prior consists of a uniform prior for k_{cat} , a badly specified log-normal prior for $[E_T]$, and a well-specified log-normal prior for K_M and ϵ . All joint priors assume independence between the model parameters and measurement error, so $p(\boldsymbol{\theta}) \equiv p(k_{\text{cat}})p([E_T])p(K_M)p(\epsilon)$. In this work, a badly specified prior for the n th parameter θ_n means that this parameter's true value $\theta_{R,n}$ has little support under the prior distribution (i.e., it lies in the tails of the prior). This is a condition referred to as a prior-data conflict that occurs when informative prior beliefs are inconsistent with the information revealed by the data (79), although the model is correctly specified as is assumed here [see (80) for discussion of appropriateness of Bayesian inference when the model is misspecified]. Alternatively, a well-specified prior in this work means that the true

parameter value $\theta_{R,n}$ is well supported by the prior distribution, i.e., it lies within the bulk of the prior distribution so that prior beliefs are consistent with the information given by the data.

In the ecological application, two different joint prior distributions are used for the 20 parameters of the ecosystem network model (table S2) and measurement error ϵ . The first joint prior distribution is chosen to be a product of vague log-normal distributions for each parameter so that this joint prior has zero covariance. Alternatively, the second joint prior distribution is chosen to be a product of well-specified log-normal distributions for parameters a_M , a_N , and a_P and vague log-normal distributions for each of the remaining parameters, including the measurement error ϵ . As discussed in Results, well-specified priors for parameters a_M , a_N , and a_P are chosen on the basis of the stiff eigenparameters, identified from the analysis of sloppiness for the case considering vague log-normal distributions for each parameter.

In the cardiac electrophysiological application, a well-specified multivariate log-normal prior distribution is used for the nine parameters of the BR model (table S4) and the SD σ . This joint prior distribution is centered at the reference parameter values and assumes zero covariance between the nine parameters and the SD σ . Stimulation conditions A_s , t_{on} , and t_{dur} ; membrane capacitance C_m ; and initial conditions $V_m(0)$, $[Ca]_i(0)$, $x_1(0)$, $m(0)$, $h(0)$, $j(0)$, $d(0)$, and $f(0)$ are set to their reference values (table S4) and are not estimated via our model-data fitting techniques.

Step iv. Fitting the model to data

We use two approaches to fit each example model to data, with the first being MLE and the second being Bayesian inference. To implement these two approaches, we conveniently rewrite the Gaussian likelihood function defined in Eq. 2 as

$$\mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta}) = \prod_{j=1}^{N_y} \prod_{i=1}^{N_x} \frac{1}{\sqrt{2\pi} \sigma_{y_j}(x_i)} \exp \left[-\frac{1}{2} \left(\frac{y_{\text{obs},j}(x_i) - y_{\text{model},j}(x_i, \boldsymbol{\theta})}{\sigma_{y_j}(x_i)} \right)^2 \right] \quad (15)$$

where $\sigma_{y_j}(x_i) = \epsilon y_{\text{model},j}(x_i, \boldsymbol{\theta})$ when heteroscedastic noise is used to generate the synthetic data and $\sigma_{y_j}(x_i) = \sigma$ when homoscedastic noise is instead used. Then, we use this Gaussian likelihood function and specified prior distributions (step iii) to approximate the joint posterior distribution $\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}})$ via Bayes' theorem (Eq. 9) by implementing the SMC sampling algorithm adapted from Adams *et al.* (17, 55). In our implementation of this posterior sampling algorithm, we use a sample size of $M = 10,000$, Metropolis-Hastings acceptance fraction of $C = 0.95$, and effective sample size reduction target of $\Delta = 0.001$. These settings were sufficient for reproducible sampler performance: Results did not vary in independent runs of the sampling algorithm using a smaller sample size of $M = 5000$ and larger effective sample size reduction target of $\Delta = 0.005$ (figs. S15 to S21).

Once the joint posterior probability distributions $\pi(\boldsymbol{\theta} | \mathbf{y}_{\text{obs}})$ are obtained for each example, we estimate the best-fit parameter values $\boldsymbol{\theta}^*$ (MLE) by minimizing the cost function $C(\boldsymbol{\theta}) = -\log \mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta})$ with $\mathcal{L}(\mathbf{y}_{\text{obs}} | \boldsymbol{\theta})$ given by Eq. 15 while using the posterior mean as the initial guess to start the optimization. Here, the sets of best-fit parameter values $\boldsymbol{\theta}^*$, $\boldsymbol{\theta}_1^*$, and $\boldsymbol{\theta}_2^*$ are only used to calculate the sensitivity matrices ($\mathbf{H}$ or $\mathbf{L}$) via the standard approach, while the already obtained prior and posterior distributions are used to calculate the sensitivity matrices ($\mathbf{P}$ and $\mathbf{G}$) based on the Bayesian approach.

Step v. Calculating the sensitivity matrix

Equations 5, 6, 10, and 14 are used to calculate the Hessian $\mathbf{H}$, the Levenberg-Marquardt Hessian $\mathbf{L}$, the PCA Hessian $\mathbf{P}$, and the LIS $\mathbf{G}$, respectively. Each of these matrices acts as a sensitivity matrix for the purpose of analyzing sloppiness. In the absence of analytical derivatives, we use central finite differences (81) to approximate first- and second-order derivatives of the log-likelihood function with respect to the logarithm of parameters, with a step size $\Delta\theta_i = \delta \times \theta_i$, $i = 1, \dots, N_\theta$, where δ is a small scalar between 10^{-4} and 10^{-2} . Finite differencing is the most widely used technique for numerical differentiation in physical applications (81), including approximations of the sensitivity matrix ($\mathbf{H}$ or $\mathbf{L}$) in standard analysis of model sloppiness (13, 15, 30). However, for more complex models than those considered here, this technique can become computationally expensive, as it requires multiple model evaluations for approximating derivatives. As an alternative, more sophisticated methods such as automatic differentiation (82) may also be used (where appropriate) in conjunction with the analysis of model sloppiness (11, 28).

Rows and columns of each sensitivity matrix characterizing the sensitivity of the model-data fit with respect to the measurement error (represented by ϵ or σ in Eq. 15) are not calculated, thus preventing the measurement error from appearing in the parameter combinations that are identified through the analysis of model sloppiness. Here, the measurement error is effectively treated as a nuisance parameter, that is, it is involved with the model-data fitting procedure but does not provide information to identify relevant parameter combinations. In addition, for likelihood functions of the form given by Eq. 15, small changes to the measurement error are expected to affect only the degree of overall curvature of the model-data fit surface but not the directions of high or low curvature. As a result, the dimension of the square symmetric sensitivity matrices $\mathbf{H}$, $\mathbf{L}$, $\mathbf{P}$, and $\mathbf{G}$ obtained here is equal to the number of model parameters, excluding the measurement error (ϵ or σ in Eq. 15), i.e., we obtain 3×3 sensitivity matrices in the Michaelis-Menten kinetic example, 20×20 in the ecological application, and 9×9 in the cardiac electrophysiological application.

Step vi. Identifying stiff eigenparameters

Eigenvalues and eigenvectors of the sensitivity matrices $\mathbf{H}$, $\mathbf{L}$, $\mathbf{P}$, and $\mathbf{G}$ are calculated via singular value decomposition (12). Then, eigenparameters $\hat{\theta}$ are obtained via Eq. 8, in which we consider the contribution of parameter θ_j to eigenparameter $\hat{\theta}_n$ only when element j of the normalized n th eigenvector $(v_n)_j$ satisfies $|(v_n)_j| \geq 0.2$ (see step 3, standard approach). We also rescale exponents $(v_n)_j$ of the bare parameters θ_j associated with each eigenparameter $\hat{\theta}_n$ so that the magnitude of the largest/smallest index $(v_n)_j$ for every eigenvector v_n is 1. Here, eigenvalues are ordered from largest to smallest so that the corresponding eigenparameters are also ordered from stiffest to sloppiest.

Trade-offs of locally and globally analyzing model sloppiness

To summarize these methods, we have proposed a unified framework to obtain locally and globally the key quantity for analyzing model sloppiness—the sensitivity matrix $\mathbf{S}$ (e.g., $\mathbf{H}$, $\mathbf{L}$, $\mathbf{P}$, and $\mathbf{G}$). This approach accurately estimates uncertainty in parameter values, constrained by the combination of prior information and data, with the key benefit of robustly identifying the relative effect of this prior information in the inference of critical parameter combinations that control the quality of the model-data fit. This is a key achievement of this work as it extends the application of the analysis of sloppiness beyond systems where there is little prior knowledge about the

model parameter values (11, 12) to those where prior information is more readily available (32, 34, 45, 55), and thus can be confidently incorporated as part of the Bayesian model-data fitting process to constrain parameter values (7, 8, 17, 25).

In the implementation of this framework, the local (standard) methods to analyzing sloppiness (matrices $\mathbf{H}$ or $\mathbf{L}$) were found to be computationally inexpensive in comparison to the Bayesian methods (matrices $\mathbf{P}$ and $\mathbf{G}$). Thus, standard methods can be very useful in model-data fitting applications where computationally expensive models make implementation of Bayesian inference impractical. Nevertheless, since local analysis of sloppiness considers a single point estimation in parameter space (i.e., the best-fit parameter values), this local approach can only accurately quantify the model sensitivity to parameter changes when the likelihood function maximum (or cost function minimum) is well defined (11, 59). Unfortunately, if the likelihood surface is relatively complicated (e.g., with ridges), this method can mislead inference of stiff eigenparameters (see, for example, Table 2) (15, 57). Careful selection of the optimization algorithm is thus needed to avoid convergence to local optima (12, 16, 57). In addition to this limitation, both standard methods to analyzing sloppiness require a closed-form likelihood function to calculate sensitivity matrices $\mathbf{H}$ and $\mathbf{L}$, such as the Gaussian likelihood functions (e.g., Eqs. 2 and 15) considered here.

Alternatively, the global (Bayesian) analysis of sloppiness looks beyond the curvature of the likelihood function surface at a single point while fully exploring the topography of this surface by using an ensemble of plausible parameter values to characterize the sensitivities of the model-data fit to changes in parameter values. As part of this global approach, we exploited PCA (36) to implement the posterior covariance method (matrix $\mathbf{P}$) that assesses the data informativity about the critical parameter combinations while accounting for (including) any prior information about parameter values. This method does not require approximating gradients of the log likelihood (Eq. 10), and so, calculating sensitivity matrix $\mathbf{P}$ is computationally inexpensive after the posterior distribution is obtained via Bayesian inference. However, since the posterior covariance method (matrix $\mathbf{P}$) assumes that the posterior structure is well captured by a covariance matrix Σ of the logarithm of parameters $\log \theta$, it is also restricted to applications where the posterior distribution for $\log \theta$ is approximately multivariate normally distributed. Despite this limitation, the posterior covariance method has the added benefit of being readily applicable to all kinds of statistical models, even to those with intractable likelihood functions where “likelihood-free” Bayesian methods, such as ABC (67) and Bayesian synthetic likelihood (68), are prevalent.

As part of the global approach, we also implemented the LIS method for Bayesian dimensionality reduction (35, 37) to analyze model sloppiness. Following its origins, the LIS method (matrix $\mathbf{G}$) adapted here assesses the data informativity about the critical parameter combinations while also acknowledging and excluding any prior information. As with the standard methods (matrices $\mathbf{H}$ or $\mathbf{L}$), the LIS method requires a closed-form likelihood function (e.g., Gaussian likelihood functions in Eqs. 2 and 15) to obtain the sensitivity matrix $\mathbf{G}$ (Eq. 14). More so, given that the LIS method also involves calculation of the Hessian matrix at a posterior sample (Eq. 12), approximating second-order derivatives for all posterior samples can become computationally expensive via finite differencing for models with time-consuming solutions. Despite these limitations, as the LIS method does not assume a given shape for the posterior distribution, it has the added benefit of being readily applicable to systems with

non-Gaussian posterior distributions. On the other hand, where the posterior distribution is close to a Gaussian, one may replace the LIS's average over Hessians (Eq. 14) with the posterior covariance

$$\mathbf{K} \approx \mathbf{L}_p \boldsymbol{\Sigma}^{-1} \mathbf{L}_p \quad \text{or} \quad \mathbf{K}^{-1} \approx \mathbf{L}_p^{-1} \boldsymbol{\Sigma} \mathbf{L}_p^{-1} \quad (16)$$

with $\mathbf{L}_p \mathbf{L}_p^\top = \boldsymbol{\Omega}$ (see LIS method), hence formulating a likelihood-free approximation to matrix $\mathbf{G}$. This makes the LIS applicable for stochastic models with intractable likelihoods and greatly reduces the extra computational cost of Hessian calculation at all posterior samples. This idea is similar to other approaches comparing prior and posterior covariance matrices to understand the posterior in the context of the prior (69, 70), arising from the generalized eigenproblem $\mathbf{H}\mathbf{v} \approx \boldsymbol{\Sigma}^{-1}\mathbf{v} = \boldsymbol{\Lambda}\boldsymbol{\Omega}^{-1}\mathbf{v}$ upon approximating the Hessian with the inverse covariance matrix $\boldsymbol{\Sigma}^{-1}$ for Gaussian settings (37). Matrix $\mathbf{K}$ can also be obtained by transforming this generalized eigenproblem into a standard eigenproblem, for which the eigenvectors are then readily interpretable for analyzing model sloppiness. As the eigenvectors of matrix $\boldsymbol{\Sigma}$ in Eq. 11 are equivalent to those of matrix $\mathbf{H}$ in Eq. 12 for a Gaussian posterior distribution, matrix $\mathbf{G}$ in Eq. 14 and matrix $\mathbf{K}$ (or $\mathbf{K}^{-1}$) in Eq. 16 share the same eigenvectors [see (37) for discussion of eigenproperties of matrices $\mathbf{H}$ and $\boldsymbol{\Sigma}$ under Gaussian settings].

Beyond the Bayesian methods discussed here, for applications in which sampling the posterior distribution is infeasible or simply impractical, forward sensitivity analysis methods such as the active subspace (AS) (83) could potentially be used as an alternative to assess sensitivities of the model-data fit function to changes in parameter values. The AS method has the advantage of evaluating a similar sensitivity matrix at a prior sample [referred to as matrix $\mathbf{C}$ by Constantine *et al.* (83)], which makes its implementation less computationally expensive than that of the LIS method since a posterior sample is not needed. Similar to the LIS, the AS identifies a set of important (stiff) directions in the space of all parameters (83, 84). However, the method has also been recently shown to be not completely analogous to the LIS method for both Gaussian and non-Gaussian settings in the context of Bayesian dimensionality reduction (85). Consequently, for analyzing model sloppiness, eigenparameters obtained from the AS method are expected to have a different interpretation than that of eigenparameters obtained from matrices $\mathbf{P}$ and $\mathbf{G}$ in relation to acknowledging the source of information (i.e., prior and/or data). Thus, exploring how the AS method compares to the Bayesian methods discussed here could be an interesting direction for future work.

Hence, given the great flexibility of the techniques discussed here to unveil sensitivities of the model-data fit to changes in parameter values, our comprehensive approach to analyzing model sloppiness does comprise a suitable set of tools to aid in understanding many of nature's systems, ranging from a single cell in the human body (7, 8) and the myriad of microorganisms found almost everywhere (3–5) to large ecosystem networks (2, 17) and beyond (9, 10), through the simultaneous usage of experimental data, mathematical models, and computer simulation.

SUPPLEMENTARY MATERIALS

Supplementary material for this article is available at <https://science.org/doi/10.1126/sciadv.abm5952>

[View/request a protocol for this paper from Bio-protocol.](#)

REFERENCES AND NOTES

1. A. Ma'ayan, Complex systems biology. *J. R. Soc. Interface* **14**, 20170391 (2017).
2. W. L. Geary, M. Bode, T. S. Doherty, E. A. Fulton, D. G. Nimmo, A. I. T. Tulloch, V. J. D. Tulloch, E. G. Ritchie, A guide to ecosystem models and their environmental applications. *Nat. Ecol. Evol.* **4**, 1459–1471 (2020).
3. A. F. Villaverde, J. R. Banga, Reverse engineering and identification in systems biology: Strategies, perspectives and challenges. *J. R. Soc. Interface* **11**, 20130505 (2014).
4. N. Mouquet, Y. Lagadeuc, V. Devictor, L. Doyen, A. Duputié, D. Eveillard, D. Faure, E. Garnier, O. Gimenez, P. Huneman, F. Jabot, P. Jarne, D. Joly, R. Julliard, S. Kéfi, G. J. Kergoat, S. Lavorel, L. L. Gall, L. Meslin, S. Morand, X. Morin, H. Morlon, G. Pinay, R. Pradel, F. M. Schurr, W. Thuiller, M. Loreau, REVIEW: Predictive ecology in a changing world. *J. Appl. Ecol.* **52**, 1293–1310 (2015).
5. C. C. Drovandi, A. N. Pettitt, Estimation of parameters for macroparasite population evolution using approximate Bayesian computation. *Biometrics* **67**, 225–233 (2011).
6. T. Schlick, *Molecular Modeling and Simulation: An Interdisciplinary Guide* (Springer, 2010), vol. 2.
7. B. A. J. Lawson, C. C. Drovandi, N. Cusimano, P. Burrage, B. Rodriguez, K. Burrage, Unlocking data sets by calibrating populations of models to data density: A study in atrial electrophysiology. *Sci. Adv.* **4**, e1701676 (2018).
8. R. H. Johnstone, E. T. Y. Chang, R. Bardenet, T. P. de Boer, D. J. Gavaghan, P. Pathmanathan, R. H. Clayton, G. R. Mirams, Uncertainty and variability in models of the cardiac action potential: Can we build trustworthy models? *J. Mol. Cell. Cardiol.* **96**, 49–62 (2016).
9. K. Veltin, *Mathematical Modeling and Simulation: Introduction for Scientists and Engineers* (Wiley-VCH, 2009).
10. M. Sundberg, Creating convincing simulations in astrophysics. *Sci. Technol. Human Values* **37**, 64–87 (2010).
11. R. N. Gutenkunst, J. J. Waterfall, F. P. Casey, K. S. Brown, C. R. Myers, J. P. Sethna, Universally sloppy parameter sensitivities in systems biology models. *PLOS Comput. Biol.* **3**, e189 (2007).
12. K. S. Brown, J. P. Sethna, Statistical mechanical approaches to models with many poorly known parameters. *Phys. Rev. E* **68**, 021904 (2003).
13. K. S. Brown, C. C. Hill, G. A. Calero, C. R. Myers, K. H. Lee, J. P. Sethna, R. A. Cerione, The statistical mechanics of complex signaling networks: Nerve growth factor signaling. *Phys. Biol.* **1**, 184 (2004).
14. A. Lewbel, The identification zoo: Meanings of identification in econometrics. *J. Econ. Lit.* **57**, 835–903 (2019).
15. L. Geris, D. Gomez-Cabrero, *Uncertainty in Biology: A Computational Modeling Approach* (Springer International Publishing, 2016).
16. M. K. Transtrum, B. B. Machta, J. P. Sethna, Geometry of nonlinear least squares with applications to sloppy models and optimization. *Phys. Rev. E* **83**, 036701 (2011).
17. M. P. Adams, S. A. Sisson, K. J. Helmstedt, C. M. Baker, M. H. Holden, M. Plein, J. Holloway, K. L. Mengersen, E. McDonald-Madden, Informing management decisions for ecological networks, using dynamic models calibrated to noisy time-series data. *Ecol. Lett.* **23**, 607–619 (2020).
18. S. Marino, I. B. Hogue, C. J. Ray, D. E. Kirschner, A methodology for performing global uncertainty and sensitivity analysis in systems biology. *J. Theor. Biol.* **254**, 178–196 (2008).
19. A. Saltelli, T. H. Andres, T. Homma, Sensitivity analysis of model output: An investigation of new techniques. *Comput. Stat. Data Anal.* **15**, 211–238 (1993).
20. I. M. Sobol', Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Math. Comput. Simul.* **55**, 271–280 (2001).
21. A. Saltelli, M. Ratto, T. Andres, F. Campolongo, J. Cariboni, D. Gatelli, M. Saisana, S. Tarantola, *Global Sensitivity Analysis: The Primer* (John Wiley & Sons Ltd, 2008).
22. M. Girolami, Bayesian inference for differential equations. *Theor. Comput. Sci.* **408**, 4–16 (2008).
23. A. Gelman, J. B. Carlin, H. S. Stern, D. B. Dunson, A. Vehtari, D. B. Rubin, *Bayesian Data Analysis* (Chapman and Hall/CRC, ed. 3, 2013).
24. D. Luengo, L. Martino, M. Bugallo, V. Elvira, S. Särkkä, A survey of Monte Carlo methods for parameter estimation. *EURASIP J. Adv. Signal Process.* **2020**, 25 (2020).
25. C. C. Drovandi, N. Cusimano, S. Psaltis, B. A. J. Lawson, A. N. Pettitt, P. Burrage, K. Burrage, Sampling methods for exploring between-subject variability in cardiac electrophysiology experiments. *J. R. Soc. Interface* **13**, 20160214 (2016).
26. M. K. Transtrum, B. B. Machta, K. S. Brown, B. C. Daniels, C. R. Myers, J. P. Sethna, Perspective: Sloppiness and emergent theories in physics, biology, and beyond. *J. Chem. Phys.* **143**, 010901 (2015).
27. A. White, M. Tolman, H. D. Thames, H. R. Withers, K. A. Mason, M. K. Transtrum, The limitations of model-based experimental design and parameter estimation in sloppy systems. *PLOS Comput. Biol.* **12**, e1005227 (2016).
28. M. K. Transtrum, A. T. Sarić, A. M. Stanković, Measurement-directed reduction of dynamic models in power systems. *IEEE Trans. Power Syst.* **32**, 2243 (2016).

29. D. R. Hagen, J. K. White, B. Tidor, Convergence in parameters and predictions using computational experimental design. *Interface Focus* **3**, 20130008 (2013).
30. J. F. Apgar, D. K. Witmer, F. M. White, B. Tidor, Sloppy models, parameter uncertainty, and the role of experimental design. *Mol. Biosyst.* **6**, 1890–1900 (2010).
31. E. Dufresne, H. A. Harrington, D. V. Raman, The geometry of sloppiness. *J. Algebraic Stat.* **9**, 30–68 (2018).
32. P. P.-Y. Wu, M. J. Caley, G. A. Kendrick, K. McMahon, K. Mengersen, Dynamic Bayesian network inferencing for non-homogeneous complex systems. *Appl. Stat.* **67**, 417–434 (2018).
33. S. L. Choy, R. O'Leary, K. Mengersen, Elicitation by design in ecology: Using expert opinion to inform priors for Bayesian statistical models. *Ecology* **90**, 265–277 (2009).
34. C. M. Baker, M. Bode, N. Dexter, D. B. Lindenmayer, C. Foster, C. M. Gregor, M. Plein, E. McDonald-Madden, A novel approach to assessing the ecosystem-wide impacts of reintroductions. *Ecol. Appl.* **29**, e01811 (2019).
35. T. Cui, J. Martin, Y. M. Marzouk, A. Solonen, A. Spantini, Likelihood-informed dimension reduction for nonlinear inverse problems. *Inverse Probl.* **30**, 114015 (2014).
36. H. Hotelling, Analysis of a complex of statistical variables into principal components. *J. Educ. Psychol.* **24**, 417–441 (1933).
37. A. Spantini, A. Solonen, T. Cui, J. Martin, L. Tenorio, Y. Marzouk, Optimal low-rank approximations of Bayesian linear inverse problems. *SIAM J. Sci. Comput.* **37**, A2451–A2487 (2015).
38. I. T. Jolliffe, J. Cadima, Principal component analysis: A review and recent developments. *Philos. Trans. R. Soc. A* **374**, 20150202 (2016).
39. T. J. Rothenberg, Identification in parametric models. *Econometrica* **39**, 577–591 (1971).
40. F. P. Casey, D. Baird, Q. Feng, R. N. Gutenkunst, J. J. Waterfall, C. R. Myers, K. S. Brown, R. A. Cerione, J. P. Sethna, Optimal experimental design in an epidermal growth factor receptor signalling and down-regulation model. *IET Syst. Biol.* **1**, 190–202 (2007).
41. L. Michaelis, M. Menten, Die Kinetik der Invertinwirkung. *Biochem. Z.* **49**, 333 (1913).
42. R. P. Pech, G. M. Hood, Foxes, rabbits, alternative prey and rabbit calicivirus disease: Consequences of a new biological control agent for an outbreaking species in Australia. *J. Appl. Ecol.* **35**, 434–453 (1998).
43. G. W. Beeler, H. Reuter, Reconstruction of the action potential of ventricular myocardial fibres. *J. Physiol.* **268**, 177–210 (1977).
44. F.-G. Wieland, A. L. Hauber, M. Rosenblatt, C. Tönsing, J. Timmer, On structural and practical identifiability. *Curr. Opin. Syst. Biol.* **25**, 60–69 (2021).
45. B. Choi, G. A. Rempala, J. K. Kim, Beyond the Michaelis-Menten equation: Accurate and efficient estimation of enzyme kinetic parameters. *Sci. Rep.* **7**, 17018 (2017).
46. G. E. Briggs, J. B. Haldane, A note on the kinetics of enzyme action. *Biochem. J.* **19**, 338–339 (1925).
47. J. M. Tomczak, E. Węglarz-Tomczak, Estimating kinetic constants in the Michaelis-Menten model from one enzymatic assay using Approximate Bayesian Computation. *FEBS Lett.* **593**, 2742–2750 (2019).
48. T. Cui, Y. Marzouk, K. Willcox, Scalable posterior approximations for large-scale Bayesian inverse problems via likelihood-informed parameter and state reduction. *J. Comput. Phys.* **315**, 363–387 (2016).
49. T. Krogh-Madsen, D. J. Christini, *Modeling and Simulating Cardiac Electrical Activity* (IOP Publishing, 2020).
50. X. Zhou, A. Bueno-Orovio, B. Rodriguez, In silico evaluation of arrhythmia. *Curr. Opin. Phys.* **1**, 95–103 (2018).
51. O. J. Britton, A. Bueno-Orovio, K. Van Ammel, H. R. Lu, R. Towart, D. J. Gallacher, B. Rodriguez, Experimentally calibrated population of models predicts and explains intersubject variability in cardiac cellular electrophysiology. *Proc. Natl. Acad. Sci. U.S.A.* **110**, E2098–E2105 (2013).
52. E. Passini, O. J. Britton, H. R. Lu, J. Rohrbacher, A. N. Hermans, D. J. Gallacher, R. J. H. Greig, A. Bueno-Orovio, B. Rodriguez, Human in silico drug trials demonstrate higher accuracy than animal models in predicting clinical pro-arrhythmic cardiotoxicity. *Front. Physiol.* **8**, 668 (2017).
53. A. Muszkiewicz, O. J. Britton, P. Gemmell, E. Passini, C. Sánchez, X. Zhou, A. Carusi, T. A. Quinn, K. Burrage, A. Bueno-Orovio, B. Rodriguez, Variability in cardiac electrophysiology: Using experimentally-calibrated populations of models to move beyond the single virtual physiological human paradigm. *Prog. Biophys. Mol. Biol.* **120**, 115–127 (2016).
54. M. Zaniboni, I. Riva, F. Cacciani, M. Groppi, How different two almost identical action potentials can be: A model study on cardiac repolarization. *Math. Biosci.* **228**, 56–70 (2010).
55. M. P. Adams, E. J. Y. Koh, M. P. Vilas, C. J. Collier, V. M. Lambert, S. A. Sisson, M. Quiroz, E. McDonald-Madden, L. J. McKenzie, K. R. O'Brien, Predicting seagrass decline due to cumulative stressors. *Environ. Modelling Software* **130**, 104717 (2020).
56. T. Cui, X. Tong, O. Zahm, Prior normalization for certified likelihood-informed subspace detection of Bayesian inverse problems. arXiv:2202.00074 [math.NA] (31 January 2022).
57. D. Fernández Slezak, C. Suárez, G. A. Cecchi, G. Marshall, G. Stolovitzky, When the optimal is not the best: Parameter estimation in complex biological models. *PLOS ONE* **5**, e13283 (2010).
58. A. Doucet, S. Godsill, C. Andrieu, On sequential Monte Carlo sampling methods for Bayesian filtering. *Stat. Comput.* **10**, 197–208 (2000).
59. M. K. Transtrum, P. Qiu, Optimal experiment selection for parameter estimation in biological differential equation models. *BMC Bioinformatics* **13**, 181 (2012).
60. B. Schölkopf, A. Smola, K.-R. Müller, Nonlinear component analysis as a kernel eigenvalue problem. *Neural Comput.* **10**, 1299–1319 (1998).
61. T. Pavlenko, A. Björkström, A. Tillander, Covariance structure approximation via gLasso in high-dimensional supervised classification. *J. Appl. Stat.* **39**, 1643–1666 (2012).
62. C. Tönsing, J. Timmer, C. Kreutz, Cause and cure of sloppiness in ordinary differential equation models. *Phys. Rev. E* **90**, 023303 (2014).
63. S. Kleinegesse, C. Drovandi, M. U. Gutmann, Sequential Bayesian experimental design for implicit models via mutual information. *Bayesian Anal.* **3**, 773–802 (2021).
64. C. Beisbart, N. J. Saam, *Computer Simulation Validation: Fundamental Concepts, Methodological Frameworks, and Philosophical Perspectives* (Springer, 2019).
65. C. E. Dangerfield, D. Kay, K. Burrage, Modeling ion channel dynamics through reflected stochastic differential equations. *Phys. Rev. E* **85**, 051907 (2012).
66. J. Grazzini, M. Richiardi, Estimation of ergodic agent-based models by simulated minimum distance. *J. Econ. Dyn. Control.* **51**, 148–165 (2015).
67. S. A. Sisson, Y. Fan, M. A. Beaumont, *Handbook of Approximate Bayesian Computation* (Chapman & Hall/CRC Press, 2018).
68. L. F. Price, C. C. Drovandi, A. Lee, D. J. Nott, Bayesian synthetic likelihood. *J. Comput. Graph. Stat.* **27**, 1–11 (2018).
69. A. Beskos, A. Jasra, K. Law, Y. Marzouk, Y. Zhou, Multilevel sequential Monte Carlo with dimension-independent likelihood-informed proposals. *SIAM/ASA J. Uncertain. Quantif.* **6**, 762 (2018).
70. U. K. Müller, Measuring prior sensitivity and prior informativeness in large Bayesian models. *J. Monetary Econ.* **59**, 581–597 (2012).
71. E. J. Milner-Gulland, K. Shea, Embracing uncertainty in applied ecology. *J. Appl. Ecol.* **54**, 2063–2068 (2017).
72. W. K. Newey, D. McFadden, *Handbook of Econometrics* (Elsevier, 1994), vol. 4, pp. 2111–2245.
73. K. M. Banner, K. M. Irvine, T. J. Rodhouse, The use of Bayesian priors in ecology: The good, the bad and the not great. *Methods Ecol. Evol.* **11**, 882 (2020).
74. C. P. Robert, G. Casella, *Monte Carlo Statistical Methods* (Springer-Verlag, 1999).
75. M. I. Jordan, Z. Ghahramani, T. S. Jaakkola, L. K. Saul, An introduction to variational methods for graphical models. *Mach. Learn.* **37**, 183–233 (1999).
76. L. Tierney, J. B. Kadane, Accurate approximations for posterior moments and marginal densities. *J. Am. Stat. Assoc.* **81**, 82–86 (1986).
77. J. B. Tenenbaum, V. de Silva, J. C. Langford, A global geometric framework for nonlinear dimensionality reduction. *Science* **290**, 2319–2323 (2000).
78. T. Cui, K. J. H. Law, Y. M. Marzouk, Dimension-independent likelihood-informed MCMC. *J. Comput. Phys.* **304**, 109–137 (2016).
79. M. Evans, H. Moshonov, Checking for prior-data conflict. *Bayesian Anal.* **1**, 893 (2006).
80. S. G. Walker, Bayesian inference with misspecified models. *J. Stat. Planning Inference* **143**, 1621–1633 (2013).
81. J. D. Hoffman, S. Frankel, *Numerical Methods for Engineers and Scientists* (CRC Press, 2018).
82. H. M. Bücker, G. Corliss, P. Hovland, U. Naumann, B. Norris, *Automatic Differentiation: Applications, Theory, and Implementations* (Springer-Verlag, 2006), vol. 50.
83. P. G. Constantine, C. Kent, T. Bui-Thanh, Accelerating Markov chain Monte Carlo with active subspaces. *SIAM J. Sci. Comput.* **38**, A2779 (2016).
84. T. Cui, X. T. Tong, A unified performance analysis of likelihood-informed subspace methods. arXiv:2101.02417 [stat.CO] (7 January 2021).
85. O. Zahm, T. Cui, K. Law, A. Spantini, Y. Marzouk, Certified dimension reduction in nonlinear Bayesian inverse problems. arXiv:1807.03712 [math.PR] (2 July 2022).
86. A. W. Bowman, A. Azzalini, *Applied Smoothing Techniques for Data Analysis: The Kernel Approach with S-Plus Illustrations* (Oxford Univ. Press Inc., 1997), vol. 18.
87. M. C. Jones, Simple boundary correction for kernel density estimation. *Stat. Comput.* **3**, 135–146 (1993).
88. S. Dokos, *Modelling Organs, Tissues, Cells and Devices: Using MATLAB and COMSOL Multiphysics* (Springer, 2017).

Acknowledgments

Funding: This work has been supported by a Research Stimulus (RS) Postdoctoral Fellowship from the University of Queensland (to G.M.M.-B.), Australian Research Council (ARC) Centre of Excellence for Mathematical and Statistical Frontiers grant no. CE140100049 (to B.A.J.L.), ARC Centre of Excellence for Plant Success in Nature and Agriculture grant no. CE200100015 (to K.B.), NSF grant MCB no. 1715342 (to K.S.B.), ARC Linkage grant no. LP160100496 (to E.M.-M.),

and ARC Discovery Early Career Researcher Award no. DE200100683 (to M.P.A.). **Author contributions:** All authors contributed to the conceptualization, original draft, and review and editing of the manuscript. G.M.M.-B., B.A.J.L., C.D., K.B., and M.P.A. contributed to the design of numerical experiments. G.M.M.-B. and B.A.J.L. performed numerical experiments. All authors contributed to the analysis of results. **Competing interests:** The authors declare that they have no competing interests. **Data and materials availability:** All data needed to evaluate the conclusions in the paper are present in the paper and/or the Supplementary Materials.

MATLAB code associated with this paper is available from the Figshare Repository: <https://doi.org/10.6084/m9.figshare.19228008>.

Submitted 1 October 2021

Accepted 10 August 2022

Published 21 September 2022

10.1126/sciadv.abm5952