

Interacting Langevin Diffusions: Gradient Structure and Ensemble Kalman Sampler*

Alfredo Garbuno-Inigo[†], Franca Hoffmann[†], Wuchen Li[‡], and Andrew M. Stuart[†]

Abstract. Solving inverse problems without the use of derivatives or adjoints of the forward model is highly desirable in many applications arising in science and engineering. In this paper we propose a new version of such a methodology, a framework for its analysis, and numerical evidence of the practicality of the method proposed. Our starting point is an ensemble of overdamped Langevin diffusions which interact through a single preconditioner computed as the empirical ensemble covariance. We demonstrate that the nonlinear Fokker–Planck equation arising from the mean-field limit of the associated stochastic differential equation (SDE) has a novel gradient flow structure, built on the Wasserstein metric and the covariance matrix of the noisy flow. Using this structure, we investigate large time properties of the Fokker–Planck equation, showing that its invariant measure coincides with that of a single Langevin diffusion, and demonstrating exponential convergence to the invariant measure in a number of settings. We introduce a new noisy variant on ensemble Kalman inversion (EKI) algorithms found from the original SDE by replacing exact gradients with ensemble differences; this defines the ensemble Kalman sampler (EKS). Numerical results are presented which demonstrate its efficacy as a derivative-free approximate sampler for the Bayesian posterior arising from inverse problems.

Key words. ensemble Kalman inversion, Kalman–Wasserstein metric, gradient flow, mean-field Fokker–Planck equation

AMS subject classifications. 65N21, 35Q84, 62F15, 65C30, 65C35, 82C80

DOI. 10.1137/19M1251655

1. Problem setting.

1.1. Background. Consider the inverse problem of finding $u \in \mathbb{R}^d$ from $y \in \mathbb{R}^K$, where

$$(1.1) \quad y = \mathcal{G}(u) + \eta,$$

$\mathcal{G} : \mathbb{R}^d \rightarrow \mathbb{R}^K$ is a known nonlinear forward operator, and η is the unknown observational noise. Although η itself is unknown, we assume that it is drawn from a known probability

*Received by the editors March 25, 2019; accepted for publication (in revised form) by Y. Bakhtin October 10, 2019; published electronically February 4, 2020.

<https://doi.org/10.1137/19M1251655>

Funding: The work of the first and fourth authors was supported by the generosity of Eric and Wendy Schmidt by recommendation of the Schmidt Futures program, by Earthrise Alliance, by the Paul G. Allen Family Foundation, and by the National Science Foundation (NSF grant AGS-1835860). The work of the fourth author was also supported by NSF grant DMS-1818977. The work of the second author was partially supported by Caltech’s von Karman postdoctoral instructorship. The work of the third author was supported by AFOSR MURI FA9550-18-1-0502.

[†]Department of Computing and Mathematical Sciences, California Institute of Technology, Pasadena, CA 91125 (agarbuno@caltech.edu, <http://agarbuno.github.io/>; fkoh@caltech.edu, <https://francahoffmann.wordpress.com/>; astuart@caltech.edu, <http://stuart.caltech.edu/>).

[‡]Department of Mathematics, UCLA, Los Angeles, CA 90095 (wcli@math.ucla.edu, <https://www.math.ucla.edu/~wcli/>).

distribution; to be concrete, we assume that this distribution is a centered Gaussian: $\eta \sim \mathcal{N}(0, \Gamma)$ for a known covariance matrix $\Gamma \in \mathbb{R}^{K \times K}$. In summary, the objective of the inverse problem is to find information about the truth $u^\dagger$ underlying the data y ; the forward map $\mathcal{G}$, the covariance Γ , and the data y are all viewed as given.

A key role in any optimization scheme for solving (1.1) is played by $\ell(y, \mathcal{G}(u))$ for some loss function $\ell : \mathbb{R}^K \times \mathbb{R}^K \mapsto \mathbb{R}$. For additive Gaussian noise the natural loss function is¹

$$\ell(y, y') = \frac{1}{2} \|y - y'\|_\Gamma^2,$$

leading to the nonlinear least squares functional

$$(1.2) \quad \Phi(u) = \frac{1}{2} \|y - \mathcal{G}(u)\|_\Gamma^2.$$

In the Bayesian approach to inversion [46] we place a prior distribution on the unknown u , with Lebesgue density $\pi_0(u)$; then the posterior density on $u|y$, denoted $\pi(u)$, is given by

$$(1.3) \quad \pi(u) \propto \exp(-\Phi(u)) \pi_0(u).$$

In this paper we will concentrate on the case when the prior is a centered Gaussian $\mathcal{N}(0, \Gamma_0)$, assuming throughout that Γ_0 is strictly positive-definite and hence invertible. If we define

$$(1.4) \quad R(u) = \frac{1}{2} \|u\|_{\Gamma_0}^2$$

and

$$(1.5) \quad \Phi_R(u) = \Phi(u) + R(u),$$

then

$$(1.6) \quad \pi(u) \propto \exp(-\Phi_R(u)).$$

Note that the regularization $R(\cdot)$ is of Tikhonov–Phillips form [30].

Our focus throughout is on using interacting particle systems to approximate Langevin-type stochastic dynamical systems to sample from (1.6). Ensemble Kalman inversion (EKI), and variants of it, will be central in our approach because these methods play an important role in large-scale scientific and engineering applications in which it is undesirable, or impossible, to compute derivatives and adjoints defined by the forward map $\mathcal{G}$. Our goals are to introduce a noisy version of EKI which may be used to generate approximate samples from (1.6) based only on evaluations of $\mathcal{G}(u)$, to exemplify its potential use, and to provide a framework for its analysis. We refer to the new methodology as *ensemble Kalman sampling* (EKS).

¹For any positive-definite symmetric matrix A we define $\langle a, a' \rangle_A = \langle a, A^{-1} a' \rangle = \langle A^{-\frac{1}{2}} a, A^{-\frac{1}{2}} a' \rangle$ and $\|a\|_A = \|A^{-\frac{1}{2}} a\|$.

1.2. Literature review. The overdamped Langevin equation provides the simplest example of a reversible diffusion process with the property that it is invariant with respect to (1.6) [79]. It provides a conceptual starting point for a range of algorithms designed to draw approximate samples from the density (1.6). This idea may be generalized to nonreversible diffusions, such as those with state-dependent noise [27], those which are higher order in time [76], and combinations of these two [33]. In the case of higher order dynamics the desired target measure is found by marginalization. There are also a range of methods, often going under the collective names Nosé–Hoover–Poincaré, which identify the target measure as the marginal of an invariant measure induced by (ideally) chaotic and mixing deterministic dynamics [54] or a mixture of chaotic and stochastic dynamics [56]. Furthermore, the Langevin equation may be shown to govern the behavior of a wide range of Monte Carlo Markov chain (MCMC) methods; this work was initiated in the seminal paper [85] and has given rise to many related works [86, 87, 6, 5, 7, 66, 80, 76]; for a recent overview see [97].

In this paper we will introduce an interacting particle system generalization of the overdamped Langevin equation and use ideas from ensemble Kalman methodology to generate approximate solutions of the resulting stochastic flow, and hence approximate samples of (1.6), without computing derivatives of the log likelihood. The ensemble Kalman filter was originally introduced as a method for state estimation and later extended as the EKI to the solution of general inverse problems and parameter estimation problems. For a historical development of the subject, the reader may consult the books [32, 70, 63, 53, 84] and the recent review [11]. The Kalman filter itself was derived for linear Gaussian state estimation problems [47, 48]. In the linear setting, ensemble Kalman based methods may be viewed as Monte Carlo approximations of the Kalman filter; in the nonlinear case ensemble Kalman based methods do not converge to the filtering or posterior distribution in the large particle limit [31]. Related interacting particle based methodologies of current interest include Stein variational gradient descent [61, 60, 24] and the Fokker–Planck particle dynamics of Reich [83, 78], both of which map an arbitrary initial measure into the desired posterior measure over an infinite time horizon $s \in [0, \infty)$. A related approach is to introduce an artificial time $s \in [0, 1]$ and a homotopy between the prior at time $s = 0$ and the posterior measure at time $s = 1$ and write an evolution equation for the measures [20, 81, 28, 52]; this evolution equation can be approximated by particle methods. There are also other approaches in which optimal transport is used to evolve a sequence of particles through a transportation map [82, 65] to solve probabilistic state estimation or inversion problems as well as interacting particle systems designed to reproduce the solution of the filtering problem [19, 98]. The paper [23] studies ensemble Kalman filters from the perspective of the mean-field process and propagation of chaos. Also of interest are the consensus-based optimization techniques, which are given a rigorous treatment in [12].

The idea of using interacting particle systems derived from coupled Langevin-type equations is introduced within the context of MCMC methods in [55]; these methods require computation of derivatives of the log likelihood. In [26], concurrent with this paper, the interacting Langevin diffusions (2.3), (2.4) below are studied, with the goal being to demonstrate that the preconditioning removes slow relaxation rates when they are present in the standard Langevin equation (2.1); such a result is proved in the case when the potential Φ_R is quadratic and the posterior measure of interest is Gaussian. A key concept underlying both [55] and

[26] is the idea of finding algorithms which converge to equilibrium at rates independent of the conditioning of the Hessian of the log posterior, an idea introduced in the affine invariant samplers of Goodman and Weare [35].

Continuous-time limits of ensemble Kalman filters for state estimation were first introduced and studied systematically in the papers [10, 81, 8, 9]; the papers [8, 9] studied the “analysis” step of filtering (using Bayes’ theorem to incorporate data) through introduction of an artificial continuous time; the papers [10, 81] developed a seamless framework that integrated the true time for state evolution and the artificial time for incorporation of data into a modeling framework. The resulting methodology has been studied in a number of subsequent papers; see [22, 23, 21, 93] and references therein. A slightly different seamless continuous-time formulation was introduced and analyzed a few years later in [53, 49]. Continuous-time limits of ensemble methods for solving inverse problems were introduced and analyzed in the paper [89]; in fact, the work in the papers [8, 9] can be reinterpreted in the context of ensemble methods for inversion and also results in similar, but slightly different, continuous-time limits. The idea of iterating ensemble methods to solve inverse problems originated in the papers [16, 29], which focused on applications to oil reservoirs; the paper [43] describes, and demonstrates the promise of, the methods introduced in those papers for quite general inverse problems. The specific continuous-time version of the methodology, which we refer to as EKI in this paper, was identified in [89].

There has been significant activity devoted to the gradient flow structure associated with the Kalman filter itself. A well-known result is that for a constant state process, Kalman filtering is the gradient flow with respect to the Fisher–Rao metric [52, 38, 71]. It is worth noting that the Fisher–Rao metric connects to the covariance matrix; see details in [3]. On the other hand, optimal transport [96] demonstrates the importance of the L^2 -Wasserstein metric in probability density space. The space of densities equipped with this metric introduces an infinite-dimensional Riemannian manifold called the density manifold [51, 75, 57]. Solutions to the Fokker–Planck equation are gradient flows of the relative entropy in the density manifold [75, 45]. Designing time-stepping methods which preserve gradient structure is also of current interest; see [78] and, within the context of Wasserstein gradient flows, [59, 94, 58]. The subject of discrete gradients for time-integration of gradient and Hamiltonian systems is developed in [40, 34, 67, 37]. Furthermore, the papers [89, 88] study continuous-time limits of EKI algorithms and, in the case of linear inverse problems, exhibit a gradient flow structure for the standard least squares loss function, preconditioned by the empirical covariance of the particles; a related structure was highlighted in [8]. The paper [39], which has inspired aspects of our work, builds on the paper [89] to study the same problem in the mean-field limit; their mean-field perspective brings considerable insight, which we build upon in this paper. The recent work [25] studies the mean-field limit of EKI and makes a connection to the appropriate Fokker–Planck equation, whose solution characterizes the distribution in the mean-field limit.

In this paper, we study a new noisy version of EKI, called the ensemble Kalman sampler (EKS), and related mean-field limits, with the aim of constructing methods which lead to approximate posterior samples, without the use of adjoints, and overcoming the issue that the standard noisy EKI does not reproduce the posterior distribution, as highlighted in [31]. We emphasize that the practical derivative-free algorithm that we propose rests on a particle-based approximation of a specific preconditioned gradient flow, as described in section 4.3

of the paper [50]; we add a judiciously chosen noise to this setting, and it is this additional noise which enables approximate posterior sampling. Related approximations are also studied in the paper [78], in which the effects of both time-discretization and particle approximation are discussed when applied to various deterministic interacting particle systems with gradient structure. In order to frame the analysis of our methods, we introduce a new metric, named the Kalman–Wasserstein metric, defined through both the covariance matrix of the mean-field limit and the Wasserstein metric. The work builds on the novel perspectives introduced in [39] and leads to new algorithms that will be useful within large-scale parameter learning and uncertainty quantification studies, such as those proposed in [90]. Since the completion of this work, several papers have appeared studying related problems: in Nüsken and Reich [69] a correction for finite particle size is introduced; Ding and Li (<https://arxiv.org/abs/1910.12923>) study the limit of an infinite number of particles; and Carrillo and Vaes (<https://arxiv.org/abs/1910.07555>) study the stability of solutions with respect to perturbations in the initial condition.

1.3. Our contributions. The contributions in this paper are as follows:

- We introduce a new noisy perturbation of the continuous-time EKS algorithm, leading to an interacting particle system in stochastic differential equation (SDE) form, the EKS.
- We also introduce a related SDE, in which ensemble differences are approximated by gradients; this approximation is exact for linear inverse problems. We study the mean-field limit of this related SDE and exhibit a novel Kalman–Wasserstein gradient flow structure in the associated nonlinear Fokker–Planck equation.
- Using this Kalman–Wasserstein structure, we characterize the steady states of the nonlinear Fokker–Planck equation and show that one of them is the posterior density (1.6).
- By explicitly solving the nonlinear Fokker–Planck equation in the case of linear $\mathcal{G}$, we demonstrate that the posterior density is a global attractor for all initial densities of finite energy which are not a Dirac measure.
- We provide numerical examples which demonstrate that the EKS algorithm gives good approximate samples from the posterior distribution for both a simple low-dimensional test problem and a PDE inverse problem arising in Darcy flow.

In section 2 we introduce the various stochastic dynamical systems which form the basis of the proposed methodology and analysis: subsection 2.1 describes an interacting particle system variant on Langevin dynamics; subsection 2.2 recaps the EKI methodology and describes the SDE arising in the case when the data is perturbed with noise; and subsection 2.3 introduces the new noisy EKS algorithm, which arises from perturbing the particles with noise rather than perturbing the data. In section 3 we discuss the theoretical properties underpinning the proposed new methodology, and in section 4 we describe numerical results which demonstrate the value of the proposed new methodology. We conclude in section 5.

2. Dynamical systems setting. This section is devoted to the various noisy dynamical systems that underpin the paper: in the three constituent subsections we introduce an interacting particle version of Langevin dynamics, the EKI algorithm, and the new EKS algorithm. In doing so, we introduce a sequence of continuous-time problems that are designed to either

maximize the posterior distribution $\pi(u)$ (EKI) or generate approximate samples from the posterior distribution $\pi(u)$ (noisy EKI and the EKS). We then make a linear approximation within part of the EKS and take the mean-field limit leading to a novel nonlinear Fokker–Planck equation studied in section 3.

2.1. Variants on Langevin dynamics. The overdamped Langevin equation has the form

$$(2.1) \quad \dot{u} = -\nabla \Phi_R(u) + \sqrt{2} \dot{\mathbf{W}},$$

where $\mathbf{W}$ denotes a standard Brownian motion in $\mathbb{R}^d$.² References to the relevant literature may be found in the introduction. A common approach to speed up convergence is to introduce a symmetric matrix $\mathbf{C}$ in the corresponding gradient descent scheme,

$$(2.2) \quad \dot{u} = -\mathbf{C} \nabla \Phi_R(u) + \sqrt{2\mathbf{C}} \dot{\mathbf{W}}.$$

The key concept behind this stochastic dynamical system is that, under conditions on Φ_R which ensure ergodicity, *an arbitrary initial distribution is transformed into the desired posterior distribution over an infinite time horizon.*

Finding a suitable matrix $\mathbf{C} \in \mathbb{R}^{d \times d}$ is of general interest. We propose to evolve an interacting set of particles $U = \{u^{(j)}\}_{j=1}^J$ according to the following system of SDEs:

$$(2.3) \quad \dot{u}^{(j)} = -\mathbf{C}(U) \nabla \Phi_R(u^{(j)}) + \sqrt{2\mathbf{C}(U)} \dot{\mathbf{W}}^{(j)}.$$

Here, the $\{\mathbf{W}^{(j)}\}$ are a collection of independent and identically distributed (i.i.d.) standard Brownian motions in the space $\mathbb{R}^d$. The matrix $\mathbf{C}(U)$ depends nonlinearly on all ensemble members and is chosen to be the empirical covariance between particles,

$$(2.4) \quad \mathbf{C}(U) = \frac{1}{J} \sum_{k=1}^J (u^{(k)} - \bar{u}) \otimes (u^{(k)} - \bar{u}) \in \mathbb{R}^{d \times d},$$

where $\bar{u}$ denotes the sample mean

$$\bar{u} = \frac{1}{J} \sum_{j=1}^J u^{(j)}.$$

This choice of preconditioning is motivated by an underlying gradient flow structure, which we exhibit in subsection 3.3. System (2.3) can be rewritten as

$$(2.5) \quad \dot{u}^{(j)} = -\frac{1}{J} \sum_{k=1}^J \langle D\mathcal{G}(u^{(j)})(u^{(k)} - \bar{u}), \mathcal{G}(u^{(j)}) - y \rangle_{\Gamma} u^{(k)} - \mathbf{C}(U) \Gamma_0^{-1} u^{(j)} + \sqrt{2\mathbf{C}(U)} \dot{\mathbf{W}}^{(j)}.$$

(We used the fact that it is possible to replace $u^{(k)}$ by $u^{(k)} - \bar{u}$ after the Γ -weighted inner-product in (2.5) without changing the equation.) We will introduce the EKS as an ensemble Kalman based methodology to approximate this interacting particle system.

²In this SDE and all that follows, the rigorous interpretation is through the Itô integral formulation of the problem.

2.2. Ensemble Kalman inversion. To understand the EKS we first recall the EKI methodology, which can be interpreted as a derivative-free optimization algorithm to invert $\mathcal{G}$ [43, 42]. The continuous-time version of the algorithm is given by [89]

$$(2.6) \quad \dot{u}^{(j)} = -\frac{1}{J} \sum_{k=1}^J \langle \mathcal{G}(u^{(k)}) - \bar{\mathcal{G}}, \mathcal{G}(u^{(j)}) - y \rangle_{\Gamma} u^{(k)}.$$

This interacting particle dynamic acts to both drive particles towards consensus and fit the data. In [16, 29] the idea of using ensemble Kalman methods to map prior samples into posterior samples was introduced (see the introduction for a literature review). Interpreted in our continuous-time setting, the methodology operates by evolving a noisy set of interacting particles given by

$$(2.7) \quad \dot{u}^{(j)} = -\frac{1}{J} \sum_{k=1}^J \langle \mathcal{G}(u^{(k)}) - \bar{\mathcal{G}}, \mathcal{G}(u^{(j)}) - y \rangle_{\Gamma} u^{(k)} + C^{up}(U) \Gamma^{-1} \sqrt{\Sigma} \dot{\mathbf{W}}^{(j)},$$

where the $\{\mathbf{W}^{(j)}\}$ are a collection of i.i.d. standard Brownian motions in the data space $\mathbb{R}^K$; different choices of Σ allow us to remove noise and obtain an optimization algorithm ($\Sigma = 0$) or to add noise in a manner which, for linear problems, creates a dynamic transporting the prior into the posterior in one time unit ($\Sigma = \Gamma$; see discussion below).

Here, the operator C^{up} denotes the empirical cross covariance matrix of the ensemble members,

$$(2.8) \quad C^{up}(U) := \frac{1}{J} \sum_{k=1}^J (u^{(k)} - \bar{u}) \otimes (\mathcal{G}(u^{(k)}) - \bar{\mathcal{G}}) \in \mathbb{R}^{d \times K}, \quad \bar{\mathcal{G}} := \frac{1}{J} \sum_{k=1}^J \mathcal{G}(u^{(k)}).$$

The approach is designed in the linear case to *transform prior samples into posterior samples in one time unit* [16]. In contrast to Langevin dynamics, this has the desirable property that it works over a single time unit rather than over an infinite time horizon. But it is considerably more rigid, as it requires initialization at the prior. Furthermore, the long time dynamics do not have the desired sampling property but rather collapse to a single point, solving the optimization problem of minimizing $\Phi(u)$. We now demonstrate these points by considering the linear problem.

To be explicit we consider the case when

$$(2.9) \quad \mathcal{G}(u) = Au.$$

In this case, the regularized misfit equals

$$(2.10) \quad \Phi_R(u) = \frac{1}{2} \|Au - y\|_{\Gamma}^2 + \frac{1}{2} \|u\|_{\Gamma_0}^2.$$

The corresponding gradient can be written as

$$(2.11) \quad \nabla \Phi_R(u) = B^{-1}u - r, \\ r := A^{\top} \Gamma^{-1} y \in \mathbb{R}^d, \quad B := \left(A^{\top} \Gamma^{-1} A + \Gamma_0^{-1} \right)^{-1} \in \mathbb{R}^{d \times d}.$$

The posterior mean is thus Br , and the posterior covariance is B .

In the linear setting (2.9) and with the choice $\Sigma = \Gamma$, the EKI algorithm defined in (2.7) has mean $\mathbf{m}$ and covariance $\mathfrak{C}$, which satisfy the closed equations

$$(2.12a) \quad \frac{d}{dt}\mathbf{m}(t) = -\mathfrak{C}(t)(A^\top \Gamma^{-1} A \mathbf{m}(t) - r),$$

$$(2.12b) \quad \frac{d}{dt}\mathfrak{C}(t) = -\mathfrak{C}(t)A^\top \Gamma^{-1} A \mathfrak{C}(t).$$

These results may be established by techniques similar to those used below in subsection 3.2. (A more general analysis of the SDE (2.7), and its related nonlinear Fokker–Planck equation, is undertaken in [25].) It follows that

$$\frac{d}{dt}\mathfrak{C}(t)^{-1} = -\mathfrak{C}(t)^{-1} \left(\frac{d}{dt}\mathfrak{C}(t) \right) \mathfrak{C}(t)^{-1} = A^\top \Gamma^{-1} A,$$

and therefore $\mathfrak{C}(t)^{-1}$ grows linearly in time. If the initial covariance is given by the prior Γ_0 , then

$$\mathfrak{C}(t)^{-1} = \Gamma_0^{-1} + A^\top \Gamma^{-1} A t,$$

demonstrating that $\mathfrak{C}(1)$ delivers the posterior covariance; furthermore, it then follows that

$$\frac{d}{dt}\{\mathfrak{C}(t)^{-1}\mathbf{m}(t)\} = r$$

so that, initializing with prior mean $\mathbf{m}(0) = 0$, we obtain

$$\mathbf{m}(t) = \left(\Gamma_0^{-1} + A^\top \Gamma^{-1} A t \right)^{-1} r t,$$

and $\mathbf{m}(1)$ delivers the posterior mean.

The resulting equations for the mean and covariance are simply those which arise from applying the Kalman–Bucy filter [48] to the model

$$\begin{aligned} \frac{d}{dt}u &= 0, \\ \frac{d}{dt}z &:= y = Au + \sqrt{\Gamma}\dot{\mathbf{W}}, \end{aligned}$$

where $\mathbf{W}$ denotes a standard unit Brownian motion in the data space $\mathbb{R}^K$. The exact closed form of equations for the first two moments, in the setting of the Kalman–Bucy filter, was established in section 4 of the paper [81] for finite particle approximations, and it transfers verbatim to this mean-field setting.

The analysis reveals interesting behavior in the large time limit: the covariance shrinks to zero, and the mean converges to the solution of the unregularized least squares problem; we thus have ensemble collapse and solution of an optimization problem rather than a sampling problem. This highlights an interesting perspective on the EKI, namely as an optimization method rather than a sampling method. A key point to appreciate is that the noise introduced in (2.7) arises from the observation y being perturbed with additional noise. In what follows we instead directly perturb the particles themselves. The benefit of introducing noise on the particles, rather than the data, was demonstrated in [50], although in that setting only optimization, and not Bayesian inversion, is considered.

2.3. The ensemble Kalman sampler. We now demonstrate how to introduce noise on the particles within the ensemble Kalman methodology, with our starting point being (2.5). This gives the EKS. In contrast to the standard noisy EKI (2.7), the EKS is based on a dynamic which *transforms an arbitrary initial distribution into the desired posterior distribution, over an infinite time horizon*. In many applications, derivatives of the forward map $\mathcal{G}$ are either not available or extremely costly to obtain. A common technique used in ensemble Kalman methods is to approximate the gradient $\nabla\Phi_R$ by differences in order to obtain a derivative-free algorithm for inverting $\mathcal{G}$. To this end, consider the dynamical system (2.5) and invoke the approximation

$$D\mathcal{G}(u^{(j)})(u^{(k)} - \bar{u}) \approx (\mathcal{G}(u^{(k)}) - \bar{\mathcal{G}}).$$

This leads to the following derivative-free algorithm to generate approximate samples from the posterior distribution:

$$(2.13) \quad \dot{u}^{(j)} = -\frac{1}{J} \sum_{k=1}^J \langle \mathcal{G}(u^{(k)}) - \bar{\mathcal{G}}, \mathcal{G}(u^{(j)}) - y \rangle_{\Gamma} u^{(k)} - \mathbf{C}(U) \Gamma_0^{-1} u^{(j)} + \sqrt{2\mathbf{C}(U)} \dot{\mathbf{W}}^{(j)}.$$

This dynamical system is similar to the noisy EKI (2.7) but has a different noise structure (noise in parameter space and not in data space) and explicitly accounts for the prior on the right-hand side (rather than having it enter through initialization). Inclusion of the Tikhonov regularization term within EKI is introduced and studied in [15].

Note that in the linear case (2.9) the two systems (2.5) and (2.13) are identical. It is also natural to conjecture that if the particles are close to one another, then (2.5) and (2.13) will generate similar particle distributions. Based on this exact (in the linear case) and conjectured (in the nonlinear case) relationship, we propose (2.13) as a derivative-free algorithm to approximately sample the Bayesian posterior distribution, and we propose (2.5) as a natural object of analysis in order to understand this sampling algorithm.

2.4. Mean-field limit. In order to write down the mean-field limit of (2.5), we define the macroscopic mean and covariance:

$$m(\rho) := \int v \rho \, dv, \quad \mathcal{C}(\rho) := \int (v - m(\rho)) \otimes (v - m(\rho)) \rho(v) \, dv.$$

Taking the large particle limit leads to the mean-field equation

$$(2.14) \quad \dot{u} = -\mathcal{C}(\rho) \nabla \Phi_R(u) + \sqrt{2\mathcal{C}(\rho)} \dot{W},$$

with corresponding nonlinear Fokker–Planck equation

$$(2.15) \quad \partial_t \rho = \nabla \cdot (\rho \mathcal{C}(\rho) \nabla \Phi_R(u)) + \mathcal{C}(\rho) : D^2 \rho.$$

Here $A_1 : A_2$ denotes the Frobenius inner-product between matrices A_1 and A_2 . The existence and form of the mean-field limit are suggested by the exchangeability of the process (existence) and by application of the law of large numbers (form). Exchangeability is exploited in a related context in [22, 23]. The rigorous derivation of the mean-field equations (2.14) and (2.15) is left for future work; for foundational work relating to mean-field limits, see [92, 44, 13, 36, 77, 95] and references therein. The following lemma states the intuitive fact that the covariance, which plays a central role in (2.15), vanishes only for Dirac measures.

Lemma 2.1. *The only probability densities $\rho \in \mathcal{P}(\mathbb{R}^d)$ at which $\mathcal{C}(\rho)$ vanishes are Diracs:*

$$\rho(u) = \delta_v(u) \text{ for some } v \in \mathbb{R}^d \quad \Leftrightarrow \quad \mathcal{C}(\rho) = 0.$$

Proof. The fact that $\mathcal{C}(\delta_v) = 0$ follows by direct substitution. For the converse, note that $\mathcal{C}(\rho) = 0$ implies $\int |u|^2 \rho \, du = \left(\int u \rho \, du \right)^2$, which is the equality case of Jensen's inequality, and therefore only holds if ρ is the law of a constant random variable. ■

3. Theoretical properties. In this section we discuss theoretical properties of (2.15) which motivate the use of (2.5) and (2.13) as particle systems to generate approximate samples from the posterior distribution (1.6). In subsection 3.1 we exhibit a gradient flow structure for (2.15) which shows that solutions evolve towards the posterior distribution (1.6) unless they collapse to a Dirac measure. In subsection 3.2 we show that in the linear case, collapse to a Dirac does not occur if the initial condition is a Gaussian with nonzero covariance, and instead convergence to the posterior distribution is obtained. In subsection 3.3 we introduce a novel metric structure which underpins the results in subsections 3.1 and 3.2.

3.1. Nonlinear problem. Because $\mathcal{C}(\rho)$ is independent of u , we may write (2.15) in divergence form, which facilitates the revelation of a gradient structure,

$$(3.1) \quad \partial_t \rho = \nabla \cdot (\rho \mathcal{C}(\rho) \nabla \Phi_R(u) + \rho \mathcal{C}(\rho) \nabla \ln \rho),$$

where we use the fact that $\rho \nabla \ln \rho = \nabla \rho$. Indeed, (3.1) is nothing but the Fokker–Planck equation for (2.2) for a time-dependent matrix $C(t) = \mathcal{C}(\rho)$. Thanks to the divergence form, it follows that (3.1) conserves mass along the flow, and so we may assume $\int \rho(t, u) \, du = 1$ for all $t \geq 0$. Defining the energy

$$(3.2) \quad E(\rho) = \int (\rho(u) \Phi_R(u) + \rho(u) \ln \rho(u)) \, du,$$

solutions to (3.1) can be written as a gradient flow,

$$(3.3) \quad \partial_t \rho = \nabla \cdot \left(\rho \mathcal{C}(\rho) \nabla \frac{\delta E}{\delta \rho} \right),$$

where $\frac{\delta}{\delta \rho}$ denotes the L^2 first variation. This will be made more explicit in subsection 3.3: see Proposition 3.6. Thanks to the gradient flow structure (3.3), stationary states of (2.15) are given either by critical points of the energy E or by choices of ρ such that $\mathcal{C}(\rho) = 0$ as characterized in Lemma 2.1. Critical points of E solve the corresponding Euler–Lagrange condition

$$(3.4) \quad \frac{\delta E}{\delta \rho} = \Phi_R(u) + \ln \rho(u) = c \quad \text{on supp } (\rho)$$

for some constant c . The unique solution to (3.4) with unit mass is given by the Gibbs measure

$$(3.5) \quad \rho_\infty(u) := \frac{e^{-\Phi_R(u)}}{\int e^{-\Phi_R(u)} \, du}.$$

Then, up to an additive normalization constant, the energy $E(\rho)$ is exactly the following relative entropy of ρ with respect to ρ_∞ , also known as the Kullback–Leibler divergence $\text{KL}(\rho(t)\|\rho_\infty)$,

$$\begin{aligned} E(\rho) &= \int (\Phi_R + \ln \rho(t)) \rho \, du \\ &= \int \frac{\rho(t)}{\rho_\infty} \ln \left(\frac{\rho(t)}{\rho_\infty} \right) \rho_\infty \, du + \ln \left(\int e^{-\Phi_R(u)} \, du \right) \\ &= \text{KL}(\rho(t)\|\rho_\infty) + \ln \left(\int e^{-\Phi_R(u)} \, du \right). \end{aligned}$$

Thanks to the gradient flow structure (3.3), we can compute the dissipation of the energy

$$\begin{aligned} \frac{d}{dt} \{E(\rho)\} &= \left\langle \frac{\delta E}{\delta \rho}, \partial_t \rho \right\rangle_{L^2(\mathbb{R}^d)} \\ (3.6) \quad &= - \int \rho \left\langle \nabla \frac{\delta E}{\delta \rho}, \mathcal{C}(\rho) \nabla \frac{\delta E}{\delta \rho} \right\rangle \, du \\ &= - \int \rho \left| \mathcal{C}(\rho)^{\frac{1}{2}} \nabla (\Phi_R + \ln \rho) \right|^2 \, du. \end{aligned}$$

As a consequence, the energy E decreases along trajectories until either $\mathcal{C}(\rho)$ approaches zero (that is, ρ collapses to a Dirac measure by Lemma 2.1) or ρ becomes the Gibbs measure with density ρ_∞ .

The dissipation of the energy along the evolution of the classical Fokker–Planck equation is known as the Fisher information [96]. We reformulate (3.6) by defining the following generalized Fisher information for any covariance matrix Λ :

$$\mathcal{I}_\Lambda(\rho(t)\|\rho_\infty) := \int \rho \left\langle \nabla \ln \left(\frac{\rho}{\rho_\infty} \right), \Lambda \nabla \ln \left(\frac{\rho}{\rho_\infty} \right) \right\rangle \, du.$$

One may also refer to $\mathcal{I}_\Lambda$ as a Dirichlet form as it is known in the theory of large particle systems, since we can write

$$\mathcal{I}_\Lambda(\rho(t)\|\rho_\infty) = 4 \int \rho_\infty \left\langle \nabla \sqrt{\frac{\rho}{\rho_\infty}}, \Lambda \nabla \sqrt{\frac{\rho}{\rho_\infty}} \right\rangle \, du.$$

For $\Lambda = \mathcal{C}(\rho)$, we name functional $\mathcal{I}_\mathcal{C}$ the relative *Kalman–Fisher information*. We conclude that the following energy dissipation equality holds:

$$\frac{d}{dt} \text{KL}(\rho(t)\|\rho_\infty) = -\mathcal{I}_\mathcal{C}(\rho(t)\|\rho_\infty).$$

To derive a rate of decay to equilibrium in entropy, we aim to identify conditions on Φ_R such that the following logarithmic Sobolev inequality holds: there exists $\lambda > 0$ such that

$$(3.7) \quad \text{KL}(\rho(t)\|\rho_\infty) \leq \frac{1}{2\lambda} \mathcal{I}_{I_d}(\rho(t)\|\rho_\infty) \quad \forall \rho.$$

By [4], it is enough to impose sufficient convexity on Φ_R , i.e., $D^2\Phi_R \geq \lambda I_d$, where $D^2\Phi_R$ denotes the Hessian of Φ_R . This allows us to deduce convergence to equilibrium as long as $\mathcal{C}(\rho)$ is uniformly bounded from below following standard arguments for the classical Fokker–Planck equation as presented, for example, in [64].

Proposition 3.1. *Assume there exists $\alpha > 0$ and $\lambda > 0$ such that*

$$\mathcal{C}(\rho(t)) \geq \alpha I_d, \quad D^2\Phi_R \geq \lambda I_d.$$

Then any solution $\rho(t)$ to (3.1) with initial condition ρ_0 satisfying $\text{KL}(\rho_0 \|\rho_\infty) < \infty$ decays exponentially fast to equilibrium: there exists a constant $c = c(\rho_0, \Phi_R) > 0$ such that for any $t > 0$,

$$\|\rho(t) - \rho_\infty\|_{L^1(\mathbb{R}^d)} \leq ce^{-\alpha\lambda t}.$$

This rate of convergence can most likely be improved using the correct logarithmic Sobolev inequality weighted by the covariance matrix $\mathcal{C}$. However, the above estimate already indicates the effect of having the covariance matrix $\mathcal{C}$ present in the Fokker–Planck equation (3.1). The properties of such inequalities in a more general setting are an interesting future avenue to explore. The weighted logarithmic Sobolev inequality that is well adapted to the setting here depends on the geometric structure of the Kalman–Wasserstein metric; see related studies in [57].

Proof. Thanks to the assumptions, and using the logarithmic Sobolev inequality (3.7), we obtain decay in entropy,

$$\frac{d}{dt} \text{KL}(\rho(t) \|\rho_\infty) \leq -\alpha \mathcal{I}_{I_d}(\rho(t) \|\rho_\infty) \leq -2\alpha\lambda \text{KL}(\rho(t) \|\rho_\infty).$$

We conclude using the Csiszár–Kullback inequality as it is mainly known to analysts, also referred to as the Pinsker inequality in probability (see [2] for more details):

$$\frac{1}{2} \|\rho(t) - \rho_\infty\|_{L^1(\mathbb{R}^d)}^2 \leq \text{KL}(\rho(t) \|\rho_\infty) \leq \text{KL}(\rho_0 \|\rho_\infty) e^{-2\alpha\lambda t}. \quad \blacksquare$$

3.2. Linear problem. Here we show that in the case of a linear forward operator $\mathcal{G}$, the Fokker–Planck equation (which is still nonlinear) has exact Gaussian solutions. This property may be seen to hold in two ways: (i) by considering the case in which the covariance matrix is an exogenously defined function of time alone and in which the observation is straightforward; and (ii) because the mean-field equation (2.14) leads to exact closed equations for the mean and covariance. Once the covariance is known, the nonlinear Fokker–Planck equation (2.15) becomes linear and is explicitly solvable if $\mathcal{G}$ is linear and the initial condition is Gaussian. Consider (2.14) in the context of a linear observation map (2.9). The misfit is given by (2.10), and the gradient of Φ_R is given in (2.11). Note that since we assume that the covariance matrix Γ_0 is invertible, it is then also strictly positive-definite. Thus it follows that B is strictly positive-definite and hence invertible too. We define $u_0 := Br$, noting that this is the solution of the regularized normal equations defining the minimizer of Φ_R in this linear case;

equivalently, u_0 maximizes the posterior density. Indeed, by completing the square we see that we may write

$$(3.8) \quad \rho_\infty(u) \propto \exp\left(-\frac{1}{2}\|u - u_0\|_B^2\right).$$

Lemma 3.2. *Let $\rho(t)$ be a solution of (2.15) with $\Phi_R(\cdot)$ given by (2.10). Then the mean $m(\rho)$ and covariance matrix $\mathcal{C}(\rho)$ are determined by $\mathbf{m}(t)$ and $\mathfrak{C}(t)$ which satisfy the evolution equations*

$$(3.9a) \quad \frac{d}{dt}\mathbf{m}(t) = -\mathfrak{C}(t)(B^{-1}\mathbf{m}(t) - r),$$

$$(3.9b) \quad \frac{d}{dt}\mathfrak{C}(t) = -2\mathfrak{C}(t)B^{-1}\mathfrak{C}(t) + 2\mathfrak{C}(t).$$

In addition, for any $\mathfrak{C}(t)$ satisfying (3.9b), its determinant and inverse solve

$$(3.10) \quad \frac{d}{dt} \det \mathfrak{C}(t) = -2 (\det \mathfrak{C}(t)) \operatorname{Tr} [B^{-1}\mathfrak{C}(t) - I_d],$$

$$(3.11) \quad \frac{d}{dt}(\mathfrak{C}(t)^{-1}) = 2B^{-1} - 2\mathfrak{C}(t)^{-1}.$$

As a consequence, $\mathfrak{C}(t) \rightarrow B$ and $\mathbf{m}(t) \rightarrow u_0$ exponentially as $t \rightarrow \infty$.

In fact, solving the ODE (3.11) explicitly and using (3.9a), we see that exponential decay immediately follows as

$$(3.12) \quad \mathfrak{C}(t)^{-1} = (\mathfrak{C}(0)^{-1} - B^{-1})e^{-2t} + B^{-1}$$

and

$$(3.13) \quad \|\mathbf{m}(t) - u_0\|_{\mathfrak{C}(t)} = \|\mathbf{m}(0) - u_0\|_{\mathfrak{C}(0)}e^{-t}.$$

Proof. We begin by deriving the evolution of the first and second moments. This is most easily accomplished by working with the mean-field flow SDE (2.14), using the regularized linear misfit written in (2.10). This yields the update

$$\dot{u} = -\mathcal{C}(\rho)(B^{-1}u - r) + \sqrt{2\mathcal{C}(\rho)}\dot{\mathbf{W}},$$

where $\dot{\mathbf{W}}$ denotes a zero mean random variable. Identical results can be obtained by working directly with the PDE for the density, namely (2.15) with the regularized linear misfit given in (2.10). Taking expectations with respect to ρ results in

$$\dot{m}(\rho) = -\mathcal{C}(\rho)(B^{-1}m(\rho) - r).$$

Let us use the auxiliary variable $e = u - m(\rho)$. By linearity of differentiation we can write

$$\dot{e} = -\mathcal{C}(\rho)B^{-1}e + \sqrt{2\mathcal{C}(\rho)}\dot{\mathbf{W}}.$$

By definition of the covariance operator, $\mathcal{C}(\rho) = \mathbb{E}[e \otimes e]$, its derivative with respect to time can be written as

$$\dot{\mathcal{C}}(\rho) = \mathbb{E}[\dot{e} \otimes e + e \otimes \dot{e}].$$

However, we must also include the Itô correction, using Itô's formula, and we can write the evolution equation of the covariance operator as

$$\dot{\mathcal{C}}(\rho) = -2\mathcal{C}(\rho)B^{-1}\mathcal{C}(\rho) + 2\mathcal{C}(\rho).$$

This concludes the proof of (3.9b). For the evolution of the determinant and inverse, note that

$$\frac{d}{dt} \det \mathcal{C}(\rho) = \text{Tr} \left[\det \mathcal{C}(\rho) \mathcal{C}(\rho)^{-1} \frac{d}{dt} \mathcal{C}(\rho) \right], \quad \frac{d}{dt} \mathcal{C}(\rho)^{-1} = -\mathcal{C}(\rho)^{-1} \left(\frac{d}{dt} \mathcal{C}(\rho) \right) \mathcal{C}(\rho)^{-1},$$

and so (3.10), (3.11) directly follow. Finally, exponential decay is a consequence of the explicit expressions (3.12) and (3.13). ■

Thanks to the evolution of the covariance matrix and its determinant, we can deduce that there is a family of Gaussian initial conditions that stay Gaussian along the flow and converge to the equilibrium ρ_∞ .

Proposition 3.3. *Fix a vector $m_0 \in \mathbb{R}^d$ and a matrix $\mathcal{C}_0 \in \mathbb{R}^{d \times d}$ and take as initial density the Gaussian distribution*

$$\rho_0(u) := \frac{1}{(2\pi)^{d/2}} (\det \mathcal{C}_0)^{-1/2} \exp \left(-\frac{1}{2} \|u - m_0\|_{\mathcal{C}_0}^2 \right)$$

with mean m_0 and covariance $\mathcal{C}_0$. Then the Gaussian profile

$$\rho(t, u) := \frac{1}{(2\pi)^{d/2}} (\det \mathfrak{C}(t))^{-1/2} \exp \left(-\frac{1}{2} \|u - \mathbf{m}(t)\|_{\mathfrak{C}(t)}^2 \right)$$

solves evolution equation (2.15) with initial condition $\rho(0, u) = \rho_0(u)$, and where $\mathbf{m}(t)$ and $\mathfrak{C}(t)$ evolve according to (3.9a) and (3.9b) with initial conditions m_0 and $\mathcal{C}_0$. As a consequence, for such initial conditions $\rho_0(u)$, the solution of the Fokker–Planck equation (2.15) converges to $\rho_\infty(u)$ given by (3.8) as $t \rightarrow \infty$.

Proof. It is straightforward to see that for $m(\rho)$ and $\mathcal{C}(\rho)$ given by Lemma 3.2,

$$\nabla \rho = -\mathcal{C}(\rho)^{-1}(u - m(\rho)) \rho,$$

since both $m(\rho)$ and $\mathcal{C}(\rho)$ are independent of u . Therefore, substituting the Gaussian ansatz $\rho(t, u)$ into the first term in the right-hand side of (2.15), we have

$$\begin{aligned} \nabla \cdot (\rho \mathcal{C}(\rho)(B^{-1}u - r)) &= (\nabla \rho) \cdot \mathcal{C}(\rho)(B^{-1}u - r) + \rho \nabla \cdot (\mathcal{C}(\rho)B^{-1}u) \\ &= (-\mathcal{C}(\rho)^{-1}(u - m(\rho)) \cdot \mathcal{C}(\rho)(B^{-1}u - r) + \text{Tr}[\mathcal{C}(\rho)B^{-1}]) \rho \\ (3.14) \quad &= \left(-\|u - m(\rho)\|_B^2 + \left\langle u - m(\rho), u_0 - m(\rho) \right\rangle_B + \text{Tr}[\mathcal{C}(\rho)B^{-1}] \right) \rho, \end{aligned}$$

where $B^{-1} = A^\top \Gamma^{-1} A + \Gamma_0^{-1}$, $r = A^\top \Gamma^{-1} y$, and $u_0 = B r$. Recall that B^{-1} is invertible. The second term on the right-hand side of (2.15) can be simplified as follows:

$$(3.15) \quad \begin{aligned} \mathcal{C}(\rho) : D^2 \rho &= \mathcal{C}(\rho) : \left(-\mathcal{C}(\rho)^{-1} + (\mathcal{C}(\rho)^{-1}(u - m(\rho))) \otimes (\mathcal{C}(\rho)^{-1}(u - m(\rho))) \right) \rho \\ &= \left(-\text{Tr}[I_d] + \|u - m(\rho)\|_{\mathcal{C}(\rho)}^2 \right) \rho. \end{aligned}$$

Thus, combining the previous two equations, we see that the right-hand side of (2.15) is given by the following expression:

$$(3.16) \quad \left[\text{Tr}[B^{-1}\mathcal{C}(\rho) - I_d] - \|u - m(\rho)\|_B^2 + \|u - m(\rho)\|_{\mathcal{C}(\rho)}^2 + \left\langle u - m(\rho), u_0 - m(\rho) \right\rangle_B \right] \rho.$$

For the left-hand side of (2.15), note that by (3.9a) and (3.9b),

$$\begin{aligned} \frac{d}{dt} \|u - m(\rho)\|_{\mathcal{C}(\rho)}^2 &= 2 \left\langle \frac{d}{dt}(u - m(\rho)), \mathcal{C}(\rho)^{-1}(u - m(\rho)) \right\rangle \\ &\quad + \left\langle (u - m(\rho)), \frac{d}{dt}(\mathcal{C}(\rho)^{-1})(u - m(\rho)) \right\rangle \\ &= -2 \left\langle u - m(\rho), u_0 - m(\rho) \right\rangle_B + 2 \|u - m(\rho)\|_B^2 - 2 \|u - m(\rho)\|_{\mathcal{C}(\rho)}^2, \end{aligned}$$

and therefore, combining with (3.10), we obtain

$$(3.17) \quad \begin{aligned} \partial_t \rho &= \left[-\frac{1}{2} (\det \mathcal{C}(\rho))^{-1} \left(\frac{d}{dt} \det \mathcal{C}(\rho) \right) - \frac{1}{2} \frac{d}{dt} \|u - u_0\|_{\mathcal{C}(\rho)}^2 \right] \rho \\ &= \left[\text{Tr}[B^{-1}\mathcal{C}(\rho) - I_d] - \|u - m(\rho)\|_B^2 + \|u - m(\rho)\|_{\mathcal{C}(\rho)}^2 + \left\langle u - m(\rho), u_0 - m(\rho) \right\rangle_B \right] \rho, \end{aligned}$$

which concludes the first part of the proof. The second part concerning the large time asymptotics is a straightforward consequence of the asymptotic behavior of $\mathfrak{m}$ and $\mathfrak{C}$ detailed in Lemma 3.2. ■

In the case of the classical Fokker–Planck equation $\mathfrak{C}(t) = I_d$ with a quadratic confining potential, the result in Proposition 3.3 follows from the fact that the fundamental solution of (2.15) is a Gaussian; see [14].

Corollary 3.4. *Let ρ_0 be a non-Gaussian initial condition for (2.15) in the case when Φ_R is given by (2.10). Assume that ρ_0 satisfies $\text{KL}(\rho_0 \| \rho_\infty) < \infty$. Then any solution of (2.15) converges exponentially fast to ρ_∞ given by (3.5) as $t \rightarrow \infty$ both in entropy and in $L^1(\mathbb{R}^d)$.*

Proof. Let $a \in \mathbb{R}^d$ have Euclidean norm 1, and define $q(t) := \langle a, \mathfrak{C}(t)^{-1} a \rangle$. From (3.11) it follows that

$$\dot{q} \leq 2\lambda - 2q,$$

where λ is the maximum eigenvalue of B^{-1} . Hence it follows that q is bounded above, independently of a , and that $\mathfrak{C}$ is bounded from below as an operator. Together with the fact that the Hessian $D^2 \Phi_R = B^{-1}$ is bounded from below, we conclude using Proposition 3.1. ■

3.3. Kalman–Wasserstein gradient flow. We introduce an infinite-dimensional Riemannian metric structure, which we name the Kalman–Wasserstein metric, in density space. It allows the interpretation of solutions to (2.15) as gradient flows in density space. To this end we denote by $\mathcal{P}$ the space of probability measures on a convex set $\Omega \subseteq \mathbb{R}^d$:

$$\mathcal{P} := \left\{ \rho \in L^1(\Omega) : \rho \geq 0 \text{ a.e.}, \int \rho(x) dx = 1 \right\}.$$

The probability simplex $\mathcal{P}$ is a manifold with boundary. For simplicity, we focus on the subset

$$\mathcal{P}_+ := \{ \rho \in \mathcal{P} : \rho > 0 \text{ a.e.}, \rho \in C^\infty(\Omega) \}.$$

The tangent space of $\mathcal{P}_+$ at a point $\rho \in \mathcal{P}_+$ is given by

$$\begin{aligned} T_\rho \mathcal{P}_+ &= \left\{ \left. \frac{d}{dt} \rho(t) \right|_{t=0} : \rho(t) \text{ is a curve in } \mathcal{P}_+, \rho(0) = \rho \right\} \\ &= \left\{ \sigma \in C^\infty(\Omega) : \int \sigma dx = 0 \right\}. \end{aligned}$$

The second equality follows since for all $\sigma \in T_\rho \mathcal{P}_+$, we have $\int \sigma(x) dx = 0$ as the mass along all curves in $\mathcal{P}_+$ remains constant. For the set $\mathcal{P}_+$, the tangent space $T_\rho \mathcal{P}_+$ is therefore independent of the point $\rho \in \mathcal{P}_+$. Cotangent vectors are elements of the topological dual $T_\rho^* \mathcal{P}_+$ and can be identified with tangent vectors via the action of the *Onsager operator* [68, 72, 73, 62, 74]

$$V_{\rho, \mathcal{C}} : T_\rho^* \mathcal{P}_+ \rightarrow T_\rho \mathcal{P}_+.$$

In this paper, we introduce the following new choice of Onsager operator:

$$(3.18) \quad V_{\rho, \mathcal{C}}(\phi) = -\nabla \cdot (\rho \mathcal{C}(\rho) \nabla \phi) =: (-\Delta_{\rho, \mathcal{C}}) \phi.$$

By Lemma 2.1, the weighted elliptic operator $\Delta_{\rho, \mathcal{C}}$ becomes degenerate if ρ is a Dirac. For points ρ in the set $\mathcal{P}_+$ that are bounded away from zero, the operator $\Delta_{\rho, \mathcal{C}}$ is well-defined, nonsingular, and invertible since $\rho \mathcal{C}(\rho) > 0$. Thus we can write

$$V_{\rho, \mathcal{C}}^{-1} : T_\rho \mathcal{P}_+ \rightarrow T_\rho^* \mathcal{P}_+, \quad \sigma \mapsto (-\Delta_{\rho, \mathcal{C}})^{-1} \sigma.$$

This provides a 1-to-1 correspondence between elements $\phi \in T_\rho^* \mathcal{P}_+$ and $\sigma \in T_\rho \mathcal{P}_+$. For general $\rho \in \mathcal{P}_+$, we can instead use the pseudoinverse $(-\Delta_{\rho, \mathcal{C}})^\dagger$; see [57]. With the above choice of Onsager operator, we can define a generalized Wasserstein metric tensor as follows.

Definition 3.5 (Kalman–Wasserstein metric tensor). Define

$$g_{\rho, \mathcal{C}} : T_\rho \mathcal{P}_+ \times T_\rho \mathcal{P}_+ \rightarrow \mathbb{R}$$

as

$$g_{\rho, \mathcal{C}}(\sigma_1, \sigma_2) = \int_\Omega \langle \nabla \phi_1, \mathcal{C}(\rho) \nabla \phi_2 \rangle \rho dx,$$

where $\sigma_i = (-\Delta_{\rho, \mathcal{C}}) \phi_i = -\nabla \cdot (\rho \mathcal{C}(\rho) \nabla \phi_i) \in T_\rho \mathcal{P}_+$ for $i = 1, 2$.

With this metric tensor, the Kalman–Wasserstein metric $\mathcal{W}_{\mathcal{C}}: \mathcal{P}_+ \times \mathcal{P}_+ \rightarrow \mathbb{R}$ can be represented by the geometric action function. Given two densities $\rho^0, \rho^1 \in \mathcal{P}_+$, consider

$$\begin{aligned} W_{\mathcal{C}}(\rho^0, \rho^1)^2 &= \inf \int_0^1 \int_{\Omega} \langle \nabla \phi_t, \mathcal{C}(\rho_t) \nabla \phi_t \rangle \rho_t \, dx \\ &\text{subject to } \partial_t \rho_t + \nabla \cdot (\rho_t \mathcal{C}(\rho_t) \nabla \phi_t) = 0, \quad \rho_0 = \rho^0, \quad \rho_1 = \rho^1, \end{aligned}$$

where the infimum is taken among all continuous density paths $\rho_t := \rho(t, x)$ and potential functions $\phi_t := \phi(t, x)$. The Kalman–Wasserstein metric has several interesting mathematical properties, which will be the focus of future work. In this paper, working in $(\mathcal{P}_+, g_{\rho, \mathcal{C}})$, we derive the gradient flow formulation that underpins the formal calculations given in subsection 3.1 for the energy functional E defined in (3.2).

Proposition 3.6. *Given a finite functional $\mathcal{F}: \mathcal{P}_+ \rightarrow \mathbb{R}$, the gradient flow of $\mathcal{F}(\rho)$ in $(\mathcal{P}_+, g_{\rho, \mathcal{C}})$ satisfies*

$$\partial_t \rho = \nabla \cdot \left(\rho \mathcal{C}(\rho) \nabla \frac{\delta \mathcal{F}}{\delta \rho} \right).$$

Proof. The Riemannian gradient operator $\text{grad} \mathcal{F}(\rho)$ is defined via the metric tensor $g_{\rho, \mathcal{C}}$ as follows:

$$g_{\rho, \mathcal{C}}(\sigma, \text{grad} \mathcal{F}(\rho)) = \int_{\Omega} \frac{\delta}{\delta \rho(u)} \mathcal{F}(\rho) \sigma(u) \, du \quad \forall \sigma \in T_{\rho} \mathcal{P}_+.$$

Thus, for $\phi := (-\Delta_{\rho, \mathcal{C}})^{-1} \sigma \in T_{\rho}^* \mathcal{P}_+$, we have

$$\begin{aligned} g_{\rho, \mathcal{C}}(\sigma, \text{grad} \mathcal{F}(\rho)) &= \int \phi(u) \text{grad} \mathcal{F}(\rho) \, du = - \int \nabla \cdot (\rho \mathcal{C}(\rho) \nabla \phi) \frac{\delta}{\delta \rho} \mathcal{F}(\rho) \, du \\ &= \int \left\langle \nabla \phi, \mathcal{C}(\rho) \nabla \frac{\delta}{\delta \rho} \mathcal{F}(\rho) \right\rangle \rho \, du \\ &= - \int \phi(u) \nabla \cdot \left(\rho \mathcal{C}(\rho) \nabla \frac{\delta}{\delta \rho} \mathcal{F}(\rho) \right) \, du. \end{aligned}$$

Hence

$$\text{grad} \mathcal{F}(\rho) = - \nabla \cdot \left(\rho \mathcal{C}(\rho) \nabla \frac{\delta}{\delta \rho} \mathcal{F}(\rho) \right).$$

Thus we derive the gradient flow by

$$\partial_t \rho = - \text{grad} \mathcal{F}(\rho) = \nabla \cdot \left(\rho \mathcal{C}(\rho) \nabla \frac{\delta}{\delta \rho} \mathcal{F}(\rho) \right). \quad \blacksquare$$

Remark 3.7. Our derivation concerns the gradient flow on the subset $\mathcal{P}_+$ of $\mathcal{P}$ for simplicity of exposition. However, a rigorous analysis of the evolution of the gradient flow (3.3) requires extending the above arguments to the full set of probabilities $\mathcal{P}$, especially because we want to study Dirac measures in view of Lemma 2.1. If ρ is an element of the boundary of $\mathcal{P}$, one may consider instead the pseudoinverse of the operator $\Delta_{\rho, \mathcal{C}}$. This will be the focus of future work; also see the more general analysis in [1], e.g., Theorem 11.1.6.

4. Numerical experiments. In this section we demonstrate that the intuition developed in the previous two sections does indeed translate into useful algorithms for generating approximate posterior samples without computing derivatives of the forward map $\mathcal{G}$. We do this by considering non-Gaussian inverse problems defined through a nonlinear forward operator $\mathcal{G}$, showing how numerical solutions of (2.13) are distributed after large time, and comparing them with exact posterior samples found from MCMC methods.

Achieving the mean-field limit requires J to be large and hence typically larger than the dimension d of the parameter space. There are interesting and important problems arising in science and engineering in which the number of parameters to be estimated is small, even though evaluation of $\mathcal{G}$ involves solution of computationally expensive PDEs; in this case, choosing $J > d$ is not prohibitive. We also include numerical results which probe outcomes when $J < d$. To this end we study two problems, the first an inverse problem for a two-dimensional vector arising from a two point boundary-value problem, and the second an inverse problem for permeability from pressure measurements in Darcy flow; in this second problem the dimension of the parameter space is, in principle, tunable from small up to infinite dimension.

4.1. Derivative-free. In this subsection we describe how to use (2.13) for the solution of the inverse problem (1.1). We approximate the continuous-time stochastic dynamics by means of a linearly implicit split-step discretization scheme given by

$$(4.1a) \quad u_{n+1}^{(*,j)} = u_n^{(j)} - \Delta t_n \frac{1}{J} \sum_{k=1}^J \langle \mathcal{G}(u_n^{(k)}) - \bar{\mathcal{G}}, \mathcal{G}(u_n^{(j)}) - y \rangle_{\Gamma} u_n^{(k)} - \Delta t_n \mathbf{C}(U_n) \Gamma_0^{-1} u_{n+1}^{(*,j)},$$

$$(4.1b) \quad u_{n+1}^{(j)} = u_{n+1}^{(*,j)} + \sqrt{2 \Delta t_n \mathbf{C}(U_n)} \xi_n^{(j)},$$

where $\xi_n^{(j)} \sim \mathcal{N}(0, I)$, Γ_0 is the prior covariance, and Δt_n is an adaptive timestep computed as in [50].

4.2. Gold standard: MCMC. In this subsection we describe the specific Random Walk Metropolis Hastings (RWMH) algorithm used to solve the same Bayesian inverse problem as in the previous subsection; we view the results as gold standard samples from the desired posterior distribution. The link between RWMH methods and Langevin sampling is explained in the literature review within the introduction, where it is shown that the latter arises as a diffusion limit of the former, as shown in numerous papers following from the seminal work in [85]. The proposal distribution is a Gaussian centered at the current state of the Markov chain with covariance given by $\Sigma = \tau \times \mathbf{C}(U^*)$, where $\mathbf{C}(U^*)$ is the covariance computed from the last iteration of the algorithm described in the preceding subsection, and τ is a scaling factor tuned for an acceptance rate of approximately 25% [85]. In our case, $\tau = 4$. The RWMH algorithm was used to get $N = 10^5$ samples, with the Markov chain starting at an approximate solution given by the mean of the last step of the algorithm from the previous subsection. For the high-dimensional problem we use the preconditioned Crank–Nicolson (pCN) variant on RWMH [18]; this too has a diffusion limit of Langevin form [80].

4.3. Numerical results: Low-dimensional parameter space. The numerical experiment considered here is the example originally presented in [31] and also used in [39]. We start

by defining the forward map, which is given by the one-dimensional elliptic boundary-value problem

$$(4.2) \quad -\frac{d}{dx} \left(\exp(u_1) \frac{d}{dx} p(x) \right) = 1, \quad x \in [0, 1],$$

with boundary conditions $p(0) = 0$ and $p(1) = u_2$. The explicit solution for this problem (see [39]) is given by

$$(4.3) \quad p(x) = u_2 x + \exp(-u_1) \left(-\frac{x^2}{2} + \frac{x}{2} \right).$$

The forward model operator $\mathcal{G}$ is then defined by

$$(4.4) \quad \mathcal{G}(u) = \begin{pmatrix} p(x_1) \\ p(x_2) \end{pmatrix}.$$

Here $u = (u_1, u_2)^\top$ is a constant vector that we want to find, and we assume that we are given noisy measurements y of $p(\cdot)$ at locations $x_1 = 0.25$ and $x_2 = 0.75$. The precise Bayesian inverse problem considered here is to find the distribution of the unknown u conditioned on the observed data y , assuming additive Gaussian noise $\eta \sim \mathcal{N}(0, \Gamma)$, where $\Gamma = 0.1^2 I_2$ and $I_2 \in \mathbb{R}^{2 \times 2}$ is the identity matrix. We use as prior distribution $\mathcal{N}(0, \Gamma_0)$, $\Gamma_0 = \sigma^2 I_2$ with $\sigma = 10$. The resulting Bayesian inverse problem is then solved, approximately, by the algorithms we now outline and with observed data $y = (27.5, 79.7)^\top$. Following [39], we consider an initial ensemble drawn from $\mathcal{N}(0, 1) \times \mathcal{U}(90, 110)$.

Figure 1 shows the results for the solution of the Bayesian inverse problem considered above. In addition to implementing the algorithms described in the previous two subsections, we also employ a specific implementation of the EKI formulation introduced in the paper of Herty and Visconti [39], and defined by the numerical discretization shown in (4.1), but with $C(U)$ replaced by the identity matrix I_2 ; this corresponds to the algorithm from equation (20) of [39], and in particular the last display of their section 5, with $\xi \sim \mathcal{N}(0, I_2)$. The blue dots correspond to the output of this algorithm at the last iteration. The red dots correspond to the last ensemble of the EKI algorithm as presented in [50]. The orange dots depict the RWMH gold standard described above. Finally, the green dots show the ensemble members at the last iteration of the proposed EKS (2.13). In this experiment, all versions of the ensemble Kalman methods were run with the adaptive timestep scheme from subsection 4.1, and all were run for 30 iterations with an ensemble size of $J = 10^3$.

Consider first the top-left panel. The true distribution, computed by RWMH, is shown in orange. Note that the algorithm of [50] collapses to a point (shown in red), unable to escape overfitting, and relating to a form of consensus formation. In contrast, the algorithm of [39], while avoiding overfitting, overestimates the spread of the ensemble members, relative to the gold standard RWMH; this is exhibited by the blue overdispersed points. The proposed EKS (green points) gives results close to the RWMH gold standard. These issues are further demonstrated in the lower panel, which shows the misfit (loss) function as a function of iterations for the three algorithms (excluding RWMH); the red line demonstrates overfitting as the misfit value falls below the noise level, whereas the other two algorithms avoid overfitting.

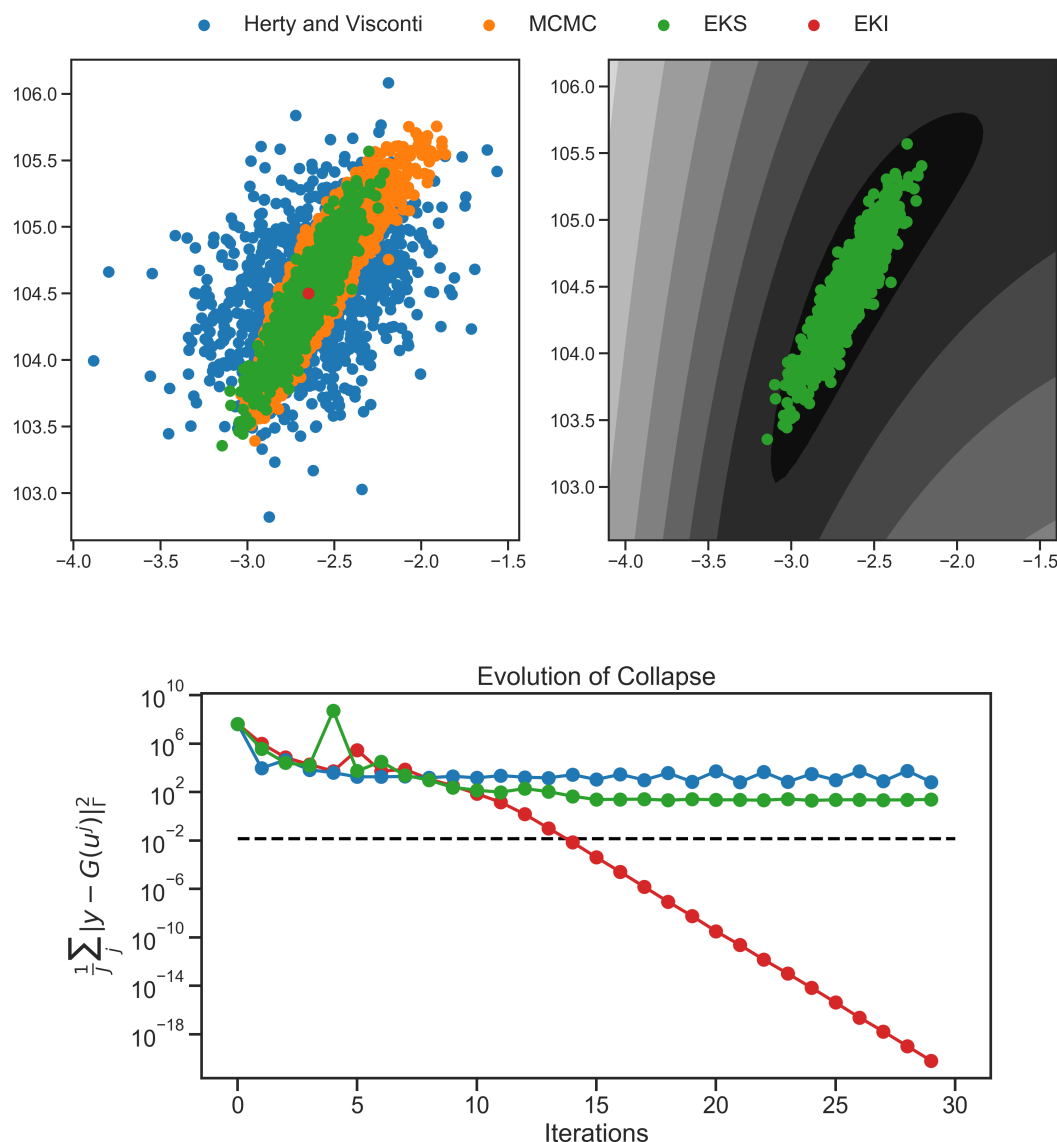


Figure 1. Results of applying different versions of ensemble Kalman methods to the nonlinear elliptic boundary problem. For comparison, an RWMH algorithm is also displayed to provide a gold standard. The proposed EKS captures approximately the true distribution, effectively avoiding overfitting or overdispersion shown with the other two implementations. Overfitting is clearly shown from the red line in the lower subfigure. The line in blue shows overdispersion exhibited by the algorithm proposed in [39]. The upper right subfigure illustrates the approximation to the posterior. Color coding is consistent among the subfigures.

We include the derivative-free optimization algorithm EKI (red points) because it gives insight into what can be achieved with these ensemble-based methods in the absence of noise (namely, derivative-free optimization); we include the noisy EKI algorithm of [39] (blue points) to demonstrate that considerable care is needed with the introduction of noise if the goal is

to produce posterior samples, and we include our proposed EKS algorithm (green points) to demonstrate that judicious addition of noise into the EKI algorithm helps to produce approximate samples from the true posterior distribution of the Bayesian inverse problem; we include true posterior samples (orange points) for comparison. We reiterate that the methods of [16, 29] also hold the potential to produce good approximate samples, though they suffer from the rigidity of needing to be initialized at the prior and integrated to exactly time 1.

4.4. Numerical results: High-dimensional parameter space. The forward problem of interest is to find the pressure field $p(\cdot)$ in a porous medium defined by permeability field $a(\cdot)$; for simplicity we assume that $a(\cdot)$ is a scalar-field in this paper. Given a scalar-field f defining sources and sinks of fluid, and assuming Dirichlet boundary conditions on the pressure for simplicity, we obtain the following elliptic PDE for the pressure:

$$(4.5a) \quad -\nabla \cdot (a(x)\nabla p(x)) = f(x), \quad x \in D,$$

$$(4.5b) \quad p(x) = 0, \quad x \in \partial D.$$

In what follows we will work on the domain $D = [0, 1]^2$. We assume that the permeability is dependent on unknown parameters $u \in \mathbb{R}^d$, so that $a(x) = a(x; u)$. The inverse problem of interest is to determine u from d linear functionals (measurements) of $p(x; u)$, subject to additive noise. Thus

$$(4.6) \quad \mathcal{G}_j(u) = \ell_j(p(\cdot; u)) + \eta_j, \quad j = 1, \dots, K.$$

We will assume that $a(\cdot) \in L^\infty(D; \mathbb{R})$ so that $p(\cdot) \in H_0^1(D; \mathbb{R})$, and thus we take the ℓ_j to be linear functionals on the space $H_0^1(D; \mathbb{R})$. In practice we will work with pointwise measurements so that $\ell_j(p) = p(x_j)$; these are not elements of the dual space of $H_0^1(D; \mathbb{R})$ in dimension 2, but mollifications of them are, and in practice mollification with a narrow kernel does not affect results of the type presented here, and so we do not use it [41]. We model $a(x; u)$ as a log-Gaussian field with the precision operator defined as

$$(4.7) \quad \mathcal{C}^{-1} = (-\Delta + \tau^2 \mathcal{I})^\alpha,$$

where the Laplacian Δ is equipped with Neumann boundary conditions on the space of spatial-mean zero functions, and τ and α are known constants that control the underlying lengthscales and smoothness of the underlying random field. In our experiments, $\tau = 3$ and $\alpha = 2$. Such parametrization yields a Karhunen–Loève (KL) expansion

$$(4.8) \quad \log a(x; u) = \sum_{\ell \in K} u_\ell \sqrt{\lambda_\ell} \varphi_\ell(x),$$

where the eigenpairs are of the form

$$(4.9) \quad \varphi_\ell(x) = \cos(\pi \langle \ell, x \rangle), \quad \lambda_\ell = (\pi^2 |\ell|^2 + \tau^2)^{-\alpha},$$

where $K \equiv \mathbb{Z}^2$ is the set of indices over which the random series is summed and the $u_\ell \sim N(0, 1)$ are i.i.d. [79]. In practice we will approximate K by $K_d \subset \mathbb{Z}^2$, a set with finite

cardinality d , and consider different d . For visualization we will sometimes find it helpful to write (4.8) as a sum over a one-dimensional variable rather than a lattice:

$$(4.10) \quad \log a(x; u) = \sum_{k \in \mathbb{Z}^+} u'_k \sqrt{\lambda'_k} \varphi'_k(x).$$

We order the indices in $\mathbb{Z}^+$ so that the eigenvalues λ'_k are in descending order by size.

We generate a truth random field by constructing $u^\dagger \in \mathbb{R}^d$ by sampling it from $N(0, I_d)$, with $d = 2^8$ and I_d the identity on $\mathbb{R}^d$ and using $u^\dagger$ as the coefficients in (4.8). We create data y from (1.1) with $\eta \sim N(0, 0.1^2 \times I_K)$. For the Bayesian inversion we choose prior covariance $\Gamma_0 = 10^2 I_d$; we also sample from this prior to initialize the ensemble for EKS. We run the experiments with different ensemble sizes to understand both strengths and limitations of the proposed algorithm for nonlinear forward models. Finally, we choose $J \in \{8, 32, 128, 512, 2048\}$, which allows the study of both $J > d$ and $J < d$ within the methodology.

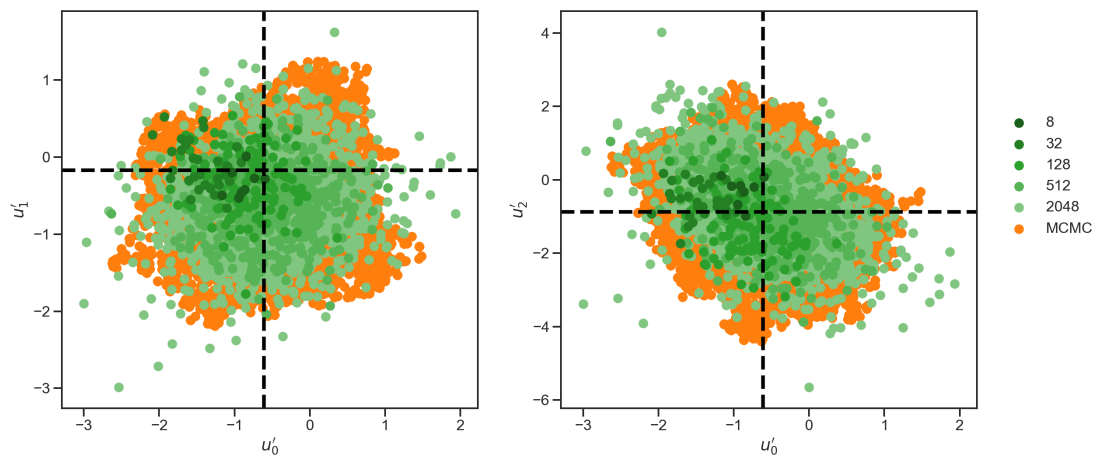
Results showing the solution of this Bayesian inverse problem by MCMC (orange dots), with 10^5 samples, and by the EKS with different J (different shades of green dots) are shown in Figures 2 and 3. For every ensemble size configuration, the EKS algorithm was run until two units of time were achieved. As can be seen from Figure 2(b) the algorithm has reached an equilibrium after this duration. The two-dimensional scatter plots in Figure 2(a) show components u'_k with $k = 0, 1, 2$. That is, we are showing the components of u which are associated to the three largest eigenvalues in the KL expansion (4.8) under the posterior distribution. We can see that sample spread is better matched to the gold standard MCMC spread as the size J of the EKS ensemble is increased. In Figures 2(b) and 2(c) we show the evolution of the dispersion of the ensemble around its mean at every timestep, $\bar{u}(t)$, and around the truth $u^\dagger$. The metrics we use to test the ensemble spread are

$$(4.11) \quad d_{H^{-2}}(\cdot) = \sqrt{\frac{1}{J} \sum_{j=1}^J \|u^{(j)}(t) - \cdot\|_{H^{-2}}^2}, \quad d_{L^2}(\cdot) = \sqrt{\frac{1}{J} \sum_{j=1}^J \|u^{(j)}(t) - \cdot\|_{L^2}^2},$$

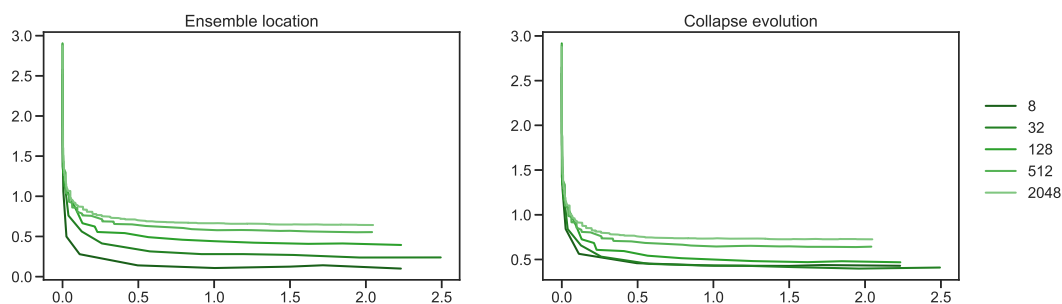
where both are evaluated at $\bar{u}(t)$ and $u^\dagger$ at every simulated time t . For these metrics we use the norms defined by

$$(4.12) \quad \|u\|_{H^{-2}} = \sqrt{\sum_{\ell \in K_d} |u_\ell|^2 \lambda_\ell}, \quad \|u\|_{L^2} = \sqrt{\sum_{\ell \in K_d} |u_\ell|^2},$$

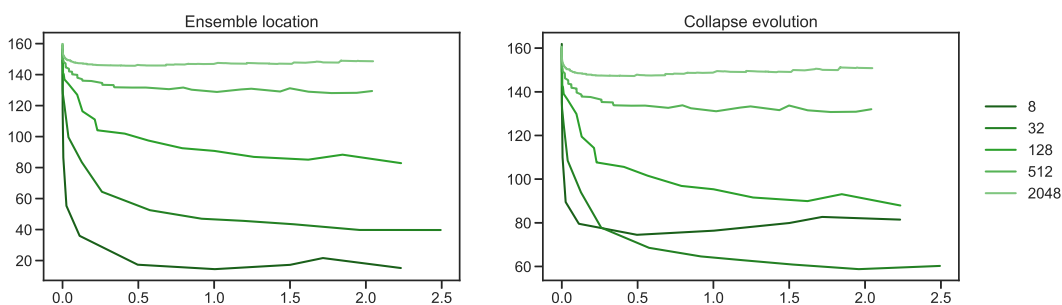
where the first is defined in the negative Sobolev space H^{-2} , while the second is defined in the L^2 space. The first norm allows for higher discrepancy in the estimation of the tail of the modes in equation (4.8), whereas the second norm penalizes equally discrepancies in the tail of the KL expansion. In Figure 2(b), we see rapid convergence of the spread around the mean and around the truth for all ensemble sizes J . The evolution in Figure 2(b) for both cases shows that the algorithm reaches its stationary distribution while incorporating higher variability with increasing ensemble size. The figures are similar because the posterior mean and the truth are close to one another. Lower values of the metrics in Figures 2(b) and 2(c) for smaller ensembles can be understood due to a mixed effect of reduced variability and overfitting to the maximum a posteriori (MAP) estimate of the Bayesian inverse problem.



(a) Bivariate scatter plots of the approximate posterior distribution on the three largest modes (as ordered by prior variance and here labeled u_0, u_1, u_2) in the KL expansion (4.8). The pCN algorithm (orange dots) is used as a reference with 10^5 samples.



(b) Evolution statistics of the EKS with respect to simulated time under the negative Sobolev norm $\|\cdot\|_{H^{-2}}$.



(c) Evolution statistics of the EKS with respect to simulated time under norm $\|\cdot\|_{L^2}$.

Figure 2. Results for the Darcy flow inverse problem in high dimensions. The top panel shows scatter plots for different combinations of the higher modes in the KL expansion (4.8). The green dots correspond to the last iteration of the EKS at every ensemble size setting as labeled in the legend. Tracking the negative Sobolev norm of the ensembles with respect to its mean $\bar{u}(t)$ and underlying truth $u^\dagger$ shows a good match with both the solution of the inverse problem and the stationary distribution of the Fokker–Planck equation (3.5).

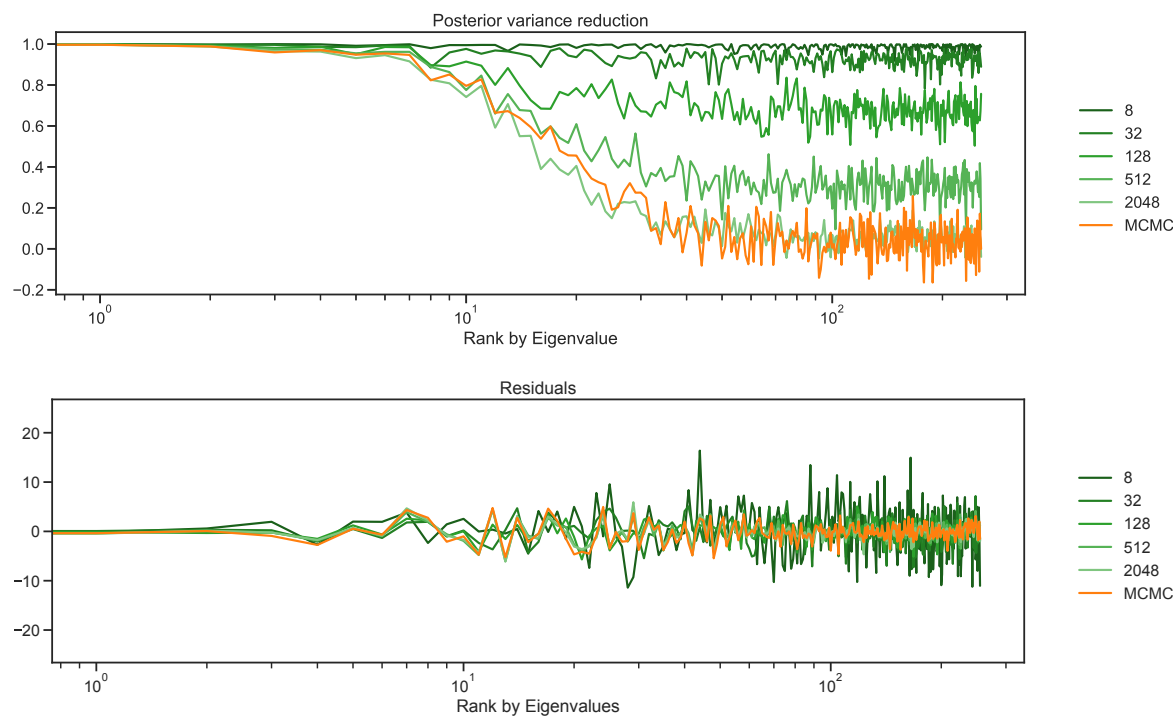


Figure 3. Results showing Darcy flow parameter identifiability. The top panel illustrates how bigger ensemble sizes are able to better capture the marginal posterior variability of each component, whereas the lower panel illustrates both variability and consistency of the approximate posterior samples from EKS.

The results using the L^2 norm in Figure 2(c) allow us to see more discrepancy between ensemble sizes. The higher metric value for larger ensembles is due to the ensemble better approximating the posterior, as will be discussed below. In summary, Figure 2 shows evidence that the EKS is generating samples from a good approximation to the posterior and that this posterior is centered close to the truth. Increasing the ensemble size improves these features of the EKS method.

Figure 3 demonstrates how different ensemble sizes are able to capture the marginal posterior variances of each component in the unknown u . The top panel in Figure 3 tracks the posterior variance reduction statistic for every component of $u' \in \mathbb{R}^d$, which, as mentioned before, now is viewed as a vector of d components rather than a function on subset K_d of the two-dimensional lattice. The posterior variance reduction is a measure of the relative decrease of variance for a given quantity of interest under the posterior with respect to the prior. It is defined as

$$(4.13) \quad \zeta_k = 1 - \frac{\mathbb{V}(u'_k|y)}{\mathbb{V}(u'_k)},$$

where $\mathbb{V}(\cdot)$ denotes the variance of a random variable. The summary statistic $\sum_k \zeta_k$ has been used in [91] to estimate the effective number of parameters in Bayesian models. When this

parameter is close to 1, the algorithm has reduced the uncertainty considerably, relative to the prior; when it is close to zero, it has reduced it very little, in comparison with the prior. By studying the figure for the MCMC algorithm (orange) and comparing with EKS for increasing J (green), we see that for J around 2000 the match between EKS and MCMC is excellent. We also see that for smaller-sized ensembles there is a regularizing effect which artificially reduces the posterior variation for larger k . On the other hand, the lower panel in Figure 3 allows us to identify the location of ensemble density by plotting the residuals $u'_k - (u_k^\dagger)'$ for every component $k = 1, \dots, d$; in particular we plot the algorithmic mean of this quantity and 95% confidence intervals. It can be seen that the ensemble is well located, as most of the components include the zero horizontal line, meaning that marginally the distribution of every component includes the correct value with high probability. Moreover, we can see two effects in this figure. First, the lower variability in the first components also shows that there is enough information on the observed data to identify these components. Second, it can be seen that for very low sized ensembles the least important components of u incorporate higher error when comparing the green EKS samples with the orange MCMC samples.

Overall, the mismatch between the results from EKS and the MCMC reference in both numerical examples can be understood from the fact that the use of the ensemble equations (2.13) introduces a linear approximation to the curvature of the regularized misfit. This effect is demonstrated clearly in Figure 1, which shows the samples from EKS against a background of the level sets of the posterior. However, despite this mismatch, the key point is that a relatively good set of approximate samples in green is computed without use of the derivative of the forward model $\mathcal{G}$ in both numerical examples; it thus holds promise as a method for large-scale nonlinear inverse problems.

5. Conclusions. In this paper we have demonstrated a methodology for the addition of noise to the basic EKI algorithm so that it generates approximate samples from the Bayesian posterior distribution; we call this method the ensemble Kalman sampler (EKS). Our starting point is a set of interacting Langevin diffusions, preconditioned by their mutual empirical covariance. To understand this system, we introduce a new mean-field Fokker–Planck equation which has the desired posterior distribution as an invariant measure. We exhibit the new Kalman–Wasserstein metric with respect to which the Fokker–Planck equation has gradient structure. We also show how to compute approximate samples from this model by using a particle approximation based on using ensemble differences in place of gradients, leading to the EKS algorithm.

In the future we anticipate that methodology to correct for the error introduced by the use of ensemble differences will be a worthwhile development from the algorithms proposed, and we are actively pursuing this [17]. Furthermore, recent interesting work of Nüsken and Reich [69] studies the invariant measures of the finite particle system (2.3), (2.4). The authors identify a simple linear correction term of order J^{-1} in (2.3) which renders the J -fold product of the posterior distribution invariant for finite ensemble number; since one of the major motivations for the use of ensemble methods is their robustness for small J , this correction is important.

We also recognize that other difference-based methods for approximating gradients may emerge and that developing theory to quantify and control the errors arising from such

difference approximations will be of interest. We believe that our proposed ensemble-based difference approximation is of particular value to the growing community of scientists and engineers who work directly with ensemble-based methods because of their simplicity and black-box nature. In the future, we will also study the properties of the Kalman–Wasserstein metric, including its duality, geodesics, and geometric structure, a line of research that is of independent mathematical interest in the context of generalized Wasserstein-type spaces. We will investigate the analytical properties of the new metric within Gaussian families. We expect these studies will bring insight to designing new numerical algorithms for the approximate solution of inverse problems.

Acknowledgments. The authors are grateful to José A. Carrillo, Greg Pavliotis, and Sebastian Reich for helpful input which improved this paper.

REFERENCES

- [1] L. AMBROSIO, N. GIGLI, AND G. SAVARÉ, *Gradient Flows: In Metric Spaces and in the Space of Probability Measures*, Birkhäuser, Basel, 2005.
- [2] A. ARNOLD, P. MARKOWICH, G. TOSCANI, AND A. UNTERREITER, *On convex Sobolev inequalities and the rate of convergence to equilibrium for Fokker-Planck type equations*, Comm. Partial Differential Equations, 26 (2001), pp. 43–100, <https://doi.org/10.1081/PDE-100002246>.
- [3] N. AY, J. JOST, H. V. LÊ, AND L. J. SCHWACHHÖFER, *Information Geometry*, Ergeb. Math. Grenzgeb. (3), A Series of Modern Surveys in Mathematics [Results in Mathematics and Related Areas], 3rd Series, A Series of Modern Surveys in Mathematics 64, Springer, Cham, 2017.
- [4] D. BAKRY AND M. ÉMERY, *Diffusions hypercontractives*, in Séminaire de Probabilités XIX, 1983/84, Lecture Notes in Math. 1123, Springer, Berlin, 1985, pp. 177–206, <https://doi.org/10.1007/BFb0075847>.
- [5] M. BEDARD, *Optimal acceptance rates for Metropolis algorithms: Moving beyond 0.234*, Stochastic Process. Appl., 118 (2008), pp. 2198–2222.
- [6] M. BÉDARD, *Weak convergence of Metropolis algorithms for non-i.i.d. target distributions*, Ann. Appl. Probab., 17 (2007), pp. 1222–1244.
- [7] M. BÉDARD AND J. S. ROSENTHAL, *Optimal scaling of Metropolis algorithms: Heading toward general target distributions*, Canad. J. Statist., 36 (2008), pp. 483–503.
- [8] K. BERGEMANN AND S. REICH, *A localization technique for ensemble Kalman filters*, Q. J. R. Meteorol. Soc., 136 (2010), pp. 701–707.
- [9] K. BERGEMANN AND S. REICH, *A mollified ensemble Kalman filter*, Q. J. R. Meteorol. Soc., 136 (2010), pp. 1636–1643.
- [10] K. BERGEMANN AND S. REICH, *An ensemble Kalman-Bucy filter for continuous data assimilation*, Meteorol. Z., 21 (2012), pp. 213–219.
- [11] A. CARRASSI, M. BOCQUET, L. BERTINO, AND G. EVENSEN, *Data assimilation in the geosciences: An overview of methods, issues, and perspectives*, WIREs Climate Change, 9 (2018), e535.
- [12] J. A. CARRILLO, Y.-P. CHOI, C. TOTZECK, AND O. TSE, *An analytical framework for consensus-based global optimization method*, Math. Models Methods Appl. Sci., 28 (2018), pp. 1037–1066.
- [13] J. A. CARRILLO, M. FORNASIER, G. TOSCANI, AND F. VECIL, *Particle, kinetic, and hydrodynamic models of swarming*, in Mathematical Modeling of Collective Behavior in Socio-economic and Life Sciences, Model. Simul. Sci. Eng. Technol., Birkhäuser Boston, Inc., Boston, MA, 2010, pp. 297–336, https://doi.org/10.1007/978-0-8176-4946-3_12.
- [14] J. A. CARRILLO AND G. TOSCANI, *Exponential convergence toward equilibrium for homogeneous Fokker-Planck-type equations*, Math. Methods Appl. Sci., 21 (1998), pp. 1269–1286, [https://doi.org/10.1002/\(SICI\)1099-1476\(19980910\)21:13%3C1269::AID-MMA995%3E3.0.CO;2-O](https://doi.org/10.1002/(SICI)1099-1476(19980910)21:13%3C1269::AID-MMA995%3E3.0.CO;2-O).
- [15] N. K. CHADA, A. M. STUART, AND X. T. TONG, *Tikhonov Regularization within Ensemble Kalman Inversion*, preprint, <https://arxiv.org/abs/1901.10382>, 2019.
- [16] Y. CHEN AND D. S. OLIVER, *Ensemble randomized maximum likelihood method as an iterative ensemble smoother*, Math. Geosci., 44 (2012), pp. 1–26.

- [17] E. CLEARY, A. GARBUNO-INIGO, S. LAN, T. SCHNEIDER, AND A. M. STUART, *Calibrate, Emulate, Sample*, preprint, <https://arxiv.org/abs/2001.03689>, 2020.
- [18] S. L. COTTER, G. O. ROBERTS, A. M. STUART, AND D. WHITE, *MCMC methods for functions: Modifying old algorithms to make them faster*, *Statist. Sci.*, 28 (2013), pp. 424–446.
- [19] D. CRISAN AND J. XIONG, *Approximate McKean–Vlasov representations for a class of SPDEs*, *Stochastics*, 82 (2010), pp. 53–68.
- [20] F. DAUM AND J. HUANG, *Particle flow for nonlinear filters*, in *Proceedings of the 2011 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, IEEE, Piscataway, NJ, 2011, pp. 5920–5923.
- [21] J. DE WILJES, S. REICH, AND W. STANNAT, *Long-time stability and accuracy of the ensemble Kalman–Bucy filter for fully observed processes and small measurement noise*, *SIAM J. Appl. Dyn. Syst.*, 17 (2018), pp. 1152–1181, <https://doi.org/10.1137/17M1119056>.
- [22] P. DEL MORAL, A. KURTZMANN, AND J. TUGAUT, *On the stability and the uniform propagation of chaos of a class of extended ensemble Kalman–Bucy filters*, *SIAM J. Control Optim.*, 55 (2017), pp. 119–155, <https://doi.org/10.1137/16M1087497>.
- [23] P. DEL MORAL AND J. TUGAUT, *On the stability and the uniform propagation of chaos properties of ensemble Kalman–Bucy filters*, *Ann. Appl. Probab.*, 28 (2018), pp. 790–850.
- [24] G. DETOMMASO, T. CUI, Y. MARZOUK, A. SPANTINI, AND R. SCHEICHL, *A Stein variational Newton method*, in *Proceedings of the 32nd Conference on Neural Information Processing Systems (NeurIPS 2018)*, Montreal, Canada, *Advances in Neural Information Processing Systems* 31, S. Bengio, ed., 2018, pp. 9169–9179.
- [25] Z. DING AND Q. LI, *Mean-Field Limit and Numerical Analysis for Ensemble Kalman Inversion: Linear Setting*, preprint, <https://arxiv.org/abs/1908.05575v1>, 2019.
- [26] A. DUNCAN AND L. SZPRUCH, *private communication*, 2019.
- [27] A. B. DUNCAN, T. LELIEVRE, AND G. PAVLIOTIS, *Variance reduction using nonreversible Langevin samplers*, *J. Statist. Phys.*, 163 (2016), pp. 457–491.
- [28] T. A. EL MOSELHY AND Y. M. MARZOUK, *Bayesian inference with optimal maps*, *J. Comput. Phys.*, 231 (2012), pp. 7815–7850.
- [29] A. A. EMERICK AND A. C. REYNOLDS, *Investigation of the sampling performance of ensemble-based methods with a simple reservoir model*, *Comput. Geosci.*, 17 (2013), pp. 325–350.
- [30] H. W. ENGL, M. HANKE, AND A. NEUBAUER, *Regularization of Inverse Problems*, *Math. Appl.* 375, Springer Netherlands, 1996.
- [31] O. G. ERNST, B. SPRUNGK, AND H.-J. STARKLOFF, *Analysis of the ensemble and polynomial chaos Kalman filters in Bayesian inverse problems*, *SIAM/ASA J. Uncertain. Quantif.*, 3 (2015), pp. 823–851, <https://doi.org/10.1137/140981319>.
- [32] G. EVENSEN, *Data Assimilation: The Ensemble Kalman Filter*, Springer-Verlag, Berlin, 2009.
- [33] M. GIROLAMI AND B. CALDERHEAD, *Riemann manifold Langevin and Hamiltonian Monte Carlo methods*, *J. Roy. Statist. Soc. Ser. B Statist. Methodol.*, 73 (2011), pp. 123–214.
- [34] O. GONZALEZ, *Time integration and discrete Hamiltonian systems*, *J. Nonlinear Sci.*, 6 (1996), 449.
- [35] J. GOODMAN AND J. WEARE, *Ensemble samplers with affine invariance*, *Commun. Appl. Math. Comput. Sci.*, 5 (2010), pp. 65–80.
- [36] S.-Y. HA AND E. TADMOR, *From particle to kinetic and hydrodynamic descriptions of flocking*, *Kinet. Relat. Models*, 1 (2008), pp. 415–435, <https://doi.org/10.3934/krm.2008.1.415>.
- [37] E. HAIRER AND C. LUBICH, *Energy-diminishing integration of gradient systems*, *IMA J. Numer. Anal.*, 34 (2013), pp. 452–461.
- [38] A. HALDER AND T. T. GEORGIU, *Gradient Flows in Filtering and Fisher-Rao Geometry*, preprint, <https://arxiv.org/abs/1710.00064>, 2017.
- [39] M. HERTY AND G. VISCONTI, *Kinetic Methods for Inverse Problems*, preprint, <https://arxiv.org/abs/1811.09387>, 2018.
- [40] A. R. HUMPHRIES AND A. M. STUART, *Runge–Kutta methods for dissipative and gradient dynamical systems*, *SIAM J. Numer. Anal.*, 31 (1994), pp. 1452–1485, <https://doi.org/10.1137/0731075>.
- [41] M. A. IGLESIAS, *Iterative regularization for ensemble data assimilation in reservoir models*, *Comput. Geosci.*, 19 (2015), pp. 177–212.
- [42] M. A. IGLESIAS, *A regularizing iterative ensemble Kalman method for PDE-constrained inverse problems*, *Inverse Problems*, 32 (2016), 025002.

- [43] M. A. IGLESIAS, K. J. LAW, AND A. M. STUART, *Ensemble Kalman methods for inverse problems*, Inverse Problems, 29 (2013), 045001.
- [44] P.-E. JABIN AND Z. WANG, *Mean field limit for stochastic particle systems*, in Active Particles, Vol. I, Model. Simul. Sci. Eng. Technol., Birkhäuser, Cham, 2017, pp. 379–402, https://doi.org/10.1007/978-3-319-49996-3_10.
- [45] R. JORDAN, D. KINDERLEHRER, AND F. OTTO, *The variational formulation of the Fokker–Planck equation*, SIAM J. Math. Anal., 29 (1998), pp. 1–17, <https://doi.org/10.1137/S0036141096303359>.
- [46] J. KAIPIO AND E. SOMERSALO, *Statistical and Computational Inverse Problems*, Appl. Math. Sci. 160, Springer-Verlag, New York, 2006.
- [47] R. E. KALMAN, *A new approach to linear filtering and prediction problems*, Trans. ASME Ser. D J. Basic Engrg., 82 (1960), pp. 35–45.
- [48] R. E. KALMAN AND R. S. BUCY, *New results in linear filtering and prediction theory*, Trans. ASME Ser. D J. Basic Engrg., 83 (1961), pp. 95–108.
- [49] D. T. KELLY, K. LAW, AND A. M. STUART, *Well-posedness and accuracy of the ensemble Kalman filter in discrete and continuous time*, Nonlinearity, 27 (2014), pp. 2579–2603.
- [50] N. B. KOVACHKI AND A. M. STUART, *Ensemble Kalman Inversion: A Derivative-Free Technique for Machine Learning Tasks*, preprint, <https://arxiv.org/abs/1808.03620>, 2018.
- [51] J. D. LAFFERTY, *The density manifold and configuration space quantization*, Trans. Amer. Math. Soc., 305 (1988), pp. 699–741.
- [52] R. S. LAUGESSEN, P. G. MEHTA, S. P. MEYN, AND M. RAGINSKY, *Poisson’s equation in nonlinear filtering*, SIAM J. Control Optim., 53 (2015), pp. 501–525, <https://doi.org/10.1137/13094743X>.
- [53] K. LAW, A. STUART, AND K. ZYGALAKIS, *Data Assimilation: A Mathematical Introduction*, Texts Appl. Math. 62, Springer International Publishing, Switzerland, 2015.
- [54] B. LEIMKUHLER AND C. MATTHEWS, *Molecular Dynamics: With Deterministic and Stochastic Numerical Methods*, Interdiscip. Appl. Math. 39, Springer International Publishing, Switzerland, 2016.
- [55] B. LEIMKUHLER, C. MATTHEWS, AND J. WEARE, *Ensemble preconditioning for Markov chain Monte Carlo simulation*, Stat. Comput., 28 (2018), pp. 277–290.
- [56] B. LEIMKUHLER, E. NOORIZADEH, AND F. THEIL, *A gentle stochastic thermostat for molecular dynamics*, J. Statist. Phys., 135 (2009), pp. 261–277.
- [57] W. LI, *Geometry of Probability Simplex via Optimal Transport*, preprint, <https://arxiv.org/abs/1803.06360>, 2018.
- [58] W. LI, A. LIN, AND G. MONTÚFAR, *Affine natural proximal learning*, in Proceedings of the International Conference on Geometric Science of Information (GSI 2019): Geometric Science of Information, Lecture Notes in Comput. Sci. 11712, Springer, Cham, 2019, pp. 705–714.
- [59] W. LI AND G. MONTUFAR, *Natural Gradient via Optimal Transport*, preprint, <https://arxiv.org/abs/1803.07033>, 2018.
- [60] Q. LIU AND D. WANG, *Stein variational gradient descent: A general purpose Bayesian inference algorithm*, in Proceedings of the 30th Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain, Advances in Neural Information Processing Systems 29, D. D. Lee et al., eds., 2016, pp. 2378–2386.
- [61] J. LU, Y. LU, AND J. NOLEN, *Scaling Limit of the Stein Variational Gradient Descent Part I: The Mean Field Regime*, preprint, <https://arxiv.org/abs/1805.04035v1>, 2018.
- [62] S. MACHLUP AND L. ONSAGER, *Fluctuations and irreversible process II: Systems with kinetic energy*, Phys. Rev. (2), 91 (1953), pp. 1512–1515.
- [63] A. J. MAJDA AND J. HARLIM, *Filtering Complex Turbulent Systems*, Cambridge University Press, Cambridge, 2012.
- [64] P. A. MARKOWICH AND C. VILLANI, *On the trend to equilibrium for the Fokker–Planck equation: An interplay between physics and functional analysis*, VI Workshop on Partial Differential Equations, Part II (Rio de Janeiro, 1999), Mat. Contemp., 19 (2000), pp. 1–29.
- [65] Y. MARZOUK, T. MOSELHY, M. PARNO, AND A. SPANTINI, *Sampling via measure transport: An introduction*, in Handbook of Uncertainty Quantification, Springer, Cham, 2016, pp. 1–41, https://doi.org/10.1007/978-3-319-11259-6_23-1.
- [66] J. C. MATTINGLY, N. S. PILLAI, AND A. M. STUART, *Diffusion limits of the random walk Metropolis algorithm in high dimensions*, Ann. Appl. Probab., 22 (2012), pp. 881–930.

- [67] R. I. MCLACHLAN, G. QUISPEL, AND N. ROBIDOUX, *Geometric integration using discrete gradients*, R. Soc. Lond. Philos. Trans. Ser. A Math. Phys. Eng. Sci., 357 (1999), pp. 1021–1045.
- [68] A. MIELKE, M. A. PELETIER, AND M. RENGIER, *A generalization of Onsager's reciprocity relations to gradient flows with nonlinear mobility*, J. Non-Equilibrium Thermodynamics, 41 (2016), pp. 141–149, <http://publications.imp.fu-berlin.de/1890/>.
- [69] N. NÜSKEN AND S. REICH, *Note on Interacting Langevin Diffusions: Gradient Structure and Ensemble Kalman Sampler by Garbuno-Inigo, Hoffmann, Li and Stuart*, preprint, <https://arxiv.org/abs/1908.10890>, 2019.
- [70] D. S. OLIVER, A. C. REYNOLDS, AND N. LIU, *Inverse Theory for Petroleum Reservoir Characterization and History Matching*, Cambridge University Press, Cambridge, 2008.
- [71] Y. OLLIVIER, *Online Natural Gradient as a Kalman Filter*, preprint, <https://arxiv.org/abs/1703.00209>, 2017.
- [72] L. ONSAGER, *Reciprocal relations in irreversible processes: I*, Phys. Rev., 37 (1931), pp. 405–426, <https://doi.org/10.1103/PhysRev.37.405>.
- [73] L. ONSAGER, *Reciprocal relations in irreversible processes: II*, Phys. Rev., 38 (1931), pp. 2265–2279, <https://doi.org/10.1103/PhysRev.38.2265>.
- [74] H. ÖTTINGER, *Beyond Equilibrium Thermodynamics*, John Wiley & Sons, Inc., 2005, <https://books.google.com/books?id=Prh9moT1WzMC>.
- [75] F. OTTO, *The geometry of dissipative evolution equations: The porous medium equation*, Comm. Partial Differential Equations, 26 (2001), pp. 101–174.
- [76] M. OTTOBRE AND G. PAVLIOTIS, *Asymptotic analysis for the generalized Langevin equation*, Nonlinearity, 24 (2011), pp. 1629–1656.
- [77] L. PARESCHI AND G. TOSCANI, *Interacting Multiagent Systems: Kinetic Equations and Monte Carlo Methods*, no. 9780199655465, OUP Catalogue, Oxford University Press, 2013, <https://ideas.repec.org/b/oxp/obooks/9780199655465.html>.
- [78] S. PATHIRAJA AND S. REICH, *Discrete Gradients for Computational Bayesian Inference*, preprint, <https://arxiv.org/abs/1903.00186>, 2019; Comput. Dyn., to appear.
- [79] G. A. PAVLIOTIS, *Stochastic Processes and Applications: Diffusion Processes, the Fokker-Planck and Langevin Equations*, Texts Appl. Math. 60, Springer-Verlag, New York, 2014.
- [80] N. S. PILLAI, A. M. STUART, AND A. H. THIÉRY, *Noisy gradient flow from a random walk in Hilbert space*, Stoch. Partial Differ. Equ. Anal. Comput., 2 (2014), pp. 196–232.
- [81] S. REICH, *A dynamical systems framework for intermittent data assimilation*, BIT, 51 (2011), pp. 235–249.
- [82] S. REICH, *A nonparametric ensemble transform method for Bayesian inference*, SIAM J. Sci. Comput., 35 (2013), pp. A2013–A2024, <https://doi.org/10.1137/130907367>.
- [83] S. REICH, *Data Assimilation: The Schrödinger Perspective*, preprint, <https://arxiv.org/abs/1807.08351>, 2018.
- [84] S. REICH AND C. COTTER, *Probabilistic Forecasting and Bayesian Data Assimilation*, Cambridge University Press, Cambridge, 2015.
- [85] G. O. ROBERTS, A. GELMAN, AND W. R. GILKS, *Weak convergence and optimal scaling of random walk Metropolis algorithms*, Ann. Appl. Probab., 7 (1997), pp. 110–120.
- [86] G. O. ROBERTS AND J. S. ROSENTHAL, *Optimal scaling of discrete approximations to Langevin diffusions*, J. R. Stat. Soc. Ser. B Stat. Methodol., 60 (1998), pp. 255–268.
- [87] G. O. ROBERTS AND J. S. ROSENTHAL, *Optimal scaling for various Metropolis-Hastings algorithms*, Statist. Sci., 16 (2001), pp. 351–367.
- [88] C. SCHILLINGS AND A. M. STUART, *Convergence analysis of ensemble Kalman inversion: The linear, noisy case*, Appl. Anal., 97 (2018), pp. 107–123.
- [89] C. SCHILLINGS AND A. M. STUART, *Analysis of the ensemble Kalman filter for inverse problems*, SIAM J. Numer. Anal., 55 (2017), pp. 1264–1290, <https://doi.org/10.1137/16M105959X>.
- [90] T. SCHNEIDER, S. LAN, A. STUART, AND J. TEIXEIRA, *Earth system modeling 2.0: A blueprint for models that learn from observations and targeted high-resolution simulations*, Geophys. Res. Lett., 44 (2017), pp. 12396–12417.
- [91] D. J. SPIEGELHALTER, N. G. BEST, B. P. CARLIN, AND A. VAN DER LINDE, *Bayesian measures of model complexity and fit*, J. R. Stat. Soc. Ser. B Stat. Methodol., 64 (2002), pp. 583–639.
- [92] A.-S. SZNITMAN, *Topics in propagation of chaos*, in Ecole d'Été de Probabilités de Saint-Flour XIX — 1989, P.-L. Hennequin, ed., 1991, Springer, Berlin, 1991, pp. 165–251.

- [93] A. TAGHVAEI, J. DE WILJES, P. G. MEHTA, AND S. REICH, *Kalman filter and its modern extensions for the continuous-time nonlinear filtering problem*, J. Dyn. Syst. Meas. Control, 140 (2018), 030904.
- [94] A. TONG LIN, W. LI, S. OSHER, AND G. MONTÚFAR, *Wasserstein Proximal of GANS*, CAM report 18-53, Department of Mathematics, University of California, Los Angeles, 2019.
- [95] G. TOSCANI, *Kinetic models of opinion formation*, Commun. Math. Sci., 4 (2006), pp. 481–496, <http://projecteuclid.org/euclid.cms/1175797553>.
- [96] C. VILLANI, *Optimal Transport: Old and New*, Grundlehren Math. Wiss. 338, Springer, Berlin, 2009.
- [97] J. YANG, G. O. ROBERTS, AND J. S. ROSENTHAL, *Optimal Scaling of Metropolis Algorithms on General Target Distributions*, preprint, <https://arxiv.org/abs/1904.12157>, 2019.
- [98] T. YANG, P. G. MEHTA, AND S. P. MEYN, *Feedback particle filter*, IEEE Trans. Automat. Control, 58 (2013), pp. 2465–2480.