

Titre: A Stochastic iteratively regularized Gauss-Newton method
Title:

Auteurs: Elhoucine Bergou, Neil K. Chada, & Youssef Diouane
Authors:

Date: 2025

Type: Article de revue / Article

Référence: Bergou, E., Chada, N. K., & Diouane, Y. (2025). A Stochastic iteratively regularized Gauss-Newton method. Inverse Problems, 41(1), 015005 (22 pages).
Citation: <https://doi.org/10.1088/1361-6420/ad9d72>

 Document en libre accès dans PolyPublie
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/61637/>
PolyPublie URL:

Version: Version officielle de l'éditeur / Published version
Révisé par les pairs / Refereed

Conditions d'utilisation: Creative Commons Attribution 4.0 International (CC BY)
Terms of Use:

 Document publié chez l'éditeur officiel
Document issued by the official publisher

Titre de la revue: Inverse Problems (vol. 41, no. 1)
Journal Title:

Maison d'édition: IOP Publishing
Publisher:

URL officiel: <https://doi.org/10.1088/1361-6420/ad9d72>
Official URL:

Mention légale: Original content from this work may be used under the terms of the Creative Commons Attribution 4.0 license (<https://creativecommons.org/licenses/by/4.0/>). Any further distribution of this work must maintain attribution to the author(s) and the title of the work, journal citation and DOI.
Legal notice:

PAPER • OPEN ACCESS

A Stochastic iteratively regularized Gauss–Newton method

To cite this article: Elhoucine Bergou *et al* 2025 *Inverse Problems* **41** 015005

View the [article online](#) for updates and enhancements.

You may also like

- [Goal oriented adaptivity in the IRGNM for parameter identification in PDEs: II. all-at-once formulations](#)
B Kaltenbacher, A Kirchner and B Vexler
- [Analysis of a generalized regularized Gauss–Newton method under heuristic rule in Banach spaces](#)
Zhenwu Fu, Yong Chen, Li Li et al.
- [Goal oriented adaptivity in the IRGNM for parameter identification in PDEs: I. reduced formulation](#)
B Kaltenbacher, A Kirchner and S Veljovi

A Stochastic iteratively regularized Gauss–Newton method

Elhoucine Bergou¹, Neil K Chada^{2,*} 
and Youssef Diouane³ 

¹ Mohammed VI Polytechnic University, Ben Guerir, Morocco

² Department of Mathematics, City University of Hong Kong, 83 Tat Chee Avenue, Kowloon Tong, Hong Kong Special Administrative Region of China, People's Republic of China

³ GERAD and Department of Mathematics and Industrial Engineering, Polytechnique Montreal, Montreal, Canada

E-mail: neilchada123@gmail.com, elhoucine.bergou@um6p.ma and youssef.diouane@polymtl.ca

Received 13 September 2024; revised 18 November 2024

Accepted for publication 11 December 2024

Published 23 December 2024



CrossMark

Abstract

This work focuses on developing and motivating a stochastic version of a wellknown inverse problem methodology. Specifically, we consider the iteratively regularized Gauss–Newton method, originally proposed by Bakushinskii for infinite-dimensional problems. Recent work have extended this method to handle sequential observations, rather than a single instance of the data, demonstrating notable improvements in reconstruction accuracy. In this paper, we further extend these methods to a stochastic framework through mini-batching, introducing a new algorithm, the stochastic iteratively regularized Gauss–Newton method (SIRGNM). Our algorithm is designed through the use randomized sketching. We provide an analysis for the SIRGNM, which includes a preliminary error decomposition and a convergence analysis, related to the residuals. We provide numerical experiments on a 2D elliptic partial differential equation example. This illustrates the effectiveness of the SIRGNM,

* Author to whom any correspondence should be addressed.



Original Content from this work may be used under the terms of the [Creative Commons Attribution 4.0 licence](https://creativecommons.org/licenses/by/4.0/). Any further distribution of this work must maintain attribution to the author(s) and the title of the work, journal citation and DOI.

through maintaining a similar level of accuracy while reducing on the computational time.

Keywords: stochastic optimization, inverse problems, regularization, Gauss–Newton method, convergence analysis, random projection

1. Introduction

The field of optimization [12, 39, 40] is of crucial importance many areas of applied mathematics, which is concerned with minimizing functions, or functionals of interest. Popular examples include control theory, calculus of variations, numerical analysis and others. This is particularly apparent with its recent success in the emerging fields of machine learning, data science and deep learning [11, 45, 49]. However a particular application, which is of interest for this work, is parameter estimation associated with differential equations. This mathematical discipline is commonly known as inverse problems [37, 44, 46], which is concerned with the recovery of some quantity, or parameter, $u^\dagger \in \mathcal{X}$ of interest from some noisy observations $y^\dagger \in \mathcal{Y}$, i.e.

$$y^\dagger = F(u^\dagger), \quad (1.1)$$

with $F: \mathcal{X} \rightarrow \mathcal{Y}$ being the forward operator associated with the targeted problem, and $\mathcal{X}$ and $\mathcal{Y}$ are assumed to be two Hilbert spaces. In practice, due to the unavoidable presence of observational noise in real applications, equation (1.1) must be replaced by

$$y^\delta = F(u^\dagger) + \eta, \quad \eta \sim \mathcal{N}(0, \delta^2 I), \quad (1.2)$$

where $y^\delta \in \mathcal{Y}$ represents the corrupt noisy data, η takes the form of additive Gaussian noise, and $\delta > 0$. Traditionally inverse problems are solved using variational techniques, motivated from regularized least-squares, where the solution of (1.2) is estimated by the following minimizer

$$u^* \in \arg \min_{u \in \mathcal{X}} \frac{1}{2} \|y^\delta - F(u)\|_{\mathcal{Y}}^2 + \alpha \|u\|_{\mathcal{X}}^2, \quad (1.3)$$

where α is a step-size rule, and $\|u\|_{\mathcal{X}}^2$ acts as a penalty term. This is a well-used, and popular choice of penalty known as Tikhonov regularization. There has been huge developments in the area of developing regularized least-squares solvers for inverse problems, which are ill-posed, known as iterative regularization methods [19, 34]. Well-known algorithms based on this include the regularized Levenberg–Marquardt method (LMM), ν -methods and the Landweber iteration [9]. Our focus in this work will be another method, which has connections to the above stated methods which is the iterated regularized Gauss–Newton method (IRGNM) proposed by Baksushinkii [2, 3], which aims to mimic the *classical* iterated Gauss–Newton method, but for ill-posed problems. Since its development it has seen significant advances both related to theory and applications. Such example are the development of error bounds in different function space settings, demonstrating convergence and applications to data assimilation-based methods, such as when one has sequential data [4, 10, 15, 26, 31, 32, 48]. Despite these important directions and works, there has been somewhat a discontinuity between classical optimization understanding, and development, of algorithms, and those related to inverse problems. In particular a recent attraction has been that of stochastic gradient methods [11, 12, 43], which have proven popular due to their cost-effectiveness, compared to full-gradient methods. This is related to not using the full derivative information, but rather a random subset of it. The most well-known example of this includes the stochastic gradient descent. In the context of least-squares optimization this has also been done, for the non-regularized Gauss–Newton method,

the LMM and evolutionary algorithms. However this has not been very well-understood in the context of iterative regularization methods for inverse problems. Recent work has looked at understanding stochastic iterative methods for inverse problems, however much of the analysis is assumed to be in the linear and finite-dimensional setting [27, 30, 33]. Therefore this acts as our primary motivation to apply these techniques to iterative regularization methods, such as the IRGNM.

Specifically our form of stochasticity we apply is based on random projection methods which have been introduced in [20, 41, 42], which is also known as sketching. These forms of projection has been shown to be easy to apply with desirable convergence properties. It is also well-known that sketching, in this form, is equivalent to using mini batching within optimization. Another motivation from the development of a stochastic version, is the derivation of convergence analysis which has not been considered before. In the classical IRGNM this has not been considered. We will consider both the IRGNM and a recent dynamic version (with multiple observations [15]), which are both presented in the infinite-dimensional setting.

1.1. Contributions

Our highlighting contributions from this paper are provided below:

- We propose a stochastic version of the IRGNM, which multiplies a projection operator by the forward operator. This technique is related to the sketching. Our method is entitled the stochastic iteratively regularized Gauss–Newton method (SIRGNM). The addition of stochastic gradients is considered for both the classic and sequential/dynamic IRGNM.
- We present an analysis of the convergence properties of SIRGNM, which extends the well-established properties of the deterministic IRGNM. Our analysis is based on a projection operator for the sketching, where we present an error decomposition and a convergence analysis, which includes the residuals.
- We evaluate the efficiency of our novel approach through experimentation on a 2D partial differential equation (PDE)-based inverse problem. Our PDE of choice is referred to as Darcy, or groundwater flow, concerned with modelling fluid flow in a porous medium underground. We demonstrate the efficiency of exploiting a stochastic version while attaining a similar order of accuracy.

1.2. Outline

The outline of this work is as follows. In section 2 we provide a primer on the necessary material required for the manuscript. This includes a review on the IRGNM, while discussing the aspects of projection, and sketching. This will lead onto section 3 where we present our new developed methodology of the stochastic method SIRGNM. We aim to carry out an analysis related to convergence, which will be presented in section 4. Numerical results are presented in section 5 demonstrating the effectiveness of the our proposed method. Finally we present an overview of our work, while mentioning fruitful areas of future work, to be conducted in section 6.

2. Background material

In this section we briefly recall our motivating algorithm, which is the iterative regularized Gauss–Newton method and its dynamic counterpart, i.e. the dIRGNM.

Without loss of generality and to avoid unnecessarily notation we assume $y^\delta = 0$ in (1.3). Hence, our target in this paper reduces to solving the following least squares problem:

$$\min_{u \in \mathcal{X}} f(u) := \frac{1}{2} \|F(u)\|_{\mathcal{Y}}^2 + \alpha \|u\|_{\mathcal{X}}^2. \quad (2.4)$$

A popular and well-known methodology for solving (2.4) is the iterated regularized Gauss–Newton method (iRGNM). The variational form based on the IRGNM is defined as:

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\|F(u_n) + F'[u_n](u - u_n)\|_{\mathcal{Y}}^2 + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right], \quad (2.5)$$

where $\alpha_n > 0$ is the step-size (which is iteration dependent). The control of the parameter α_n of the iterations is often essential for the convergence of the iterative process. In our setting, $u_0 \in \mathcal{X}$ represents the apriori information on u_n . If $u_0 = u_n$, it corresponds to the classical LMM. u_0 can be also set to some initial guess that does not depend on the iterative process [14]. $F'[u]$ is the Fréchet derivative of F at u . Without loss of generality, a covariance operator can be included to scale the term $\|u - u_0\|_{\mathcal{X}}^2$. Namely, the latter term can be replaced with $\|C^{-1/2}(u - u_0)\|_{\mathcal{X}}^2$ where C is a covariance operator, which we will discuss in more details in the numerical section.

In the work of [14], the authors designed an alternative version for (2.5), which takes into consideration random sequential data, motivated from the field of data assimilation. Specifically, the observational model with *sequential noisy observations* has the form

$$Y_i = F(u^\dagger) + \sigma \xi_i, \quad i = 1, 2, \dots, \quad (2.6)$$

where now the averaged observations, Z_n , of (2.6) is defined as

$$Z_n = n^{-1} \sum_{i=1}^n Y_i = F(u^\dagger) + \frac{\sigma}{n} \sum_{i=1}^n \xi_i. \quad (2.7)$$

The motivation behind considering (2.7) is that the covariance of the data decreases as i increasing, unlike the fixed data choice of Y which is σ , i.e.

$$\text{Cov} Z_n = \text{Cov} \frac{\sigma}{n} \sum_{i=1}^n \xi_i = \frac{\sigma^2}{n^2} \sum_{i=1}^n \text{Cov} \xi_i = \frac{\sigma^2}{n} \text{id}_{\mathcal{Y}}.$$

For the negative log-likelihood functional of the normal distribution, it replaces the term $g := F(u_n) + F'[u_n](u - u_n)$ by

$$\begin{aligned} \mathcal{S}_n(F(u_n) + F'[u_n](u - u_n)) &:= \frac{1}{2} \|F(u_n) + F'[u_n](u - u_n)\|_{\mathcal{Y}}^2 \\ &\quad - \langle F(u_n) + F'[u_n](u - u_n), Z_n \rangle, \end{aligned} \quad (2.8)$$

where $Z_n \in \mathcal{Y}$ is observation dependent term, typically Z_n is set as the average of a selected subset of sequential observations.

Hence, instead of (2.5), one uses the following method modification:

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\mathcal{S}_n(F(u_n) + F'[u_n](u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right]. \quad (2.9)$$

Remark 2.1. An important question is why do we consider the negative log-likelihood defined in (2.8) compared to $\|g\|_{\mathcal{Y}}^2$? The reason for this is because, in infinite dimensions, the quantity $\|g\|_{\mathcal{Y}}^2$ is almost-surely infinite, with respect to the data Y . This is because Gaussian measures

Algorithm 1. Iterated regularized Gauss–Newton method.

inputs: $u_0 \in \mathcal{X}$, $\alpha_0 > 0$ (an initial value for the regularisation), and M (a maximum number of iterations).

for $n = 0, \dots, M - 1$ **do**

Select $Z_n \in \mathcal{Y}$.

Set the $\mathcal{S}_n$ as defined in (2.8).

Compute u_{n+1} , i.e.,

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\mathcal{S}_n (F(u_n) + F'[u_n](u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right].$$

Update the regularization parameter α_{n+1} .

end

output: u_M .

in infinite dimensions do not lie in the associated Cameron–Martin space $\text{Im}(\mathcal{Y}^{1/2})$. This is also the case because Y as a function rather than a finite-dimensional vector. Therefore what we do is subtract off the ‘infinite’ part in (2.8).

Alternatively we can express the minimization procedure of (2.9) in terms of the first order optimality condition as

$$u_{n+1} = u_n + (F'[u_n]^* F'[u_n] + \alpha_n \text{Id}_{\mathcal{X}})^{-1} (F'[u_n]^* (Z_n - F(u_n)) + \alpha_n (u_0 - u_n)), \quad (2.10)$$

where $F'[u_n]^* : \mathcal{Y} \rightarrow \mathcal{X}$ is the adjoint operator of $F'[u_n] : \mathcal{X} \rightarrow \mathcal{Y}$. For computational efficiency, we can then use Woodbury lemma for (2.10) yielding

$$u_{n+1} = u_0 + F'[u_n]^* (F'[u_n] F'[u_n]^* + \alpha_n \text{Id}_{\mathcal{X}})^{-1} (Z_n - F(u_n) - F'[u_n](u_0 - u_n)). \quad (2.11)$$

An overview of the iterated regularized Gauss–Newton method (IRGNM) is provided in algorithm 1. We note that we now take the notation for the maximum number of iterations to M .

3. Stochastic IRGNM

The recent success and popularity of variational, or optimization schemes, has been evident in the areas of machine learning and data science, where one aims to optimize a function or objective function. In particular this has been the case primarily due to stochastic optimization methods, which are known to have faster convergence properties while taking less computational time than evaluating the full gradient. The disadvantage of these methodologies is generally the accuracy is slightly worse. Given this, our aim is to incorporate ideas from stochastic gradient descent applied to the IRGNM, in the context of solving (1.1).

3.1. Towards a general stochastic framework for IRGNM

To initiate this we can take two particular approaches in doing so. The first is based on a sketching operator, that projects the problem on a *lower-dimensional* space, i.e. one uses $\mathcal{P}[F(u)]$ as a stochastic estimation of $F(u)$, where $\mathcal{P}$ is the stochastic sketch operator that we assume is unbiased

Algorithm 2. Stochastic iterated regularized Gauss–Newton method.

Inputs: $u_0 \in \mathcal{X}$, $\alpha_0 > 0$ (an initial value for the regularisation), and M (a maximum number of iterations).

for $n = 0, \dots, M - 1$ **do**

Select $\mathcal{P}_n$ a stochastic sketching operator.

Select $Z_n \in \mathcal{Y}_p$.

Set the $\tilde{\mathcal{S}}_n$ as defined in (3.13).

Compute u_{n+1} , i.e.,

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\tilde{\mathcal{S}}_n (F(u_n) + F'[u_n](u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right].$$

Update the regularization parameter α_{n+1} .

end

Output: u_M .

As a result we can include the modified operator $\mathcal{P}_n$ (depending on the iteration index n) in the variational problem as

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\tilde{\mathcal{S}}_n (F(u_n) + F'[u_n](u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right]. \quad (3.12)$$

The operator $\tilde{\mathcal{S}}_n$ is defined as follows:

$$\tilde{\mathcal{S}}_n(g) := \frac{1}{2} \|\mathcal{P}_n[g]\|_{\mathcal{Y}_p}^2 - \langle Z_n, \mathcal{P}_n[g] \rangle. \quad (3.13)$$

where $\mathcal{Y}_p$ is set as the projected observations space. The vector point $Z_n \in \mathcal{Y}_p$ is selected by the user. The proposed framework is detailed in algorithm 2. We note that the update of the regularization parameter is purposefully left unspecified at this stage. From now on, we will refer to our proposed approach as SIRGNM.

3.2. On the randomized sketching operator

In what comes next, we use the following notations to easy the readability of the paper:

$$F_n^\dagger := F(u_n) + F'[u_n](u^\dagger - u_n); \quad (3.14)$$

$$F_n^{n+1} := F(u_n) + F'[u_n](u_{n+1} - u_n). \quad (3.15)$$

Note that in the case $\mathcal{P}_n$ is the identity (meaning that we use full batch in the case of sampling), we have $\tilde{\mathcal{S}}_n(\cdot) = \mathcal{S}_n(\cdot)$. The sketching operator $\mathcal{P}_n$ plays a key role in our proposed framework. For a better illustration, we can consider the case where the computation of $F'(u_n)$ is very expensive, not feasible numerically or only a part of it can be computed. In this case, the operator $\mathcal{P}_n$ can be represented by a diagonal matrix P_n with ones and zeros elements randomly distributed over its diagonal. We can also consider the sketched version of the Jacobean, in the definition of $\mathcal{P}_n$, i.e, $P_n F'(u_n)$. Note that for the new matrix $P_n F'(u_{n-1})$, we need to compute only the derivatives of the function F component for which we have one in the corresponding position in P_n . Therefore, by controlling the ‘number of ones’ in P_n , we control the number of components of F for which we compute the derivative. Therefore, in this case, the sketching is equivalent to consider a random mini batch of the sum instead of considering the whole sum at each iteration.

Another example is where the precise evaluation of the operator F and its Jacobian are unattainable. This situation is particularly evident when, given a specific input x , the operator $F(x)$ is defined as the expected value over a random variable ξ , for instance $F(x) = \mathbb{E}_\xi(\tilde{F}(x, \xi))$. In such cases, the only viable approach for estimating F and its Jacobian is through sampling. In this work we will make the following assumption on the sketching operator. The assumption is inspired by the standard mini-batch assumption in the stochastic gradient method [23].

Assumption 3.1. The randomized sketching operator satisfies the following properties:

- $\mathbb{E}_n[\mathcal{P}_n[F_n^\dagger]] = F_n^\dagger$.
- There exists $\sigma \geq 0$, such that $\mathbb{E}_n[\|\mathcal{P}_n[F_n^{n+1}] - F_n^{n+1}\|] \leq \sigma$.
- There exists a constant $\delta > 0$, such that $\|\mathcal{P}_n[F_n^{n+1}] - Z_n\| \leq \delta$.

where $\mathbb{E}_n[X] := \mathbb{E}[X | \mathcal{P}_1, \dots, \mathcal{P}_{n-1}]$ designates the conditional expectation of the random variable X with respect to the filtration generated by the sigma-algebra of $\mathcal{P}_1, \dots, \mathcal{P}_{n-1}$.

3.2.1. Discussion on assumption 3.1. We note that for the first part of this assumption, the sketching operator when it is applied to $F_n^\dagger$ is leading to unbiased stochastic estimates for all objectives. These type of assumption are very common in the machine learning community where the sketching operator is often related to the use of subsampling over small fractions of the data [23]. The second part of the assumption, is related to the fact that the variance of the sketching operator $\mathcal{P}_n$ is bounded which is also a very standard assumption similar when we consider the particular case where the residual function is scalar and sketching operator operates via subsampling [23]. The third part of the assumption indicates that the gap between the $\mathcal{P}_n[F_n^{n+1}]$ and Z_n has to be uniformly bounded. In practice, Z_n is often selected as the average of observations over F_n^{n+1} . One example of $\mathcal{P}_n$ that one may consider is the diagonal matrix including 0's and 1's that are chosen randomly according to the Bernoulli distribution with some parameter $0 < p \leq 1$. This is a common example, which we will use in our numerical experiments in section 5.

4. Convergence analysis

In this section we present our main theoretical contribution which is a convergence analysis of our SIRGNM method. This will require a preliminary error decomposition, based on different error terms. We will then discuss our main result, which provides a bound on the expected residual, independent of α_n .

4.1. Preliminary error decomposition

In this section, we firstly begin by introducing some notation we will throughout. Specifically, we introduce the effective noise level as follows:

$$\text{err}_n(g) := \frac{1}{2} \|g - y^\dagger\|_{\mathcal{Y}}^2 - (\mathcal{S}_n(g) - \mathcal{S}_n(y^\dagger)), \quad g \in \mathcal{Y},$$

and its stochastic version

$$\begin{aligned} \widetilde{\text{err}}_n(g) &:= \frac{1}{2} \|g - y^\dagger\|_{\mathcal{Y}}^2 - (\mathcal{S}_n(g) - \mathcal{S}_n(y^\dagger)) + (\mathcal{S}_n(g) - \tilde{\mathcal{S}}_n(g)) \\ &= \text{err}_n(g) + (\mathcal{S}_n(g) - \tilde{\mathcal{S}}_n(g)), \quad g \in \mathcal{Y}. \end{aligned}$$

Let also

$$\partial \widetilde{\text{err}}_n := \widetilde{\text{err}}_n(F_n^{n+1}) - \widetilde{\text{err}}_n(F_n^\dagger), \quad (4.16)$$

$$\partial \text{err}_n := \text{err}_n(F_n^{n+1}) - \text{err}_n(F_n^\dagger), \quad (4.17)$$

$$\partial \mathcal{S}_n^{n+1} := \mathcal{S}_n(F_n^{n+1}) - \tilde{\mathcal{S}}_n(F_n^{n+1}), \quad (4.18)$$

$$\partial \mathcal{S}_n^\dagger := \mathcal{S}_n(F_n^\dagger) - \tilde{\mathcal{S}}_n(F_n^\dagger). \quad (4.19)$$

Then, we have

$$\partial \widetilde{\text{err}}_n = \partial \text{err}_n + \partial \mathcal{S}_n^{n+1} + \partial \mathcal{S}_n^\dagger,$$

The first term in $\partial \widetilde{\text{err}}_n$ is due to the noise in the data and in case of noise-free observations it vanishes. The rest is due to sampling.

Now, using assumption 3.1, we introduce the following lemma, which follows naturally.

Lemma 4.1. *Let assumption 3.1 hold. Then the randomized error $\widetilde{\text{err}}_n$ evaluated at $F_n^\dagger$ is an unbiased estimator of err_n at $F_n^\dagger$, i.e.*

$$\mathbb{E}_n[\widetilde{\text{err}}_n(F_n^\dagger)] = \text{err}_n(F_n^\dagger).$$

Then next auxiliary results is usefull for our convergence analysis.

Lemma 4.2. *Under our working assumptions, one has*

$$\alpha_n \left[\|u_{n+1} - u_0\|_{\mathcal{X}}^2 - \|u^\dagger - u_0\|_{\mathcal{X}}^2 \right] + \frac{1}{2} \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2 \leq \partial \widetilde{\text{err}}_n + \frac{1}{2} \|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2.$$

Proof. Indeed, one has,

$$\begin{aligned} & \tilde{\mathcal{S}}_n(F_n^\dagger) - \tilde{\mathcal{S}}_n(F_n^{n+1}) \\ &= \left(\tilde{\mathcal{S}}_n(F_n^\dagger) - \mathcal{S}_n(F_n^\dagger) \right) + \left(\mathcal{S}_n(F_n^\dagger) - \mathcal{S}_n(y^\dagger) \right) \\ & \quad + \left(\mathcal{S}_n(y^\dagger) - \mathcal{S}_n(F_n^{n+1}) \right) + \left(\mathcal{S}_n(F_n^{n+1}) - \tilde{\mathcal{S}}_n(F_n^{n+1}) \right) \\ &= \left(\text{err}_n(F_n^\dagger) - \widetilde{\text{err}}_n(F_n^\dagger) \right) + \left(\frac{1}{2} \|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2 - \text{err}_n(F_n^\dagger) \right) \\ & \quad - \left(\frac{1}{2} \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2 - \text{err}_n(F_n^{n+1}) \right) + \left(\widetilde{\text{err}}_n(F_n^{n+1}) - \text{err}_n(F_n^{n+1}) \right). \end{aligned}$$

Hence,

$$\tilde{\mathcal{S}}_n(F_n^\dagger) - \tilde{\mathcal{S}}_n(F_n^{n+1}) = \frac{1}{2} \|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2 - \frac{1}{2} \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2 + \widetilde{\text{err}}_n(F_n^{n+1}) - \widetilde{\text{err}}_n(F_n^\dagger). \quad (4.20)$$

On the other hand, the minimality condition of (3.12) implies that

$$\tilde{\mathcal{S}}_n(F_n^{n+1}) + \alpha_n \|u_{n+1} - u_0\|_{\mathcal{X}}^2 \leq \alpha_n \|u^\dagger - u_0\|_{\mathcal{X}}^2 + \tilde{\mathcal{S}}_n(F_n^\dagger).$$

Hence,

$$\alpha_n \left[\|u_{n+1} - u_0\|_{\mathcal{X}}^2 - \|u^\dagger - u_0\|_{\mathcal{X}}^2 \right] \leq \tilde{\mathcal{S}}_n(F_n^\dagger) - \tilde{\mathcal{S}}_n(F_n^{n+1}).$$

Now, by using the equality (4.20), one gets

$$\alpha_n \left[\|u_{n+1} - u_0\|_{\mathcal{X}}^2 - \|u^\dagger - u_0\|_{\mathcal{X}}^2 \right] + \frac{1}{2} \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2 \leq \partial \widetilde{\text{err}}_n + \frac{1}{2} \|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2.$$

□

To proceed further, we need the following variational source condition, which has been first formulated in [24] and has become a standard assumption in the analysis of variational regularization methods. In many situations it turns out that variational source conditions are necessary and sufficient for convergence rates [25]. Note that—as typical for source conditions in general—the smoothness of $u^\dagger$ is therein measured relative to the smoothing properties of F .

Assumption 4.3 (Variational source condition). There exists a concave index function φ (i.e. $\varphi(0) = 0$ and $\varphi \nearrow \infty$) such that for all $u \in D(F)$ it holds

$$\|u - u^\dagger\|_{\mathcal{X}}^2 \leq \|u - \hat{u}_0\|_{\mathcal{X}}^2 - \|u^\dagger - \hat{u}_0\|_{\mathcal{X}}^2 + \varphi \left(\frac{1}{2} \|F(u) - F(u^\dagger)\|_{\mathcal{Y}}^2 \right). \quad (4.21)$$

In order to further treat the nonlinearity, we employ also the following assumption.

Assumption 4.4 (Tangential cone condition). There exists constants $C_{\text{tc}} \geq 1$ and $\eta > 0$ sufficiently small such that, for all $u, v \in \mathcal{X}$,

$$\begin{aligned} \frac{1}{C_{\text{tc}}} \|F(v) - y^\dagger\|_{\mathcal{Y}}^2 - \eta \|F(u) - y^\dagger\|_{\mathcal{Y}}^2 &\leq \|F(u) + F'[u](v - u) - y^\dagger\|_{\mathcal{Y}}^2 \\ &\leq C_{\text{tc}} \|F(v) - y^\dagger\|_{\mathcal{Y}}^2 + \eta \|F(u) - y^\dagger\|_{\mathcal{Y}}^2. \end{aligned}$$

Remark 4.5. This tangential cone condition follows from the standard tangential cone condition with some C_{tc} , see [26, lemma 5.2]. If $\varphi \geq \sqrt{t}$ as $t \rightarrow 0$, it can—using the techniques from [48]—be replaced by a Lipschitz-type assumption.

In what comes next, we will use the following notation:

$$d_n := \mathbb{E} \left[\|u_n - u^\dagger\|_{\mathcal{X}}^2 \right], \quad (4.22)$$

$$t_n := \frac{1}{2} \mathbb{E} \left[\|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2 \right]. \quad (4.23)$$

Therewith, we have proven the following:

Lemma 4.6 (Preliminary error estimate). Let assumptions 4.3 and 4.4 hold and assume that $u_n \in D(F)$ is well defined. Then u_{n+1} satisfies the following error decomposition

$$\alpha_n d_{n+1} + \frac{1}{2C_{\text{tc}}} t_{n+1} \leq \mathbb{E} [\partial \widetilde{\text{err}}_n] + \alpha_n \Psi(2C_{\text{tc}}\alpha_n) + 2\eta t_n, \quad (4.24)$$

where we abbreviate $\Psi(2C_{\text{tc}}\alpha_n)$ as

$$\Psi(2C_{\text{tc}}\alpha_n) := (-\varphi)^* \left(-\frac{1}{2C_{\text{tc}}\alpha_n} \right) = \sup_{\tau \geq 0} \left[\varphi(\tau) - \frac{\tau}{2C_{\text{tc}}\alpha_n} \right], \quad (4.25)$$

with $(-\varphi)^*$ being the Fenchel conjugate of the convex function $-\varphi$.

Proof. Indeed, plugging assumption 4.3 into lemma 4.2 with $u = u_{n+1}$ yields

$$\begin{aligned} & \alpha_n \|u_{n+1} - u^\dagger\|_{\mathcal{X}}^2 + \frac{1}{2} \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2 \\ & \leq \widetilde{\partial \text{err}}_n + \alpha_n \varphi \left(\frac{1}{2} \|F(u_{n+1}) - F(u^\dagger)\|_{\mathcal{Y}}^2 \right) + \frac{1}{2} \|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2. \end{aligned} \quad (4.26)$$

Now, using $v = u_{n+1}$ and $u = u_n$ in the tangential cone condition (from assumption 4.4) gives for the second term on the left-hand side of (4.26) that

$$\frac{1}{C_{\text{tc}}} \|F(u_{n+1}) - y^\dagger\|_{\mathcal{Y}}^2 - \eta \|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2 \leq \|F_n^{n+1} - y^\dagger\|_{\mathcal{Y}}^2,$$

and for the third term on the right-hand side, with (1.1), that

$$\|F_n^\dagger - y^\dagger\|_{\mathcal{Y}}^2 \leq C_{\text{tc}} \|F(u^\dagger) - y^\dagger\|_{\mathcal{Y}}^2 + \eta \|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2 = \eta \|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2.$$

By inserting the last two inequalities (above) into (4.26), we obtain the recursive error estimate

$$\begin{aligned} & \alpha_n \|u_{n+1} - u^\dagger\|_{\mathcal{X}}^2 + \frac{1}{2C_{\text{tc}}} \|F(u_{n+1}) - y^\dagger\|_{\mathcal{Y}}^2 \\ & \leq \widetilde{\partial \text{err}}_n + \alpha_n \varphi \left(\frac{1}{2} \|F(u_{n+1}) - F(u^\dagger)\|_{\mathcal{Y}}^2 \right) + \eta \|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2. \end{aligned} \quad (4.27)$$

Now by taking the expectation on both sides of the last inequality (and then apply the Jensen's inequality), one gets

$$\begin{aligned} & \alpha_n \mathbb{E} \left[\|u_{n+1} - u^\dagger\|_{\mathcal{X}}^2 \right] + \frac{1}{2C_{\text{tc}}} \mathbb{E} \left[\|F(u_{n+1}) - y^\dagger\|_{\mathcal{Y}}^2 \right] \\ & \leq \mathbb{E} [\widetilde{\partial \text{err}}_n] + \alpha_n \varphi \left(\frac{1}{2} \mathbb{E} \left[\|F(u_{n+1}) - F(u^\dagger)\|_{\mathcal{Y}}^2 \right] \right) + \eta \mathbb{E} \left[\|F(u_n) - y^\dagger\|_{\mathcal{Y}}^2 \right]. \end{aligned}$$

Equivalently,

$$\alpha_n d_{n+1} + \frac{1}{C_{\text{tc}}} t_{n+1} \leq \mathbb{E} [\partial \text{err}_n] + \alpha_n \varphi(t_{n+1}) + 2\eta t_n. \quad (4.28)$$

On the other hand, by using the properties of the Fenchel conjugate of $-\varphi$,

$$\Psi(2C_{\text{tc}}\alpha_n) = (-\varphi)^* \left(-\frac{1}{2C_{\text{tc}}\alpha_n} \right) = \sup_{\tau \geq 0} \left[\varphi(\tau) - \frac{\tau}{2C_{\text{tc}}\alpha_n} \right] \geq \varphi(t_{n+1}) - \frac{t_{n+1}}{2C_{\text{tc}}\alpha_n}.$$

Thus,

$$\alpha_n \varphi(t_{n+1}) \leq \alpha_n \Psi(2C_{\text{tc}}\alpha_n) + \frac{t_{n+1}}{2C_{\text{tc}}}. \quad (4.29)$$

Hence, by combining inequalities (4.28) and (4.29), the proof is concluded. $\square$

Lemma 4.7. Under assumption 3.1, we have that

$$\mathbb{E} [\widetilde{\partial \text{err}}_n] \leq \mathbb{E} [\partial \text{err}_n] + \frac{\sigma^2}{2} + \delta \sigma.$$

Proof. We have

$$\partial \widetilde{\text{err}}_n = \partial \text{err}_n + \partial \mathcal{S}_n^{n+1} + \partial \mathcal{S}_n^\dagger.$$

Since $F_n^\dagger$ does not depend on u_{n+1} , thus by lemma 4.1, one has

$$\mathbb{E}_n [\partial \mathcal{S}_n^\dagger] = \mathbb{E}_n [\tilde{\mathcal{S}}_n (F_n^\dagger) - \mathcal{S}_n (F_n^\dagger)] = 0.$$

Note that for the remaining term (4.17), i.e. $\partial \mathcal{S}_n^{n+1}$, we can not use lemma 4.1 because of the dependence between F_n^{n+1} , u_{n+1} and the sketching operator $\mathcal{P}_n$. In this case, one has

$$\begin{aligned} \partial \mathcal{S}_n^{n+1} &= \mathcal{S}_n (F_n^{n+1}) - \tilde{\mathcal{S}}_n (F_n^{n+1}) \\ &= \frac{1}{2} \|F_n^{n+1}\|_{\mathcal{Y}_p}^2 - \langle Z_n, F_n^{n+1} \rangle - \frac{1}{2} \|\mathcal{P}_n [F_n^{n+1}]\|_{\mathcal{Y}_p}^2 + \langle Z_n, \mathcal{P}_n [F_n^{n+1}] \rangle \\ &= \frac{1}{2} \|F_n^{n+1} - \mathcal{P}_n [F_n^{n+1}]\|_{\mathcal{Y}_p}^2 + \langle \mathcal{P}_n [F_n^{n+1}] - Z_n, F_n^{n+1} - \mathcal{P}_n [F_n^{n+1}] \rangle \\ &\leq \frac{1}{2} \|F_n^{n+1} - \mathcal{P}_n [F_n^{n+1}]\|_{\mathcal{Y}_p}^2 + \|\mathcal{P}_n [F_n^{n+1}] - Z_n\|_{\mathcal{Y}_p} \|F_n^{n+1} - \mathcal{P}_n [F_n^{n+1}]\|_{\mathcal{Y}_p}. \end{aligned}$$

Hence using assumption 3.1, one gets

$$\mathbb{E}_n [\partial \mathcal{S}_n^{n+1}] \leq \frac{\sigma^2}{2} + \delta \sigma.$$

Which, by using to total expectation theorem, concludes the proof. $\square$

For the sake of readability and simplicity of our final results on the convergence rates and without loss of generality, we assume that $\varphi(t) = \sqrt{t}$, therefore $\psi(t) = t/2$.

The following theorem provides the main bound on the expected residual at iteration $n + 1$. This bound comprises three terms: the first term accounts for the noise error, the second term can be managed by selecting an appropriate step size rule, and the third term is controlled by choosing a sufficiently small parameter η .

Theorem 4.8. *Let assumptions 3.1, 4.3 and 4.4 hold. Assume that $\eta < \frac{1}{4C_{\text{tc}}}$ then we have the following convergence rates*

$$t_{n+1} \leq (2C_{\text{tc}}e_{\max} + C_{\text{tc}}\sigma^2 + 2C_{\text{tc}}\delta\sigma) \sum_{i=0}^n (4C_{\text{tc}}\eta)^i + 2C_{\text{tc}}^2 \sum_{i=0}^n (4C_{\text{tc}}\eta)^i \alpha_{n-i}^2 + (4C_{\text{tc}}\eta)^n t_0,$$

where $e_{\max} = \max_{n \geq 0} \mathbb{E} [\partial \text{err}_n]$.

Proof. From lemmas 4.6 and 4.7, we have

$$\begin{aligned} t_{n+1} &\leq 2C_{\text{tc}}\mathbb{E} [\partial \widetilde{\text{err}}_n] + C_{\text{tc}}^2 \alpha_n^2 + 4C_{\text{tc}}\eta t_n \\ &\leq 2C_{\text{tc}}\mathbb{E} [\partial \text{err}_n] + C_{\text{tc}}\sigma^2 + 2C_{\text{tc}}\delta\sigma + 2C_{\text{tc}}^2 \alpha_n^2 + 4C_{\text{tc}}\eta t_n. \end{aligned}$$

By unrolling the recurrence we get

$$t_{n+1} \leq (2C_{\text{tc}}e_{\max} + C_{\text{tc}}\sigma^2 + 2C_{\text{tc}}\delta\sigma) \sum_{i=0}^n (4C_{\text{tc}}\eta)^i + 2C_{\text{tc}}^2 \sum_{i=0}^n (4C_{\text{tc}}\eta)^i \alpha_{n-i}^2 + (4C_{\text{tc}}\eta)^n t_0.$$

$\square$

The following corollary gives the convergence rates for the constant and the geometrically decreasing regularization.

Corollary 4.9. *Let assumptions 3.1, 4.3 and 4.4 hold. Assume also that we have a constant or linear decreasing step size, i.e. there exists $\gamma \leq 1$ such that $\alpha_n = \gamma \alpha_{n-1}$ and that $\eta < \frac{\gamma^2}{4C_{tc}}$ then we have the following convergence rates*

(i) (General case)

$$t_{n+1} \leq \frac{(2C_{tc}e_{\max} + C_{tc}\sigma^2 + 2C_{tc}\delta\sigma)}{1 - 4C_{tc}\eta} + \mathcal{O}(\alpha_n^2),$$

(ii) (Noise free case, i.e. $\partial \text{err}_n = 0$ and full sampling)

$$t_{n+1} = \mathcal{O}(\alpha_n^2) \text{ and } d_{n+1} = \mathcal{O}(\alpha_n).$$

Proof. (i) From theorem 4.8 we have

$$\begin{aligned} t_{n+1} &\leq (2C_{tc}e_{\max} + C_{tc}\sigma^2 + 2C_{tc}\delta\sigma) \sum_{i=0}^n (4C_{tc}\eta)^i \\ &\quad + 2C_{tc}^2 \sum_{i=0}^n (4C_{tc}\eta)^i \alpha_{n-i}^2 + (4C_{tc}\eta)^n t_0. \end{aligned}$$

Therefore by using the fact that $\alpha_n = \gamma \alpha_{n-1}$ we get

$$\begin{aligned} t_{n+1} &\leq (2C_{tc}e_{\max} + C_{tc}\sigma^2 + 2C_{tc}\delta\sigma) \sum_{i=0}^n (4C_{tc}\eta)^i \\ &\quad + 2C_{tc}^2 \alpha_0^2 \gamma^{2n} \sum_{i=0}^n \left(\frac{4C_{tc}\eta}{\gamma^2} \right)^i + (4C_{tc}\eta)^n t_0 \\ &\leq \frac{(2C_{tc}e_{\max} + C_{tc}\sigma^2 + 2C_{tc}\delta\sigma)}{1 - 4C_{tc}\eta} + \frac{2C_{tc}^2 \alpha_0^2 \gamma^{2n}}{1 - \frac{4C_{tc}\eta}{\gamma^2}} + (4C_{tc}\eta)^n t_0 \\ &= \frac{(2C_{tc}e_{\max} + C_{tc}\sigma^2 + 2C_{tc}\delta\sigma)}{1 - 4C_{tc}\eta} + \mathcal{O}(\alpha_n^2). \end{aligned}$$

(ii) This is the noise free case, i.e. $\partial \text{err}_n = 0$ for all $n \geq 0$ so $e_{\max} = 0$. We have also full sampling so $\sigma = \delta = 0$. Injecting this in the previous result we get $t_{n+1} = \mathcal{O}(\alpha_n^2)$.

Finally, for d_{n+1} from lemma 4.6 we have

$$d_{n+1} \leq C_{tc}^2 \alpha_n + \frac{4C_{tc}\eta t_n}{\alpha_n} - \frac{t_{n+1}}{\alpha_n C_{tc}} = \mathcal{O}(\alpha_n).$$

□

The previous corollary establishes that the residuals exhibit convergence toward a neighborhood of zero, and this convergence occurs at a rate proportional to α_n^2 . In other words, as the sequence progresses, the residuals progressively approach a neighborhood of zero with quadratic rate. In the noise free case, the theorem asserts the dual convergence of both the residual and iterate sequences toward zero. Specifically, the residual sequence converges with a rate proportional to α_n^2 , indicating a quadratic decrease, while the iterate sequence exhibits a convergence rate of α_n , implying a linear decrease.

5. Numerical experiments

In this section we provide numerical experiments demonstrating, and highlighting the benefit, of our stochastic IRGNM. We apply this in the context of both noise-free and noisy data, where we test our algorithm on a random PDE example, taking motivation from geophysical sciences. We will compare our methodology to previously discussed methods which demonstrate an improvement on accuracy.

5.1. Overview of algorithms

We briefly provide an overview of our algorithms before discussing our numerical simulations. Throughout the experiments we use algorithms 1, 2 and 4. Algorithm 1 consists of the version of the vanilla iterative scheme, referred to as the IRGNM. For this algorithm we consider the ‘classical’ version which contains no modified noise, where $Z_n = Y$ as from (2.6), and the dynamic version with sequential observations from (2.7). This takes motivation from the work of Chada *et al* [15], which takes sequential observations. We then consider a stochastic version of this which we refer it as the sIRGNM in algorithm 2. Finally our last algorithm is algorithm 4 which is the stochastic version of the dynamic algorithm, referred to as SdIRGNM.

5.2. Darcy flow model

We consider the inversion of the Darcy flow model [22], which models groundwater flow in a porous medium. For this we consider the problem of solving an elliptic PDE, where the forward problem is to solve the PDE

$$-\nabla \cdot (\kappa \nabla p) = f, \quad x \in \Omega, \quad (5.30)$$

with mixed boundary conditions

$$p(x_1, 0) = 100, \quad \frac{\partial p}{\partial x_1}(6, x_2) = 0, \quad -\kappa \frac{\partial p}{\partial x_1}(0, x_2) = 500, \quad \frac{\partial p}{\partial x_2}(x_1, 6) = 0, \quad (5.31)$$

and the source term f defined as

$$f(x_1, x_2) = \begin{cases} 0, & \text{if } 0 \leq x_2 \leq 4, \\ 137, & \text{if } 4 \leq x_2 \leq 5, \\ 274, & \text{if } 5 \leq x_2 \leq 6, \end{cases}$$

where $\kappa \in L^\infty(\Omega)$ denotes the permeability, $f \in L^\infty(\Omega)$ is a source term and $p \in H^1(\Omega)$ is the pressure, representing the forward solution. For our problem the domain is specified as $\Omega = [0, 6]^2$. The inverse problem associated with (5.30) is to determine the permeability $\kappa \in L^\infty(\Omega)$ from pointwise measurements of the PDE solution $p \in H^1(\Omega)$. The specific form of the PDE model (5.30), (in relation to the boundary conditions and f), is tested by Hanke [22] whom motivated a regularized LM algorithm. More information on this PDE model can be found in [8]. To treat a more general setting, we consider the following weighted L^2 space

$$\mathcal{H} \equiv \left\{ u \in L^2(\Omega) \mid \|C^{-1/2}u\|_{L^2(\Omega)} \leq \infty \right\}, \quad (5.32)$$

such that our covariance operator $\mathcal{C}$ induced by a correlation function given by

$$\mathcal{C}(u)[x] = \int_{\Omega} u(x') c(x, x') dx dx'. \quad (5.33)$$

Specifically we choose a Matérn covariance function which takes the form

$$\mathcal{C}(x, x') = \frac{\Gamma(\nu)}{2^{1-\nu}} K_\nu \left(\frac{|x - x'|}{\ell} \right) \left(\frac{|x - x'|}{\ell} \right)^\nu,$$

where $(\nu, \ell) \in \mathbb{R}^+ \times \mathbb{R}$ are associated hyperparameters which define the smoothness and lengthscale, and $K_\nu(\cdot)$ represents a Bessel function of the second kind, where we use the same notation of $\mathcal{C}$ as the covariance operator, and its corresponding function. We use this function, as it is common when aiming to simulate Gaussian random fields.

We compare the performances of our approaches against their corresponding deterministic versions. The equivalent updates of (2.10) of the deterministic algorithm cIRGNM [15] that take into consideration weighting of $\mathcal{C}$ is given by

$$u_{n+1} = u_n + (F'[u_n]^* F'[u_n] + \alpha_n \mathcal{C}^{-1})^{-1} (F'[u_n]^* (y - F(u_n)) + \alpha_n \mathcal{C}^{-1} (u_0 - u_n)), \quad (5.34)$$

and similarly for the dIRGNM [15]. For computational efficiency we use the Woodbury lemma to compute the updates, for instance (5.34) is equivalent to

$$u_{n+1} = u_0 + \mathcal{C} F'(u_n)^* (F'(u_n) \mathcal{C} F'(u_n)^* + \alpha_n I)^{-1} (y - F(u_n) - F'(u_n) (u_0 - u_n)).$$

We consider two experiments with two different groundtruths to recover, the first being a smooth function defined as,

$$u^\dagger(x, y) = \exp \left[-100 \left((x - 0.3)^2 + (y - 0.7)^2 \right) \right] + \frac{1}{2} \exp \left[-100 \left((x - 0.7)^2 + (y - 0.35)^2 \right) \right],$$

over the domain $\Omega = [0, 6]^2$. The second ground truth is a piecewise constant function based on a level set thresholding of a Gaussian random field.

5.2.1. Parameter choices. We choose 64 pointwise measurements within our domain Ω , and run our experiment for $n = 10^4$ iterations. For the regularization parameter, we set $\alpha_0 = 0.5$. We modify the choice of the regularization parameter, which is explained in the next subsection, and within the figures. The variance of the data is chosen as $\delta = 0.1$. For the Matérn covariance function, we specify $(\nu, \ell) = (2.4, 15)$. $u_0 = 1$ is chosen to be the vector of ones.

For our stochastic algorithm variants, we choose $\mathcal{P}_n$ to be the ‘minibatch’ operator, in the sense that it selects only a part of the vector. Within our experiment, at each iteration, for our stochastic variants (SIRGNM and SdIRGNM) we keep 32 components selected randomly at uniform. In the numerical experiments, as a stopping criteria, we use the of relative error which is defined as

$$\text{err}_{\text{relative}} = \frac{\|u - u^\dagger\|_{L^2(\Omega)}}{\|u^\dagger\|_{L^2(\Omega)}},$$

which considers the percentage of difference with respect to the $L^2(\Omega)$ norm. For the initial guess of the discontinuous truth we set $u_0 = 1$, and $u_0 = 0.1$ for the smooth ground truth. As discussed in [15], this choice is not important which was tested for different choices previously.

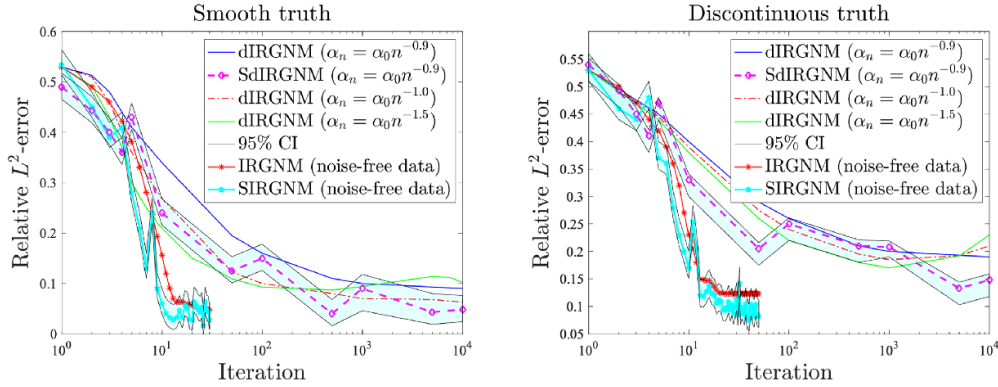


Figure 1. Stochastic IRGNM and dIRGNM for a discontinuous truth $u^\dagger$.

5.2.2. Discussion on assumptions. We emphasize with the PDE model we have chosen, that demonstrating our assumptions, regarding assumption 4.3 and 4.4, hold is beyond the scope of this paper. It is not entirely obvious to show this, for the case of the stochastic method, and the fact that we are using dynamic data. These assumptions in general are difficult to show, and are sensitive to choice of F . This is also discussed in the work of [15]. However, despite this we will show that our numerical simulations work well and follow the theory, in terms of demonstrating convergence, despite the fact we cannot exactly verify certain assumptions.

5.3. Summary from figures

Our first set of plots relate to comparing SIRGNM and SdIRGNM to the deterministic algorithms (IRGNM and dIRGNM), where we consider the relative L_2 error. Figure 1 presents a comparison of different plots where for the noisy-data, we consider alternative choices for the adaptive regularization parameter. As we can observe in both instances the stochastic version performs better than its deterministic counterpart, where we see an improvement in accuracy. The fluctuations that arise, are induced from the uncertainty. We also provide 95% confidence intervals, where we notice a low variance where the upper level usually is as good as its deterministic method.

The results obtained both both ground truth are similar and the stochastic algorithms seem to perform the best. We further show the effect if we consider a constant α_n , in comparison with noise-free data cases in figure 2 which further highlights the benefit of the stochastic method, and the lack of significant effect on the constant choice for the regularization parameter. Finally, to get an understanding on the performance of these methods, we compare the final reconstruction of some of these algorithms where we plot the SIRGNM in figures 3 and 4 and in all cases we see a similar performance, where the noise-free algorithms seem to perform the best. We note the early termination of the noise-free data is based on a threshold, which is attained in both experiments. We also terminate the noise-free simulations early as we aim to demonstrate the convergence speed difference, even though we are primarily interested in noisy observations.

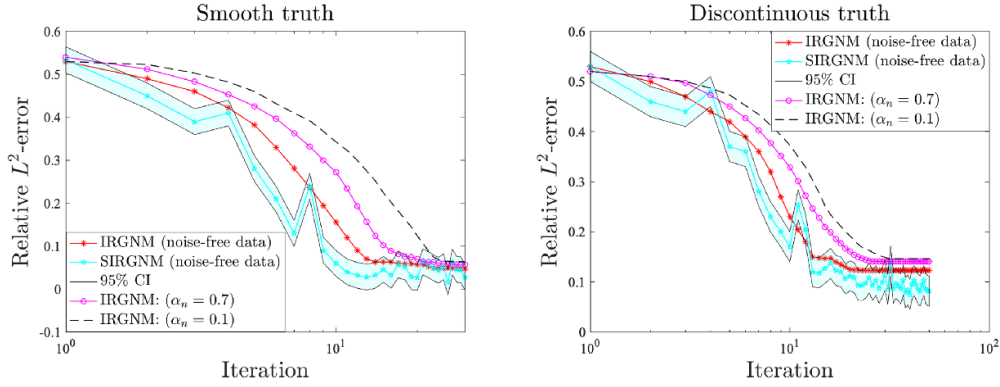


Figure 2. IRGNM and SIRGNM with constant regularization term.

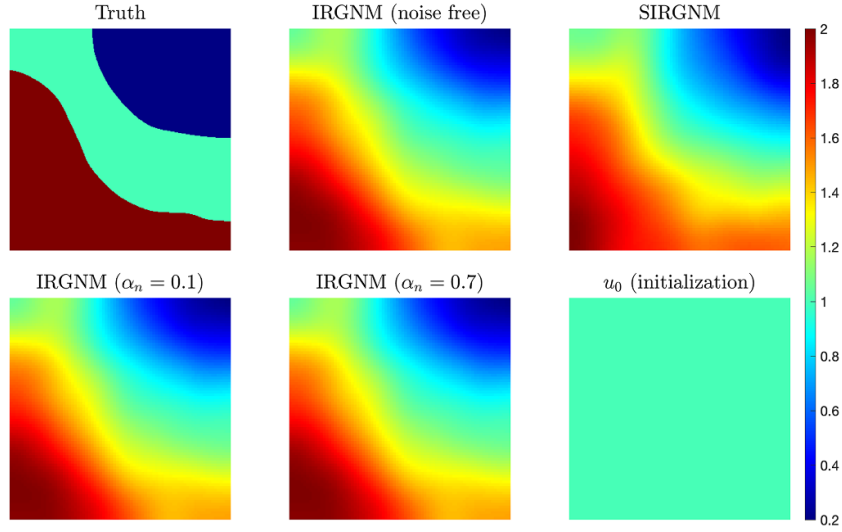


Figure 3. Final reconstruction plots of various methods on the discontinuous groundtruth.

5.4. Effect of mini-batching on the convergence speed

Another question, in the context of stochastic optimization algorithms, is the effect of mini-batching. In particular we repeated the experiments above with the effect of mini-batching. Table 1 provides results of this, where we change the size of mini-batches corresponding to the data set, and monitor the final reconstruction error at the end of each experiment. Our results here are for both the smooth and non-smooth truths. The table is considered for the adaptive choice of $\alpha_n = \alpha_0 n^{-0.9}$.

As demonstrated in the table, as we increase the batch-size, the CPU time and iterations required to achieve a relative error of 0.1 increases, which is expected. However the relative errors, independent of the batch-size, does not differ significantly, which improves slightly as we increase the batch-size.

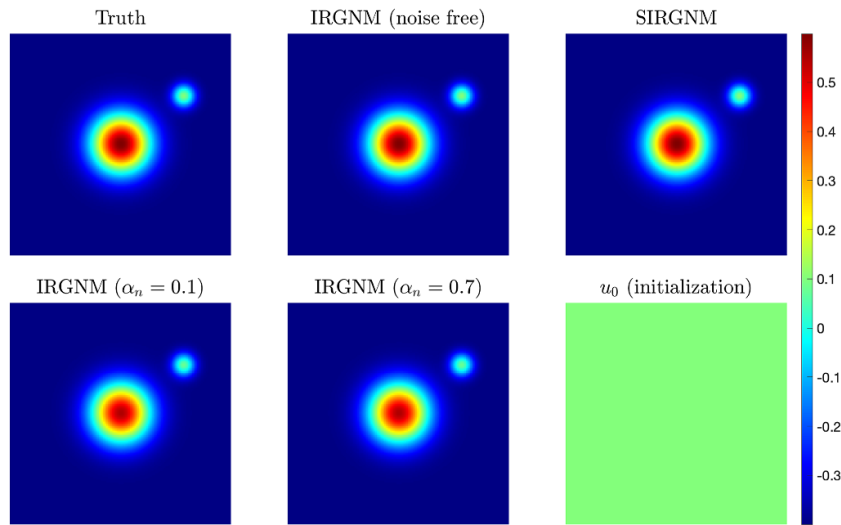


Figure 4. Final reconstruction plots of various methods on the smooth groundtruth.

Table 1. The effect of mini-batching and truth type on the convergence speed (measured by the number of iterations and related CPU time) of the proposed method. The stopping criterion is set to achieve $\text{err}_{\text{relative}} \leq 0.1$.

Truth $u^\dagger$	Mini-batch size	Relative error (final)	# of iterations	CPU time
Discontinuous	16	0.198	1925	4.5 h
	32	0.173	2269	4.4 h
	64	0.172	2741	4.5 h
	128	0.166	3034	4.6 h
Smooth	16	0.092	613	4.4 h
	32	0.083	697	4.4 h
	64	0.077	723	4.5 h
	128	0.076	856	4.7 h

To emphasize again, as we observe from table 1, using a small min-batch size of 16, or 32, achieves a good final reconstruction in terms of a relative error, without the use of many iterations. As expected as we increase the size of the mini-batch, this increases the numbers of iterations and CPU time. This is especially the case of the smooth truth, where the accuracy is higher, independent (relatively) of the mini-batch size. To save on both iterations and CPU time, we can consider a small batch size, while retaining a good level of accuracy.

6. Conclusion

Stochastic gradient-based algorithms have become very important in the fields of applied mathematics and machine learning, where exact gradient-evaluations can be cumbersome and costly. This argument extends to this work where we develop a stochastic version of the inverse-problem methodology which is the iteratively regularized Gauss–Newton method. In

particular we also extend this to recent work of Chada *et al* [15] which proposed a dynamics version which allows sequential observations. We provided the error analysis related to our new approach. We then provided numerical experiments on a nonlinear 2D elliptic PDE example, which demonstrate the improved performance of the SIRGNM method, and its dynamic variant. This is compared to various methods with varying choices of regularization parameters.

In terms of further directions to consider for future work, we list a number of these below.

- (i) This work considered a version of the Gauss–Newton method. It would be of interest to consider to the extension to the regularized LMM, which follows very similarly with a key difference being the penalty term. A more thorough investigation is needed in order to develop sound theory.
- (ii) Another potential direction would be to consider a further analysis for infinite-dimensional optimization such as gradient flow systems.
- (iii) Using these ideas for sequential observations, and uncertainty could be potentially applied, and used, in understanding the ensemble Kalman filter algorithm for inverse problems [13, 14, 16, 17, 47]. Despite the algorithm being derivative-free, it has many connections with both the Gauss–Newton and LMMs.
- (iv) Another potential work is the setting where one has a corrupt operator, or biased operator with respect to the forward operator $F(\cdot)$. Such a consideration is common in the stochastic optimization which is covered in various pieces of literature [1, 18].
- (v) Finally one could aim to understand various stopping criterion's in the context of stochastic iterative regularization. This has been achieved in the non-dynamic version, however this is still very much in the dynamic setting. One starting point may be the use of the Lepskii principle [36, 38].

Data availability statement

No new data were created or analysed in this study.

Acknowledgments

N K C is supported by an EPSRC-UKRI AI for Net Zero Grant: ‘Enabling CO2 Capture And Storage Projects Using AI’, (Grant EP/Y006143/1). NKC is also supported by a City University of Hong Kong startup grant.

Appendix. Algorithms

Algorithm 3. Stochastic iteratively regularized Gauss–Newton method (SIRGNM).

inputs: $u_0 \in \mathcal{X}$, $\alpha_0 > 0$ (an initial value for the regularisation), and M (a maximum number of iterations).

for $n = 0, \dots, M - 1$ **do**

- Select $\mathcal{P}_n$ a stochastic sketching operator.
- Select $\alpha_n \in \mathbb{R}$.
- Select $Z_n = Y \in \mathcal{Y}_p$.
- Set the $\tilde{S}_n$ as defined in (3.13).
- Compute u_{n+1} , i.e.,

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\tilde{S}_n (F(u_n) + F' [u_n] (u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right].$$

end

output: u_M .

Algorithm 4. Stochastic dynamic iterated regularized Gauss–Newton method (SdIRGNM).

inputs: $u_0 \in \mathcal{X}$, $\alpha_0 > 0$ (an initial value for the regularisation), and M (a maximum number of iterations).

for $n = 0, \dots, M - 1$ **do**

- Select $\mathcal{P}_n$ a stochastic sketching operator.
- Select $\alpha_n \in \mathbb{R}$.
- Set

$$Z_n = \frac{1}{n} \sum_{i=1}^n Y_i \in \mathcal{Y}_p,$$

i.e. the averaged observations.

- Set the $\tilde{S}_n$ as defined in (3.13).
- Compute u_{n+1} , i.e.,

$$u_{n+1} := \arg \min_{u \in \mathcal{X}} \left[\tilde{S}_n (F(u_n) + F' [u_n] (u - u_n)) + \alpha_n \|u - u_0\|_{\mathcal{X}}^2 \right].$$

end

outputs: u_M .

ORCID iDs

Neil K Chada  <https://orcid.org/0000-0002-2180-0985>

Youssef Diouane  <https://orcid.org/0000-0002-6609-7330>

References

- [1] Ajalloeian A and Stich S U 2020 On the convergence of stochastic gradient descent with biased gradients (arXiv:2008.00051)
- [2] Bakushinskii A B 1992 The problem of the convergence of the iteratively regularized Gauss–Newton method *Comput. Math. Math. Phys.* **32** 1353–9
- [3] Bakushinskii A B and Kokurin M Y 2004 *Iterative Methods for Approximate Solution of Inverse Problems* (Springer)
- [4] Bauer F, Hohage T and Munk A 2009 Iteratively regularized Gauss–Newton method for nonlinear inverse problems with random noise *SIAM J. Numer. Anal.* **47** 1827–46

- [5] Bell B M 1994 The iterated Kalman smoother as a Gauss-Newton method *SIAM J. Optim.* **4** 626–36
- [6] Bergou E, Diouane Y, Kungurtsev V and Royer C W 2022 A stochastic Levenberg-Marquardt method using random models with complexity results *SIAM/ASA J. Uncertain. Quantification* **10** 507–36
- [7] Bergou E, Diouane Y and Kungurtsev V 2020 Convergence and complexity analysis of a Levenberg-Marquardt algorithm for inverse problems *J. Optim. Theory Appl.* **185** 927–44
- [8] Carera J and Neuman S P 1986 Estimation of aquifer parameters under transient and steady state conditions: application to synthetic and field data *Water Resour. Res.* **22** 228–42
- [9] Bergou E, Gratton S and Vicente L N 2016 Levenberg-Marquardt methods based on probabilistic gradient models and inexact subproblem solution, with application to data assimilation *SIAM/ASA J. Uncertain. Quantification* **4** 924–51
- [10] Blaschke B, Neubauer A and Scherzer O 1997 On convergence rates for the iteratively regularized Gauss-Newton method *IMA J. Numer. Anal.* **17** 421–36
- [11] Bottou L, Curtis F E and Nocedal J 2018 Optimization methods for large-scale machine learning *SIAM Rev.* **60** 223–311
- [12] Boyd S and Vandenberghe L 2004 *Convex Optimization* (Cambridge University Press)
- [13] Chada N K, Schillings C and Weissmann S 2019 On the incorporation of box-constraints for ensemble Kalman inversion *Found. Data Sci.* **1** 433–56
- [14] Chada N K, Chen Y and Sanz-Alonso D 2021 Iterative ensemble Kalman methods: a unified perspective with some new variants *Found. Data Sci.* **3** 331–69
- [15] Chada N K, Iglesias M A, Lu S and Werner F 2023 On a dynamic variant of the iteratively regularized Gauss-Newton method with sequential data *SIAM J. Sci. Comput.* **45** 3020–46
- [16] Chada N K, Iglesias M A, Roininen L and Stuart A M 2018 Parameterizations for ensemble Kalman inversion *Inverse Problems* **34** 055009
- [17] Chada N K and Tong X T 2022 Convergence acceleration of ensemble Kalman inversion in non-linear settings *Math. Comput.* **91** 1247–80
- [18] Demidovich Y, Malinovsky G, Sokolov I and Richtarik P 2023 A guide through the zoo of biased SGD *Advances in Neural Information Processing Systems* vol 36
- [19] Engl H W, Hanke K and Neubauer A 1996 *Regularization of Inverse Problems (Mathematics and its Applications* vol 375) (Kluwer Academic Publishers Group)
- [20] Gower R M and Richtárik P 2015 Randomized iterative methods for linear systems *SIAM J. Matrix Anal. Appl.* **36** 1660–90
- [21] Gratton S, Lawless A S and Nichols N K 2007 Approximate Gauss-Newton methods for nonlinear least squares problems *SIAM J. Optim.* **18** 106–32
- [22] Hanke M 1997 A regularizing Levenberg-Marquardt scheme, with applications to inverse groundwater filtration problems *Inverse Problems* **13** 79–95
- [23] Hardt M, Recht B and Singer Y 2016 Train faster, generalize better: stability of stochastic gradient descent *Proc. 33rd Int. Conf. on Machine Learning* vol 48 (PMLR) pp 1225–34
- [24] Hofmann B, Kaltenbacher B, Pöschl C and Scherzer O 2007 A convergence rates result for Tikhonov regularization in Banach spaces with non-smooth operators *Inverse Problems* **23** 987–1010
- [25] Hohage T and Weidling F 2017 Characterizations of variational source conditions, converse results and maxisets of spectral regularization methods *SIAM J. Numer. Anal.* **55** 598–620
- [26] Hohage T and Werner F 2013 Iteratively regularized Newton-type methods for general data misfit functionals and applications to Poisson data *Numer. Math.* **123** 745–79
- [27] Hong Y, Bergou E, Doucet N, Zhang H, Cranney J, Ltaief H, Gratadour D, Rigaut F and Keyes D 2021 Outsmarting the atmospheric turbulence for ground-based telescopes using the stochastic Levenberg-Marquardt method *Eur. Conf. on Parallel Processing* pp 565–79
- [28] Iglesias M A, Law K J H and Stuart A M 2013 Ensemble Kalman methods for inverse problems *Inverse Problems* **29** 045001
- [29] Iglesias M A, Lin K, Lu S and Stuart A M 2017 Filter based methods for statistical linear inverse problems *Commun. Math. Sci.* **15** 1867–96
- [30] Jin Q and Chen L 2024 Stochastic variance reduced gradient method for linear ill-posed inverse problems (arXiv:2403.12460)
- [31] Jin Q 2000 On the iteratively regularized Gauss-Newton method for solving nonlinear ill-posed problems *Math. Comput.* **69** 1603–23
- [32] Jin Q 2011 A general convergence analysis of some Newton-type methods for nonlinear inverse problems *SIAM J. Numer. Anal.* **49** 549–73

- [33] Jin Q and Liu Y 2024 Convergence analysis of a stochastic heavy-ball method for linear ill-posed problems (arXiv:[2406.16814](#))
- [34] Kaltenbacher B, Neubauer A and Scherzer O 2008 *Iterative Regularization Methods for Nonlinear ill-Posed Problems (Radon Series on Computational and Applied Mathematics vol 6)* (Walter de Gruyter GmbH & Co. KG)
- [35] Larson J, Menickelly M and Wild S M 2019 Derivative-free optimization methods *Acta Numer.* **20** 287–404
- [36] Lepski O V and Spokoiny V G 1997 Optimal pointwise adaptive methods in nonparametric estimation *Ann. Stat.* **25** 2512–46
- [37] Lu S and Pereverzev S V 2013 *Regularization Theory for Ill-Posed Problems (Inverse Ill-Posed Probl. Ser. 58)* (De Gruyter)
- [38] Mathé P 2006 The Lepskii principle revisited *Inverse Problems* **22** L11
- [39] Nesterov Y 2004 *Introductory Lectures on Convex Optimization (Applied Optimization vol 87)* (Kluwer Academic Publishers)
- [40] Nocedal J and Wright S J 2006 *Numerical Optimization* 2nd edn (Springer)
- [41] Raskutti G and Mahoney M W 2016 A statistical perspective on randomized sketching for ordinary least-squares *J. Mach. Learn. Res.* **17** 1–31
- [42] Richtárik P and Takáč M 2020 Stochastic reformulation of linear systems: algorithms and convergence theory *SIAM J. Matrix Anal. Appl.* **41** 487–524
- [43] Robbins H and Monro S 1951 A stochastic approximation method *Ann. Math. Stat.* **22** 400–7
- [44] Stuart A M 2010 Inverse problems: a Bayesian perspective *Acta Numer.* **19** 451–559
- [45] Sun S, Cao Z, Zhu H and Zhao J 2020 A survey of optimization methods from a machine learning perspective *IEEE Trans. Cybern.* **50** 3668–81
- [46] Tarantola A 1987 *Inverse Problem Theory and Methods for Model Parameter Estimation* (Elsevier)
- [47] Weissmann S, Chada N K, Schillings C and Tong X T 2022 Adaptive Tikhonov strategies for stochastic ensemble Kalman inversion *Inverse Problems* **38** 045009
- [48] Werner F 2015 On convergence rates for iteratively regularized Newton-type methods under a Lipschitz-type nonlinearity condition *J. Inverse Ill-Posed Problems* **23** 75–84
- [49] Wibisono A, Wilson A C and Jordan M I 2016 A variational perspective on accelerated methods in optimization *Proc. Natl Acad. Sci.* **113** 7351–8